

ТРУДЫ

**ИНСТИТУТА СИСТЕМНОГО
ПРОГРАММИРОВАНИЯ РАН**

**PROCEEDINGS OF THE INSTITUTE
FOR SYSTEM PROGRAMMING OF THE RAS**

ISSN Print 2079-8156
Том 33 Выпуск 2

ISSN Online 2220-6426
Volume 33 Issue 2

Институт системного
программирования
им. В.П. Иванникова РАН

Москва, 2021

ИСП **РАН**

Труды Института системного программирования РАН Proceedings of the Institute for System Programming of the RAS

Труды ИСП РАН – это издание с двойной анонимной системой рецензирования, публикующее научные статьи, относящиеся ко всем областям системного программирования, технологий программирования и вычислительной техники. Целью издания является формирование научно-информационной среды в этих областях путем публикации высококачественных статей в открытом доступе. Издание предназначено для исследователей, студентов и аспирантов, а также практиков. Оно охватывает широкий спектр тем, включая, в частности, следующие:

- операционные системы;
- компиляторные технологии;
- базы данных и информационные системы;
- параллельные и распределенные системы;
- автоматизированная разработка программ;
- верификация, валидация и тестирование;
- статический и динамический анализ;
- защита и обеспечение безопасности ПО;
- компьютерные алгоритмы;
- искусственный интеллект.

Журнал издается по одному тому в год, шесть выпусков в каждом томе.

Поддерживается открытый доступ к содержанию издания, обеспечивая доступность результатов исследований для общественности и поддерживая глобальный обмен знаниями.

Труды ИСП РАН реферируются и/или индексируются в:

Proceedings of ISP RAS are a double-blind peer-reviewed journal publishing scientific articles in the areas of system programming, software engineering, and computer science. The journal's goal is to develop a respected network of knowledge in the mentioned above areas by publishing high quality articles on open access. The journal is intended for researchers, students, and practitioners. It covers a wide variety of topics including (but not limited to):

- Operating Systems.
- Compiler Technology.
- Databases and Information Systems.
- Parallel and Distributed Systems.
- Software Engineering.
- Software Modeling and Design Tools.
- Verification, Validation, and Testing.
- Static and Dynamic Analysis.
- Software Safety and Security.
- Computer Algorithms.
- Artificial Intelligence.

The journal is published one volume per year, six issues in each volume.

Open access to the journal content allows to provide public access to the research results and to support global exchange of knowledge. **Proceedings of ISP RAS** is abstracted and/or indexed in:



Редколлегия

Главный редактор - [Аветисян Арутюн Ишханович](#), академик РАН, доктор физико-математических наук, профессор, ИСП РАН (Москва, Российская Федерация)

Заместитель главного редактора - [Кузнецов Сергей Дмитриевич](#), д.т.н., профессор, ИСП РАН (Москва, Российская Федерация)

Члены редколлегии

[Воронков Андрей Анатольевич](#), доктор физико-математических наук, профессор, Университет Манчестера (Манчестер, Великобритания)

[Вирбицкайте Ирина Бонавентуровна](#), профессор, доктор физико-математических наук, Институт систем информатики им. академика А.П. Ершова СО РАН (Новосибирск, Россия)

[Коннов Игорь Владимирович](#), кандидат физико-математических наук, Технический университет Вены (Вена, Австрия)

[Ластовецкий Алексей Леонидович](#), доктор физико-математических наук, профессор, Университет Дублина (Дублин, Ирландия)

[Ломазова Ирина Александровна](#), доктор физико-математических наук, профессор, Национальный исследовательский университет «Высшая школа экономики» (Москва, Российская Федерация)

[Новиков Борис Асенович](#), доктор физико-математических наук, профессор, Санкт-Петербургский государственный университет (Санкт-Петербург, Россия)

[Петренко Александр Федорович](#), доктор наук, Исследовательский институт Монреаля (Монреаль, Канада)

[Черных Андрей](#), доктор физико-математических наук, профессор, Научно-исследовательский центр CICESE (Энсенада, Баха Калифорния, Мексика)

[Шустер Ассаф](#), доктор физико-математических наук, профессор, Технион — Израильский технологический институт Technion (Хайфа, Израиль)

Адрес: 109004, г. Москва, ул. А. Солженицына, дом 25.

Телефон: +7(495) 912-44-25

E-mail: proceedings@ispras.ru

Сайт: <https://ispranproceedings.elpub.ru/>

Editorial Board

Editor-in-Chief - [Arutyun I. Avetisyan](#), Academician of RAS, Dr. Sci. (Phys.–Math.), Professor, Ivannikov Institute for System Programming of the RAS (Moscow, Russian Federation)

Deputy Editor-in-Chief - [Sergey D. Kuznetsov](#), Dr. Sci. (Eng.), Professor, Ivannikov Institute for System Programming of the RAS (Moscow, Russian Federation)

Editorial Members

[Igor Konnov](#), PhD (Phys.–Math.), Vienna University of Technology (Vienna, Austria)

[Alexey Lastovetsky](#), Dr. Sci. (Phys.–Math.), Professor, UCD School of Computer Science and Informatics (Dublin, Ireland)

[Irina A. Lomazova](#), Dr. Sci. (Phys.–Math.), Professor, National Research University Higher School of Economics (Moscow, Russian Federation)

[Boris A. Novikov](#), Dr. Sci. (Phys.–Math.), Professor, St. Petersburg University (St. Petersburg, Russian Federation)

[Alexandre F. Petrenko](#), PhD, Computer Research Institute of Montreal (Montreal, Canada)

[Assaf Schuster](#), Ph.D., Professor, Technion - Israel Institute of Technology (Haifa, Israel)

[Andrei Tchernykh](#), Dr. Sci., Professor, CICESE Research Centre (Ensenada, Baja California, Mexico).

[Irina B. Virbitskaite](#), Dr. Sci. (Phys.–Math.), The A.P. Ershov Institute of Informatics Systems, Siberian Branch of the RAS (Novosibirsk, Russian Federation)

[Andrew Voronkov](#), Dr. Sci. (Phys.–Math.), Professor, University of Manchester (Manchester, United Kingdom)

Address: 25, Alexander Solzhenitsyn st., Moscow, 109004, Russia.

Tel: +7(495) 912-44-25

E-mail: proceedings@ispras.ru

Web: <https://ispranproceedings.elpub.ru/>

Современное состояние методов расчёта глобальной освещённости в задачах реалистичной компьютерной графики <i>Фролов В.А., Волобой А.Г., Еришов С.В., Галактионов В.А.</i>	7
Разработка Web-приложений с учетом элементов качества данных – DQAWA <i>Гуерра-Гарсия С.А., Перес-Гонсалес Г.Х., Рамирес-Торрес М.Т., Хуарес-Рамирес Р.</i>	49
Регулярные выражения для обнаружения Web-рекламы на основе автоматического скользящего алгоритма <i>Рианьо-Энрикес Д., Пинон-Аяла Р., Молеро-Кастильо Г., Барсенас Э., Веласкес-Мена А.</i>	65
Оценка пользовательских историй на основе декомпозиции сложности с использованием байесовских сетей <i>Дуран М., Хуарес-Рамирес Р., Хименес С., Тона К.</i>	77
Новая интеллектуальная система для обнаружения сахарного диабета 2-го типа с модифицированной функцией потерь и регуляризацией <i>Г.К.М., Альсадун А., Фам Д.Т.Х., Абдулла С.Х., Май Х.Т., Прасад П.В.Ч., Нгуен Ч.К.В.</i>	93
Классификация депрессивных эпизодов на основе ночных измерений: многомерный и одномерный анализ данных <i>Родригес-Руиз Д.Г., Гальван-Техада К.Э., Васкес-Рейес С., Гальван-Техада Х.И., Гамбоа-Росалес Х.</i>	115
Стратегии управления спросом в умных сетях электроснабжения для центров данных <i>Муранья-Сильвера Д., Несмачнов-Кановас С.Э., Итурриага-Фабра С.Д., Монтес Де Ока С., Белкредди Г., Монзон-Рангелофф П.А., Шепелёв В.Д., Черных А.Н.</i>	125
Классификатор изображений транспортных средств для определения корреляции их воздействия со смещением координат моста <i>Флорес-Фуэнтес В.</i>	137
Обнаружение объектов в аэронавигации с использованием вейвлет-преобразования и сверточных нейронных сетей: первый подход <i>Фортуна-Сервантес Х.М., Рамирес-Торрес М.Т., Мартинес-Карранса Х., Мургуия-Ибарра Х.С., Мехия-Карлос М.</i>	149
Решение проблемы обеспечения качества дерева многоадресной рассылки услуг <i>Риссо-Монтальдо К.Э., Робледо-Амоза Ф.Р., Несмачнов-Кановас С.Э.</i>	163
Исследование задачи обеспечения безопасности при хранении и обработке конфиденциальных данных <i>Мартиншин С.А., Храпченко М.В., Шокуров А.В.</i>	173

Выполнимость мю-исчисления с арифметическими ограничениями
Лимон-Приего Й., Барсенас-Патиньо И.Э., Бенетес-Герреро Э.И., Молеро-Кастильо
Г.Х., Веласкес-Мена А. 191

Table of Contents

The current state of the methods for calculating global illumination in tasks of realistic computer graphics <i>Frolov V., Voloboy A.G., Ershov S.V., Galaktionov V.A.</i>	7
Developing Web Applications with Awareness of Data Quality Elements – DQAWA <i>Guerra-García C.A., Perez-Gonzalez H.G., Ramírez-Torres M.T., Ramírez R.J.</i>	49
Regular Expressions for Web Advertising Detection based on an Automatic Sliding Algorithm <i>Riaño Enriquez D., Pinon-Ayala R., Molero-Castillo G., Barcenás E., Velázquez-Mena A.</i>	65
User Story Estimation based on the Complexity Decomposition using Bayesian Networks <i>Durán M., Juárez-Ramírez R., Jiménez S., Tona C.</i>	77
A Novel Intelligent System for Detection of Type 2 Diabetes with Modified Loss Function and Regularization <i>G.C. M., Alsadoon A., Pham D.T.H., Abdullah S.H., Mai H.T., Prasad P.W. C., Nguyen T.Q.V.</i>	93
Classification of Depressive Episodes Using Nighttime Data: Multivariate and Univariate Analysis <i>Rodríguez-Ruiz J.G., Galván-Tejada C.E., Vázquez-Reyes S., Galván-Tejada J.I., Gamboa-Rosales H.</i>	115
Smart grid demand response strategies for datacenters <i>Muraña-Silvera J., Nesmachnow-Cánovas S.E., Iturriaga-Fabra S.D., Montes De Oca S., Belcredi G., Monzón-Rangeloff P.A., Shepelev V.D., Tchernykh A.N.</i>	125
Vehicle Image Classifier for Bridge Displacement Correlation <i>Flores-Fuentes W.</i>	137
Object Detection in Aerial Navigation using Wavelet Transform and Convolutional Neural Networks: A first Approach <i>Fortuna-Cervantes J.M., Ramírez-Torres M.T., Martínez-Carranza J., Murguía-Barra J.S., Mejía-Carlos M.</i>	149
Solving the Quality of Service Multicast Tree Problem <i>Risso-Montaldo C.E., Robledo-Amoza F.R., Nesmachnow-Cánovas S.E.</i>	163
Study of the problem of ensuring security in the storage and processing of confidential data <i>Martishin S.A., Khrapchenko M.V., Shokurov A.V.</i>	173
Mu-Calculus Satisfiability with Arithmetic Constraints <i>Limón-Priego Y., Bárcenas-Patiño I.E., Benítez-Guerrero E.I., Molero-Castillo G.G., Velázquez-Mena A.</i>	191

DOI: 10.15514/ISPRAS-2020-33(2)-1



Современное состояние методов расчёта глобальной освещённости в задачах реалистичной компьютерной графики

^{1,2} В.А. Фролов, ORCID: 0000-0001-8829-9884 <vfrolov@graphics.cs.msu.ru>

¹ А.Г. Волобой, ORCID: 0000-0003-1252-8294 <voloboy@gin.keldysh.ru>

¹ С.В. Ершов, ORCID: 0000-0002-5493-1076 <sergey_65@mail.ru>

¹ В.А. Галактионов, 0000-0001-6460-7539 <vlgal@gin.keldysh.ru>

¹ Институт прикладной математики им. М.В. Келдыша РАН,
125047, Россия, Москва, Миусская пл., д. 4

² Московский государственный университет имени М.В. Ломоносова,
119991, Россия, Москва, Ленинские горы, д. 1

Аннотация. Современная реалистичная компьютерная графика базируется на физически корректном моделировании распространения света. Одной из основных и трудно вычислимых задач при этом является расчет глобальной освещённости, т.е. распределения света в виртуальной сцене, учитывающий множественные отражения и рассеяния света и всевозможные виды взаимодействия его с объектами сцены. Этой проблеме посвящены сотни публикаций, описывающие десятки методов вычисления глобальной освещённости и их модификации. В данной обзорной статье мы бы хотели не просто перечислить и кратко описать эти методы, но и дать некоторую «карту» существующих работ, которая позволит читателю сориентироваться, понять их достоинства и недостатки и, тем самым, выбрать для себя подходящий базовый метод. Особое внимание уделяется таким характеристикам методов как надёжность и универсальность в отношении используемых моделей, прозрачность их верификации, возможность эффективной реализации на GPU, а также накладываемые на сцену или феномены освещённости ограничения. В отличие от существующих обзорных работ анализируется не только эффективность методов, но также их ограничения и сложность программной реализации. Кроме того, мы предоставляем результаты собственных численных экспериментов с различными методами, служащих иллюстрациями к выводам.

Ключевые слова: расчёт освещённости, глобальная освещённость, трудновычислимые феномены освещённости

Для цитирования: Фролов В.А., Волобой А.Г., Ершов С.В., Галактионов В.А. Современное состояние методов расчёта глобальной освещённости в задачах реалистичной компьютерной графики. Труды ИСП РАН, том 33, вып. 2, 2021 г., стр. 7-48. DOI: 10.15514/ISPRAS-2021-33(2)-1

Благодарности. Авторы выражают благодарность Кириллу Гаранже (K. Garanzha) за ряд ценных критических замечаний, высказанных в процессе подготовки статьи.

The current state of the methods for calculating global illumination in tasks of realistic computer graphics

^{1,2} V.A. Frolov, ORCID: 0000-0001-8829-9884 <vfrolov@graphics.cs.msu.ru>

¹ A.G. Voloboy, ORCID: 0000-0003-1252-8294 <voloboy@gin.keldysh.ru>

¹ S.V. Ershov, ORCID: 0000-0002-5493-1076 <sergey_65@mail.ru>

¹ V.A. Galaktionov, 0000-0001-6460-7539 <vlgal@gin.keldysh.ru>

¹ Keldysh Institute of Applied Mathematics Russian Academy of Science,
4, Miusskaya sq., Moscow, 125047, Russia

² Lomonosov Moscow State University,
GSP-1, Leninskie Gory, Moscow, 119991, Russia

Abstract. Modern realistic computer graphics are based on light transport simulation. In this case, one of the main and difficult to calculate tasks is to calculate the global illumination, i.e. distribution of light in a virtual scene, taking into account multiple reflections and scattering of light and all kinds of its interaction with objects in the scene. Hundreds of publications and describing dozens of methods are devoted to this problem. In this state-of-the-art review, we would like not only to list and briefly describe these methods, but also to give some “map” of existing works, which will allow the reader to navigate, understand their advantages and disadvantages, and, thereby, choose a right method for themselves. Particular attention is paid to such characteristics of the methods as robustness and universality in relation to the used mathematical models, the transparency of the method verification, the possibility of efficient implementation on the GPU, as well as restrictions imposed on the scene or illumination phenomena. In contrast to the existing survey papers, not only the efficiency of the methods is analyzed, but also their limitations and the complexity of software implementation. In addition, we provide the results of our own numerical experiments with various methods that serve as illustrations for the conclusions.

Keywords: light transport simulation; global illumination; hard sampling lighting phenomena

For citation: Frolov V.A., Voloboy A.G., Ershov S.V., Galaktionov V.A. The current state of the methods for calculating global illumination in tasks of realistic computer graphics. *Trudy ISP RAN/Proc. ISP RAS*, vol. 33, issue 2, 2021, pp. 7-48 (in Russian). DOI: 10.15514/ISPRAS-2021-33(2)-1.

Acknowledgments. The authors are grateful to Kirill Garanzha for a number of valuable critical comments during the preparation of this article.

1. Введение

Расчёт освещённости (другие названия области – глобальное освещение, global illumination, lighting simulation, light transport) в реалистичной компьютерной графике – это бездонный колодец, в котором необходимость в повышении производительности не иссякает. Ускорение расчетов – доминирующий мотив большинства научных работ в данной области. Причина этого заключается в том, что моделирование явлений реального мира – это задача, не имеющая предела. Большая производительность на практике, как ни странно, не всегда приводит к уменьшению времени расчёта кадра. Она приводит к тому, что пользователь за отведённое время может посчитать более сложную 3D сцену или промоделировать более сложные явления.

За последнее десятилетие компьютерная графика добилась впечатляющих успехов в области моделирования и расчёта глобальной освещённости. Сегодня системы моделирования умеют рассчитывать такие явления и конфигурации оптических систем, которые ранее было практически невозможно моделировать из-за того, что методы просто не сходились к точному решению за разумное время. Кроме того, с развитием вычислительных систем стало возможным моделировать освещение для сцен с массивной геометрией, а также рассчитывать точное освещение в анимации для киноиндустрии.

Однако на данный момент в области глобальной освещённости так много разрозненных научных работ, что прикладному исследователю или разработчику крайне сложно выбрать

правильный метод для конкретной задачи. Эта проблема усугубляется несколькими трудностями.

Во-первых, сравнение производительности существующих методов – довольно сложная задача. В компьютерном зрении, например, существует большое количество фиксированных наборов данных, на которых проверяется эффективность и точность алгоритмов классификации. Они являются *de facto* стандартами, по отношению к которым оценивается тот или иной алгоритм. В графике ситуация обстоит иначе. Не существует открытых наборов сцен, в которых разные исследователи могли бы получить совпадающие изображения (в основном из-за отсутствия необходимых для этого стандартов). Это приводит к практике так называемого «cherry picking» – аккуратного подбора сцен и условий освещения таким образом, чтобы продемонстрировать преимущества алгоритма, разработанного авторами. Cherry picking в общем случае не является недостатком работы, т.к. новые алгоритмы, как правило, разрабатываются с тем, чтобы решить определённый класс проблем предыдущих методов. Если эти проблемы решены, то другие проблемы могут оставаться за рамками исследования. Однако разработчикам практических приложений, перед которыми стоит непростой вопрос выбора метода, от этого легче не становится.

Вторая проблема заключается в том, что некоторые современные и эффективные двунаправленные методы не просто более сложны, но и значительно более ограничены условиями, в которых эти методы работают правильно. Кроме того, нет общей методологии верификации, которая гарантировала бы корректность метода на любой сцене. А это означает, что во многих случаях эти методы не могут быть использованы для инженерных целей, как, например, проектирование оптических устройств, где корректность важна в первую очередь.

В результате, в каждом конкретном случае выбор базового расчётного метода и его развитие становится нетривиальной задачей. Мы полагаем, что наша работа поможет исследователям и разработчикам в области компьютерной графики и оптического моделирования сделать обоснованный выбор базового метода и грамотно определить собственное направление развития.

2. Используемые сокращения и термины

Далее мы в алфавитном порядке расшифруем основные сокращения из нашей работы и рис.

1. Большинство из них являются общепринятыми.

- BDPM – Bidirectional Photon Mapping [37], метод двунаправленных фотонных карт;
- BSDF – Bidirectional Scattering Distribution Function или двунаправленная функция отражения-рассеяния; именно эта функция описывает взаимодействие света с поверхностью, т.е. определяет модель материала;
- BPT – Bidirectional Path Tracing [3], двунаправленная трассировка путей.
- CC-BPT – Caustic Connection Strategies for Bidirectional Path Tracing [14];
- CMIS – Continuous Multiple Importance Sampling [57], метод многократной выборки по значимости, расширенный на случай континуума стратегий;
- ERPT – Energy Redistribution Path Tracing [105];
- FG – Final Gathering, финальный сбор, метод, откладывающий сбор из фотонной карты на одно переотражение;
- HHMC – Hessian Hamiltonian Monte Carlo light transport [94];
- HMC – Hamiltonian Monte Carlo [91];
- HSLT – Half Space Light Transport [77];
- IBPT – Instant BPT [6], урезанная версия BPT, которая может рассматриваться как оптимизация;

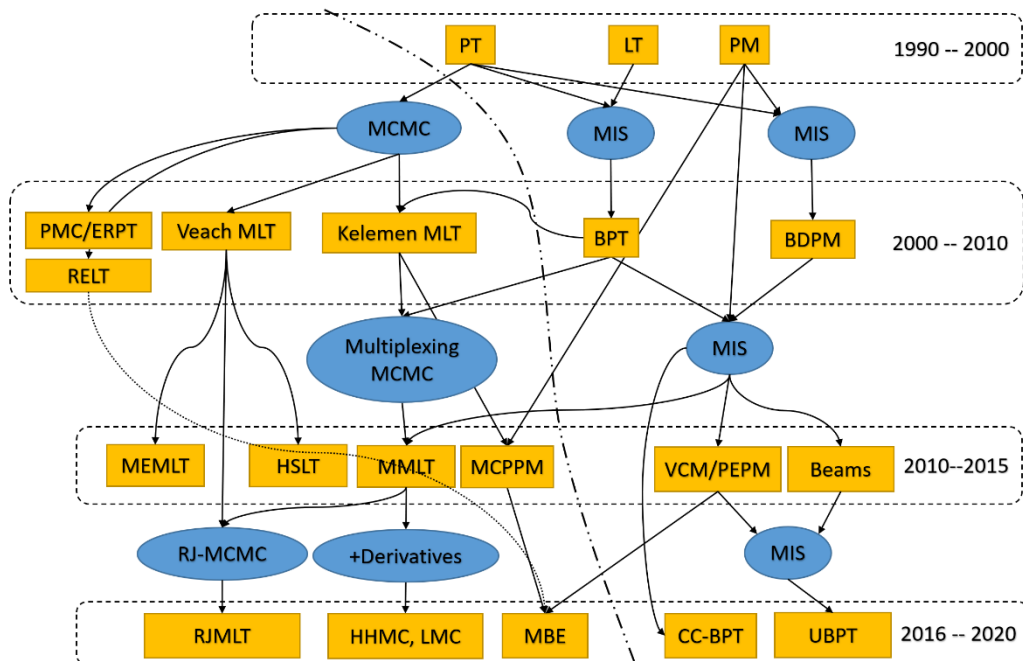


Рис. 1. Приблизительная карта современных методов интегрирования освещённости. В правой части изображены методы, основанные на обычном Монте-Карло, в левой части – на основе Монте-Карло по схеме марковских цепей. Левая и правая части условно разделены пунктирной линией с двумя точками. Прямоугольники представляют собой названия конкретного метода или класса методов расчёта освещения. Овалы представляют собой базовый математический инструмент, на основе которого строятся методы. Стрелки показывают, что одни методы построены на основе других методов или определённых математических инструментов

Fig. 1. An approximate map of modern methods of light transport. Methods based on ordinary Monte Carlo are shown on the right side, on the left side – based on Markov Chain Monte Carlo. The left and right parts are conditionally separated by a dotted line with two dots. The rectangles represent the names of a particular method or class of light transport methods. Ovals represent the basic mathematical toolkit on the basis of which methods are built. Arrows indicate that some methods are built on the basis of other methods or certain mathematical tools

- Kelemen MLT (или Primary Sample Space MLT (PSSMLT)) – явное указание на то, что MLT реализован в первичном пространстве путей, как в работе Келемена (Csaba Kelemen) и др. [75];
- LMC – Langevin Monte Carlo light transport [97];
- leap frog – традиционный способ реализации HMC, требующий большого количества промежуточных шагов, на каждом из которых необходимо вычислять целевую функцию [91].
- LT – Light Tracing (Forward Monte Carlo), прямая Монте-Карло трассировка;
- MALA – Metropolis-adjusted Langevin algorithm [103].
- MBE – Metropolized Bidirectional Estimator [41];
- MCMC – Markov Chain Monte Carlo, Монте-Карло по схеме марковских цепей;
- MEMLT – Manifold Exploration Metropolis Light Transport [15];
- MCPPM – Markov Chain PPM (Progressive Photon Mapping) [39];
- MIS – Multiple Importance Sampling [3], многократная выборка по значимости;

- MIS PT – MIS Path Tracing, обратная трассировка путей с использованием многократной выборки по значимости;
- MLT – Metropolis Light Transport [73];
- MMLT – Multiplexed Metropolis Light Transport, Multiplexed MLT [78];
- OMC – Ordinary Monte Carlo, обыкновенный метод Монте-Карло интегрирования;
- PDF – Probability Density Function, плотность вероятности;
- PCBPT – Probability Connection Bidirectional Path Tracing [7], оптимизация BPT;
- PCLT – Pixel Cache Light Tracing [30];
- PEPM – Path space Extension for Photon Mapping [25];
- PM – Photon Mapping, фотонные карты [17];
- PMC – Population Monte Carlo [106], Монте-Карло на основе отбора популяции выборок;
- PT – Path Tracing [2] (Backward Monte Carlo), обратная трассировка путей;
- PPM – Progressive Photon Mapping [18], прогрессивные фотонные карты;
- RELT – Replica Exchange Light Transport [80];
- RJMLT – Reversible Jump Metropolis Light Transport [88-90];
- RIS – Re-sampling for Importance Sampling [58];
- SDS-пути – Specular Diffuse Specular, вид путей, в которых между двумя зеркальным отражениями встречается одно диффузное;
- SPPM – Stochastic Progressive Photon Mapping [19], стохастические прогрессивные фотонные карты;
- Startup Bias – начальное смещение. Ошибка в изображении, проявляющаяся в виде неправильной оценки яркости отдельных областей изображения (например, недостаточно яркий каустик или наоборот яркая область, которая в процессе расчёта будет темнеть). Проблема свойственна методам на основе MLT.
- strMCMC-LT -- Stratified Markov Chain Monte Carlo Light Transport [83];
- SVBSDF – Spatial Varying Bidirectional Scattering Distribution Function [109], вид BSDF, когда свойства поверхности заданы в текстурах, т.е. могут различаться для разных точек. Например, маска смещения двух материалов или параметр "glossiness", заданный в текстуре, позволяют классифицировать BSDF как SVBSDF.
- UBPT – Unifying points, Beams and Paths in volumetric light transport simulation [26];
- VCM – Vertex Connection and Merging [24];
- Veach MLT (или Path Space MLT) – явное указание на то, что MLT реализован в мировом пространстве путей, как в оригинальной работе Вича [73].

Кроме того, нам потребуются ещё несколько определений:

Надёжность. Алгоритм расчёта освещения называется *надёжным (robust)* [3] на определённом сценарии освещения, если при расчёте интеграла отсутствуют выбросы – редкие Монте-Карло выборки с крайне большими значениями, препятствующие сходимости расчёта за приемлемое время. Надёжность является крайне важной характеристикой, т.к. более надёжные методы позволяют посчитать более сложные сценарии освещения, которые мы будем иногда называть *трудновычислимыми*. Таким образом, надёжность является краеугольным камнем эффективности расчёта освещения.

Сходимость Монте-Карло метода мы будем называть функцию $C(\dots)$, обратно пропорционально которой убывает ошибка. Например, сходимость $C(N) = \frac{1}{\sqrt{N}}$, где N – число выборок, означает, что ошибка убывает пропорционально $\frac{1}{\sqrt{N}}$. В этом случае, если мы хотим увеличить точность метода в 10 раз по сравнению с определённым значением, придётся увеличить количество выборок в 100 раз.

Эффективностью метода расчёта освещения будем называть далее:

- i. процент Монте-Карло выборок, вносящих существенный вклад в изображение, для методов на основе ОМС; *существенным* вкладом будем называть такой вклад, яркость которого сопоставима по порядку со средней яркостью изображения или значительно больше её;
- ii. среднюю вероятность принятия предложения перехода (acceptance rate) для методов на основе МСМС.

3. Общая классификация

Развитие методов вычисления освещённости шло, в основном, двумя путями на основе:

- i. обыкновенного Монте-Карло интегрирования (Ordinary Monte Carlo, ОМС);
- ii. Монте-Карло интегрирования по схеме марковских цепей (Markov Chain Monte Carlo, МСМС) [1].

Обе группы методов в настоящее время успешно применяются для вычисления интеграла освещённости и решения так называемого уравнения рендеринга [2] и имеют свои плюсы и минусы. На рис. 1 представлена приблизительная карта методов и их аббревиатуры, которые мы будем раскрывать по ходу статьи.

Кроме рассмотренной классификации на ОМС и МСМС возможна, как минимум, ещё одна независимая классификация:

- i. методы, работающие с тонкими лучами в терминах яркости;
- ii. методы, работающие с элементами конечного размера в терминах светового потока.

Такое разделение возможно как для методов на основе ОМС, так и для методов на основе МСМС. Однако, поскольку методы, работающие в терминах светового потока, используются в основном в обыкновенном Монте-Карло, мы будем рассматривать их в разделе методов на основе ОМС, а позже сделаем ремарку про применение этих методов в МСМС (подраздел 5.9, Markov Chain PPM).

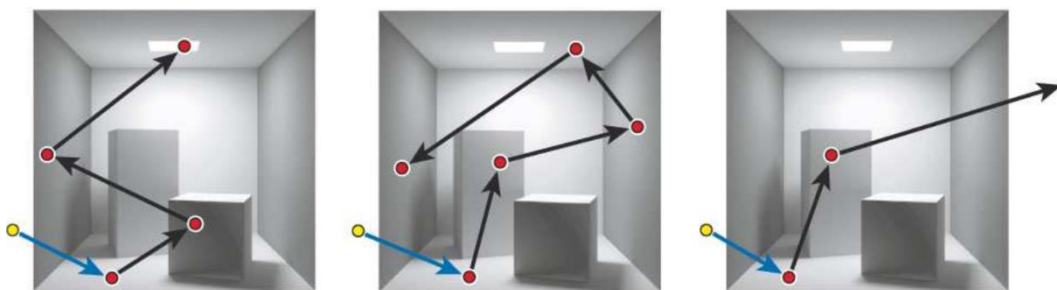


Рис. 2. Трассировка путей (Path Tracing)
Fig. 2. Path Tracing

4. Методы на основе ОМС

4.1 Методы, работающие в терминах яркости

Ключевым механизмом, на основе которого строятся все современные методы интегрирования освещённости на основе ОМС, является многократная выборка по значимости (Multiple Importance Sampling, MIS) [3]. Для того чтобы объяснить её суть, мы подробно рассмотрим применение MIS на примере самого базового метода – трассировки путей (Path Tracing, PT) [2].

4.1.1 Path Tracing (PT)

В простейшем варианте трассировки путей луч путешествует по сцене случайно до тех пор, пока не попадёт в источник света, не выйдет за пределы сцены или не будет достигнута заданная максимальная глубина рекурсии (рис. 2).

На рис. 3 такой алгоритм обозначен как «Simple PT» (или Naive PT). Как не трудно догадаться, основная проблема этого метода в том, что на типичных сценах вероятность случайного попадания в источник света крайне мала. Поэтому обычно в трассировку путей добавляют теневые лучи (shadow rays или visibility tests), которые при каждом переотражении летят напрямую в источник («Shadow PT» на рис. 3). Для потолочных светильников такой вариант, как правило, существенно улучшает ситуацию. Однако, на самом деле одного лишь «Shadow PT» недостаточно, даже если мы рассматриваем только прямое освещение. Shadow PT не работает, когда источник света имеет крупный размер, и/или материал освещаемой поверхности имеет глянцевые отражения (glossy reflections). В этом случае лишь небольшой участок источника будет вносить вклад в освещение, в то время как Shadow PT генерирует выборки по всей поверхности источника. Классический пример такой сцены: стол с глянцевой поверхностью, освещаемый окном. Оказывается, что в этом случае, как ни странно, Simple/Naive PT работает существенно лучше, поскольку выпускаемые им лучи учитывают свойство материала и почти всегда попадают в источник, внося ненулевой вклад в изображение.

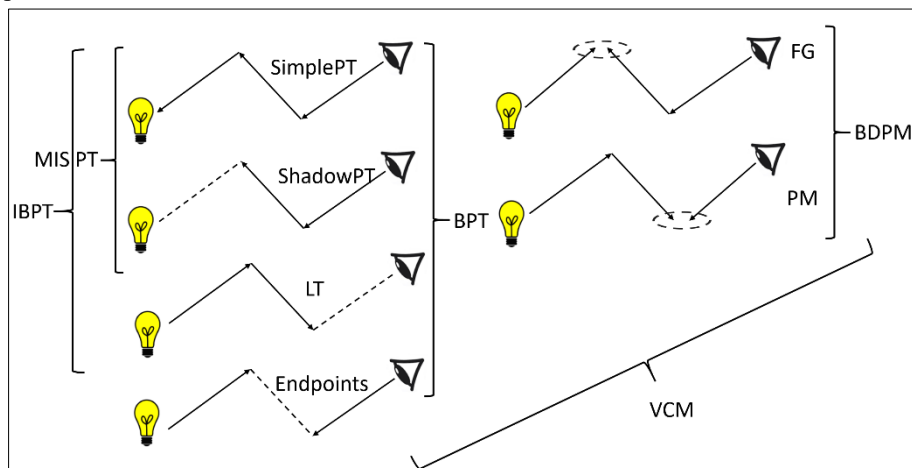


Рис. 3. Отдельные стратегии генерации выборок и методы, которые строятся на основе их комбинации

Fig. 3. Separate strategies for generating samples and methods that are based on their combination

Рассмотренные нами алгоритмы используют две различные стратегии создания выборок: Simple PT использует *неявную* стратегию, а Shadow PT – *явную*, выпуская лучи непосредственно в источник света. При этом следует отметить, что даже для прямого освещения недостаточно использования какой-либо одной стратегии: явной или неявной. Поэтому в трассировке путей используют обе, и такой алгоритм называется MIS PT (рис. 3). Он комбинирует вклады от явной (теневых лучей) и неявной (лучей, случайно попавших в источник) стратегий при помощи многократной выборки по значимости. Это, однако, требует вычисления весов выборки в многократной выборке по значимости, что на самом деле является нетривиальным решением по двум причинам.

- (а) Вычисление весов выборок в многократной выборке по значимости требует от реализации функций генерирования выборок для материалов (неявной стратегии) и источников (явной стратегии) корректного вычисления плотностей вероятности для каждой выборки. Это существенно усложняет дизайн функций генерирования выборок

материалов и источников и делает его доступным только для узких специалистов [4], поскольку требует знания о плотностях вероятности всех стратегий. Причём весьма нетривиальным образом: когда бы ни была вычислена некоторая выборка какой-то одной стратегией, для этой конкретной выборки необходимо уметь вычислять плотности всех других стратегий, то есть *как если бы* именно эта выборка *была бы* вычислена другими стратегиями.

- (б) Усложняется верификация MIS PT, т.к. для многих источников и материалов (точечный и направленный источники, источник в виде панорамы окружения и так называемые «небесные порталы», зеркальный материал, смесь материалов и др.) в коде появляется обработка специальных случаев. Причина этого в том, что, например, для идеального зеркального отражения невозможно корректно вычислить плотность вероятности [5]. В результате приходится прибегать к хитрости: в MIS PT зеркальное отражение форсированно делает вес неявной стратегии равным единице, а явной равным нулю. В двунаправленных методах, рассматриваемых далее, считают, что какой бы ни была плотность вероятности при зеркальном отражении, в прямом и обратном направлении она будет одинакова. Поэтому её можно задавать равной единице, но при этом нужно не забыть превратить её в ноль, если встречается явный способ соединения вершин.

Рассмотренный алгоритм трассировки путей с применением MIS используется в подавляющем большинстве индустриальных систем расчёта освещения в компьютерной графике (в архитектуре, кино и мультипликации). К сожалению, он даёт возможность сделать надёжный расчёт лишь прямого освещения. Как только в 3D сцене появляется существенное влияние непрямого (вторичного) освещения, возникает необходимость в добавлении новых стратегий построения выборок.

4.1.2 Bidirectional Path Tracing (BPT)

Рассмотрим сначала частный случай BPT – алгоритм *усечённой* двунаправленной трассировки, который в английском языке называется Instant Bidirectional Path Tracing (IBPT) [6]. Идея IBPT состоит в том, чтобы к двум существующим в MIS PT стратегиям добавить ещё одну: световую стратегию (LT на рис. 3). Эта стратегия работает аналогично Shadow PT, но в прямом направлении: луч стартует на источнике света, и при каждом переотражении от поверхности производится явное соединение этой точки поверхности с камерой. За счёт этой стратегии IBPT на удивление хорошо справляется со сложным вторичным освещением на ламбертовских поверхностях, но не может эффективно рассчитывать каустики, видимые через зеркала и стёкла (SDS-пути) [6].

Полноценный BPT устроен сложнее. Из источника и камеры трассируются два пути глубиной N и M переотражений соответственно. После чего делается $N \times M$ соединений между вершинами этих путей. Далее, используя полученные $N \times M$ соединений, формируются все полные пути от источника до камеры. Для каждого полного пути его вклад учитывается на основе многократной выборки по значимости.

Следует отметить, что вычисление вклада конкретного полного пути на основе веса многократной выборки по значимости в BPT никак не связано с переиспользованием вершин во время выполнения $N \times M$ соединений. То есть веса вычисляются для каждого полного пути независимо от других полных путей. Таким образом, BPT можно было бы строить иначе, соединяя лишь конечные точки: сначала выбрать случайно глубину от 0 до N и протрассировать путь от источника на глубину N (нулевое значение кодирует стратегию SimplePT, в которой мы будем пытаться поймать источник камерным лучём); затем протрассировать путь от камеры на случайно выбранную глубину от 0 до M (нулевое значение соответствует стратегии Light Tracing); наконец, соединить конечные точки путей от камеры и источника. Однако такой алгоритм (мы называем его End-Points BPT) сам по себе неэффективен вследствие более низкой вероятности получить удачный путь, априорно

выбирая глубину трассировки N и M . Тем не менее, именно End-Points BPT является базовым расчётным методом в MMLT [78], где его эффективность меняется кардинальным образом благодаря марковским цепям.

Таким образом, IBPT [6] и РСВРТ [7] могут рассматриваться как оптимизации оригинального алгоритма двунаправленной трассировки путей (BPT [3]). Они достигают лучшей производительности за счёт того, что реже используют (РСВРТ) или совсем не используют (IBPT) стратегии с промежуточными соединениями («End-points» на рис. 3), которые редко являются удачными. Все три метода тем не менее имеют следующие недостатки.

- (а) Двунаправленные методы требуют выполнения симметрии моделей. В действительности, это довольно существенное ограничение, поскольку источником ассиметрии могут быть входные данные (раздел «The Sources of Non-Symmetric Scattering» [3]): особенности геометрической модели поверхности [10], необходимость учёта поляризации, определённой только в одном направлении [11-12], а также ассиметрия, возникающая при расчёте преломлений в специфических средах [13]. Кроме того, расчёт такого явления как глубина резкости (Depth of Field, DOF) в BPT для LT стратегии затруднён из-за нарушения симметрии в силу отсутствия эффекта «огибания объектов» при проектировании точки на экранную плоскость или поверхность линзы объектива, что приводит к артефактам в виде тёмных краёв объектов. Это происходит из-за того, что не все 100% выборки, проецируемых в камеру с некоторой «дальней поверхности», в действительности достигнут её объектива, т.к. теневой луч может быть перекрыт близлежащим объектом. При расчёте же в обратном направлении все 100% лучей так или иначе достигнут какой-либо из двух рассматриваемых поверхностей – либо близлежащую, либо упомянутую ранее дальнюю поверхность.
- (б) Верификация двунаправленных методов (особенно BPT и РСВРТ, содержащих стратегии с промежуточными соединениями) становится ещё более сложной. Необходимо тестировать большое количество случаев, когда разные стратегии вносят существенный вклад в изображение. Тестирование при помощи покрытия кода не помогает, т.к. важно, чтобы рассматриваемый участок кода не просто выполнялся достаточно большое число раз, но и статистически внёс существенный вклад в изображение. А это не просто гарантировать из-за того, что вес в многократной выборке по значимости может обнулить вклад от той или иной стратегии в другом месте кода, никак не связанной с текущей стратегией (то есть важно тестировать стратегии ещё и в сочетании друг с другом). Кроме того, в двунаправленных методах необходимо все плотности вероятности вычислять в площадной мере (вероятность / m^2) [5], что увеличивает количество упомянутых ранее специальных случаев.

Мы полагаем, что метод IBPT в целом более практичен чем BPT или РСВРТ в первую очередь из-за того, что он компактный по памяти. А это существенно для реализации расчёта на GPU. Многие существующие реализации используют именно этот тип двунаправленной трассировки путей.

4.1.3 Проблема множества источников

Большое число источников света обычно является проблемой для однонаправленных методов и требует применения специальных алгоритмов для реализации эффективной выборки по значимости [8-9]. Однако интересно отметить, что алгоритмы IBT, BPT и РСВРТ могут эффективно вычислять освещение при большом числе источников благодаря использованию прямой стратегии: из какого бы источника света луч не был выпущен, он, как правило, вносит не нулевой вклад в изображение. Исключением будут являться сцены, состоящие из большого количества закрытых комнат, в которых необходимо применять те же методы повышения эффективности выборки источников что и для однонаправленных методов [8-9].

4.1.4 SDS каустики

Ранее мы упомянули, что каустики, видимые в стекле или зеркале (называемые нами *каустики второго типа*), не могут быть эффективно вычислены при помощи ВРТ. Однако, в недавней работе [14] была предложена стратегия, позволяющая решить эту проблему. Для этого авторы [14] используют исследование многообразий [15-16]. Ядром этого метода является эффективный расчёт так называемого обобщённого геометрического члена, который связывает две диффузные (или глянцевые) поверхности через последовательность зеркальных переотражений. В [14] предложена аппроксимация расчёта обобщённого геометрического члена, благодаря чему метод обретает практическую значимость.

Наиболее существенной проблемой метода [14] (как и [15-16]) является то, что для его работы требуется наличие инструментария дифференциальной геометрии, что является сильным ограничением, т.к. непосредственно влияет на то, как именно геометрические модели представлены внутри рендер-системы. Таким образом, это является классическим примером того, что более эффективный метод глобального освещения налагает существенно больше ограничений и, таким образом, оказывается трудноприменим в определённых условиях на практике (например, при наличии карт смещений и карт нормалей, о чём упоминают авторы [14-16]).

4.2 Методы, работающие в терминах светового потока

Второй класс методов – это фотонные карты. Он переходит от несмещенной оценки интеграла к смещенной (в стандартном определении фотонных карт) [17] и состоятельной для прогрессивных алгоритмов, таких как SPPM [19]. Фотонные карты делают это путем введения радиуса сбора конечного размера, в пределах которого близко расположенные точки световых путей могут быть объединены вместе, как если бы это была одна точка. Такое свойство алгоритма (смещенная / состоятельная оценка) на практике приводит к тому, что приближенное решение может быть получено быстрее чем несмещёнными методами, но точные вычисления становятся проблемой. В итоге фотонные карты часто используются в демонстрационных приложениях, но на практике есть существенные недостатки.

- (а) Более медленная сходимостъ: $O(\frac{1}{\sqrt[3]{N}})$ по сравнению с $O(\frac{1}{\sqrt{N}})$ для ВРТ [21], что происходит за счет уменьшения радиуса сбора в прогрессивных методах.
- (б) Дорогостоящая операция сбора (особенно в произвольном случае, когда требуется оценка BSDF для каждого фотона), которая становится критической из-за наличия ступок фотонов. Контроль плотности [22] амортизирует эту проблему, но его (как и метод релаксации фотонов [20]) можно использовать лишь для ограниченного случая: сбор фотонов для ламбертовской поверхности.
- (с) Фотонные карты должны применяться с осторожностью на поверхностях с микрорельефом, поскольку могут исказить вид микрорельефа.
- (д) Фотонные карты требуют построения структур пространственного разбиения (ускоряющих структур) для процедуры сбора освещённости, что не является серьезным недостатком, но значительно усложняет параллельную реализацию и требует дополнительный объем памяти по сравнению с подходами на основе ВРТ.

Тем не менее, несмотря на упомянутые недостатки фотонные карты широко используются в промышленных и научных системах расчёта по следующим причинам.

- (а) Фотонные карты могут рассчитывать довольно сложные случаи глобальной освещённости и быстро получать для них приближённое решение.
- (б) Реализация фотонных карт проста по сравнению, например, с ВРТ и сама по себе не требует вычисления плотностей вероятностей для каждого материала и источника.

4.3 Объединение ВРТ и фотонных карт

Преимущества фотонных карт (РМ) побудили исследователей к созданию методов, объединяющих в себе ВРТ и РМ. На сегодняшний день существует два достаточно стандартных способа объединения.

- (a) Выполнить сбор освещённости на первом незеркальном переотражении луча во время обратной трассировки (рис. 3, РМ). Этот способ широко используется для расчёта каустиков, для которых, как правило, заводится отдельная фотонная карта [23-24].
- (b) Выполнить сбор освещённости после второго незеркального переотражения луча во время обратной трассировки (рис. 3, FG). Этот способ называется финальным сбором (Final Gathering, FG). Фактически, фотонные карты в финальном сборе используются как аппроксимация третьего переотражения, и в таком качестве работают очень хорошо для многих сцен, поскольку, как мы обсуждали ранее, фотонные карты позволяют получить приближённое решение быстро.

4.3.1 BDPM

Однако же комбинации рассмотренных способов недостаточно, если в сцене встречаются материалы с глянцевыми отражениями. Кроме того, финальный сбор даёт артефакты в углах геометрических объектов, где его откладывают как минимум ещё на 1 переотражение (вторичный финальный сбор [23]). Подобные проблемы привели к появлению методов, делающих выбор номера переотражения на основе анализа оптических свойств текущей поверхности [33, 35-36]. К сожалению, легко показать, что в общем случае такого анализа недостаточно, т.к. заранее неизвестно насколько далеко будет располагаться следующая поверхность. Поэтому единственным решением, гарантирующим надёжность в данном случае, является сбор на каждом переотражении и комбинация результатов от разных переотражений при помощи многократной выборки по значимости. Именно так поступает алгоритм двунаправленных фотонных карт (Bidirectional Photon Mapping, BDPM) [37] (рис. 3, справа). Однако за повышение надёжности здесь, как и в ВРТ, приходится расплачиваться понижением скорости в среднем, т.к. дорогостоящий сбор освещённости теперь выполняется на каждом переотражении. Поэтому на наш взгляд имеет смысл также иметь ввиду варианты комбинирования ВРТ и фотонных карт без многократной выборки по значимости, таких как PCLT [30] и упомянутых методов на основе анализа свойств поверхности [33, 35-36]. Хотя по сравнению с полноценным BDPM они не столь универсальные и надёжные, что на практике приводит к настройке параметров алгоритма пользователем.

Заметим, что метод BDPM, несмотря на название, не является объединением классических [17, 23] и обратных фотонных карт [31-32], поскольку в обратных фотонных картах в действительности инвертируется лишь геометрическая задача поиска ближайших фотонов, а сами вычисления интеграла и стратегии генерации выборок при этом не изменяются.

4.3.2 VCM

Наконец, по-настоящему промежуточное положение между первым (ВРТ) и вторым (РМ) классами занимают методы, которые интегрируют фотонные карты в ВРТ на основе MIS: VCM [24], PERM [25], UBPT [26]. При этом многократная выборка по значимости в этих методах не решает проблему дорогой операции сбора, т.к. работает апостериорно, т.е. уже после того, как сбор освещённости был выполнен. Проблема меньшей скорости сходимости также решается не полностью, поскольку для трудновычислимых феноменов освещённости, с которыми не справляется ВРТ, будет наблюдаться ухудшенная сходимость фотонных карт $-O(\frac{1}{\sqrt[3]{N}})$.

Отдельным пунктом необходимо упомянуть метод Path Space Regularization [27], который является аналогом VCM, но вместо классического сбора в пространстве использует «сбор по

углам», так называемую угловую регуляризацию. Если говорить упрощённо, Path Space Regularization заменяет идеально-зеркальные BSDF на глянцевые и постепенно уменьшает глянцевость, двигаясь с каждым новым проходом обратно к зеркалу.

На сегодняшний день VCM является чрезвычайно популярным методом. Однако, хотя с математической точки зрения VCM и его аналоги делают большой шаг вперёд и повышают надёжность расчёта (что позволяет применять метод на более широком классе сцен), на практике зачастую эти методы (например, VCM, реализованный в системе Corona [28]) являются лишь небольшим улучшением простого смешивания изображений от BPT и SPPM по маске [29]. Это является вполне очевидным следствием механизма работы многократной выборки по значимости – сначала вычислить результат двумя разными способами, и потом скомбинировать оба с весами, которые часто вырождаются в пару (0,1) или (1,0). С другой стороны, трудоёмкость реализации на основе MIS достаточно высока [38] («... so much so that the authors also released a technical report and source code explaining how to implement the VCM algorithm...»). Верифицировать VCM или тем более UBPT крайне сложно в силу огромного количества рассматриваемых вариантов комбинирования стратегий в многократной выборке по значимости.

4.4 Адаптивная генерация выборок

Идея управляемых путей (Path Guiding) строится на аппроксимации входящего освещения в различных точках поверхности каким-либо из существующих методов (фотонные карты, функции специального вида или машинное обучение) с тем, чтобы использовать эту аппроксимацию как подсказку для выборки по значимости [42-46]. Path Guiding уменьшает уровень шума в целом по изображению, однако легко пропускает мелкие геометрические детали, которые функция аппроксимации не может учесть. Из-за этого на итоговом изображении присутствует разрозненный импульсный шум. Тем не менее Path Guiding становится популярным в индустрии, т.к. считается, что его относительно просто интегрировать в существующие системы [46], и, в отличие от MLT, он может быть использован для вычисления первичного освещения.

Path Guiding можно рассматривать как частный случай адаптивной генерации выборок, которая, в свою очередь, является частным случаем области компьютерной графики, называемой Sampling and Reconstruction [47]. Следуя классификации [47], все методы Sampling and Reconstruction делятся на априорные и апостериорные (хотя существуют методы, занимающие промежуточное положение [48] между этими двумя классами). Априорные методы строят адаптивные стратегии генерации выборок, помещая больше выборок в более сложные области изображения или пространства интегрирования. Апостериорные методы удаляют шум, работая на выходе системы расчёта освещённости, и не влияют на сам процесс построения изображения.

Однако, несмотря на достаточно высокую степень развитости, большинство априорных методов, рассматриваемых в [47], не могут быть использованы для построения эффективных алгоритмов рендеринга. Причина этого в том, что они используют один базовый алгоритм расчета. Если этот базовый метод неэффективен в некоторой области изображения или многомерного пространства интегрирования, априорные методы стараются помещать больше выборок в эту область, чтобы «задавить» шум количеством. Такой подход, очевидно, имеет определенный недостаток – при наличии трудно вычисляемых феноменов освещённости адаптивная генерация выборок начинает сосредотачивать большую часть вычислительных ресурсов в областях пространства, где расчёт неэффективен. При этом страдает качество в остальных участках изображения, а трудно вычисляемые области по-прежнему остаются шумными. Исключением является работа [49] и упомянутый ранее Path Guiding, в которых адаптивная генерация выборок производится в пространстве высокой размерности. Однако размерность пространства в [49] всё же ограничена небольшой

величиной (4D–5D), и отмечается, что для более высокой размерности возникнут проблемы, которые ещё только предстоит решить. Методы на основе MCMC, рассматриваемые в разд. 5, выполняют адаптивную генерацию выборок в многомерном пространстве автоматически. Интересно отдельно обозначить метод RIS [58], суть которого в том, что можно использовать пространственную (соседние пиксели) или темпоральную (соседние кадры) когерентность пространства интегрирования. RIS объединяет использованные распределения от нескольких соседних пикселей (и/или кадров) изображения для того, чтобы сформировать новое распределение выборок для данного пиксела. Таким образом, хорошее распределение выборок в RIS получается как взвешенная сумма большого числа распределений из соседних пикселей или кадров. В настоящий момент метод применяется только для первичного освещения и работает в пространстве экрана [58].

4.5 Рендеринг в пространстве градиентов

Идея рендеринга в пространстве градиентов состоит в том, чтобы одновременно с обычным изображением генерировать выборки ещё и для изображения градиентов. Полагая далее, что градиент посчитать проще (т.к. он разреженный по сравнению с обычным изображением), можно восстанавливать итоговое изображение по градиентам. В [50] был представлен алгоритм, названный Gradient-domain Path Tracing (GDPT), и показано улучшение относительно трассировки путей.

Metropolis Light Transport, к которому мы будем неоднократно возвращаться, обычно строит распределение пропорционально некоторой целевой функции освещения, которая имеет линейную зависимость с изображением. В работе [54] эта идея была расширена на нелинейную зависимость между изображением и целевой функцией. Генерация выборок производилась пропорционально градиентам изображения, а само изображение реконструировалось при помощи решения уравнения Пуассона. Впоследствии в работах [51–52] рендеринг в пространстве градиентов был расширен на двунаправленную трассировку путей, а в [53] и на фотонные карты.

На наш взгляд преимущества и недостатки рендеринга в пространстве градиентов в целом схожи с преимуществами и недостатками, получаемыми от простого применения апостериорных методов шумоподавления (рассматриваемых в [47] к изображению или отдельным выборкам, поскольку в обоих случаях используется идея «меньше считать, больше восстанавливать»). Даже характер артефактов схожий. С другой стороны, реализация рендеринга в пространстве градиентов существенно усложняет программную систему по сравнению с применением методов шумоподавления. По этой же причине на наш взгляд не стоит подробно рассматривать методы кэширования освещённости [55–56].

4.6 Заключение по ОМС

Таким образом, для методов на основе ОМС существует много эвристических подходов, пытающихся улучшить расчёт отдельных феноменов и хорошо работающих на определённом классе сцен, обычно, в условиях сильных ограничений. Однако можно сказать, что эти подходы на практике трудно применять, поскольку пространство интегрирования многомерное и крайне сложное: в таком пространстве любое предположение, сделанное эвристикой рано или поздно перестанет выполняться. Единственным методом, обеспечивающим надёжность расчёта в таком случае, остаётся многократная выборка по значимости. Но чем более метод надёжен, тем он сложнее в реализации и тестировании и, как правило, медленнее в среднем. Что касается ограничений, постановка задачи зачастую не позволяет их ввести. Это приводит к необходимости использовать разные методы в разных случаях и чрезвычайно усложняет жизнь как разработчикам, так и пользователям систем расчёта освещения.

На сегодняшний день итогом развития методов на основе многократной выборки по значимости (MIS) можно считать метод CMIS [57], который позволяет расширить MIS на континуум стратегий генерации выборок (представленный, как правило, набором параметрических функций), что улучшает сходимость в определённых случаях. Однако проблема оптимального выбора стратегий в CMIS не решается сама по себе. Алгоритм, который решает эту проблему, называется Multiplexed Metropolis Light Transport и рассматривается нами разд. 5.

4.7 Реализации на GPU

Мы полагаем, что исследования эффективности алгоритмов на GPU крайне важны для современных практических приложений, особенно с учётом появления аппаратного ускорения трассировки лучей в современных GPU [59-60]. Поэтому мы не можем обойти вниманием эти работы.

Для эффективной реализации методов расчёта освещения на GPU важно уделять внимание нескольким вещам.

- (a) Важно уменьшать потребление памяти на 1 поток, поскольку число запускаемых потоков (которые выполняют несколько вычислительных ядер последовательно) на GPU может исчисляться сотнями тысяч. Например, в [61] особое внимание уделяется экономии памяти для BPT. Для этого предлагается специальный механизм расчёта весов многократной выборке по значимости, для которого не нужно хранить все вершины пути в памяти. А в [6] предлагается использовать усечённую двунаправленную трассировку путей на GPU.
- (b) Необходимо стремиться уменьшить количество дивергентных потоков и решать проблему нерегулярного распределения работы, что обычно делается при помощи регенерации, уплотнения и сортировки потоков, а также разбиения сложного кода на набор простых вычислительных ядер [62-64].
- (c) Для методов на основе фотонных карт нужно использовать быстрые методы построения пространственных структур данных на GPU для поиска ближайших фотонов на основе хэш-таблиц или деревьев, использующих коды Мортонa [65-69]. Отдельно следует отметить здесь работу [70], где в отличие от предыдущих работ все уровни дерева строятся параллельно.

Интересно отметить работу [71], где на GPU был реализован алгоритм VCM. Для эффективной реализации VCM в сочетании с BPT предлагается специальная структура данных, названная Light Vertex Cache, которая позволяет одновременно решать первую и третью проблему из данного списка.

Наконец, нельзя обойти вниманием работу [72], в которой предлагается механизм эффективной диспетчеризации GPU ядер, поддержка рекурсии при помощи сохранения данных для ядер в отдельных очередях (для каждого ядра своя очередь) и использования механизма создания работы на GPU. Важно, что механизм, предлагаемый в [72], решает проблему дивергентных потоков, поскольку диспетчеризация фактически производится без ветвления: сначала для программы/ядра каждого типа в очереди набирается достаточно большое число потоков, после чего они все выполняют одну программу/ядро.

5. Методы на основе MCMC

Монте-Карло по схеме марковских цепей (Markov Chain Monte Carlo, MCMC) можно рассматривать как обобщение обыкновенного Монте-Карло. Если в обыкновенном Монте-Карло выборки независимы, то в MCMC они наоборот коррелированы. Благодаря этому в MCMC в отличие от ОМС информацию об областях функции с высокой значимостью удаётся переиспользовать. Один из наиболее распространённых вариантов MCMC – алгоритм Метрополиса и его более общий вариант, алгоритм Метрополиса-Гастингса. Цель всех 20

МСМС алгоритмов заключается в том, чтобы построить распределение выборок пропорционально произвольной целевой функции.

Таким образом, в противовес методам, основанным на многократной выборке по значимости, алгоритм Метрополиса для расчёта освещённости (Metropolis Light Transport, MLT) [73] генерирует выборки не пропорционально какой-либо одной части подынтегральной функции (что делает каждая из стратегий ВРТ), а пропорционально всему интегралу освещённости как многомерной функции, спроецированной на множество пикселей изображения: $F(x, y, r_0, r_1, \dots, r_n) \xrightarrow{\text{project}} F(x, y)$. MLT автоматически помещает больше выборок в более сложные участки функции F , значительно уменьшая дисперсию [76]. Вопрос сходимости для MLT более сложный. В частности, в [74] приводится оценка $O(\gamma^N)$, где $\gamma \in (0; 1]$.

Следует, однако, отметить, что многократная выборка по значимости и алгоритм Метрополиса *могут и должны использоваться вместе*. Именно так и было сделано в самой первой работе по MLT [73], где алгоритм Метрополиса был предложен для ВРТ, и использованы переходы в *пространстве путей* (path space) как небольшие изменения позиций вершин пути.

5.1 Основы MLT

Прежде чем перейти к дальнейшему рассмотрению методов на основе МСМС, необходимо оговорить условия, в которых эти методы позволяют получить корректный результат. Здесь есть несколько принципиальных моментов.

- (a) Ни один из МСМС методов не позволяет напрямую посчитать интеграл освещённости. Эти методы строят многомерное распределение выборок и накапливают в двумерном массиве гистограмму, которая пропорциональна целевому изображению, но не равна ему. Для того чтобы получить итоговое изображение, эту гистограмму нужно скалировать – то есть умножить на некоторую константу нормализации. Оценить константу нормализации можно, например, если посчитать среднюю яркость изображения – средняя яркость скалированной гистограммы должна быть такой же. Это можно сделать при помощи расчёта грубого изображения любым из методов на основе ОМС. Таким образом, ОМС и МСМС всегда работают вместе.
- (b) Для того чтобы МСМС методы работали корректно, марковский процесс должен быть эргодичным. Это означает, что всё пространство состояний должно быть достижимо во время случайных блужданий марковской цепи (для алгоритма Метрополиса, где прямые и обратные переходы равновероятны). Все предложения переходов разбиваются на 2 части: большие и маленькие шаги. Маленькие шаги являются той самой «рабочей лошадкой», которая обеспечивает эффективную генерацию выборок, используя тем или иным образом локальность интегрируемой функции в многомерном пространстве. Большие же шаги, как правило, призваны лишь обеспечить эргодичность. На практике, однако, низкая вероятность принятия большого шага может приводить к росту начального смещения (startup bias). Отдельно взятые большие шаги в MLT фактически являются базовым расчётным методом на основе ОМС и могут служить, например, для оценки константы нормализации.
- (c) Каждый шаг марковской цепи добавляет близкий к единичному (либо даже строго единичный) вклад в гистограмму изображения (пункт 1). Из этого необходимо сразу сделать вывод: при помощи Metropolis Light Transport не стоит считать первичное освещение, т.к. он просто «уткнётся» в блики и другие яркие области (например, источник), и практически никогда не будет из них выходить, набирая значения в ярких областях единицами.

5.2 Paths Space vs Primary Sample Space

МСМС методы в рендеринге можно условно разделить на два класса по тому, в каком пространстве они работают. К первому классу относятся методы, работающие в *мировом пространстве путей* (Path Space; мы будем также называть такой тип VeachMLT) – пространстве, составленном конкатенацией всех координат вершин путей в мировом пространстве (например, это Veach MLT [73], MEMLT [15], HSLT [77] на рис. 1). Ко второму классу можно отнести методы, работающие в так называемом *первичном пространстве путей* (Primary Sample Space; мы будем также называть такой тип KelemenMLT) – пространстве всех случайных чисел, используемых Монте-Карло выборкой, т.е. многомерном единичном кубе (например, это Kelemen MLT [75] и MMLT [78] на рис. 1). На рис. 4 приведена иллюстрация обоих пространств. В практических приложениях, как правило, используются методы второго класса, т.к. они более универсальны по отношению к различным феноменам освещённости. Path Space методы, напротив, более специфичны. Например, MEMLT хорошо работает с SDS-путями, а HSLT разработан для многократных glossy-отражений.

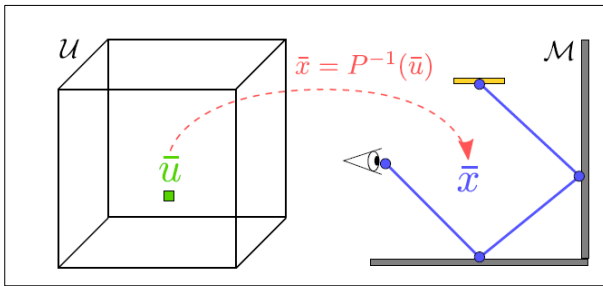


Рис. 4. Иллюстрация перехода из первичного пространства путей U (слева, представленное многомерным единичным кубом) в мировое пространство путей M (справа, представленное конкатенацией всех позиций вершин пути в мировом пространстве). Отметим, что в работах [88-90] необходимо уметь строить обратный переход $\bar{u} = P(\bar{x})$, что является нетривиальной операцией.

Использован рисунок из работы [78]

Fig. 4. An illustration of the transition from the primary path space U (on the left, represented by the multidimensional unit cube) to the world path space M (on the right, represented by the concatenation of all positions of the vertices of the path in world space). Note that in [88-90] it is necessary to be able to construct the inverse transition $\bar{u} = P(\bar{x})$, which is a nontrivial operation. Used drawing from work [78]

Наиболее существенным недостатком MLT в пространстве путей является то, что под каждый феномен освещённости необходимо не только тщательно проектировать специфические стратегии пертурбаций, но и уметь вычислять плотности вероятности переходов – $T(x \leftarrow y), T(x \rightarrow y)$, для правила Метрополиса (формула 5.1):

$$a(x \rightarrow y) = \frac{f(y)T(x \leftarrow y)}{f(x)T(x \rightarrow y)}. \quad (5.1)$$

Современные системы синтеза реалистичных изображений содержат десятки различных типов материалов и источников освещения, что приводит к сотням всевозможных типов взаимодействий (феноменам освещённости). Это чрезвычайно усложняет процесс разработки вышеупомянутых стратегий и процесс вычисления плотностей вероятности $T(x \rightarrow y)$ и $T(x \leftarrow y)$, поэтому MLT в пространстве путей практически не применяется на практике, а используется в основном в исследовательских и демонстрационных приложениях. В Primary Sample Space MLT, с другой стороны, прямая и обратная вероятности (вернее, плотности вероятности) одинаковы. Из этого следует, что $T(x \rightarrow y)$ и $T(x \leftarrow y)$ в формуле 5.1 сокращаются и, следовательно, их не нужно вычислять.

5.3 Multiplexed MLT (MMLT)

Как уже было отмечено, многократная выборка по значимости и алгоритм Метрополиса **могут и должны использоваться вместе**. Однако это можно сделать разными способами – Veach MLT [73], Kelemen MLT [75], MMLT [78]. Ключ к успеху MMLT лежит в построении такого пространства интегрирования, в котором алгоритм Метрополиса и многократная выборка по значимости не конкурируют, а усиливают друг друга.

Для того чтобы комбинировать алгоритм Метрополиса и многократную выборку по значимости более эффективно, в MMLT [78] строится «мультиплексированное» пространство интегрирования. Это происходит путём добавления двух степеней свободы: глубины трассировки d и стратегии генерации выборки пути s . Далее марковская цепь статистически находит оптимальный способ построения пути в ВРТ, варьируя параметр s (рис. 5, формула 5.2). Благодаря этому алгоритм Метрополиса автоматически перераспределяет вычислительные ресурсы таким образом, что малозначимые стратегии и соединения в ВРТ считаются редко. Причём это происходит в том числе и с учётом функции видимости, т.к. алгоритм Метрополиса строит распределение пропорционально итоговому ответу.

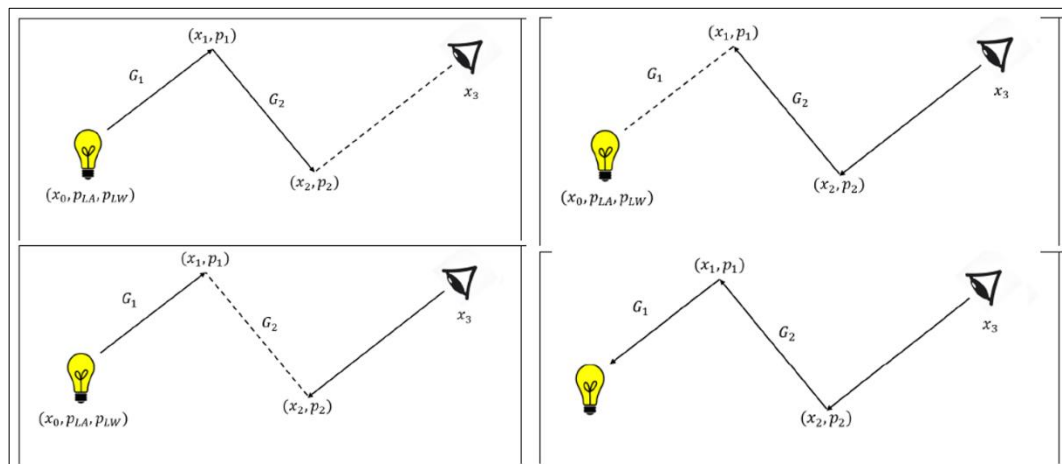


Рис. 5. Пример четырех стратегий для глубины трассировки, равной 3. Пунктирная черта изображает явное соединение вершин. Оптимальный выбор стратегии происходит автоматически алгоритмом Метрополиса, т.к. функция вклада в MMLT построена в виде взвешенной суммы различных стратегий (формула 5.2)

Fig. 5. An example of four strategies for a tracing depth of 3. The dashed line represents an explicit vertex connection. The optimal choice of strategy occurs automatically by the Metropolis algorithm, since the contribution function in MMLT is built as a weighted sum of various strategies (formula 5.2)

$$C(\bar{u}) = \sum_{i=1}^4 \hat{\omega}_i \hat{C}_i^*(\bar{u}). \quad (5.2)$$

Здесь $\hat{C}_i^*$ –результат выборки i -ой стратегии, а $\hat{\omega}_i$ – соответствующий ей вес.

Более передовые методы, как правило, построены поверх MMLT и нацелены на его улучшение или устранение его проблем. Эти проблемы происходят из-за того, что мультиплексированное пространство интегрирования в MMLT, благодаря которому алгоритм Метрополиса и многократная выборка по значимости хорошо работают в связке, само по себе более сложно для ОМС (т.е. вероятность попадания в существенную область пространства равномерно случайной выборкой в нём меньше), чем первичное пространство путей в Kelemen MLT [75]. Это происходит из-за априорного разбиения

(«мультиплексирования») путей по глубине и стратегиям, что, в свою очередь, создаёт две проблемы: (1) низкая эффективность равномерно-случайной генерации выборок, и как следствие, низкая вероятность принятия большого шага и (2) невозможность смены стратегии генерации выборки во время маленького шага.

Одним из возможных путей решения первой проблемы (т.е. повышение эффективности большого шага) являются работы [80-81], где используется «обмен репликами» (Replica Exchange). К этой работе мы вернёмся в подразделе 5.7.

5.4 Veach и Kelemen MLT вместе

Невозможность смены стратегии генерации выборки в MMLT во время маленького шага служит источником дополнительных артефактов на изображении, поскольку отдельные участки изображения (где MIS-веса далеки от нуля и единицы) строятся разными стратегиями и зачастую, как следствие, разными цепями Маркова. Для решения этой проблемы в работах [88-90] были разработаны гибриды MMLT с методами, работающими в мировом пространстве путей. В этих работах строится обратимый переход между мировым пространством путей и первичным пространством путей, за счёт чего удаётся добиться смены стратегии генерации выборки во время *маленького шага* при помощи Reversible-Jump MCMC, который способен работать сразу в нескольких доменах/пространствах интегрирования.

Улучшения, достигнутые таким способом, на наш взгляд являются не очень значительными. Однако при этом они существенно усложняют расчётный метод, поскольку не только требуют реализации пертурбаций в Path Space и вычисления транзитивных вероятностей перехода, но и дополнительно требуют реализации нетривиального преобразования самого пути обратно из Path Space в Primary Sample Space: для любого пути в сцене необходимо найти все возможные наборы случайных чисел, по которым такой путь мог бы быть получен (рис. 4).

5.5 Гибридный метод Монте-Карло

Гибридный метод Монте-Карло [92] – это довольно большой, постепенно расширяющийся класс методов. Эти методы генерируют выборки при помощи траектории динамической системы: гамильтоновой механики (НМС) или броуновского движения в вязкой среде на основе уравнения Ланжевена (ЛМС). Поскольку мы считаем эти методы перспективным направлением в рендеринге, остановимся на них более подробно.

5.5.1 Общая идея гибридного МС

В сложных сценах фазовое пространство заполнено несвязанными областями («островками») весьма сложной формы (пример на рис. 6). Целевая интегрируемая функция (значимость) равна нулю в большей части пространства и отлична от нуля только в пределах этих островков, а на их границах обычно испытывает разрыв. Поэтому если мы будем выбирать точки случайным образом, то почти все они попадут в область нулевой значимости и никакого вклада в интегрирование не внесут. Тем не менее, по островкам мы все же иногда попадаем, и потому эргодичность марковского процесса не нарушается.

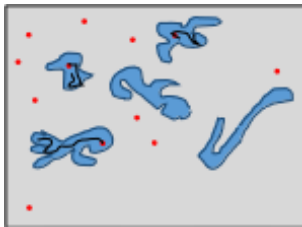


Рис. 6. Гибридный метод Монте-Карло. Прямоугольник – фазовое пространство, контуры – области, где функция значимости отлична от нуля (т.е. точка является допустимой траекторией, соединяющей камеру с источником света), отдельные точки – первичная случайная выборка, ломаные линии внутри контуров – траектории, построенные из попавших в области первичных точек

Fig. 6. Hybrid Monte Carlo method. Rectangle – phase space, contours – areas where the significance function is nonzero (i.e., a point is a valid trajectory connecting the camera to the light source), individual points – primary random sampling, broken lines inside the contours – trajectories constructed from the primary points that are in the area

Однако с ростом размерности любая тонкая область в многомерном пространстве будет «ощущаться» интегратором ещё тоньше, поскольку объём части (относительно целого) стремительно падает с ростом размерности пространства. В результате изотропное предложение перехода МСМС перестает работать: находясь в пределах островка, маловероятно попасть в пределы того же самого островка при помощи случайного блуждания. Чтобы это исправить, случайное блуждание точки в этот момент заменяется на *направленный* процесс обхода островка многомерного пространства при помощи упомянутой динамической системы.

После попадания в островок строим из этой точки *траекторию* $x(t)$, которая не покидает островок и исследует его с плотностью $\rho(x)$, пропорциональной интегрируемой функции (рис. 6). x – вектор координат трассы в первичном пространстве путей для трассы некоторой фиксированной длины. Траектория строится *фиксированное время* t_n (параметр метода), и ее конечные положения $x(t_n)$ включаются в выборку для интегрирования методом Монте-Карло.

Для реализации такого *направленного* обхода островка нужна динамическая система, траектория которой будет, во-первых, ограничена островком, а во-вторых, в пределах него будет двигаться так, чтобы плотность точек была максимально близка к нормированной функции значимости. Небольшое отличие вполне допустимо, его несложно скомпенсировать, например, методом Гастингса, сравнивая значимость для $x(t_{n+1})$ и $x(t_n)$ и выбирая, которую из них добавить к выборке.

5.5.2 Типы динамических систем

Придумать динамическую систему, которая будет порождать траектории с заданным распределением, весьма непросто. Поэтому, естественно взять аналогию из статистической физики с *каноническим распределением* $\rho(x) = Ze^{-E(x)/T}$, где E – «энергия состояния», T – температура в подходящих единицах (т.е. умноженная на постоянную Больцмана), Z – просто нормировка. В качестве энергии берем логарифм функции значимости $f(x)$: $E(x) = const - T \log f(x)$. Если моделировать динамику и брать результат в некоторые случайные моменты времени, то мы будем получать искомое распределение точек x , пропорционально целевой функции $f(x)$ [91].

Среди методов расчета глобального освещения существует две реализации гибридного Монте-Карло с динамическими системами, восходящими к механике: консервативная гамильтонова динамика [94] и броуновское движение в вязкой среде, т.е. один из типов уравнения Ланжевена [97].

5.5.3 Hamiltonian MC

В работе [94] в качестве динамической системы выбирается гамильтонова динамика

$$\frac{dx}{dt} = \frac{\partial H}{\partial p}, \quad \frac{dp}{dt} = -\frac{\partial H}{\partial x},$$

где обобщенные координаты – это и есть x , обобщенный импульс p вводим, исходя из простейшего выражения кинетической энергии $K = \sum_i \frac{p_i^2}{2m_i}$. Потенциальная же энергия будет тем самым логарифмом функции значимости (формула 5.3):

$$U(\bar{x}) = -\log(f(\bar{x})). \quad (5.3)$$

Из-за ее обнуления на границе островка получаем бесконечно высокий потенциальный барьер, пересечь который наша траектория не сможет. Коль скоро полная энергия H есть сумма (формула 5.4) потенциального члена, зависящего только от обобщенных координат, и кинетического, зависящего только от обобщенного импульса, то каноническое распределение будет *произведением* $\rho_x(x)\rho_p(p)$. Поэтому мы просто игнорируем импульс и рассматриваем распределение только координат, что удобно.

$$H(\bar{x}, \bar{p}) = U(\bar{x}) + K(\bar{p}) = -\log(f(\bar{x})) + \sum_i \frac{p_i^2}{2m_i}. \quad (5.4)$$

Сама по себе гамильтонова динамика часто эргодичностью не обладает. Это видно хотя бы уже из того, что полная энергия не меняется вдоль траектории, так что она ограничена подобластью, определяемой начальными условиями (начальной энергией). Поэтому надо добавить к нему взаимодействие с внешним стохастическим миром. Традиционно это делается так. В начальной точке задаются обобщенные координаты (это вектор координат трассы в первичном пространстве путей) и импульсы (которые в трассировке лучей не существуют). Импульсы определяют, как случайный вектор (например, с гауссовым распределением) с нулевым средним. С этой точки траекторию строят некоторое время. Затем импульсы в конечной точке заменяют на случайный вектор, внося в траекторию «излом», и ведут интегрирование дальше. Грубо говоря, это «случайный удар» по нашей обобщенной частице. Он-то и обеспечивает переходы между состояниями с разной полной энергией, поскольку в точке «удара» мы меняем кинетическую энергию. В каком-то смысле эта приходящая извне случайность может быть ассоциирована с *температурой*, которая будет мерой средней энергии возмущения. Температура в гамильтоновом Монте-Карло обычно выбирается единичной, а распределение случайных моментов – с единичной дисперсией.

Реализация этого метода достаточно сложна и содержит эмпирические решения. Скажем, кинетическую энергию можно ввести как произвольную положительную квадратичную форму от импульса (p, Ap) :

$$K(\bar{p}) = \frac{1}{2} \bar{p}^T A \bar{p}.$$

Работать метод будет при любой симметричной положительно-определенной матрице A , однако эффективность будет заметно разной [93]. Поэтому матрицу эту лучше выбирать под конкретный случай.

Другой принципиально важной особенностью практической реализации гамильтонова Монте-Карло является применение метода сравнения Гастингса к следующей точке траектории. При численном решении гамильтониан на траектории сохраняется лишь приближенно, и необходимо компенсировать это явление. Наконец, весьма нетривиален вопрос о том, как часто «ломать» траекторию, заменяя импульс на случайный вектор, и с какой дисперсией. Все эти тонкости весьма серьезно влияют на эффективность метода.

5.5.4 ННМС

Основная сложность с реализацией НМС для рендеринга связана с использованием метода интегрирования второго порядка *leap frog* [91], выполняющего симуляцию гамильтоновой динамики. Он требует производной от правой части, которая в гамильтоновой динамике сама является производной от гамильтониана, так что нам в результате нужна матрица вторых производных от логарифма функции значимости (т.н. гессиан). Вычисления матрицы вторых производных от гамильтониана довольно дороги, а делать их надо на каждом шаге. Поэтому в работе [94] предложен вариант, где вместо реального гамильтониана берется его аппроксимация квадратичной формой (ряд Тэйлора). Такая аппроксимация позволяет в течение нескольких шагов по времени вычислять вторые производные аналитически, что заметно снижает стоимость шага.

5.5.5 DRMLT

Возможная оптимизация ННМС заключается в том, чтобы на каждом шаге сначала попробовать более простое изотропное предложение перехода, и если только оно будет отвергнуто, тогда уже использовать анизотропное предложение на основе ННМС. Так было сделано в работе [102], где был предложен гибридный алгоритм, названный Delayed Rejection MLT (DRMLT).

5.5.6 Идея Langevin MC

Работа [96] рассматривает броуновское движение в качестве динамики в пределах островка. Пусть у нас есть частица, движущаяся в потенциальном поле в вязкой среде. Вязкость приводит к тому, что частица остановится в одном из локальных минимумов потенциальной энергии. Приложим теперь к ней случайную силу, она будет выталкивать нашу частицу из потенциальной ямы и заставлять обходить целевую функцию в различных местах, не застревая в локальном минимуме. Если это случайное возмущение не слишком велико, то чем глубже локальный минимум потенциала, тем дольше частица будет из него выбираться. Равновесное распределение, определяемое «борьбой» градиента потенциала и случайной силы, будет, тем самым, иметь максимум там, где потенциал имеет минимум.

Математически такое движение выражается стохастическим дифференциальным уравнением

$$d = -\nabla U dt + \sqrt{2T} dw,$$

где $U(x)$ – потенциал, а w – винеровский процесс, т.е. такой, что приращение $dw = w(t + dt) - w$ есть случайная величина с нулевым средним и дисперсией $\sqrt{t}$ [98]. Данное уравнение – один из вариантов уравнения Ланжевена. Из-за винеровского члена в этом уравнении есть перемешивание, сходимост и эргодичность. Среднее по траектории сходится к равновесному распределению $\rho(x) = Ze^{-U(x)/T}$. Заметим, что температура тут возникает достаточно естественно как амплитуда винеровского члена (случайной силы). Обычно ее берут единичной. Таким образом, как и в случае гамильтоновой механики, в качестве x берем вектор переменных узлов полной трассы луча (без источника и камеры), $T = 1$, а потенциал равен логарифму функции значимости $f(x)$: $U(x) = -\log f(x)$. На границе островка имеем опять бесконечный потенциальный барьер, который наша траектория преодолеть не может. А внутренность его она заполняет как раз с нужной нам плотностью.

В отличие от НМС, в Langevin MC нет математических проблем с перемешиванием, эргодичностью и сходимостью распределений. Температура здесь не является чужеродным фактором. С математической точки зрения этот вариант выглядит привлекательнее. Он должен работать просто при реализации в соответствии с математической идеей. Тем не менее, и для Ланжевена подчас используют метод сравнения Гастингса, а также необходимо как-то выбрать длину траектории и пр.

5.5.7 Langevin MC в практике

Работа [96] не только использует более простое представление для производных (градиент), но и кэширует их вычисление, благодаря чему удаётся существенно снизить накладные расходы на их вычисление. Основная идея кэширования заключается в том, чтобы для алгоритма MALA не вычислять градиенты целиком для каждого шага, а обновлять их при помощи ограниченной порции вычислений по мере движения цепи (эта идея в работе [96] называется «online adaptation»). При этом необходимо отметить, что алгоритм MALA, используемый в [96], является частным случаем полноценного НМС, когда leap frog [91] делает один шаг, и этот шаг сразу берётся в качестве следующего предложения перехода (формула 5.5):

$$\bar{u}^{t+1} = \bar{u}^t + \epsilon \nabla L(\bar{u}^t) + \sigma \sqrt{\epsilon} \mathcal{N}(0,1), \quad (5.5)$$

где $\bar{u}$ – вектор состояния цепи, $L(u) = -\log(F(\bar{u}))$, ϵ и σ – параметры алгоритма, а $\mathcal{N}(0,1)$ – единичный многомерный гауссиан.

К сожалению, наивное применение MALA к задаче рендеринга не даёт существенного улучшения, что демонстрирует рис. 3 работы [96]. Как там отмечается, MALA всё же предполагает, что случайные величины в векторе L имеют приблизительно одинаковую дисперсию. Если это не так, сходимость ухудшается. Существующие решения этой проблемы через так называемую матрицу предварительной обработки (preconditioning matrix) [104] оказываются слишком дорогими для рендеринга в силу дороговизны вычисления целевой функции. Работа [96] фокусируется на оптимизации вычисления этой матрицы с использованием *адаптивной предварительной обработки* (adaptive preconditioning) и реализации таким образом алгоритма MALA с *адаптацией* (MALA with adaptation). Она предлагает 3 варианта оптимизации:

- (a) MALA с подкреплением (online adaptation);
- (b) MALA на основе кэширования (cache-driven adaptation);
- (c) гибрид двух предыдущих подходов (hybrid adaptation).

5.5.8 Заключение по гибриднему MC

К сожалению, все три упомянутые работы требуют применения методов автоматического дифференцирования [99], что оставляет вопрос практической применимости открытым, т.к. делает добавление новых математических моделей материалов и источников освещения в расчётную систему крайне сложным и трудоёмким процессом. Например, в системе Mitsuba2 [101] предлагается расширяемая система со встроенными возможностями дифференцирования. Тем не менее, Mitsuba2 не реализует НМС, а использует дифференцирование для решения задачи так называемого обратного рендеринга, когда по заданному изображению требуется восстановить некоторые характеристики сцены или отдельных моделей.

Важный нюанс, на который должны обратить внимание разработчики расчётного ПО в работах [94] и [96], – это дифференцирование при помощи шаблонного мета-программирования C++ [99]. Такой код крайне трудно развивать и отлаживать, а его эффективность зачастую невозможно оценить. Кроме того, это налагает определённые ограничения на используемый компилятор и затрудняет перенос на GPU. Хотя применение инструментов на основе source-to-source преобразований исходного кода при помощи специального компилятора выглядит перспективным направлением с практической точки зрения [100] и позволяет надеяться на внедрение НМС в расчётные системы в будущем.

Кроме того, в работах [94] и [96] отмечается, что при наличии высокочастотных карт нормалей или карт смещений (имитация микрорельефа поверхности) дифференцирование не будет корректно работать. При этом очевидно, что поддержка Spatial Varying BSDF (SVBSDF [109]) также находится под вопросом, поскольку для них невозможно получить производную

аналитически. На наш взгляд, для современных приложений это даже более критичная функциональность, чем карты смещений.

Мы полагаем, что в настоящее время высокоэффективная реализация этих методов на GPU выглядит вряд ли возможной в силу того, что НМС [94] имеет высокие затраты памяти на 1 поток из-за необходимости вычисления гессииана. Аналогичные проблемы наблюдаются и в работе [96], где на каждый поток нужно хранить *preconditioning matrix*.

Тем не менее, методы расчёта освещения на основе гибридного Монте-Карло являются перспективным направлением исследований. Эти методы достаточно универсальны, их математическое обоснование хорошо проработано, они обладают лучшей сходимостью при росте размерности пространства интегрирования: $d^{5/4}$ для НМС против d^2 для МСМС, где d – размерность ([91], стр. 29, конец раздела «Scaling of HMC and random-walk Metropolis»). Численное сравнение можно найти в работе [95], раздел 7.2. Это происходит благодаря анизотропному предложению перехода, который строится на основе производных. НМС существенно выигрывает у МСМС потому, что с ростом размерности любая «тонкая» область пространства становится ещё тоньше, из-за чего у изотропного предложения перехода в МСМС очень быстро снижается шанс остаться в области функции с высокой значимостью.

5.6 Population Monte Carlo

Обособленно от других находится класс методов, основанный на отборе популяции выборок (поэтому они называются Population Monte Carlo, или PMC) [105-106]. Основная идея этих методов состоит в том, чтобы запоминать информацию о выборках во времени, и затем переиспользовать наиболее удачные выборки как стартовое состояние для небольших шагов (мутаций). Схожую идею использует упомянутая нами ранее работа [80], которая переиспользует части состояний марковской цепи для повышения вероятности принятия большого шага в стандартном МСМС с длинными цепями. Можно сказать, что эти три метода близки по способу исследования пространства к генетическим алгоритмам (что на наш взгляд в целом нельзя говорить по отношению к другим МСМС методам).

Огромное преимущество PMC по сравнению с традиционным МСМС заключается в большей параллельности метода за счёт коротких цепей, что чрезвычайно важно для реализации МСМС на графических процессорах. Например, известно, что коммерческая система расчёта освещения на GPU Octane использует PMC [107]. Кроме того, рассматриваемый класс методов обладает большей интерактивностью (т.е. в начале расчёта изображение выглядит более правильно) по сравнению с традиционными МСМС алгоритмами с длинными цепями. Также PMC может вычислять прямой свет, что категорически не рекомендуется делать во всех видах MLT. Среди недостатков, как правило, меньшая точность в пределе по сравнению с MLT [87].

5.7 Replica Exchange

Метод обмен репликами (RELT) [79-81] можно считать расширением обычного MLT в первичном пространстве путей. RELT подразумевает, что несколько марковских цепей (число K в алгоритме 1, работа [80]) работают параллельно. Отличие от обычного MLT заключается в том, что после каждого шага цепи могут обмениваться состояниями с некоторой вероятностью. Для того, чтобы такой обмен не нарушал стационарность распределения, вероятность должна вычисляться по формуле 5.6 [1]:

$$p(i, j) = \frac{F_i(u_j) \times F_j(u_i)}{F_i(u_i) \times F_j(u_j)}. \quad (5.6)$$

Здесь u_i и u_j – это состояния, между которыми предполагается произвести обмен, а F – целевая функция, пропорционально которой необходимо построить распределение выборок. Необычным здесь является то, что функций F на самом деле несколько. В методе обмена

репликами используется так называемая процедура закалки функций (function tempering [82]). Суть этой процедуры в том, чтобы применить к целевой функции преобразование следующего вида (формула 5.7):

$$F(x) = F(x)^{1/T_l}. \quad (5.7)$$

Здесь $T_l \geq 1$ – параметр, называемый температурой. Чем больше температура, тем более плоской является преобразованная функция. Цепи упорядочиваются в соответствии с их температурой, чтобы обмен производился только между соседними цепями, что повышает вероятность принятия обмена.

Преимущества RELT (как и PMC) – лучшее исследование пространства интегрирования в глобальном смысле и лучшая параллелизуемость. Основным недостатком RELT в [1] считают тот факт, что сходимость сильно зависит от количества цепей, а это количество трудно подобрать правильно. С другой стороны, при реализации на GPU большое количество цепей не является проблемой. RELT используется во многих работах [40-41, 83] как метод, улучшающий глобальное исследование пространства Марковскими цепями.

Стоит отметить, что применение идеи обмена реплик на практике имеет нюансы. Ключевой вопрос заключается в выборе стратегий генерации выборок для каждой цепи (реплики). В [80] было использовано 4 типа реплик: (1) равномерная стратегия генерации выборок (большой шаг), (2) реплика «полного пути», (3) реплика «непрямого освещения» и (4) специальная реплика для дельта-функций. При этом вторая и третья реплики на самом деле являются комбинациями различных типов путей (получаемых разными стратегиями генерации выборок), а четвёртая реплика используется исключительно для SDS каустиков при помощи неявной стратегии. Таким образом, практическое применение RELT тесно переплетается с использованием различных стратегий в многократной выборке по значимости в BPT или PT, и при этом вклад от различных реплик между собой также учитывался на основе многократной выборки по значимости.

Причина повышения эффективности RELT (по сравнению, например, с Kelemen MLT) заключается в том, что каждая реплика работает со специфическим типом путей. Благодаря этому вероятность сделать удачное предложение перехода внутри реплики повышается (можно, например, сделать шаг меньше для четвёртого типа реплик), а фактическая размерность пространства интегрирования понижается.

5.8 Стратификация (Tiled MLT)

Впервые идея стратификации была применена к MCMC в рендеринге в работе [83]. Изображение разбивалось на перекрывающиеся участки, каждый из которых обрабатывался отдельной марковской цепью независимо. Для каждого из участков оценивалась своя константа нормализации, после чего можно было собрать изображение из участков, зная константу нормализации для всего изображения. Благодаря стратификации в [83] достигался более равномерный обход изображения марковскими цепями, за счёт чего изображение сходится к эталону более предсказуемо.

5.9 Markov Chain PPM

Существует подход, объединяющий в себе фотонные карты и марковские цепи. Так в [39-40] (Markov Chain PPM, MCPPM) марковские цепи были применены к стохастическим прогрессивным фотонным картам (SPPM [19]), а в [41] (MBE) к алгоритму VCM [24]. MCPPM позволяет решить проблему плохого распределения фотонов на сценах с большой площадью. Однако MCPPM, как и оригинальный SPPM, плохо справляется с многократными глянцевыми отражениями.

Объединение фотонных карт и марковских цепей может улучшить расчётные системы, которые основаны на фотонных картах. Однако нам не кажется это направление

30

перспективным по сравнению с полностью несмещёнными методами. Причина нашего скепсиса заключается в том, что и Метрополис (марковские цепи), и фотонные карты накапливают значения интеграла единичными выборками и потому имеют схожие недостатки, которые будут усиливаться в комбинации: в тёмных областях изображение получается шумным (именно для устранения этого недостатка в MBE [41] был применён метод RELT). Отметим, что низкочастотный шум от фотонных карт крайне тяжело удалять [48].

Другой недостаток, который может усиливаться в комбинированном методе, – это образование сгустков фотонов, которые тяжелы для обработки в вычислительном плане. Как мы отмечали ранее, контроль плотности [22] возможен не для всех типов BRDF, что существенно ограничивает применимость Markov Chain PPM в компьютерной графике.

Что касается MBE [41], этот метод прежде всего является комбинацией других методов: VCM, RELT и предыдущих MCPPM подходов. Для лучей из камеры он использует обыкновенную трассировку путей, а для путей из источника использует MCPPM. По опыту реализации гибридных подходов реальная трудоёмкость его реализации и верификации должна многократно превышать сумму соответствующих величин отдельных методов. Поэтому, прежде чем приступить к его реализации, на наш взгляд придётся в том или ином виде сначала реализовать базовые методы. После этого можно будет приступить к более детальному изучению работы [41], чтобы построить гибридный подход. Интересно также отметить, что Multiplexed MBE метода, комбинирующего идеи MMLT с VCM и RELT, в настоящий момент ещё не существует.

5.10 Реализации MLT на GPU

Существенное отличие GPU реализации от CPU реализации для MCMC методов состоит в том, что длина Марковских цепей на GPU будет на несколько порядков меньше. Хотя при этом общее количество используемых цепей намного больше. Несущественное на первый взгляд, это отличие на самом деле является критическим для MCMC, о чём говорится во всех рассматриваемых далее работах.

В [85] на GPU впервые был реализован метод Kelemen MLT [75] на основе традиционной двунаправленной трассировки путей (BPT). В качестве проблем в [85] отмечены высокие траты памяти (обусловленные двунаправленной трассировкой путей и необходимостью хранить весь вектор случайных чисел на каждый поток) и высокое значение начального смещения (start-up bias). В качестве основного направления улучшения исследуются метод регенерации путей (при помощи уплотнения оставшихся активных потоков на GPU после некоторого количества переотражений) и параллельное вычисление соединений в BPT, что в сумме повышает скорость на 20-30% по сравнению с наивной реализацией.

В [86] для Kelemen MLT было предложено решение проблемы большого начального смещения на GPU при помощи отбора начальных состояний марковских цепей обыкновенным Монте-Карло (так называемого параллельного прожига). Также в [86] рассматриваются методы уменьшения занимаемой памяти, а для повышения когерентности лучей используется сортировка потоков. Это даёт примерно такой же прирост в 20-30% производительности, как и в [85].

В [84] впервые сообщается об успешной реализации MMLT на графических процессорах. Предлагаемые методы – Speculative MLT (SMLT) и Rejection Chain MLT (RCMLT) – позволяют вычислять несколько предложений переходов в MCMC параллельно, однако ухудшают ошибку в целом. Идея метода RCMLT состоит в том, чтобы вычислить сразу несколько предложений перехода с запасом. Тогда если первое предложение будет отвергнуто, то можно сразу попробовать следующее. SMLT является развитием RCMLT. Сначала параллельно вычисляется дерево всех возможных переходов, после чего производится фактический переход в один из листов.

Однако следует отметить, что MMLT – это последний из рассматриваемых нами МСМС алгоритмов, для которых известны реализации на GPU. То есть для более новых методов реализации на GPU ещё не существует, что на наш взгляд является результатом увеличения сложности подходов и возрастающих ограничений, к чему мы ещё вернёмся. При этом в [84, 87] отмечается, что из-за мультиплексированного пространства интегрирования проблема начального смещения для MMLT более существенна, чем, например, для более простого Kelemen MLT. Если для Kelemen MLT в существующих реализациях [85-86] проблема начального смещения может быть амортизирована при помощи параллельного прожига [86], то для MMLT этого метода оказывается недостаточно [87].

5.11 Дальнейшее изучение МСМС

Мы рекомендуем также ознакомиться с работой [1] для углублённого анализа методов расчёта освещения на основе МСМС. Эта работа подробно рассматривает принципы работы методов на основе Монте-Карло по схеме марковских цепей, и может служить в качестве отправной точки для понимания их на более глубоком уровне с целью практической реализации. Интересным результатом этой работы на наш взгляд является оценка эффективности МСМС методов по 3 типам и постулирование открытых проблем в области.

- (a) Local exploration. Эффективность локального исследования пространства. Это тип эффективности можно назвать основным, т.к. именно он определяет точность расчёта в пределе.
- (b) Global Exploration. Эффективность глобального исследования пространства. Этот параметр влияет на то, насколько быстро из изображения уходит начальное смещение, что важно, например, для реализации методов на GPU. В [1] отмечается, что на сегодняшний день эффективное глобальное исследование пространства является первой нерешённой проблемой.
- (c) Uniform image error. Равномерность ошибки на изображении. Этот параметр является ключевым для анимации и отчасти зависит от предыдущего. В [1] говорится что почти все МСМС алгоритмы страдают от временной нестабильности, и эта вторая нерешённая в настоящее время проблема. Кроме того, для МСМС методов отмечается отсутствие надёжной метрики оценки ошибки на изображении, которая могла бы быть использована для остановки счёта. Это – третья проблема.

Мы не считаем перечисленные проблемы критичными по следующим причинам.

- (a) Современные методы обработки изображений и шумоподавления (в том числе на основе нейросетей) по нашему опыту в состоянии амортизировать многие из артефактов, вызванных, например, второй проблемой.
- (b) Для многих применений на практике временная стабильность алгоритма не критична. Более важными параметрами могут быть надёжность (robustness) и эффективность расчёта.
- (c) При реализации МСМС методов на GPU количество марковских цепей большое, но сами цепи при этом успевают сделать меньше шагов [86-87]. Это приводит к совершенно иному влиянию обозначенных проблем на изображение, чем при расчёте на CPU. Проблема временной стабильности перестаёт быть критичной, т.к. при таком большом количестве цепей мельтешение из анимации уходит. Однако проблема начального смещения начинает проявляться в виде недооценённой яркости некоторых участков изображения, в которых марковские цепи не успели накопить в гистограмме высокие значения яркости [87]. Хотя это также является серьёзной проблемой, в некоторых приложениях (например, визуализация архитектурных проектов) на наш взгляд такой тип артефактов является менее критичным.

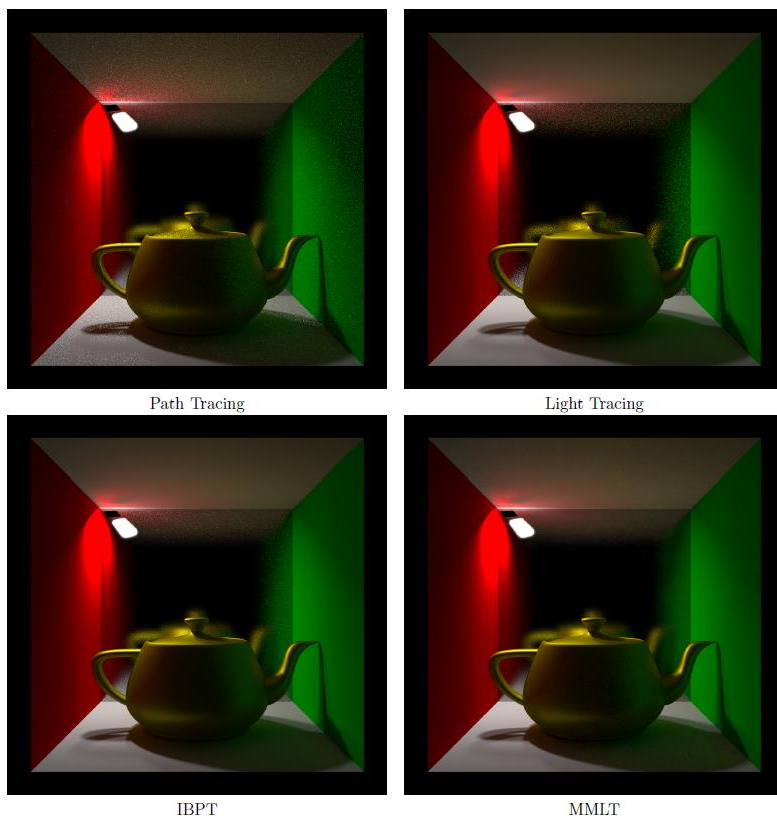


Рис. 7. Сравнение различных алгоритмов интегрирования освещённости. Время рендеринга – 10 минут на процессоре Intel Core i7 3770, 3.4 GHz. В данной сцене источник в виде площадки повёрнут в направлении диффузной стены слева и задней глянцевой стены, благодаря чему значительная часть света – вторичное освещение, а первичное представлено большим бликом на стене слева. Благодаря этому Path Tracing производит значительный шум на всём изображении. Light Tracing неплохо справляется со вторичным освещением, но по сравнению с обратной трассировкой существенно ухудшает точность на глянцевой стене. IBPT работает значительно лучше двух предыдущих алгоритмов, но шум на глянцевой стене всё ещё остаётся. Кроме того, нетрудно заметить, что этот шум больше чем в PT. Это демонстрирует классический недостаток MIS: чем больше стратегий, тем хуже могут вычисляться отдельные элементы сцены, т.к. за то же самое время они получают меньше удачных выборов. MMLT надёжно и эффективно посчитал все видимые эффекты на данной сцене

Fig. 7. Comparison of various algorithms for integrating illumination. Rendering time - 10 minutes on an Intel Core i7 3770 processor, 3.4 GHz. In this scene, the source in the form of a platform is turned towards the diffuse wall on the left and the back glossy wall, due to which a significant part of the light is secondary lighting, and the primary is represented by a large glare on the wall to the left. As a result, Path Tracing produces a significant amount of noise throughout the image. Light Tracing does a good job with secondary lighting, but compared to back tracing, it significantly degrades accuracy on a glossy wall. IBPT performs significantly better than the previous two algorithms, but there is still noise on the glossy wall. In addition, it is easy to see that this noise is greater than in PT. This demonstrates the classic disadvantage of MIS: the more strategies, the worse the individual scene elements can be calculated, since in the same time, they get fewer hits. MMLT reliably and efficiently counted all visible effects on a given scene

6. Эксперименты

В процессе работы с различными методами расчёта освещённости мы проводили экспериментальные сравнения различных методов, используя наши собственные реализации.

Здесь мы не ставим задачу объективного сравнения, поскольку считаем, что такое сравнение невозможно без создания соответствующего набора сцен, который бы объективно отражал некоторую усреднённую реальность. А это является работой, выходящей за рамки данной статьи. Кроме того, некоторые методы были реализованы нами на CPU, а некоторые на GPU. Однако, мы можем использовать эксперименты для того, чтобы зафиксировать некоторые выводы и обозначить занимаемую нами позицию по отношению к тем или иным методам.

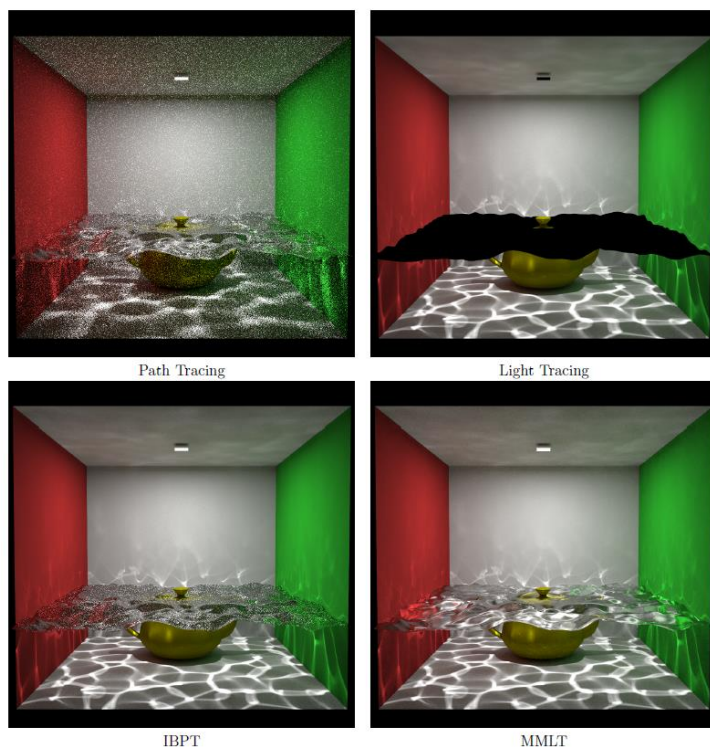


Рис. 8. Сравнение различных алгоритмов интегрирования освещённости. Время рендеринга – 10 минут на процессоре Intel Core i7 3770, 3.4 GHz. Классический пример трудновычислимых каустики. В данной сцене представлены два типа каустики: видимые камерой напрямую (колонка Ind.) на дне и стенах коробки, и каустики, видимые через зеркальное переотражение в воде (колонка SDS, само изображение поверхности воды). На данных изображениях видно, что хотя эффективность расчёта SDS каустики на поверхности воды у MMLT всё же ниже, чем эффективность расчёта более простой каустики на дне и стенах коробки, метод всё ещё можно считать надёжным

Fig. 8. Comparison of various algorithms for integrating illumination. Rendering time – 10 minutes on an Intel Core i7 3770 processor, 3.4 GHz. A classic example of hard-to-calculate caustics. There are two types of caustics in this scene: directly visible by the camera (Ind. Column) on the bottom and walls of the box, and caustics visible through specular reflection in water (SDS column, the actual image of the water surface). These images show that although the efficiency of calculating SDS caustics on the water surface of MMLT is still lower than the efficiency of calculating simpler caustics on the bottom and walls of the box, the method can still be considered reliable.

Эксперимент №1: Хотя двунаправленные методы более надёжны чем однонаправленные (рис. 7-10), нельзя сказать, что они более быстрые. Это хорошо видно из рис. 9, где прямая Монте-Карло трассировка (Light Tracing) лучше посчитала каустику от тора на полу, чем двунаправленная (IBPT). Недостаток MIS проявляется здесь в полной мере: увеличивая надёжность метода, мы понижаем скорость расчёта отдельных явлений. Метод Multiplexed MLT (MMLT) лишён данного недостатка, поскольку стохастически оптимально выбирает стратегии генерации выборок и показывает лучшие результаты на всех 4 сценах (рис. 7-10).

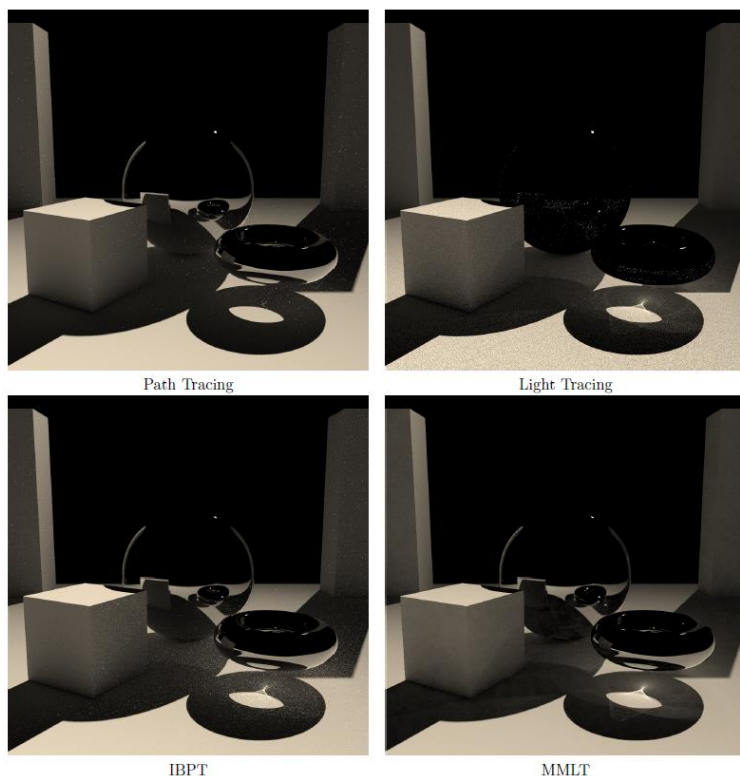


Рис. 9. Сравнение различных алгоритмов интегрирования освещённости. Время рендеринга – 5 минут на процессоре Intel Core i7 3770, 3.4 GHz. На данном изображении демонстрируется важность марковских цепей, поскольку видимая напрямую каустика внизу под тором в MMLT получилась точнее, чем в Light Tracing. Отметим, что на рассматриваемых ранее сценах это не наблюдалось, т.к. ранее объём сцены был крайне ограничен, и Light Tracing, как и фотонные карты, на подобных сценах обладают высокой эффективностью. На данной же сцене эффективность прямого Монте-Карло снижена из-за необходимости генерации случайных лучей по всей площади сцены для источника, имитирующего солнце

Fig. 9. Comparison of various algorithms for integrating illumination. Rendering time – 5 minutes on an Intel Core i7 3770 processor, 3.4 GHz. This image demonstrates the importance of Markov chains, since the caustics seen directly below the torus in MMLT is more accurate than in Light Tracing. Note that this was not observed in the previously considered scenes, since Previously, the volume of the scene was extremely limited, and Light Tracing, like photon cards, is highly efficient in such scenes. On the same scene, the efficiency of the direct Monte Carlo is reduced due to the need to generate random rays over the entire area of the scene for a source simulating the sun.

Эксперимент №2: Фотонные карты хорошо показывают себя при небольшом времени расчёта и в ограниченном наборе сценариев, где преобладают ярко выраженные диффузные и зеркальные поверхности. Но в чистом виде (т.е. без применения ряда усовершенствований) они проигрывают методам на основе Марковских цепей (Kelemen MLT и MMLT) при более длительном расчёте и при увеличении сложности освещения (рис. 11). К счастью, фотонные карты и Марковские цепи можно комбинировать [39-41].

Эксперимент №3: Применение полностью несмещённых методов предпочтительнее использованию смещённых по целому ряду причин: большая точность в пределе, универсальность по отношению к BSDF, отсутствие необходимости в построении ускоряющих структур (что облегчает реализацию на GPU).

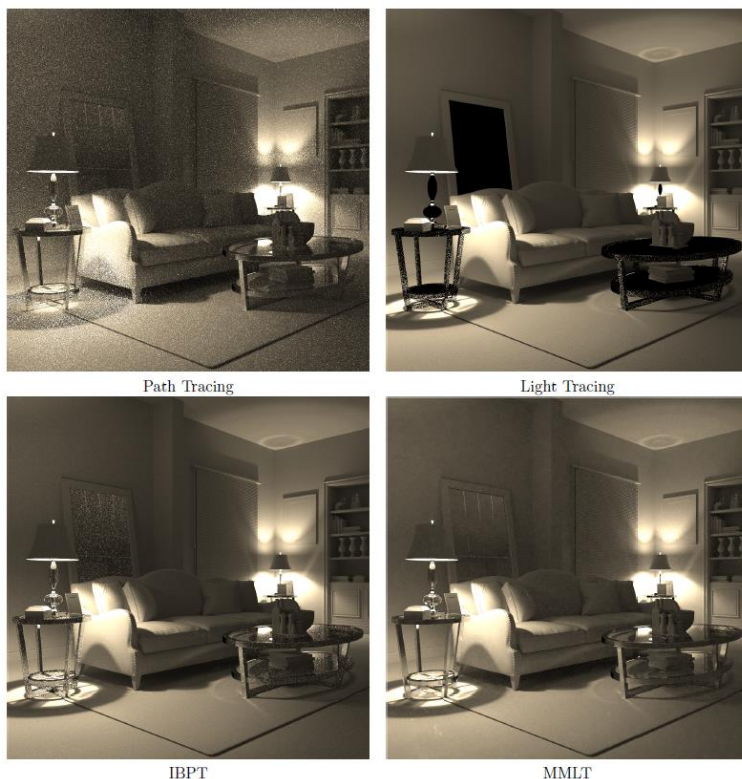


Рис. 10. Сравнение различных алгоритмов интегрирования освещённости. Время рендеринга – 60 минут на процессоре Intel Core i7 3770, 3.4 GHz. Более сложная сцена, содержащая небольшие геометрические детали и три типа материалов: ламбертовские, идеально-зеркальные и материал с глянцевым отражением (ножки столов). Алгоритм MMLT демонстрирует надёжность (robustness) в отличие от других методов. Если сопоставить данный рисунок с рис. 7-9, можно заметить, что при росте сложности сцены преимущество MMLT становится более заметным

Fig. 10. Comparison of various algorithms for integrating illumination. Rendering time – 60 minutes on an Intel Core i7 3770 processor, 3.4 GHz. A more complex scene containing small geometrical details and three types of materials: Lambertian, perfect mirror, and glossy reflective material (table legs). The MMLT algorithm demonstrates robustness in contrast to other methods. If we compare this figure with Fig. 7-9, you can see that as the complexity of the scene increases, the advantage of MMLT becomes more noticeable.

Мы использовали фотонные карты на раннем этапе нашей работы, однако впоследствии заменили их на прямую Монте-Карло трассировку лучей (Light Tracing) в целях изучения вкладов от различных стратегий в BPT. Метод PCLT [30] похожим образом заменяет сбор для первично видимых поверхностей на проекцию точек на плоскость изображения, но при этом выполняет сбор для остальных случаев. Если бы мы добавили фотонные карты к сравнениям на рис. 7-10, то увидели бы улучшенные версии изображений для Light Tracing (если добавлять их без MIS) или IBPT (если добавлять их с MIS, как это делается в VCM) на зеркальных поверхностях. Однако изображения поверхностей с многократными глянцевыми отражениями не получают улучшения от фотонных карт. Кроме того, во всех случаях, где видно преимущество MMLT над Light Tracing на диффузных поверхностях, фотонные карты также не принесут каких-либо улучшений. Скорее наоборот, результат ухудшится, т.к. высокочастотный шум от метода Light Tracing (который легко фильтровать при необходимости) перейдёт в низкочастотный трудноустраняемый шум в виде цветных пятен [48].

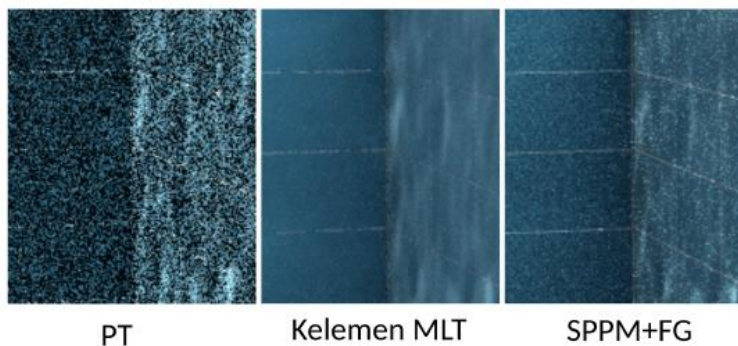


Рис. 11. Сравнение фотонных карт (SPPM+FG) против Kelemen MLT на сцене бассейна с каустиками. Время рендеринга – 10 минут на GPU AMD RX580. Низкая эффективность фотонных карт на этой сцене обусловлена тем, что камера видит не более 10% площади всей сцены (угол бассейна). Из-за этого значительная часть фотонов не долетает до видимой области сцены, и вычисления по их трассировке производятся впустую. Прямоугольник на рис. 12 отмечает участок изображения, который был увеличен

Fig. 11. Comparison of photon maps (SPPM + FG) against Kelemen MLT on the stage of a pool with caustics. Rendering time is 10 minutes on AMD RX580 GPU. The low efficiency of photon maps in this scene is due to the fact that the camera sees no more than 10% of the area of the entire scene (the corner of the pool). Because of this, a significant part of the photons do not reach the visible area of the scene, and calculations for their tracing are wasted. The rectangle in fig. 12 marks the area of the image that has been enlarged

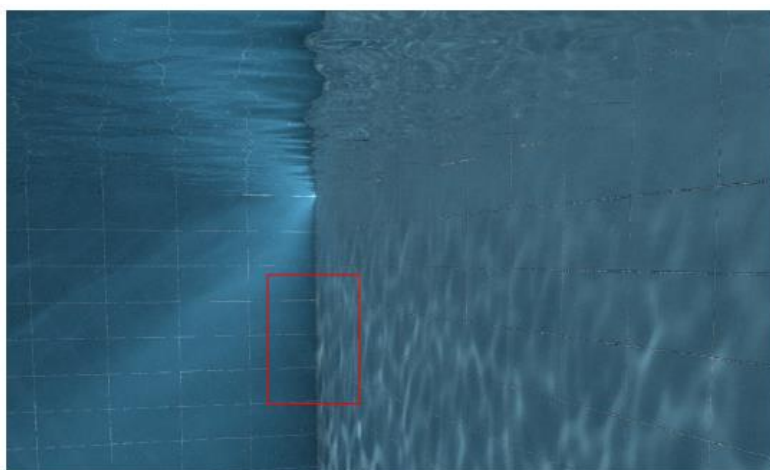


Рис. 12. Сцена с бассейном (использованная в сравнении на рис. 11)

Fig. 12. Pool scene (used in comparison in fig. 11)

Эксперимент №4: Методы на основе Марковских цепей в целом существенно лучше себя показывают на сложном вторичном освещении, чем любые методы на основе обычного Монте-Карло (рис. 7-11). В условиях сложного вторичного освещения их однозначно следует выбирать при реализации алгоритмов на CPU. Однако при реализации алгоритмов на GPU можно столкнуться с эффектом неожиданно возросшего начального смещения, что особенно сильно проявляется в MMLT (рис. 13). Поэтому мы считаем, что для GPU больше подходит Kelemen MLT [75] и PMC [105-106].

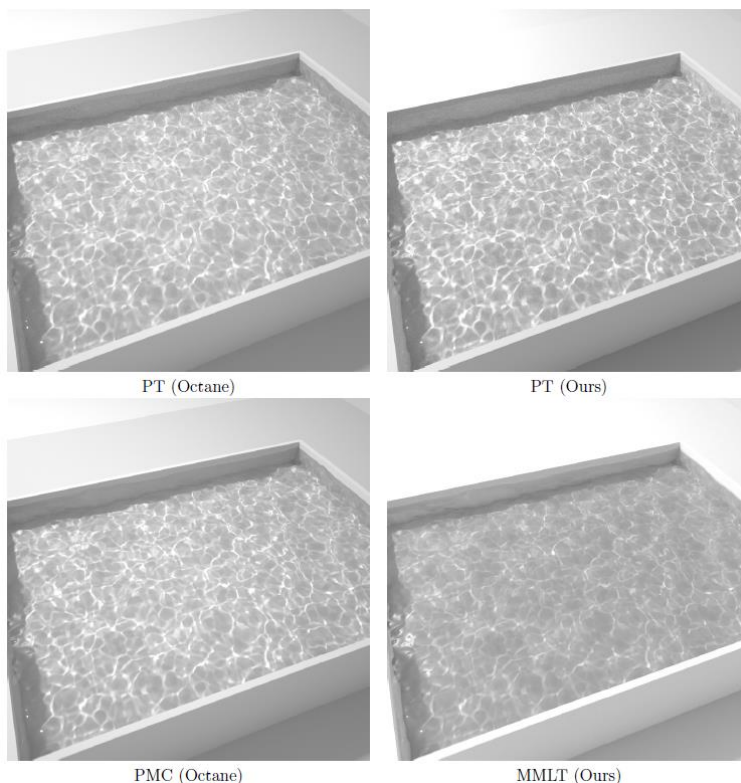


Рис. 13. Демонстрация проблемы большого начального смещения в MMLT при реализации на GPU, которое проявляется в виде недооценённой яркости каустиков. Время рендеринга – 30 минут на GTX1070. Следует отметить, что подобная сцена могла бы быть эффективно посчитана при помощи фотонных карт

Fig. 13. Demonstration of the problem of a large initial displacement in MMLT when implemented on a GPU, which manifests itself in the form of underestimated brightness of caustics. Rendering time is 30 minutes on GTX1070. It should be noted that such a scene could be efficiently computed using photon maps

Эксперимент №5: Сильно недооценённым является метод РМС, который хотя и проигрывал в среднем нашей реализации MMLT [108] (рис. 14), но не катастрофически. Кроме того, надо принять во внимание, что реализация РМС в Octane однонаправленная, в то время как MMLT – двунаправленный алгоритм, что отчасти и обуславливает его преимущества на рис 14. С учётом того, что РМС использует короткие цепи, его параллельная реализация на GPU представляется нам разумным решением, которое и сделали разработчики Octane.

7. Выбор метода

На сегодняшний день нельзя выделить универсальный или лучший метод. Если это допускает постановка задачи, следует всегда принимать во внимание специфику типичных сцен, для которых создается рендерер. Так может оказаться излишним применение развитых и более сложных методов расчёта освещённости. Кроме того, специфика сцен может очень сильно изменить баланс эффективности в сторону того или иного метода. Например, для расчёта освещения от панорамы на открытом пространстве в большинстве случаев подойдёт применение обычного MIS PT, т.к. почти всё освещение в таком случае будет первичным. Если в High Dynamic Range (HDR) панораме, например, встречаются отдельные сверхяркие области, то могут помочь простые МСМС методы, вроде Kelemen MLT или РМС поверх MIS PT. Применять двунаправленные методы в этом случае бессмысленно по причине низкой эффективности прямой стратегии. С другой стороны, для закрытых помещений и с большим

числом SDS каустик хорошо подходят методы на основе фотонных карт. Если же постановка задачи не допускает априорных предположений о сцене и требуется построение универсального и надёжного решения, тогда следует использовать передовые методы на основе MCMC. Исходя из нашей практики, мы бы рекомендовали метод MMLT.

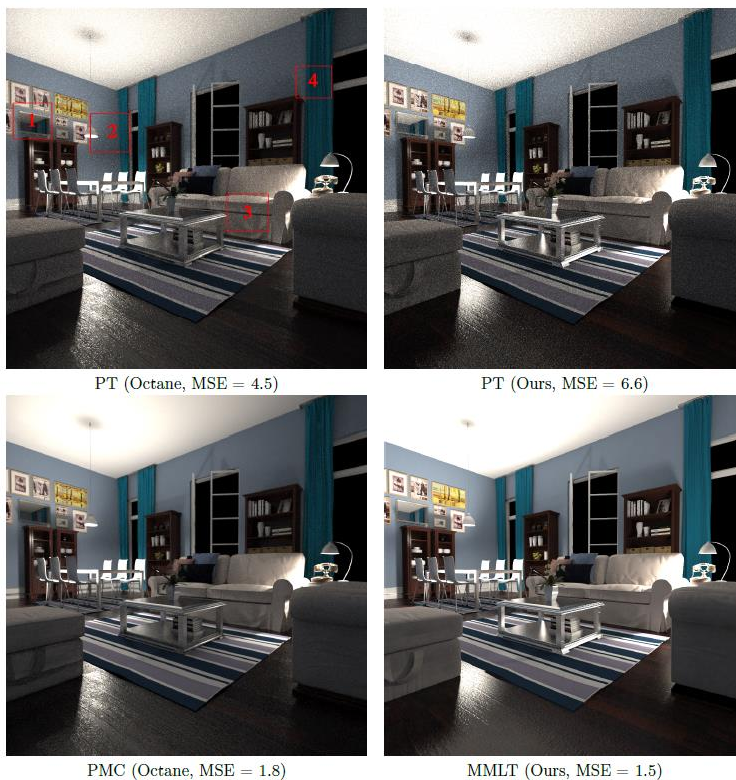


Рис. 14. Сравнение нашей реализации PT и MMLT с Octane на GPU RTX2070 (без аппаратной поддержки трассировки лучей). Время рендеринга – 30 минут. В данной сцене вторичное освещение получено путём отражения от глянцевой поверхности, близкой к идеальному зеркалу (как стол посередине комнаты). На полу присутствует микрорельеф. Сравнение производилось между методом PMC, реализованным в системе Octane, и MMLT, реализованным в системе Hydra Renderer. На рис. 15 показаны увеличенные части изображений, демонстрирующие преимущество метода MMLT. Прямоугольники с номерами на левом верхнем рисунке обозначают участки изображений, которые представлены на рис. 15

Fig. 14. Comparison of our implementation of PT and MMLT with Octane on GPU RTX2070 (no hardware support for ray tracing). Rendering time is 30 minutes. In this scene, the secondary lighting is obtained by reflecting off a glossy surface close to an ideal mirror (like a table in the middle of a room). There is a micro-relief on the floor. The comparison was made between the PMC method implemented in the Octane system and the MMLT method implemented in the Hydra Renderer system. In fig. 15 shows enlarged portions of images showing the advantage of the MMLT method. The numbered rectangles in the upper left figure denote the image areas shown in Fig. 15

Мы полагаем, что для сложного освещения наиболее эффективными могут быть признаны методы на основе современных работ: VCM и аналоги [24, 25, 27], HMC [94, 96, 102], MMLT [78, 84], RJMLT [88-90] и MBE [41]. Например, RJMLT, задуманный его авторами как улучшение MMLT, является более эффективным классом методов, чем MMLT. Однако ключевой вопрос в цене, которую приходится платить за это улучшение. В этом разделе мы хотели бы дать рекомендации для читателя с учётом целевой задачи, нашего опыта, сложности реализации метода (в первую очередь упоминая простые методы и лишь затем более сложные) и объективных ограничений, которые метод имеет. Мы будем считать, что

при переходе к следующему подразделу все феномены освещённости предыдущих подразделов также должны вычисляться методом эффективно, если обратное не оговорено. То есть сложность сцен и освещения растёт кумулятивно.

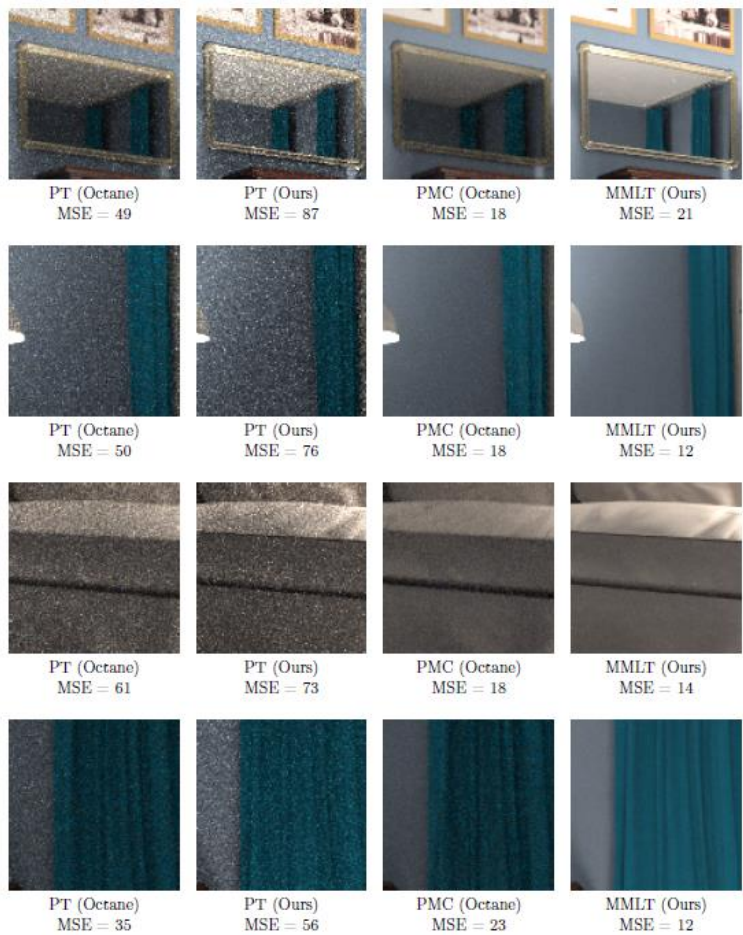


Рис. 15. Сравнение увеличенных фрагментов из рис. 14. Вырезки расположены сверху вниз в порядке их номеров от 1 до 4

Fig. 15. Comparison of enlarged fragments from fig. 14. The clippings are arranged from top to bottom in numerical order from 1 to 4

7.1 Прямое освещение

Для эффективного расчёта прямого освещения необходима реализация как минимум MIS PT с явной (для «маленьких» источников) и неявной (для «больших» источников) стратегией генерации выборок. При большом количестве источников рекомендуется дополнить MIS PT до IBPT [6], добавив ещё одну стратегию генерации выборок – Light Tracing. Ранее мы отмечали, что эта стратегия позволяет эффективно рассчитывать освещение при большом числе источников. К сожалению, у Light Tracing есть определённые сложности с моделированием объектива камеры (тёмные края объектов), что следует иметь в виду при выборе этого метода. Логичным решением в данном случае является замена на фотонные карты (что ухудшит сходимость и восприятие изображения), либо использование методов из работы [9] для эффективного выбора источника в MIS PT.

Необходимо упомянуть, что методы на основе MCMC не могут эффективно рассчитывать первичное освещение, которое, как правило, вычисляется отдельно.

7.2 Каустики первого рода

Каустиками первого рода мы называем каустики, образующиеся на диффузной поверхности и видимые непосредственно камерой. Методы IBPT [6], BPT [3], PCBPT [7] хорошо справляются с подобными каустиками, как и с любым другим типом вторичного освещения, непосредственно видимого камерой (т.е. чтобы проекция на плоскость камеры могла быть использована) и появляющегося на преимущественно диффузных поверхностях. Если применение двунаправленных методов невозможно по каким-либо причинам, нарушающим симметрию, рекомендуется использовать простые методы на основе MCMC поверх MIS PT: Kelemen MLT [75], RELT [80], PMC [105-106]. Их реализация (особенно Kelemen MLT) является достаточно простой и существенно повышает эффективность. PMC и Kelemen MLT хорошо работают на GPU, но для Kelemen MLT всё ещё может быть заметна проблема начального смещения [86].

7.3 Каустики второго рода

Каустиками второго рода мы называем SDS каустики, видимые в зеркале, стекле или под водой. То есть это такие каустики, которые не могут быть спроецированы в камеру. Во-первых, такие каустики не могут быть вычислены при помощи двунаправленных методов IBPT, BPT, PCBPT. За исключением CC-BPT [14], который является существенным усложнением BPT, т.к. требует наличия инструментария дифференциальной геометрии.

Упомянутые простые методы на основе MCMC в принципе с ними справляются, и RELT [80] лучше, чем Kelemen MLT [75]. При этом расчёт всё ещё остаётся несмещённым. Поэтому при длительном расчёте (для получения эталонного изображения) эти методы являются хорошим выбором.

Однако, если необходимо получить приближённое решение быстро, рекомендуется применять методы на основе фотонных карт (в порядке увеличения сложности реализации): SPPM [80], PCLT [75], вариации BDPM без MIS [33, 35-36], с MIS [25, 37], полноценный VCM [24]. Дополнительное применение марковских цепей к фотонным картам способно улучшить их эффективность: MCSPM [39-40] и MBE [41] при длительном расчёте.

7.4 Многократные глянцевые отражения

Ахиллесовой пятой фотонных карт являются многократные глянцевые отражения и сложные составные материалы (кроме MBE [41], который является в значительной степени гибридным методом). Поскольку в промышленных системах компьютерной графики материал может быть образован из нескольких BSDF в виде графа, вычисление BSDF становится дорогой операцией. Особенно при наличии процедурных текстур, использующих различные шумовые функции. В гибридных методах [33, 35-36] и MBE [41] приходится разделять BSDF на ламбертовскую часть, для которой используются фотонные карты, и не-ламбертовскую, для которой они не используются. Это в значительной степени усложняет разработку. Вычислять BSDF на каждый фотон является непозволительно дорогой (для составных материалов) и неэффективной (для глянцевых отражений) операцией, поскольку большая часть фотонов приходит с направлений, не участвующих в формировании изображения.

Упомянутые простые методы на основе MCMC (Kelemen MLT [75] и PMC [105-106]) прекрасно справляются с многократными глянцевыми отражениями и не вносят существенной дополнительной сложности при работе с составными материалами (для RELT [80] это не совсем так). Метод MMLT [78] особенно хорошо подходит для данного феномена освещённости.

7.5 Микрорельеф и разномасштабные сцены

Использование карт нормалей для имитации микрорельефа можно рассматривать как развитие функциональности BSDF, являющееся в некотором смысле следующим пунктом после глянцевых отражений. Во-первых, популярные микрофасетные модели BSDF используют в своей основе вариацию нормали. Во-вторых, при увеличении частоты изменения нормали (например, путём умножения текстурной координаты на некоторое большое число), зеркальная поверхность с микрорельефом начинает выглядеть как глянцевая. Однако, на практике мелкие детали могут встречаться не только из-за имитации микрорельефа, но и в действительности возникать из-за различного масштаба отдельных элементов сцены.

Использование MLT вместе с BPT или MMLT [78] повышает устойчивость расчёта для разномасштабных сцен (см. разд. 6) и сцен с картами нормалей за счёт использования стратегий с промежуточными соединениями. К сожалению, как BPT, так и MMLT являются существенно более сложными и ограниченными (в первую очередь симметрией) методами, чем однонаправленные Kelemen MLT [75] и PMC [105-106].

Фотонные карты и методы на их основе для подобных сцен следует использовать аккуратно, т.к. для них трудно подобрать правильный радиус сбора [21] и получить корректный вид поверхности при наличии микрорельефа.

7.6 Устойчивость, несмещённость, математическая обоснованность

Мы рассмотрели несколько ключевых феноменов освещённости, влияющих на сложность сцен. Однако мы не учли ещё многие феномены, такие как интегрирование объема [112], спектральный расчёт, поляризацию света и многое другое. Такие феномены будут увеличивать размерность пространства интегрирования на каждом переотражении, что в сочетании с ростом желаемой глубины просчёта (максимально допустимое количество переотражений) может доводить размерность пространства интегрирования до 100 и более.

Если важно сделать устойчивое и универсальное решение, необходимо обращать внимание на свойства алгоритма, которые это гарантируют. Для методов на основе ОМС – это, прежде всего, многократная выборка по значимости: BPT/PCBPT [7], VCM [24], CMIS [57]. Для методов на основе MCMC – это теория Монте-Карло методов при помощи марковских цепей и удачное использование особенностей пространства интегрирования: RELT [80], MMLT [78], NHMC [94] и LangevinMC [96].

7.7 Использование GPU

Графические процессоры вводят два существенных ограничения. Во-первых, объём памяти, который метод использует на 1 поток, необходимо сокращать. Поэтому здесь мы отдаём предпочтение методу IBPT [6] против, например, BPT/PCBPT или VCM, для которых необходимо существенно изменять схему вычисления весов многократной выборки по значимости только лишь для того, чтобы сэкономить память [61, 71] (Recursive MIS). Кроме того, с учётом появления аппаратного ускорения трассировки лучей в современных GPU, мы отдаём предпочтение полностью несмещённым методам против фотонных карт: трассировка луча становится существенно дешевле, чем сбор освещённости из фотонной карты, а общее количество Монте-Карло выборок, которые мы можем сделать, возрастает в десятки раз. Для такого количества более медленная сходимость фотонных карт $O(\frac{1}{\sqrt{N}})$ [21] становится заметна.

Во-вторых, для методов на основе MCMC необходимо обращать внимание на то, что длина марковских цепей на GPU будет существенно меньше, чем на CPU. Поэтому проблема начального смещения в MLT на GPU более заметна [86]. Лучше всего здесь подходят методы с короткими цепями PMC [106] и ERPT [105]. Терпимо справляется Kelemen MLT при

помощи параллельного прожига [86] (RELТ должен обладать теми же свойствами), а вот в MMLT из-за более сложного пространства интегрирования эта проблема уже становится критической [87].

7.8 Анимация

В [1] говорится, что почти все MСМС алгоритмы страдают от временной нестабильности (что недопустимо в анимации), и эта проблема в настоящее время ещё не решена. В киноиндустрии действительно в основном используются методы на основе ОМС, а вместо марковских цепей для повышения эффективности предпочитают использовать Path Guiding [42-46]. Безусловно, MСМС методы имеют существенные недостатки при расчёте анимации. Однако, в целом ситуация не столь драматична.

Во-первых, в киноиндустрии редко требуется вычислять сложные феномены освещённости с высокой точностью. Достаточно добиться правдоподобности освещения. Поэтому можно, например, просто обрезать выбросы, присутствующие в Path Guiding, или применить более прогрессивный подход [110]. Иногда, если это требуется, применяют фотонные карты как простой и быстрый способ приближённо посчитать сложные феномены освещённости с SDS путями.

Во-вторых, для повышения регулярности обхода пространства изображения всегда можно просто увеличить вероятность большого шага до 50% и смешивать ОМС (в качестве которого можно брать большие шаги) и MСМС при помощи многократной выборки по значимости, что обычно и делается в Kelemen MLT [75]. С этой точки зрения у Path Guiding нет существенных преимуществ перед PMC или Kelemen MLT, особенно при расчёте на GPU, где обновление динамических структур является нетривиальным действием.

В-третьих, можно применять стратификацию [83].

В-четвертых, временной стабильности в MLT можно добиться, если добавить временные мутации к цепи и рассчитывать несколько кадров одновременно [111].

Наконец, стоит учитывать возможности применения современных методов фильтрации и шумоподавления, особенно с учётом соседних кадров.

8. Заключение

Таким образом, за последние 10 лет появилось много новых расчётных методов, и некоторые из них можно считать чрезвычайно интересными в плане эффективности решения фундаментально трудной задачи. Заявка, сделанная методами на основе НМС [94, 96, 102], выглядит многообещающе. В этом смысле по сравнению с 2010 годом в науке расчёта глобального освещения произошёл настоящий прорыв. Мы полагаем, что за этим прорывом должен последовать соответствующий прорыв и в прикладной науке, поскольку вопрос применения многих из новых методов на практике остаётся открытым в силу большого количества ограничений. Например, система с аппаратом дифференцирования существует (Mitsuba2 [101]), однако НМС в ней не реализован.

В сегодняшней практике имеет смысл отдавать предпочтение тем методам, которые не имеют большого набора ограничений. Например, мы считаем крайне важным возможность эффективной реализации метода на GPU с учётом всё возрастающей поддержки аппаратной трассировки лучей в современных картах. По этой же причине мы полагаем, что полностью несмещённые методы более ценны, т.к. стоимость трассировки лучей по отношению к сбору освещённости из фотонной карты при использовании аппаратного ускорения существенно падает (в 5-10 раз на последних моделях GPU [59-60]). С этой точки зрения более перспективными выглядят методы IBPT [6], Kelemen MLT [75], PMC [105-106], RELT [80-81]. С другой стороны, если важно уметь получать приближённое решение быстро, то можно использовать различные методы на основе фотонных карт: SPPM [19], VCM [24], регуляризации путей [77], PCLT [30].

Список литературы / References

- [1] Sik M. and Krivanek J. Survey of Markov Chain Monte Carlo Methods in Light Transport Simulation. *IEEE Transactions on Visualization and Computer Graphics*, vol. 26, issue 4, 2018, pp. 1821-1840.
- [2] Kajiya J.T. The rendering equation. *ACM SIGGRAPH Computer Graphics*, 1986, vol. 20, no. 4, pp. 143-150.
- [3] Veach E. Robust monte carlo methods for light transport simulation. Ph.D. Dissertation, Stanford University, 1998, 406 p.
- [4] Ershov S.V., Voloboy A.G. Calculation of MIS weights for bidirectional path tracing with photonmaps in presence of direct illumination. *Mathematica Montisnigri*, vol. 48, 2020, pp. 86-102.
- [5] Pharr M., Jakob W., and Humphreys G. *Physically Based Rendering: From Theory to Implementation* (3rd. ed.). Morgan Kaufmann Publishers, 2016, 1266 p.
- [6] Bogolepov D., Ulyanov D. GPU-Optimized Bidirectional Path Tracing. In *Proc. of the 21th International Conference in Central Europe on Computer Graphics, Visualization and Computer Vision*. 2013, pp. 1-15.
- [7] Popov S., Ramamoorthi R., Durand F., Drettakis G. Probabilistic Connections for Bidirectional Path Tracing. *Computer Graphics Forum*, vol. 34, no. 4, 2015, pp. 75–86.
- [8] Walter B., Fernandez S., Arbre A. et al. A Scalable Approach to Illumination. *ACM Transactions on Graphics*, vol. 24, no. 3, 2005, pp. 1098-1107.
- [9] Moreau P., Clarber P. Importance Sampling of Many Lights on the GPU. In *Ray Tracing Gems: High-Quality and Real-Time Rendering with DXR and Other APIs*. Apress, 2019, pp. 255-283.
- [10] Schussler V., Heitz E., Hanika J. and Dachsbacher C. Microfacet-based normal mapping for robust Monte Carlo path tracing. *ACM Transactions on Graphics*, vol. 36, no. 6, 2017, pp. 1-12.
- [11] Baek S.H., Jeon D.S., Tong X., Kim M.H. Simultaneous acquisition of polarimetric SVBRDF and normal. *ACM Transactions on Graphics*, vol. 37, no. 6, 2018, pp. 1-15.
- [12] Bar C., Alterman M., Gkioulekas I., Levin A. A Monte Carlo framework for rendering speckle statistics in scattering media. *ACM Transactions on Graphics*, vol. 38, no. 4, 2019, pp. 1-22.
- [13] Zhdanov A., Zhdanov D., Sokolov V. et al. Problems of the realistic image synthesis in media with a gradient index of refraction. *Proc. of the SPIE 11548, Optical Design and Testing X*. 2020, 115480W.
- [14] Speierer S., Hery C., Villemin R., Jako W. Caustic Connection Strategies for Bidirectional Path Tracing. *Pixar Technical Memo no. 18-01*, 2018, 6 p.
- [15] Wenzel J. Light Transport on Path-Space Manifolds. Ph.D. Dissertation. Cornell University, 2013, 153 p.
- [16] Johannes H., Droske M., and Fascione L. Manifold next event estimation. *Computer Graphics Forum*, vol. 34, no. 4, 2015, pp. 87-97.
- [17] Jensen H.W. Global Illumination using Photon Maps. In *Proc. of Eurographics Workshop on Rendering*, 1996, pp. 21-30.
- [18] Hachisuka T., Ogaki S., and Jensen H.W. Progressive photon mapping. *ACM Transactions on Graphics*, vol. 27, no. 5, 2008, pp. 1-8.
- [19] Hachisuka T., and Jensen H.W. Stochastic progressive photon mapping. *ACM Transactions on Graphics*, vol. 28, no. 5, 2009, pp. 1-8.
- [20] Spencer B. and Jones M.W. Progressive photon relaxation. *ACM Transactions on Graphics*, vol. 32, no. 1, 2013, PP. 1-11.
- [21] Kaplanyan A.S., Dachsbacher C. Adaptive progressive photon mapping. *ACM Transactions on Graphics*, vol. 32, no. 2, 2013, pp. 1-13.
- [22] Suykens F., Willems Y.D. Density Control for Photon Maps. In *Proc. of the Eurographics Workshop on Rendering*, 2000, pp. 23-34.
- [23] Jarosz W., Jensen H.W., and Donner C. Advanced global illumination using photon mapping. In *ACM SIGGRAPH 2008 classes (SIGGRAPH '08)*, 2008, pp. 1-112.
- [24] Georgiev I., Krivanek J., Davidovic T., Slusallek P. Light Transport Simulation with Vertex Connection and Merging. *ACM Transactions on Graphics*, vol. 31, no. 6, 2012, pp. 1-10.
- [25] Hachisuka T., Pantaleoni J., Jensen H.W. A path space extension for robust light transport simulation. *ACM Transactions on Graphics*, vol. 31, no. 6, 2012, pp. 1-10.
- [26] Krivanek J., Georgiev I., Hachisuka T. et al. Unifying Points, Beams, and Paths in Volumetric Light Transport Simulation. *ACM Transactions on Graphics*, vol. 33, no. 4, 2014, pp. 1-13.
- [27] Kaplanyan A.S., Dachsbacher C. Path space regularization for holistic and robust light transport. *Computer Graphics Forum*, vol. 32, no. 2, 2013, pp. 63–72.

- [28] Render Legion. Corona Render System. URL: <https://corona-renderer.com/>, accessed 26.01.2021.
- [29] Изотов И. Уроки 3Ds Max. Caustics CORONA. Видеоурок на Youtube, 2018 г. / Izotov I, Lessons 3Ds Max. Caustics CORONA. Pool water with caustic in the crown. Video tutorial on Youtube. 2018.URL: <https://www.youtube.com/watch?v=Nv1ULR8sMZY>, accessed 26.01.2021 (in Russian).
- [30] Jendersie J., Rohmer K., Brüll F., Grosch T. Pixel cache light tracing. *Proceedings of the Conference on Vision, Modeling and Visualization*. 2017. pp. 137-144.
- [31] Havran V., Herzog R., Seidel H. P. Fast final gathering via reverse photon mapping. *Computer Graphics Forum*, vol. 24, no. 3, 2005, pp. 323-332.
- [32] Zhdanov A., Zhdanov D. The Backward Photon Mapping for the Realistic Image Rendering. In *Proc. 30th Conf. on Computer Graphics and Machine Vision (GraphiCon 2020)*, *CEUR Workshop Proceedings*, 2020, vol. 2744, pp. 1-12.
- [33] Жданов Д.Д., Ершов С.В., Волобой А.Г. Адаптивный выбор глубины трассировки обратного луча в методе двунаправленной стохастической трассировки лучей. *Труды 25-й Международной конференции GraphiCon*, 2015 г., стр. 44–49 / Zhdanov D.D., Ershov S.V., Voloboy A.G. Adaptive estimation of backward ray trace depth for the stochastic bi-directional ray trace method. In *Proc. of the 25th Anniversary International Conference GraphiCon*, 2015, pp. 44-49 (in Russian).
- [34] Жданов А.Д., Жданов Д.Д., Бирюков Е.Д. Реалистичный рендеринг на основе прямых и обратных фотонных карт. Препринт ИПИМ № 77, 2020 г., 22 стр. / Zhdanov A.D., Zhdanov D.D., Birukov E.D. The realistic rendering with forward and backward photon mapping. *KIAM Preprint № 77*, 2020, 22 p. (in Russian).
- [35] Ershov S.V., Zhdanov D.D., Voloboy A.G., Deryabin N.B. The method of quasi-specular elements to reduce stochastic noise in illumination simulation. *Light and Engineering*, no. 5, 2020, pp. 39-47.
- [36] Zhdanov A., Zhdanov D., Galaktionov V. Realistic image synthesis with hybrid photon maps. *Proceedings of the SPIE 11550, Optoelectronic Imaging and Multimedia Technology VII*, 2029, 115500G.
- [37] Vorba J. Bidirectional Photon Mapping. In *Proc. of the 15th Central European Seminar on Computer Graphics*. 2011, pp. 1-8.
- [38] Schutte J. Vertex Connection and Merging. *Rendering Equations Blog*, 2018. URL: <https://schuttejoe.github.io/post/vertexconnectionandmerging/>, accessed 26.01.2021.
- [39] Chen J., Wang B., and Yong J.-H. Improved stochastic progressive photon mapping with metropolis sampling. In *Proc. of the 22nd Eurographics Conference on Rendering*, 2011, pp. 1205–1213.
- [40] Hachisuka T. and Jensen H.W. Robust adaptive photon tracing using photon path visibility. *ACM Transactions on Graphics*, vol. 30, no. 5, 2011, pp. 1-11.
- [41] Sik M., Otsu H., Hachisuka T., and Krivanek J. Robust light transport simulation via metropolised bidirectional estimators. *ACM Transactions on Graphics*, vol. 35, no. 6, 2016, pp. 1-12.
- [42] Vorba J., Hanika J., Herholz S. et al. Path guiding in production. In *ACM SIGGRAPH Courses*, 2019, pp. 1-77.
- [43] Herholz S., Elek O., Vorba J. et al. Product importance sampling for light transport path guiding. *Computer Graphics Forum*, vol. 35, no. 4, 2016, pp. 67-77.
- [44] Guo J., Bauszat P., Bikker J., Eisemann E. Primary sample space path guiding. In *Proc. of the Eurographics Symposium on Rendering: Experimental Ideas & Implementations*, 2018, pp. 73–82.
- [45] Rath A., Grittmann P., Herholz S. et al. Variance-aware path guiding. *ACM Transactions on Graphics*, vol. 39, no. 4, 2020, pp. 1-12.
- [46] Muller T., Gross M., and Novik J. Practical Path Guiding for Efficient Light-Transport Simulation. *Computer Graphics Forum*, vol. 36, no. 4, 2017, pp. 91-100.
- [47] Zwicker M., Jarosz W., Lehtinen J. et al. Recent Advances in Adaptive Sampling and Reconstruction for Monte Carlo Rendering. *Computer Graphics Forum*, vol. 34, no. 2, 2015, pp. 667-681.
- [48] Ershov S.V., Zhdanov D.D., Voloboy A.G., Galaktionov V.A. Two denoising algorithms for bi-directional Monte Carlo ray tracing. *Mathematica Montisnigri*, vol. 43, 2018, pp. 78-100.
- [49] Hachisuka T., Jarosz W., Weistroffer R.P. et al. Multidimensional adaptive sampling and reconstruction for ray tracing. *ACM Transactions on Graphics*, vol. 27, no. 3, 2008, pp. 1-10.
- [50] Kettunen M., Manzi M., Aittala M. et al. Gradient-domain path tracing. *ACM Transactions on Graphics*, vol. 34, no. 4, 2015, pp. 1-13.
- [51] Manzi M., Kettunen M., Aittala M. et al. Gradient-Domain Bidirectional Path Tracing. In *Proc. of the Eurographics Symposium on Rendering - Experimental Ideas & Implementations*, 2015, pp. 65-74.
- [52] Manzi M., Kettunen M., Durand F. et al. Temporal gradient-domain path tracing. *ACM Transactions on Graphics*, vol. 35, no. 6, 2016, pp. 1-9.

- [53] Hua B.S., Gruson A., Nowrouzezahrai D., Hachisuka T. Gradient-domain Photon Density Estimation. *Computer Graphics Forum*, vol. 36, no. 2, 2017, pp. 31-38.
- [54] Lehtinen J., Karras T., Laine S. et al. Gradient-domain metropolis light transport. *ACM Transactions on Graphics*, vol. 32, no. 4, 2013, pp. 1-12.
- [55] Krivanek J., Bouatouch K., Pattanaik S., Žára J. Making Radiance and Irradiance Caching Practical: Adaptive Caching and Neighbor Clamping. In *Proc. of the 17th Eurographics Conference on Rendering Techniques*, 2006, pp. 127–138.
- [56] Schwarzhaupt J., Jensen H.W., Jarosz W. Practical Hessian-based error control for irradiance caching. *ACM Transactions on Graphics*, vol. 31, no. 6, 2012, pp. 1-10.
- [57] West R., Georgiev I., Gruson A., and Hachisuka T. Continuous multiple importance sampling. *ACM Transactions on Graphics*, vol. 39, no. 4, 2020, pp. 1-12.
- [58] Bitterli B., Wyman C., Pharr M. et al. Spatiotemporal reservoir resampling for real-time ray tracing with dynamic direct lighting. *ACM Transactions on Graphics*, vol. 39, no. 4, 2020, pp. 1-17.
- [59] Санжаров В.В., Фролов В.А., Галактионов В.А. Исследование технологии Nvidia RTX. Программирование, том 46, no. 4, 2020 г., стр. 65-72 / Sanzharov V.V., Frolov V.A., and Galaktionov V.A. Survey of Nvidia RTX Technology. *Programming and Computer Software*, vol. 46, no. 4, 2020, pp. 297–304.
- [60] Meister D., Boksansky J., Guthe M., Bittner J. On Ray Reordering Techniques for Faster GPU Ray Tracing. In *Proc. of the Symposium on Interactive 3D Graphics and Games*, 2020, pp. 1-9.
- [61] Van Antwerpen D. Recursive MIS Computation for Streaming BDPT on the GPU. Technical report, Delft University of Technology, 2011, 12 p.
- [62] Nocak J., Havran V., and Daschbacher C. Path regeneration for interactive path tracing. In *Proc. of the 31st Annual Conference of the European Association for Computer Graphics – Short Papers*, 2010, pp.61-64.
- [63] Van Antwerpen D. Unbiased physically based rendering on the GPU. M.S. Thesis, Delft University of Technology, 2011, 180 p.
- [64] Фролов В.А., Галактионов В.А. Регенерация путей с низкими накладными расходами. Программирование, том 42, no. 6, 2016 г., стр. 67-74 / Frolov V.A., Galaktionov V.A. Low overhead path regeneration. *Programming and Computer Software*, vol. 42, no. 6, 2016, pp. 382-387.
- [65] Fabianowski B., Dingliana J. Interactive Global Photon Mapping. *Computer Graphics Forum*, vol. 28, no. 4, 2009, pp. 1151–1159.
- [66] Garanzha K., Pantaleoni J., McAllister D. Simpler and Faster HLBVH with Work Queues. In *Proc. of the ACM SIGGRAPH Symposium on High Performance Graphics*, 2011, pp. 59–64.
- [67] Фролов В.А., Харламов А.А., Галактионов В.А., Востряков К.А. Окто-деревья со множественными ссылками в применении к реализации фотонных карт и кэша освещенности на GPU. Программирование, том 40, no. 4, 2014 г., стр. 64-73 / Frolov V.A., Kharlamov A.A., Galaktionov V.A., Vostriyakov K.A. Multiple reference octrees for a GPU photon mapping and irradiance caching. *Programming and Computer Software*, vol. 40, no. 4, 2014, pp. 208–214.
- [68] Hachisuka T. and Jensen H.W. Parallel progressive photon mapping on GPUs. In *Proc. of the 3rd ACM SIGGRAPH Conference and Exhibition on Computer Graphics and Interactive Techniques in Asia, Sketches*, 2010, pp. 1-1.
- [69] Garanzha K., Pantaleoni J., and McAllister D. Simpler and faster HLBVH with work queues. In *Proc. of the ACM SIGGRAPH Symposium on High Performance Graphics*, 2011, pp. 59–64.
- [70] Karras T. Maximizing Parallelism in the Construction of BVHs, Octrees, and k-d Trees. In *Proc. of the 4th ACM SIGGRAPH / Eurographics conference on High-Performance Graphics*, 2012, pp. 33-37.
- [71] Davidovic T., Krivanek J., Hasan M., and Slusallek P. Progressive Light Transport Simulation on the GPU: Survey and Improvements. *ACM Transactions on Graphics*, vol. 33, no. 3, Article 29, 2014, pp. 1-19.
- [72] Laine S., Karras T., and Aila T. Megakernels considered harmful: wavefront path tracing on GPUs. In *Proc. of the 5th High-Performance Graphics Conference*, 2013, pp. 137-143.
- [73] Veach E., Guibas L.J. Metropolis Light Transport. In *Proc. of the of the 24th Annual Conference on Computer Graphics and Interactive Techniques*, 1997, pp. 65-76.
- [74] Kiivanek J, Georgiev I., Kaplanyan A.S., Canada J. Recent Advances in Light Transport Simulation: Theory and Practice. In *ACM SIGGRAPH Courses*, 2013, pp. 1-5.
- [75] Kelemen C, Szirmay-Kalos L., Antal G., Csonka F. A Simple and Robust Mutation Strategy for the Metropolis Light Transport Algorithm. *Computer Graphics Forum*, vol. 21, no. 3, 2002, pp. 531-540.

- [76] Ashikhmin M., Premoze S., Shirley P., Smits B. A Variance Analysis of the Metropolis Light Transport Algorithm. *Computers and Graphics*, vol. 25, issue 2, 2001, pp. 287-294.
- [77] Kaplanyan A.S., Hanika J., Dachsbacher C. The Natural-constraint Representation of the Path Space for Efficient Light Transport Simulation. *ACM Transactions on Graphics*, vol. 33, no. 4, 2014, pp. 1-13.
- [78] Hachisuka T., Kaplanyan A.S., Dachsbacher C. Multiplexed Metropolis Light Transport. *ACM Transactions on Graphics*, vol. 33, no. 4, 2014, pp. 1-10.
- [79] Swendsen R.H. and Wang J.-S. Replica Monte Carlo simulation of spin-glasses. *Physical Review Letters*, vol. 57, issue 21, 1986, pp. 2607–2609.
- [80] Kitaoka S., Kitamura Y., Kishino F. Replica exchange light transport. *Computer Graphics Forum*, vol. 28, no. 8, 2009, pp. 2330-2342.
- [81] Otsu H., Yue Y., Hou Q. et al. Replica exchange light transport on relaxed distributions. In *ACM SIGGRAPH Posters*, 2013, pp. 1-1.
- [82] Earl D.J. and Deem M.W. Parallel tempering: Theory, applications, and new perspectives. *Physical Chemistry Chemical Physics*, vol. 7, issue 23, 2005, pp. 3910-3916.
- [83] Gruson A., West R., Hachisuka T. Stratified Markov Chain Monte Carlo Light Transport. *Computer Graphics Forum*, vol. 39, no. 2, 2020, pp. 351-362.
- [84] Schmidt M., Lobachev O., Guthe M. Coherent Metropolis Light Transport on the GPU using Speculative Mutations. *Journal of WSCG*, vol. 24, no. 1, 2016, pp. 1-9.
- [85] Antwerpen D. Improving SIMD efficiency for parallel Monte Carlo light transport on the GPU. In *Proc. of the ACM SIGGRAPH Symposium on High Performance Graphics*, 2011, pp. 41-50.
- [86] Фролов В.А., Галактионов В.А. Компактная по памяти реализация алгоритма Metropolis light transport на графических процессорах. *Программирование*, том 43, no. 3, 2017 г., стр. 83-92 / Frolov V.A., Galaktionov V.A. Memory compact Metropolis light transport on GPUs. *Programming and Computer Software*, vol. 43, no. 3, 2017, pp. 196-203.
- [87] Фролов В.А. Исследование алгоритма Multiplexed Metropolis Light Transport на графических процессорах. Препринт ИПМ № 267, 2018 г., 47 стр. / Frolov V.A. Investigation of Multiplexed Metropolis Light Transport on GPUs. *KIAM Preprint № 267*, 2018, 47 p. (in Russian).
- [88] Bitterli B., Jakob W., Novak J., and Jarosz W. Reversible Jump Metropolis Light Transport Using Inverse Mappings. *ACM Transactions on Graphics*, vol. 37, no. 1, 2017, pp. 1-12.
- [89] Otsu H., Kaplanyan A. S., Hanika J. et al. Fusing state spaces for markov chain Monte Carlo rendering. *ACM Transactions on Graphics*, vol. 36, no. 4, 2017, pp. 1-10.
- [90] Pantaleoni J. Charted metropolis light transport. *ACM Transactions on Graphics*, vol. 36, no. 4, 2017, pp. 1-14.
- [91] Brooks S., Gelman A., Jones G., Meng X. MCMC using Hamiltonian dynamics. In *Handbook of Markov Chain Monte Carlo*. Chapman and Hall/CRC, 2011, pp. 113-163.
- [92] Duane S., Kennedy A.D., Pendleton B.J., and Roweth D. Hybrid Monte Carlo. *Physics Letters*, vol. 195, issue 2, 1987, pp. 216–222.
- [93] Michael Betancourt. A Conceptual Introduction to Hamiltonian Monte Carlo. *arXiv:1701.02434*, 2017, 60 p.
- [94] Li T.-M., Lehtinen J., Ramamoorthi R. et al. Anisotropic Gaussian mutations for metropolis light transport through Hessian-Hamiltonian dynamics. *ACM Transactions on Graphics*, vol. 34, no. 6, 2015, pp. 1-13.
- [95] Girolami M., Calderhead B. Riemann manifold Langevin and Hamiltonian Monte Carlo methods. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, vol. 73, issue 2, 2011, pp. 123-214.
- [96] Luan F., Zhao S., Bala K., and Gkioulekas I. Langevin Monte Carlo rendering with gradient-based adaptation. *ACM Transactions on Graphics*, vol. 39, no. 4, 2020, pp. 1-16.
- [97] Roberts G.O., Tweedie R.L. Exponential convergence of Langevin distributions and their discrete approximations. *Bernoulli*, vol. 2, no. 4, 1996, pp.341-363.
- [98] Duong M.H., Lamacz A., Peletier M.A. et al. Quantification of coarsegraining error in Langevin and overdamped Langevin dynamics. *arXiv:1712.09920*, 2017, 48 p.
- [99] Leal R.J. CppADCodeGen. Source Code Generation for Automatic Differentiation using Operator Overloading. *GitHub*. 2011. URL: <https://github.com/joaoleal/CppADCodeGen>, accessed 26.01.2021.
- [100] Vassilev V., Vassilev M., Penev A. et al. Clad – Automatic Differentiation Using Clang and LLVM. *Journal of Physics: Conference Series*, vol. 608, no. 1, 2015, pp. 1-10.
- [101] Nimier-David M., Vicini D., Zeltner T., and Jakob W. Mitsuba 2: a retargetable forward and inverse renderer. *ACM Transactions on Graphics*, vol. 38, no. 6, 2019, pp. 1-17.

- [102] Rioux-Lavoie D., Litalien J., Gruson A., Hachisuka T. Delayed Rejection Metropolis Light Transport. *ACM Transactions on Graphics*, vol. 39, no. 3, 2020, pp. 1-14.
- [103] Atchade Y.F. An adaptive version for the Metropolis adjusted Langevin algorithm with a truncated drift. *Methodology and Computing in Applied Probability*, vol.8, no. 2, 2006, pp. 235-254.
- [104] Roberts G.O., Stramer O. Langevin diffusions and Metropolis-Hastings algorithms. *Methodology and Computing in Applied Probability*, vol. 4, no. 4, 2002, pp. 337-357.
- [105] Cline D., Talbot J., Egbert P. Energy redistribution path tracing. *ACM Transactions on Graphics*, vol. 24, no. 3, 2005, pp. 1186-1195.
- [106] Lai Y., Fan S.H., Chenney S., Dyer C. Photorealistic image rendering with population Monte-Carlo energy redistribution. In *Proc. of the 18th Eurographics Conference on Rendering Techniques*, 2007, pp. 287–295.
- [107] OTOY. Octane Render. URL: <https://home.otoy.com/render/octane-render/>, accessed 26.01.2021.
- [108] Hydra Renderer. Frolov V., Snazharov V., Trofimov M., Galaktionov V. Hydra Renderer. Open Source GPU Rendering System. GitHub, 2019. URL: <https://github.com/Ray-Tracing-Systems/HydraAPI>, accessed 10.02.2021.
- [109] Haindl M., Filip J. Spatially Varying Bidirectional Reflectance Distribution Functions. In *Visual Texture: Accurate Material Appearance Measurement, Representation and Modeling*, Springer, 2013, pp 119-145.
- [110] Zirr T., Hanika J., Dachsbacher C. Re-Weighting Firefly Samples for Improved Finite-Sample Monte Carlo Estimates. *Computer Graphics Forum*, vol. 37, no. 6, pp. 410-421.
- [111] Van de Woestijne J., Frederickx R., Billen N., and Dutre P. Temporal coherence for metropolis light transport. In *Proc. of the Eurographics Symposium on Rendering: Experimental Ideas and Implementations*, 2017, pp. 55–63.
- [112] Novák J. et al. Monte Carlo methods for volumetric light transport simulation // *Computer Graphics Forum*, vol. 37. no. 2, 2018, pp. 551-576

Информация об авторах / Information about authors

Владимир Александрович ФРОЛОВ – кандидат физико-математических наук, старший научный сотрудник ИПМ РАН, научный сотрудник факультета ВМК МГУ. Сфера научных интересов: реалистичная компьютерная графика, моделирование освещённости, разработка программных систем оптического моделирования, параллельные и распределённые вычисления.

Vladimir FROLOV – PhD in computer graphics, senior researcher in the KIAM and researcher in computer graphics of Moscow State University. Research interests: realistic computer graphics, light transport simulation, elaboration of optical simulation software systems, GPU computing.

Алексей Геннадьевич ВОЛОБОЙ, ведущий научный сотрудник, доктор физико-математических наук, доцент. Научные интересы: компьютерная графика, оптика, оптическое моделирование.

Alexey Gennadievich VOLOBOY, Leading Researcher, Doctor of Physical and Mathematical Sciences, Associate Professor. Research interests: computer graphics, optics, optical modeling.

Сергей Валентинович ЕРШОВ – кандидат физико-математических наук, старший научный сотрудник, доцент. Научные интересы: компьютерная графика, моделирование освещённости, научная визуализация.

Sergey Valentinovich ERSHOV, PhD, Senior Researcher, Associate Professor. Research interests: computer graphics, light transport simulation, scientific visualization.

Владимир Александрович ГАЛАКТИОНОВ, главный научный сотрудник, доктор физико-математических наук, профессор. Научные интересы: компьютерная графика, вычислительная оптика, компьютерная лингвистика, научная визуализация.

Vladimir Alexandrovich GALAKTIONOV, Chief Researcher, Doctor of Physical and Mathematical Sciences, Professor. Research interests: computer graphics, computational optics, computer linguistics, scientific visualization.



Разработка Web-приложений с учетом элементов качества данных – DQAWA

¹ Ц. Гуерра-Гарсия, ORCID: 0000-0002-9290-6170 <cesar.guerra@uaslp.mx>

¹ Г. Перес Гонсалес, ORCID: 0000-0003-3331-2230 <hectorgerardo@uaslp.mx>

¹ М. Рамирез-Торрес, ORCID: 0000-0002-7457-7318 <tulio.torres@uaslp.mx>

² Р. Хуарез-Рамирез, ORCID: 0000-0002-5825-2433 <reyesjua@uabc.edu.mx>

¹ Автономный университет Сан-Луис-Потоси,

Мексика, 78000, Сан-Луис-Потоси, Альваро Обрегон, 64

² Автономный университет Нижней Калифорнии (UABC),

Мексика, 21100, Нижняя Калифорния, Энсенада

Аннотация. В настоящее время обеспечение достаточного уровня качества данных является первостепенной задачей для поддержки успешной деятельности любой организации или предприятия. Поэтому внедрение практик управления качеством данных необходимо в ситуациях, требующих обеспечения определенного уровня качества данных для конкретной функциональности приложения или сервиса. Подобные специализированные практики управления качеством данных должны внедряться в процесс разработки программного обеспечения как можно раньше. На основе анализа существующих предложений в данной области было установлено, что все еще заметна нехватка методологических и технологических решений при разработке приложений, ориентированных на качество данных. Основываясь на достижениях области Web-проектирования под управлением моделей, данная работа представляет частичный результат наших исследований в этой области: предлагается использование метамодели и профиля UML как артефактов качества данных на стадии проектирования Web-приложений. Основной целью работы является предоставление инструментов для предотвращения проблем с качеством данных при проектировании Web-приложений.

Ключевые слова: качество данных; Web-приложения; проектирование под управлением моделей

Для цитирования: Гуерра-Гарсия Ц., Перес Гонсалес Г., Рамирез-Торрес М., Хуарез-Рамирез Р. Разработка Web-приложений с учетом элементов качества данных – DQAWA. Труды ИСП РАН, том 33, вып. 2, 2021 г., стр. 49-64. DOI: 10.15514/ISPRAS-2021-33(2)-2

Developing Web Applications with Awareness of Data Quality Elements – DQAWA

- ¹ C. Guerra-García, ORCID: 0000-0002-9290-6170 <cesar.guerra@uaslp.mx>
¹ H. Pérez González, ORCID: 0000-0003-3331-2230 <hectorgerardo@uaslp.mx>
¹ M. Ramírez-Torres, ORCID: 0000-0002-7457-7318 <tulio.torres@uaslp.mx>
² R. Juárez-Ramírez, ORCID: 0000-0002-5825-2433 <reyesjua@uabc.edu.mx>

¹ Universidad Autónoma de San Luis Potosí,
Álvaro Obregón #64, Col. Centro, San Luis Potosí, C.P. 78000, México

² Autonomous University of Baja California (UABC),
Ensenada, Mexico, 21100

Abstract. An acceptable level of quality in data is nowadays a paramount for any kind of organization or enterprise that wishes its business processes to prosper. Thus, introducing activities focused in the data quality management is a crucial requirement for the analysts if the level of quality of data for the functionality or service at hand is to be ensured. Such specialized data quality management activities should be presented as early as possible during the software development process. So far and having done a search for proposals in this field, there is still a lack of either methodological or technological proposals with which a developer could be able to design data quality aware applications in the specific field of Web application development. Considering the benefits offered in the field of Model Driven Web Engineering, this work presents a partial outcome of our research in this novel field: a metamodel and a UML profile, both able to be used as data quality artefacts during the design stage of Web applications. The main objective is to provide the designer with the tools needed to design Web applications, in order to prevent data quality issues.

Keywords: data quality; Web applications; Model Driven Web Engineering

For citation: Guerra-García C., Pérez González H., Ramírez-Torres M., Juárez-Ramírez R. Developing Web Applications with Awareness of Data Quality Elements – DQAWA. *Trudy ISP RAN/Proc. ISP RAS*, vol. 33, issue 2, 2021, pp. 49-64 (in Russian). DOI: 10.15514/ISPRAS-2021-33(2)-2.

1. Введение

С точки зрения пользователя можно было бы утверждать, что Информационная Система не может быть качественнее, чем используемые ею данные [1]. Общая проблема поиска информации в Сети не изменилась – требуется найти документы и данные, релевантные запросу, то есть документы, соответствующие информационной потребности пользователя [2]. Учитывая огромное количество данных в Интернете, для организаций как никогда важно иметь возможность собирать, хранить, обмениваться и отображать данные с адекватным уровнем качества через Web-приложения [3-5]. Одной из основных трудностей при извлечении данных с Web-страниц является автоматическое обнаружение интересующих объектов и их свойств и сохранение их в базе данных в едином формате [6]. В некоторых организациях использовались отдельные базы данных для каждого подразделения, несмотря на то что они содержали информацию об одних и тех же объектах, но различались схемами данных. Кроме того, данные часто содержат типографские ошибки, неточности в их описаниях и устаревшие значения [7]. Описание требований должно включать структуру, содержание и значения результатов, ожидаемых пользователем, объектом или системой в определенных условиях и при определенных исходных данных [8].

Наиболее распространенные требования к качеству данных (Data Quality, DQ) включают: полноту, корректность, полезность и доступность для систем [4]. Эти требования к DQ должны быть введены как можно раньше в процесс разработки Web-приложений для обеспечения работы приложения в соответствии с ними. Для этого важно изучить, как Web-приложения обеспечивают сохранность данных и как они могут адаптировать свое содержание под различные конфигурации визуализации [9].

В данной работе проблема рассматривается как с методологической, так и с технологической точки зрения в рамках Web-проектирования под управлением моделей (Model-Driven Web Engineering, MDWE) [10]. Большая часть исследовательских работ в этой области была ориентирована преимущественно на этапы анализа и проектирования [11]. Кроме того, почти все существующие работы предлагают конкретные процессы для систематической и полуавтоматической разработки приложений. Но после анализа литературы [12-15] мы обнаружили, что ни одна из этих работ в области модельно-ориентированного проектирования не затрагивает вопросы, связанные с управлением DQ при проектировании и разработке Web-приложений.

Чтобы устранить выявленную проблему, в рамках данной работы предлагается *расширенная метамодель и соответствующий ей профиль, заданный с помощью языка UML, которые позволяют на стадии проектирования процесса Web-разработки объединить определенные функции и артефакты, совместно удовлетворяющие требованиям управления качеством данных*. Кроме того, предлагается набор правил отображения, способствующих улучшению качества проектирования Web-приложений, ориентированных на качество данных.

2. Родственные области

2.1. Разработка требований к управлению качеством данных

Читатели могут обратиться к нашей предыдущей работе [14], посвященной *требованиям к DQ и их анализу* для ознакомления с необходимыми элементами управления качеством данных на стадии проектирования. Введение элементов управления DQ позволяет учитывать требования к DQ и способствует их реализации на стадии проектирования.

В данной работе под DQ понимается *пригодность данных к использованию* [16-18]. С этой точки зрения, когда от пользователя требуется оценить уровень качества какого-либо элемента данных для решения конкретной задачи в конкретном контексте, необходимо построить модель DQ и определить несколько критериев DQ.

В данном исследовании ставится следующий вопрос: как должна разрабатываться Информационная Система (ИС), чтобы обеспечить адекватный уровень качества хранящихся в ней данных? Для этого пользователям должны быть предоставлены модели DQ с подходящими критериями DQ, которые подходят под их собственные требования к DQ. В данной работе используется стандарт ИСО/МЭК 25012 [19]. В нем определены характеристики DQ, значимых при оценке уровня DQ в информационной системе. Выделяются пятнадцать характеристик DQ (*точность, полнота, непротиворечивость, достоверность, актуальность, доступность и т.д.*), которые делятся на две группы: *неотъемлемые и зависящие от системы*.

Одной из целей данной работы является проектирование конкретных функций веб-приложения, позволяющих оценить характеристики, представленные в стандарте. В настоящее время не существует исследований в области перехода от анализа к проектированию, посвященных конкретным аспектам или характеристикам управления DQ на стадии проектирования Web-приложения. Единственной связанной с данной проблемой работой является [14], в которой для *спецификации требований к DQ в Web-разработке* были предложены расширенная метамодель и UML-профиль (названный DQ_WebRE).

2.2. Веб-проектирование

Наиболее важными методологиями, используемыми на этапе моделирования при проектировании Web-приложений, являются OOWS [20], UWE [21], WebML [22], W2000 [23], WSDM [24], SOD-M [25], WebSA [26], FACPL [27] и HiLA [28].

Все эти методологии в основном посвящены выявлению и определению функциональных аспектов, которые в основном связаны с навигационными, концептуальными и

презентационными моделями. Однако ни одна из данных методологий не применяется для обеспечения DQ, за исключением подхода, предложенного Сери (Stefano Ceri) и др. (WebML) [22], где упоминаются некоторые конкретные цели и соображения, которые следует учитывать при разработке Web-приложений.

3. Расширение метамодели UWE для поддержки проектирования веб-приложений, ориентированных на качество данных

Проведя сравнительное исследование различных подходов, мы решили взять в качестве основы подход UWE. Подход UWE был впервые представлен Кох (Nora Koch) и Краусом (Andreas Kraus) в [21], а затем обновлен в [29-31]. Подход также был улучшен благодаря добавлению новых моделей, учитывающих соображения безопасности [32, 33]. В подходе UWE предлагается метамодель для описания концептов и отношений в Web-проектировании. Она определяет конкретные представления, графически изображаемые диаграммами классов UML, которые сгруппированы в четыре разных пакета: *Navigation*, *Presentation*, *Process* и *Content*. Мы решили использовать подход UWE в качестве основы, поскольку он предоставляет возможность расширения, в данном случае с помощью конкретных элементов DQ.

3.1 Спецификация стереотипов для элементов качества данных в проектировании

В этом подразделе описывается предлагаемый подход, который является расширением метамодели UWE с добавлением концепций, необходимых для управления DQ на стадии проектирования. Во-первых, мы проанализируем основные подходы, представленные в [34] и связанные с моделированием DQ, а затем рассмотрим следующие ключевые понятия пакета *Process*: *DQProcessClass*, *DQProcessProperty* и *DQLink*, вместе с тремя новыми элементами в пакете *Content* – *DQDim*, *DQDimProperty* и *DQ_Constraint*, а также один элемент пакета *Presentation* – *UIE_DQVerifier*. Ниже представлено подробное описание данных метаклассов.

- ***DQProcessClass***. Представляет собой процесс, ответственный за управление метаданными DQ и связанный с данными, которые обрабатываются в каждом *ProcessClass*. Другими словами, *DQProcessClass* может использовать класс *DQDim* и определяет метаданные DQ каждого критерия DQ.
- ***DQProcessProperty***. Указывается в *DQProcessClass* и используется для определения конкретных критериев, связанных с DQ. Каждый элемент данных, управляемый соответствующим *ProcessClass*, должен охватывать какой-либо критерий.
- ***DQLink***. Специальная связь, соединяющая элементы *ProcessClass* и *DQProcessClass*.
- ***DQDim***. Этот метакласс представляет собой элемент контента; он соответствует всем критериям DQ: он будет в основном отвечать за сохранение различных метаданных DQ, называемых *DQDimProperty*.
- ***DQDimProperty***. Принадлежит критерию DQ и используется для определения всех метаданных DQ, связанных с каждым критерием DQ (*DQDim*).
- ***DQ_Constraint***. Представляет собой элемент контента. Соответствует всем *Constraints*, которые должны быть определены и связаны с различными критериями DQ вместе с их соответствующими границами (*upper_bound* и *lower_bound*).
- ***UIE_DQVerifier***. Этот элемент представляет класс, ответственный за проверку соответствия данных, управляемых каждым элементом пользовательского интерфейса различным характеристикам типовых критериев DQ (*DQDim*).

В табл.1 представлен предлагаемый набор новых стереотипных элементов DQ в соответствии со спецификацией UML [35].

Табл. 1. Спецификация стереотипов для элементов DQ в проектировании

Table 1. Specification of stereotypes for DQ design elements

Имя	UIE_DQVerifier
Базовый класс	Класс
Описание	Класс, ответственный за определение и проверку соответствия данных, управляемых каждым элементом пользовательского интерфейса, указанным критериям DQ (DQDim).
Ограничения	Должен быть связан хотя бы с одним элементом <i>UIElement</i>
Используется в модели	Модель <i>Presentation</i>
Имя	DQProcessClass
Базовый класс	navigationNode::Class
Описание	Процесс, ответственный за контроль метаданных DQ, связанных с данными, обрабатываемыми в каждом элементе <i>ProcessClass</i>
Ограничения	Должен быть связан хотя бы с одним элементом типа <i>ProcessClass</i>
Используется в модели	Модели <i>Navigation / Process</i>
Имя	DQProcessProperty
Базовый класс	navigationProperty::Property
Описание	<i>DQProcessProperty</i> принадлежит классу <i>DQProcessClass</i> и используется для указания характеристик DQ, которым должен удовлетворять каждый элемент данных, управляемый в соответствующем классе <i>ProcessClass</i> .
Используется в модели	Модели <i>Navigation / Process</i>
Имя	DQLink
Базовый класс	link::Association
Описание	Специальная связь, соединяющая два типа конкретных элементов <i>ProcessClass</i> и <i>DQProcessClass</i>
Используется в модели	Модель <i>Navigation</i>
Имя	DQDim
Базовый класс	Класс
Описание	Соответствует всем критериям DQ: он будет в основном отвечать за сохранение различных метаданных DQ, называемых <i>DQDimProperty</i> .
Используется в модели	Модель <i>Content</i>
Имя	DQDimProperty
Базовый класс	Property
Описание	<i>DQDimProperty</i> принадлежит критерию DQ и используется для определения метаданных DQ
Используется в модели	Модель <i>Content</i> .
Имя	DQ_Constraint
Базовый класс	Класс
Описание	Соответствует набору всех <i>Constraints</i> , которые должны быть связаны с различными критериями DQ.
Используется в модели	Модель <i>Content</i> .

3.2 Правила отображения

Поскольку при получении диаграмм для проектирования на основе диаграмм требований важно поддерживать высокий уровень трассируемости, был создан набор правил отображения, приведенных в табл. 2, обеспечивающий корректный процесс создания диаграмм. Предложенные правила отображения должны быть учтены при моделировании соответствующей диаграммы (*Content, Navigation, Presentation*) на основе результатов, полученных на этапе анализа требований.

Табл. 2. Правила отображения *DQ_WebRE* в *DQAWA*

Table 2. Mapping rules *DQ_WebRE* to *DQAWA*

DQ_WebRE	DQAWA	Описание	Используется в модели
Add_DQ_Metadata	DQProcessClass	Каждый элемент <i>Add_DQ_Metadata</i> будет преобразован в <i>DQProcessClass</i> . Отношение элемента типа <i>UserTransaction</i> к другому элементу типа <i>Add_DQ_Metadata</i> будет изменено на конкретное отношение типа <i>DQLink</i> между элементами типа “ <i>ProcessClass</i> ” и <i>DQProcessClass</i> . Имена элементов типа <i>DQ_Metadata</i> и <i>DQ_Validator</i> , связанных с <i>Add_DQ_Metadata</i> , будут преобразованы в атрибуты типа <i>DQDim</i> в элементе <i>DQProcessClass</i> .	<i>Navigation</i>
DQ_Metadata	DQDim	Каждый элемент <i>DQ_Metadata</i> будет преобразован в <i>DQDim</i> , а его соответствующие определенные атрибуты (свойства <i>DQD Content im</i>) будут скопированы. Каждое отношение между элементами типа <i>DQ_Metadata</i> и <i>Contents</i> будет преобразовано в отношения между элементом <i>DQDim</i> и элементом <i>ContentClass</i> .	<i>Content</i>
DQ_Validator	UIE_DQVerifier	Каждый элемент <i>DQ_Validator</i> будет преобразован в элемент <i>UIE_DQVerifier</i> . Отношение между элементами типа <i>DQValidator</i> и <i>WebIU</i> будет изменено на отношение между <i>UIE_DQVerifier</i> и <i>PresentationPage</i> .	<i>Presentation</i>
DQConstraint	DQ_Constraint	Каждый <i>DQConstraint</i> преобразуется в другой <i>DQ_Constraint</i> . Этот элемент будет отвечать за хранение конкретных данных, с которыми будут связаны ограничения, а также за переменные, используемые для определения максимальной и минимальной границ (<i>upper_bound</i> и <i>lower_bound</i>). Отношение между <i>DQ_Validator</i> и <i>DQConstraint</i> изменяется на отношение между <i>DQDim</i> и <i>DQ_Constrain</i> ”.	<i>Content</i>

Для поддержки нашего подхода мы предлагаем новый UML-профиль для проектирования Web-приложений с ориентацией на DQ (*DQAWA*). Профиль был реализован с помощью инструмента *Enterprise Architect*, который поддерживает определение профилей на основе UML. После того, как *профиль DQAWA* был реализован, мы получили подходящую рабочую среду, в которой можно моделировать диаграммы с новыми элементами (рис. 1).

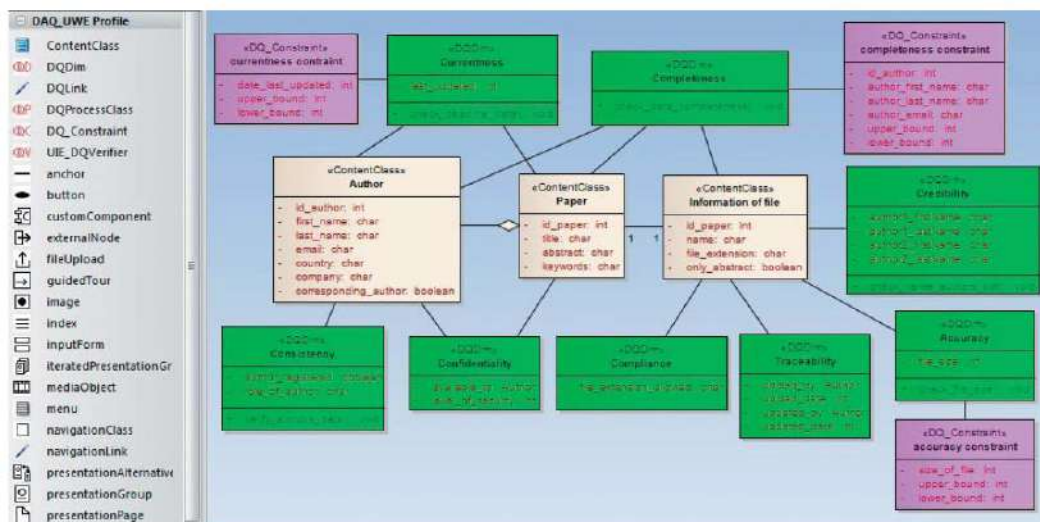


Рис. 1. Модель контента с элементами управления качеством данных
Fig. 1. Content model with elements for DQ management

4. Практическое применение

Чтобы проиллюстрировать использование предложенного профиля, в этом разделе рассматривается пример системы управления конференциями Easy Chair [36]. Это приложение позволяет подавать документы, управлять и контролировать членов Программного комитета (ПК) и назначать статьи рецензентам, а также имеет некоторые другие функции. Хотя данная работа посвящена стадии проектирования, для создания более полного понимания всего процесса ниже приведены результаты, полученные на предыдущих стадиях. Поэтому мы начнем с моделирования конкретной диаграммы использования (рис. 2) для каждого субъекта, использующего веб-приложение, а также смоделируем соответствующие конкретные требования к качеству данных с помощью элементов, которые определены в *профиле DQ_WebRE*, предложенном в [14] (подраздел 4.1). В подразделе 4.2 мы создаем соответствующие диаграммы для проектирования, используя элементы, определенные в *профиле DQAWA*, и соответствующие правила отображения.

4.1 Сбор и моделирование требований к качеству данных с помощью DQ_WebRE

На диаграмме использования представлены функциональные возможности каждого типа пользователя при использовании профиля *DQ_WebRE* (рис. 2). Таким образом, сначала необходимо смоделировать основные сценарии использования Автором, а затем с помощью диаграмм деятельности дать подробное описание каждого сценария использования. Для простоты описания было решено сосредоточиться исключительно на сценарии использования *Новое представление (New Submission)*, для которого *Автору* понадобится использовать приложение Easy Chair.

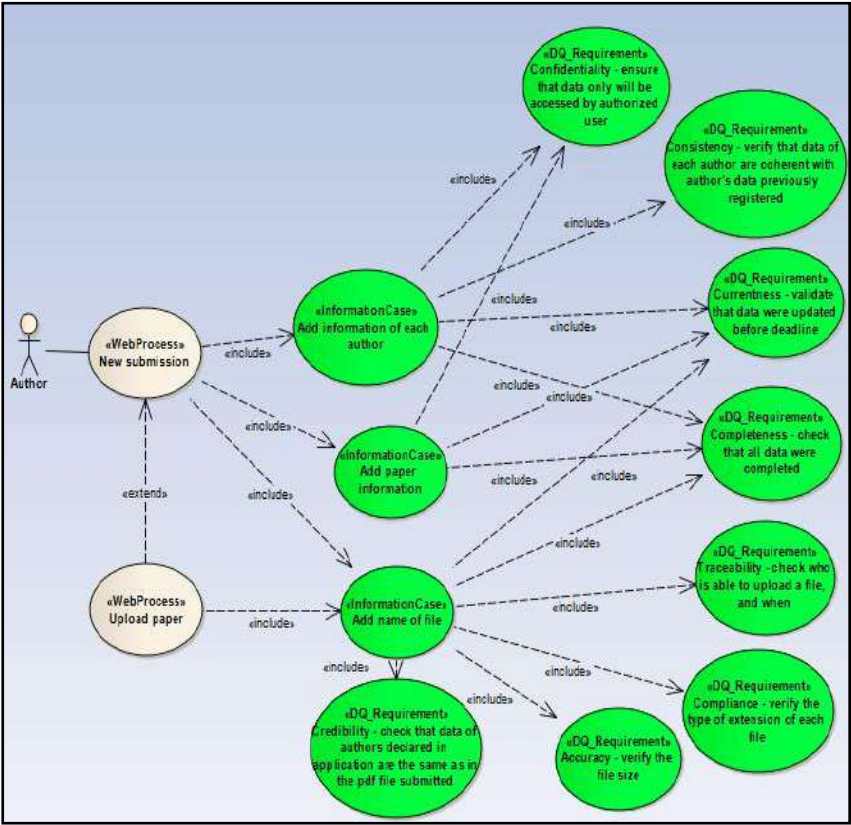


Рис. 2. Диаграмма использования с указанием требований к DQ
Fig. 2. Use case diagram specifying DQ requirements

Основываясь на работе, представленной ранее в [37], мы сформулировали требования к DQ, которые, по всей видимости, связаны с каждой из функциональных возможностей, реализованных в приложении. Учитывая, что сценарий использования *Новое представление* может быть отнесен к категории функциональности типа *Управление содержанием (Content Management)*, можно видеть, что эта функциональность может быть связана со следующими характеристиками DQ: *точность, полнота, согласованность, достоверность и т.д.* Принимая во внимание эти соображения, смоделируем диаграмму использования (рис. 2) и обратим внимание на спецификацию конкретных сценариев использования *Добавить информацию о каждом авторе (Add information of each author)*, *Добавить информацию о статье (Add paper information)* и *Добавить имя файла (Add name of file)*, стереотипизированных как *InformationCase*, которые будут отвечать за управление всеми данными в сценарии использования *Новое представление* (см. комментарий на рис. 3).

метаданные, связанные с актуальностью (*add metadata related to Currentness*), которые стереотипизируются как *Add_DQ_Metadata*, будут отвечать за сбор метаданных, связанных с конкретной характеристикой *согласованности* (*author_registered*, *role_of_author*) и *актуальности* (*last_updated*). Эти метаданные будут храниться в экземпляре класса *DQ_Metadata* и использоваться для удовлетворения требований к DQ по критериям *согласованности* и *актуальности* соответственно.

Кроме того, активность *добавить метаданные, связанные с конфиденциальностью* (*add metadata related to Confidentiality*) будет собирать соответствующие метаданные, относящиеся к *available_to* и *level_of_security*, которые хранятся в конкретном классе *DQ_Metadata*. Заметим, что все классы *DQ_Metadata* связаны с данными, управляемыми предыдущими действиями *добавить данные авторов* и *добавить данные статьи* (стереотипизированные как *UserTransaction*). Активность *Проверить данные файла* (*Check data of file*) аналогично отвечает за сбор метаданных, связанных с выполнением конкретных требований к DQ по критерию *трассируемости* (*upload_by*, *upload_date* и т.д.), *соответствия* (*fileextension_allowed*), *точности* (*file_size*) и *достоверности* (*author1_firstName*, *author1_lastName* и т.д.).

Наконец, действия *Проверить полноту данных* (*Verify Completeness of data*), *Проверить актуальность введенных данных* (*Validate the Currentness of entered data*), *Проверить размер загруженного файла* (*Verify the size of uploaded file*) и *Проверить имена авторов и организации в поданном PDF-файле* (*Check authors names and affiliations in the submitted pdf file*) будут отвечать за добавление конкретных операций *check_data_completeness()*, *check_deadline_date()*, *check_file_size()* и *check_name_of_authors_in_pdf_file()*, чтобы проверить *полноту*, *актуальность*, *точность* и *достоверность* управляемых данных в элементе *Webpage New_Submission*.

Данные, управляемые в элементе *Ограничение по полноте* (*Completeness constraint*), стереотипизированном как *DQConstraint*, будут использоваться для указания данных, необходимых для завершения представления, и спецификации границ критерия полноты, которые варьируются от *lower_bound* до *upper_bound*. Данные, хранящиеся в элементе *Ограничение по актуальности* (*Currentness constraint*), будут служить для определения допустимых границ и подтверждения актуальности даты обновленных данных *Автора* (*date_last_updated*). Действие *Проверить согласованность данных каждого автора* (*Currentness constraint*) будет отвечать за проверку этих элементов данных, введенных с помощью *Webpage New_Submission*. Наконец, данные, хранящиеся в элементе *Ограничение по точности* (*Accuracy constraint*), будут использоваться для проверки того, что поданный файл не является пустым, а также для ограничения максимального допустимого размера (*upper_bound* и *lower_bound*).

4.2 Моделирование диаграмм для проектирования с помощью DQAWA

Следует отметить, что авторы метамодели UWE предложили набор правил преобразования [38, 39], с помощью которых можно получить различные элементы, используемые в моделях проектирования из элементов, которые определены в подходе *WebRE* для уточнения требований [40]. Если взять за основу стереотипные элементы диаграммы деятельности (рис. 3) и применить наши *правила отображения*, становится возможным создать соответствующие модели для проектирования. Диаграмма, содержащая *модель Content*, показана на рис. 1 вместе с классами, определенными для хранения основной информации, касающейся *Автора*, *Статьи* и *Информации* файла (стереотипизированной как *ContentClass*).

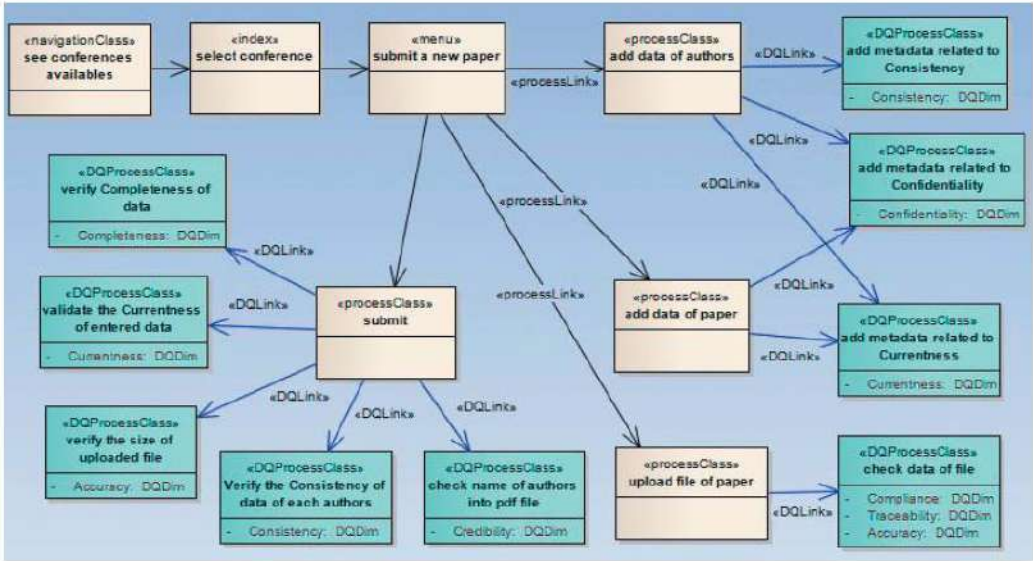


Рис. 4. Навигационная модель с управлением DQ
Fig. 4. Navigation model with DQ management

На этой диаграмме также показаны классы, связанные с определением каждого критерия качества данных (стереотипизированного как *DQDim*): полноты, соответствия, конфиденциальности и т.д.; а также их атрибуты (типа *DQDimProperty*) и соответствующие им операции. Также можно смоделировать классы *ограничение по актуальности*, *ограничение по точности* и *ограничение по полноте* (*DQ_Constraint*), которые будут хранить данные, используемые для определения ограничений и границ различных критериев качества данных.

Также возможно смоделировать гипертекстовую структуру с помощью *модели Navigation* (рис. 4). Она показывает узлы и связи (узлы именуются обозначенные как *узлами процессов* (*Process nodes*) и интегрированные в навигационный поток с помощью связей процессов). Эта диаграмма демонстрирует навигационный класс (*смотреть доступные конференции – see conferences available*), индексный класс (*выбрать конференцию – select conference*) и класс меню *представить новую статью – submit a new pape*). Все данные, вводимые в веб-приложение, моделируются с помощью классов *добавить данные авторов*, *добавить данные статьи* и *загрузить файл статьи* типа *ProcessClass*.

Управление DQ моделируется с помощью элементов типа *DQProcessClass* (например, *добавить метаданные, относящиеся к согласованности – add metadata related to Consistency*, *добавить метаданные, связанные с конфиденциальностью – add metadata related to confidentiality* и т.д.), относящихся к разным классам типа *ProcessClass* с использованием конкретного отношения *DQLink*. Элементы *DQProcessClass* отвечают за указание характеристик DQ для данных в каждом *узле процесса*.

На рис. 5 показана диаграмма *модели Presentation* для нашего сквозного примера. Страница выдачи *представить новую статью* содержит три группы выдачи: *добавить данные авторов*, *добавить данные статьи* и *загрузить файл статьи*. Для каждой группы в модели выдачи показаны различные элементы. Например, в группе выдачи *добавить данные авторов* моделируются следующие элементы: *first_name*, *last_name*, *email*, *country*, *company* (тип *textInput*) и *corresponding_author* (типа *selection*), а группа *добавить данные статьи* содержит поля *title*, *abstract* и *keywords* типа *textInput*. В этой модели выдачи мы видим, что класс *Полнота* (*Completeness*), стереотипизированный как *UIE_DQVerifier*, относится к

странице выдачи *представить новую статью*. Это означает, что этот класс будет отвечать за проверку полноты данных (с помощью операции *check_data_completeness*), введенных в каждый элемент, содержащийся в этой группе.

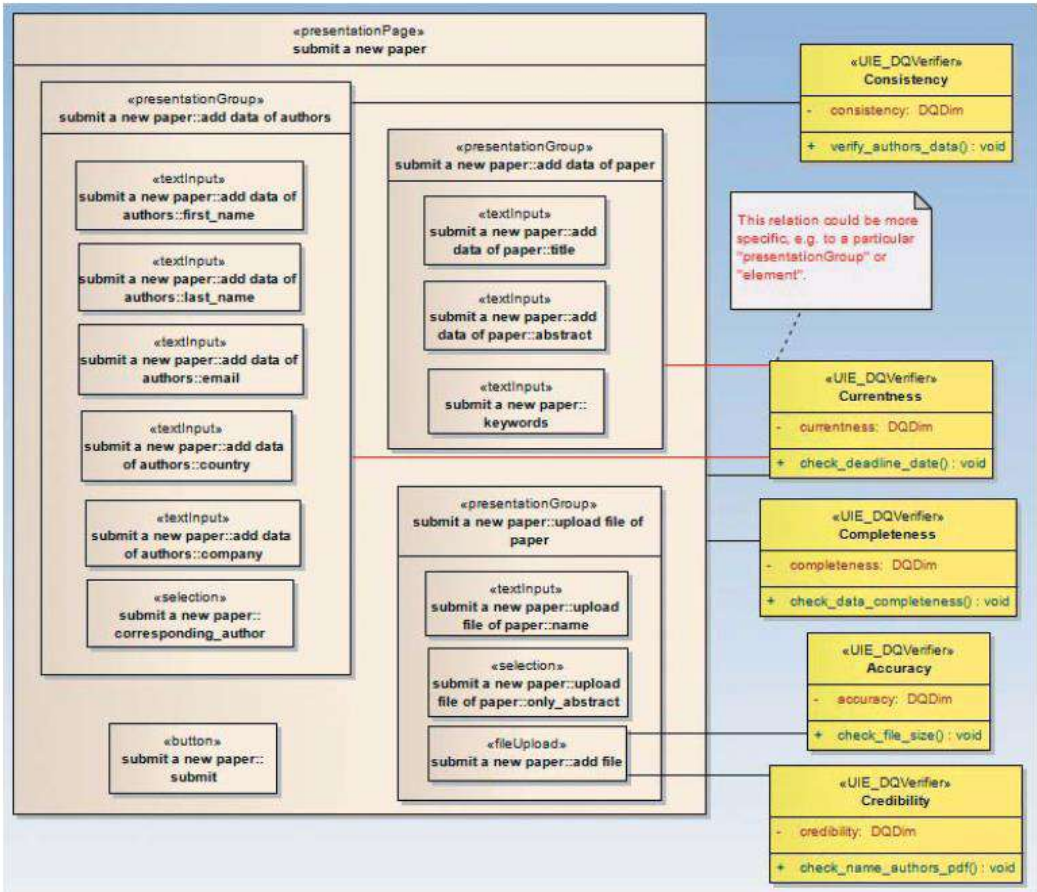


Рис. 5. Модель Presentation с элементами управления качеством данных
Fig. 5. Presentation model including elements for DQ management

Класс *Актуальность* (Currentness), стереотипизированный как *UIE_DQVerifier*, аналогичным образом будет отвечать за выполнение конкретной операции *check_deadline_date* для проверки актуальности каждого из элементов на странице выдачи *представить новую статью*. Класс *Согласованность* (Consistency) будет отвечать за проверку того, что все вводимые данные, принадлежащие каждому автору, поддерживаются согласованным образом с помощью операции *verify_authors_data*. Класс *Точность* (Accuracy) будет отвечать за выполнение операции *check_file_size* для проверки того, что размер файла находится в разрешенных границах. Наконец, класс *Достоверность* (Credibility) будет отвечать за проверку соответствия имени автора в файле формата pdf с именем, введенным в приложении.

5. Заключение

За последние два десятилетия Web-приложения стали первостепенными информационными ресурсами для большинства организаций. Поэтому важно, чтобы приложения могли обеспечивать адекватный уровень предоставляемых данных. Управление элементами

качества данных на этапе проектирования поможет избежать некоторых возможных проблем с данными, а также позволит внедрить соответствующие механизмы для требуемого уровня управления данными. Для решения этой проблемы на основе подхода MDWE представлены расширенная метамодель и соответствующий ей UML-профиль (DQAWA) для управления элементами качества данных при проектировании Web-приложений. В настоящее время данный подход реализуется в различных реальных проектах для проверки предложенных положений и определения его достоинств и недостатков, а также для получения количественных данных в управлении проектами (например, время, необходимое для внедрения требований к качеству данных). Результаты практического использования помогут оценить предлагаемый подход, как упоминалось в [41].

Список литературы / References

- [1] Strong D.M., Y.W. Lee, and R.Y. Wang. Data Quality in Context. Communications of the ACM, vol. 40, no. 5, 1997, pp. 103-110.
- [2] Некрестьянов И.С., Пантелеева Н.В. Системы текстового поиска для Веб. Программирование, том 28, no. 4, 2002 г., стр. 33-67 / Nekrestyanov I.S. and N.V. Panteleeva. Text Retrieval Systems for the Web. Programming and Computer Software, vil. 28, no. 4, 2002, pp. 207 - 225.
- [3] Akoka J., Berti-Equille L., Boucelma O et al. A Framework for Quality Evaluation in Data Integration Systems. In Proc. of the I9th ernational Conference on Enterprise Information Systems, ICEIS. 2007, pp.170-175.
- [4] Bertino E., Maurino A., and Scannapieco M. Guest Editors' Introduction: Data Quality in the Internet Era. IEEE Internet Computing, vol. 14, no. 4, 2010. p. 11-13.
- [5] ISO/TS--8000-1, ISO/TS 8000-1:2011 Data Quality - Part 1: Overview. 2011.
- [6] Варламов М.И., Турдаков Д.Ю. Обзор методов извлечения информации из Веб-ресурсов. Программирование, том 42, no. 5, 2016 г., стр. 39-48 / Varlamov M.I. and Turdakov D.Yu. A survey of methods for the extraction of information from Web resources. Programming and Computer Software, vol. 42, no. 5, 2016, p. 279-291.
- [7] Недумов Я.Р., Турдаков Д.Ю., Майоров В.Д., Овчинников П.Е. Автоматизация процесса нормализации информации при внедрении систем управления основными данными. Программирование, том 39, no. 3, 2013 г., стр. 3-14 / Nedumov Y.R., D.Yu. Turdakov, V.D. Maiorov, and Ovchinnikov P.E. Automation of data normalization for implementing master data management systems. Programming and Computer Software, vol. 39, no. 3, 2013, pp. 115-123.
- [8] Липаев В.В. Методология верификации и тестирования крупномасштабных программных средств. Программирование, том 29, no. 6, 2003 г., стр. 7-24 / Lipaev V.V. A Methodology of Verification and Testing of Large Software Systems. Programming and Computer Software, vol. 29, no. 6, 2003, pp. 298-309.
- [9] Aguilar J.A., Garrigós I., Mazón J.-N., Trujillo J. An MDA Approach for Goal-oriented Requirement Analysis in Web Engineering. Universal Computer Science, vol. 16, issue 17, 2010, pp. 2475-2494.
- [10] Moreno N. and A. Vallecillo. Towards interoperable Web engineering methods. Journal of American Society for Information Science and Technology, vol. 59, no. 7, 2008, pp. 1073-1092.
- [11] Escalona M.J., Torres J., Mejías M. et al., The treatment of navigation in web engineering. Advances in Engineering Software, vol. 38, no. 4, 2007, pp. 267-282.
- [12] Guerra-García C., Juárez-Ramírez R., Menéndez-Domínguez V. et al. Improving the Project Planning Process Considering Artifacts with Quality. In Proc. of the 4th. International Conference in Software Engineering Research and Innovation, 2016, pp. 15-20.
- [13] Guerra-García C., Llamas R. and Montaña-Rivas O. et al. QUACOP: An approach to Increase the Quality of Artifacts considered in a Project Planning Process. In Software Engineering: Methods, Modeling and Teaching. Universidad San Buenaventura Medellin, 2017.
- [14] Guerra-García C., Caballero I., and Piatini M. Capturing data quality requirements for Web applications by means of DQ_WebRE. Information Systems Frontiers, vol. 15, issue 3, 2013, pp. 433-445.
- [15] Guerra-García C., Caballero I., and Piatini M. A Survey on How to Manage Specific Data Quality Requirements during Information System Development. Communications in Computer and Information Science, vol. 230, 2011, pp. 16-30.
- [16] Ge M. and Helfert M. A Review of Information Quality Research - Develop a Research Agenda. In Proc. of the 12th International Conference on Information Quality, 2007, pp. 76-91.

- [17] Pipino L.L., Wang R.Y., Funk J.D., Lee Y.W. Journey to Data Quality. The MIT Press, 2006, 240 p.
- [18] Wang R., Pierce E., Madnick S. et al., eds. Information Quality. Advances in Management Information Systems. Routledge, 2005, 265 p.
- [19] ISO-25012, ISO/IEC 25012: Software Engineering-Software Product Quality Requirements and Evaluation (SQuaRE)-Data Quality Model. 2008.
- [20] Fons J., et al. Development of Web Applications from Web Enhanced Conceptual Schemas. Lecture Notes in Computer Science, vol. 2813, 2003, pp. 232-245.
- [21] Koch N. and A. Kraus. The Expressive Power of UML-based Web Engineering. In Proc. of the 2nd International Workshop on Web-oriented Software Technology (IWWOST '02), 2002, pp. 105-119.
- [22] Ceri S., Fraternali P., and Bongio A. Web Modeling Language (WebML): a modeling language for designing Web sites. Computer Networks, vol. 33, issues 1-6, 2000, pp. 137-157.
- [23] Baresi L., Garzotto F., Mainetti L., Paolini P. Meta-modeling Techniques Meet Web Application Design Tools. Lecture Notes in Computer Science, vol. 2306, 2002, pp. 182-206.
- [24] De Troyer O.M.F. and Leune C.J. WSDM: a user centered design method for Web sites. Computer Networks and ISDN Systems, vol. 30, issues 1-7, 1998, pp. 85-94.
- [25] De Castro V. and Marcos E. Towards a Service-Oriented MDA-Based Approach to the Alignment of Business Process with IT Systems: from the Business Model to a Web Service Composition Model. International Journal of Cooperative Information Systems, vol. 18, no. 2, 2009, pp. 225-260.
- [26] Meliá S. and Gómez J. Applying Transformations to Model Driven Development of Web applications. Lecture Notes in Computer Science, vol. 3770, 2005, pp. 63-73.
- [27] Busch M., Koch N., Masi M. et al. Towards Model-Driven Development of Access Control Policies for Web Applications. In Proc. of the Workshop on Model-Driven Security, 2012, pp. 1-6.
- [28] Zhang G. and Hölzl M. Aspect-Oriented Modeling of Web Applications with HiLA. Lecture Notes in Computer Science, vol. 7059, 2012, pp. 211-222.
- [29] Koch N., Knapp A., Zhang G., and Baumeister H. Uml-Based Web Engineering. In Web Engineering: Modelling and Implementing Web Applications, Springer, 2008, pp. 157-191.
- [30] Koch N. and Kraus A. Towards a Common Metamodel for the Development of Web Applications. Lecture Notes in Computer Science, vol. 2722, 2003, pp. 497-506.
- [31] Koch N. and Kroib C. UWE Metamodel and Profile. User Guide and Reference. Technical Report 0802, Institute for Informatics. Ludwig-Maximilians-Universität München (LMU), 2008, 35 p.
- [32] Busch M. Evaluating and Engineering: an Approach for the Development of Secure Web Applications. Dissertation, Ludwig-Maximilians-Universität München, 2016, 215 p.
- [33] Busch M., Koch N., and Suppan S. Modeling Security Features of Web Applications. Lecture Notes in Computer Science, vol. 8431, 2014, pp. 119-139.
- [34] Guerra-García C., Caballero I., and Piattini M. A Systematic Literature Review of How to Introduce Data Quality Requirements into a Software Product Development. In Proc. of the 5th International Conference on Evaluation of Novel Approaches to Software Engineering, 2010, pp. 12-19.
- [35] OMG. Unified Modeling Language: Superstructure. Versión 2.0. 2005.
- [36] EasyChair. EasyChair Conference System. Available at: <http://www.easychair.org/>.
- [37] Guerra-García C., Caballero I., and Piattini M. DQ-VORD: A Methodology for Managing and Integrating Data Quality Requirements into Software Requirement Specification. In Proc. of the IADIS International Conference on WWW/INTERNET, 2009, pp. 392–399.
- [38] Koch N., Zhang G., and Escalona M.J. Model transformations from requirements to Web system design. In Proc. of the 6th international conference on Web engineering, 2006, pp. 281–288.
- [39] Kraus A., Knapp A., and Koch N. Model-Driven Generation of Web Applications in UWE. In Proc. of the 3rd International Workshop on Model-Driven Web Engineering, 2007, pp. 1-16.
- [40] Escalona M.J. and Koch N. Metamodeling the Requirements of Web Systems. Lecture Notes in Business Information Processing book series, vol. 1, 2006, pp. 267-280.
- [41] Batini C. and Scannapieco M. Data Quality: Concepts, Methodologies and Techniques. Springer, 2006, 281 p.

Информация об авторах / Information about authors

Сезар Артуро ГУЕРРА-ГАРСИЯ, Ph.D., профессор. Область научных интересов: программная инженерия, качество данных и информации, инженерия требований.

César Arturo GUERRA GARCÍA, Doctor of Computer Science, Full Time Professor. Research interests include Software Engineering, Data and Information Quality, Requirements Engineering.

Гектор Херардо ПЕРЕС-ГОНСАЛЕС, кандидат наук, профессор. Область научных интересов: программная инженерия, проектирование программного обеспечения, инженерия требований.

Hector Gerardo PEREZ-GONZALEZ, Ph.D., Full Time Professor. Research interests include Software Engineering, Software Design, Requirements Engineering.

Марко Тулио РАМИРЕС-ТОРРЕС – кандидат прикладных наук, профессор. Область научных интересов: криптография, вейвлет-анализ, виртуальные приборы.

Marco Tulio RAMÍREZ-TORRES, Ph.D. in Applied Sciences, Full time Professor. Research interests include Cryptography, Wavelet Analysis, Virtual Instrumentation.

Рейес ХУАРЕС-РАМИРЕС, кандидат наук, профессор. Область научных интересов: разработка программного обеспечения, оценка неопределенности программного обеспечения и взаимодействие человека с компьютером.

Reyes JUÁREZ-RAMÍREZ, Ph.D., Full Time Professor. Research interests include Software Engineering, software uncertainty estimation, and human-computer interaction.

DOI: 10.15514/ISPRAS-2020-33(2)-3



Регулярные выражения для обнаружения Web-рекламы на основе автоматического скользящего алгоритма

*Д. Рианьо, ORCID: 0000-0001-8998-0129 <donovan20@comunidad.unam.mx>
Р. Пиньон, ORCID: 0000-0002-4819-5703 <rodrigo_pinon@comunidad.unam.mx>
Г. Молеро-Кастильо, ORCID: 0000-0002-6330-6408 <gmolero@fi-b.unam.mx>
Э. Барсена, ORCID: 0000-0002-1523-1579 <ebarcenas@unam.mx>
А. Веласкес-Мена, ORCID: 0000-0003-3509-6236 <mena@fi-b.unam.mx>*

*Национальный автономный университет Мексики,
Мексика, 04510, Мехико, Мэрия Койоакана*

Аннотация. Представлена реализация алгоритма распознавания Web-рекламы с использованием регулярных выражений. Сегодня при разработке программного обеспечения важную роль играет использование регулярных выражений, оптического распознавания символов, баз данных и автоматизированного тестирования. Тесты проводились в трех веб-браузерах. Результатом явилось выявление рекламы на испанском языке, которая отвлекает внимание пользователей, а прежде всего, позволяет получать информацию о них. Основная особенность алгоритма заключается в том, что его автоматическое и настраиваемое выполнение не требует доступа к коду рассматриваемой страницы, и в будущем может появиться приложение, работающее в фоновом режиме. Кроме того, при поддержке оптического распознавания символов мы получаем приемлемую эффективность при выявлении рекламы.

Ключевые слова: цифровой маркетинг; оптическое распознавание символов; регулярные выражения; интернет-реклама

Для цитирования: Рианьо Д., Пиньон Р., Молеро-Кастильо Г., Барсена, Э., Веласкес-Мена А. Регулярные выражения для обнаружения Web-рекламы на основе автоматического скользящего алгоритма. Труды ИСП РАН, том 33, вып. 2, 2021 г., стр. 65-76. DOI: 10.15514/ISPRAS-2021-33(2)-3

Regular Expressions for Web Advertising Detection based on an Automatic Sliding Algorithm

*D. Riaño, ORCID: 0000-0001-8998-0129 <donovan20@comunidad.unam.mx>
R. Piñon, ORCID: 0000-0002-4819-5703 <rodrigo_pinon@comunidad.unam.mx>
G. Molero-Castillo, ORCID: 0000-0002-6330-6408 <gmolero@fi-b.unam.mx>
E. Bárcenas, ORCID: 0000-0002-1523-1579 <ebarcenas@unam.mx>
A. Velázquez-Mena, ORCID: 0000-0003-3509-6236 <mena@fi-b.unam.mx>*

*National Autonomous University of Mexico,
Coyoacan Mayor's Office, Mexico City, CDMX, C.P. 04510, Mexico*

Abstract. This paper presents the automation of a Web advertising recognition algorithm, using regular expressions. Currently, the use of regular expressions, optical character recognition, Databases, and automation tests have been critical for multiple Software implementations. The tests were carried out in three Web browsers. As a result, the detection of advertisements in Spanish, that distract attention and that above all extract

information from users was achieved. The main feature of the algorithm is that automatic and versatile execution does not require access to the code of the page in question and that in the future it can be an application with background operation. In addition, being supported by optical character recognition gives us acceptable efficiency in detecting advertising.

Keywords: Digital Marketing; Optical Character Recognition; Regular Expressions; Web Advertising

For citation: Riaño D., Piñon R., Molero-Castillo G., Bárcenas E., Velázquez-Mena A. Regular Expressions for Web Advertising Detection based on an Automatic Sliding Algorithm. *Trudy ISP RAN/Proc. ISP RAS*, vol. 33, issue 2, 2021, pp. 65-76 (in Russian). DOI: 10.15514/ISPRAS-2021-33(2)-3.

1. Введение

В прошлом маркетинг не существовал в режиме онлайн, и его главная цель заключалась в привлечении средств массовой информации, чтобы сформировать у людей или даже компаний положительное мнение о рекламируемой продукции или о планируемых к продажам идеях. Но теперь этому наступил конец; имея под рукой все инструменты, которые дала нам современная технология, люди используют поисковые системы, чтобы найти все необходимое и более того, они могут получить доступ также и к критике и комментариям сообщества.

С появлением цифрового маркетинга основной целью становится пользователь [1], поэтому изменилась парадигма методов маркетинга. Сегодня цифровая стратегия должна включать все подходящие пространства для взаимодействия с «целью», выискивая людей, могущих повлиять на мнение других пользователей, чтобы привлечь их в свою сеть и усилить позиции идей или продуктов. Эта стратегия также может быть направлена на совершенствование поисковых систем, которые, как показывает опыт, становятся все более назойливыми в своих попытках проникновения в умы пользователей.

Для лучшего понимания текущих тенденций цифрового маркетинга на основе динамических Web-страниц, была разработана семантическая сеть. Важно отметить, что семантическая сеть – это форма представления знаний через взаимосвязи в виде графа [2]. На рис. 1 показаны взаимосвязи семантической сети, в которой рекламные объявления на веб-страницах можно разделить на три группы: а) сайты, которые существуют за счет рекламы; б) информационные, корпоративные сайты и интернет-магазины; в) социальные сети или нечто подобное.

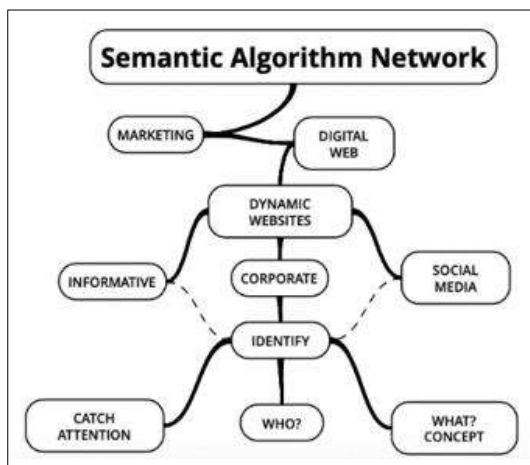


Рис. 1. Семантическая сеть современной Web-рекламы
Fig. 1. Semantic network of current Web advertising

Эта семантическая сеть связывает следующее: i) кто предлагает, ii) какой продукт или услуга предлагается, iii) каким образом привлекается внимание (предложение или раскрутка). Эти

типы рекламы распознаются при просмотре Интернета, некоторые из них были предметом анализа в настоящей исследовательской работе.

Веб-сайты, существующие за счет рекламы, предлагают краткие новостные, образовательные или развлекательные материалы, и их рекламный контент предоставляется Google. В них реклама постоянно меняется. Информационные, корпоративные сайты и интернет-магазины представлены такими сайтами, как MSN, Amazon, Sanborns, Adidas и др., целью которых является предложение текущих трендов продуктов или услуг. Социальные сети и их вариации, такие как Facebook, Instagram, LinkedIn, Twitter и пр. являются сайтами, которые предлагают взаимодействие с пользователями и рекламу, привязанную к профилям пользователей.

В настоящей статье представлена реализация алгоритма распознавания Web-рекламы с использованием регулярных выражений. Тесты проводились в трёх браузерах: Chrome, Firefox и Safari. Особенностью алгоритма является его автоматическое и настраиваемое выполнение, поскольку он не требует доступа к коду просматриваемой Web-странице и может считаться приложением, работающим в фоновом режиме.

2. Предварительная информация

При наличии высокого потенциала и развитых стратегий современного цифрового маркетинга эта деятельность используется все чаще как важная часть деятельности кампаний по обеспечению приверженности потребителей к торговым маркам, поскольку обеспечивается коммуникационный канал для взаимодействия между потенциальными клиентами и торговой маркой. К числу действий, ориентированных на интернет-маркетинг, относятся [3]: обеспечение потребительской лояльности; повышение имиджа торговой марки и продаж продуктов; проведение акций по раскрутке и тестированию продуктов; поощрение повторной покупки продукта той же торговой марки; проведение прямых и персонализированных коммуникационных кампаний.

Одним из способов борьбы с рекламой и обеспечения конфиденциальности в Web-браузерах является использование регулярных выражений (regular expression, RegExp, Regex, или RE). Эти выражения представляют собой способ представления символьных строк, которые соответствуют некоторому шаблону [4] [5]. Приложения RegExp разнообразны, например, валидация полей форм, идентификация текстовых строк в социальных сетях, поисковые команды и т.д. [6] [7].

2.1 Управление конфиденциальностью пользователей

В настоящее время управление информацией и конфиденциальностью находится в критической ситуации [8]. Такие компании, как Facebook, Twitter, Google, Amazon и др. включают в свои веб-сайты, приложения или поисковые системы средства накопления информации. Эти компании предоставляют услуги «бесплатно», но используют информацию о пользователях по своему усмотрению, в рекламных целях.

Увеличение количества рекламы в Web-браузерах ставит под угрозу конфиденциальность пользователей. Web-сайты, предлагаемые по результату поиска определенных видов товаров и услуг, перенасыщены информацией о клиентах. Другой важный аспект заключается в том, что сами поисковые запросы пользователей позволяют собрать большой объем данных. Как следствие, эти данные при определенной обработке могут также предоставить явную информацию для более эффективного управления целевой рекламой [9] [10].

Для примера на рис. 2 показано, как поисковая система Google автоматически создает профили пользователей на основе просмотра ими веб-страниц. Например, релевантным является пользователь мужского рода, возраст которого колеблется от 18 до 24 лет, заинтересованный в покупке онлайн, проведении открытых мероприятий, разработке мобильных приложений, просмотре фильмов и прочее. Например, может быть актуален

пользователь мужского пола в возрасте от 18 до 24 лет, который заинтересован в совершении покупок в Интернете, занятиях активным отдыхом, разработке мобильных приложений, просмотре фильмов и т.д. Данные об таких пользователях собираются через файлы cookie, трекеры и информацию из профиля Gmail, и эти данные могут быть использованы для группировки профилей с общими характеристиками с целью рекламы различных предложений и промоакций.

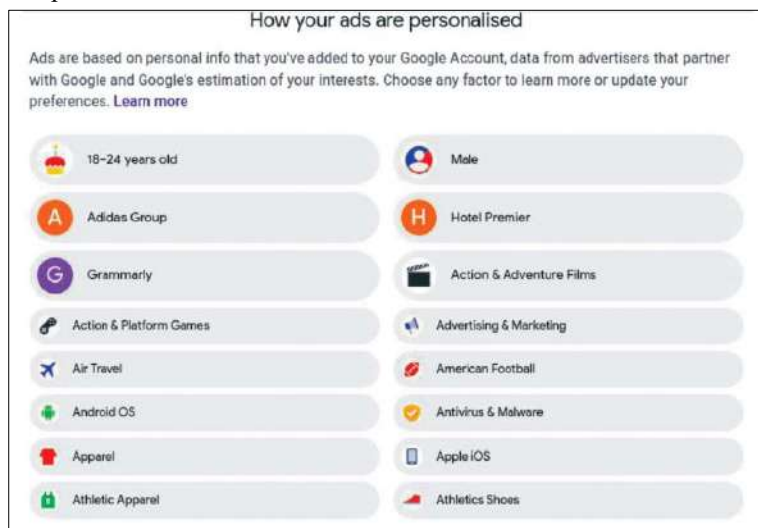


Рис. 2. Поисковая система Google для автоматического создания профилей пользователей на основе просмотра Web-страниц

Fig. 2. Google search engine for automatic creation of user profiles based on Web browsing

Поэтому при постоянном взаимодействии со службами Google стоит почаще знакомиться с изменениями в политике конфиденциальности в Маунтин-Вью¹, чтобы иметь лучшее представление о безопасности и использовании нашей информации.

2.2 Блокировщики рекламы веб-браузеров

В последние годы были увеличены усилия по внедрению блокировщиков веб-рекламы [11]. Первоначально они были реализованы для браузера Firefox, который годами совершенствовался, чтобы получился браузер без рекламы. Но компании искали способы продолжить рассылку рекламы. В скором времени возникла полемика [1] в связи со значительными потерями возможной прибыли [12].

Еще один термин, получивший распространение в настоящее время, – это трекеры в маркетинге, которые являются средствами отслеживания и индикаторами эффективности таргетированных рекламных кампаний. Эти инструменты сохраняют информацию с помощью файлов cookie, а те, в свою очередь, предоставляют местоположение выполняемых поисковых запросов. По мере роста спроса и расширения возможностей блокировки рекламы в большинстве Web-браузеров были реализованы новые блокировщики рекламы.

Необходимо учитывать тот важный факт, что в последние годы Google контролирует 85% мирового рекламного бизнеса в поисковых системах и около 50% всей онлайн-рекламы [13]. Люди и общество, как правило, считают Google сервисом [14], но за этим стоит технология, которая поддерживает специальные функции и включает уникальные и ограничительные расширения.

¹ В Маунти-Вью (Mountain View), Калифорния находится штаб-квартира компании Google.

2.3 Управление Web-браузерами на основе Selenium

Selenium, называемый также Selenium Webdriver, – это инструмент для автоматизации процессов в различных Web-браузерах [15]. Его цель – улучшить поддержку обнаружения проблем в любом Web-браузере [16]. Этот инструмент позволяет тестировать любой веб-браузер с получением HTML-кода, изменять и открывать вкладки окон браузера, а также и перемещаться между вкладками, изменять размеры окон, создавать скриншоты, заполнять поля форм многое другое.

Selenium поддерживает языки программирования Java, Python, C#, Ruby, Perl и JavaScript. Операционные системы – Windows, Mac OS и Linux, каждая со своими соответствующими пакетами и интегрированными средами разработки (Integrated Development Environment, IDE) [17].

С другой стороны, в настоящее время существует интерес к использованию оптического распознавания символов (OCR) при обнаружении Web-рекламы [14], однако все еще требуются дополнительные усилия для охвата всех стратегий цифрового маркетинга.

2.4 Родственные работы

Одной из работ, посвященных обнаружению Web-сайтов, пригодных для таргетированной рекламы, была статья [18], где представлен алгоритм, выполняющий Web-краулинг. Метод состоит в получении информации с Web-сайтов путем обучения системы на заранее размеченном наборе страниц (для этого использовались страницы MSN). Классификация производилась по предполагаемым ключевым словам, встречающимся в начале фразы, в середине фразы, в конце фразы и вне фразы.

Работа [19] описывает анализ контекстной рекламы с помощью PageSense, который направлен на ассоциирование рекламы со стилем Web-страниц. С помощью этой платформы выявляются пустые области Web-страниц и выбирается не назойливое место для размещения рекламы, не нарушая оригинальный стиль Web-страницы. Для анализа использовались байесовские комбинации и вероятности, которые отражают процентное соотношение рекламных объявлений для различных видов товаров или услуг.

В работе [20] анализ производился на основе евклидовых расстояний. Эти расстояния позволяют оценить, какой интерес представляет реклама для пользователей, насколько она помогает при поиске продуктов, а также способствует адаптации профилей пользователей, позволяя разбить его на разделы, такие как здоровье, спорт, бизнес-общество, образование, искусство, наука, компьютер и т.д.

Наконец, в работе [21] описан анализ, проведенный примерно на 500 веб-страницах и направленный на выявление типов (не содержания) рекламы. Среди проанализированных типов рекламы выделяются всплывающие окна, карусели, видео, GIF-файлы, игры, стикеры или текст. Также были проанализированы страны, откуда поступают рекламные объявления, их частота, размер и происхождение URL.

3. Методика

Алгоритм был реализован на языке Python с применением библиотек Tesseract [22], Pillow и OpenCV [23], а также пакета Tesseract-OCR. Реализация начинается с точки, предшествующей процессу OCR. Кроме того, мы использовали библиотеку Selenium для выполнения автоматического скольжения в Web-браузере.

При реализации также учитывалась разница в высоте мониторов различных размеров, представленных на рынке, поэтому алгоритм динамически определяет высоту окон, подстраиваясь под любой размер экрана. По этой причине перемещение информации в настоящей работе является вертикальным.

3.1 Скольжение по Web-сайту

Через Selenium открывается новое окно для браузера, совместимого с этим инструментом. Терминал запрашивает у пользователя поле URL. Далее тестируемый браузер расширяется на весь экран для быстрого и полного сканирования Web-сайта.

С целью приостановки процесса на несколько мгновений используются программные потоки управления (thread). Первая пауза делается для того, чтобы дать странице загрузиться, так как в зависимости от скорости Интернета и даже от состояния страницы требуется время для полной ее загрузки, чтобы можно было сделать правильный снимок экрана для его последующего анализа с помощью OCR.

На следующем этапе окно разворачивается до максимума, а затем скользит для захвата и сохранения экранов. Это вызывает вторую паузу, которая составляет 0,5 секунды. Последовательно запускается процесс захвата экрана, пока не будет обработано все вертикальное содержимое Web-страницы. После этого окно Web-браузера автоматически закрывается.

3.2 Оптическое распознавание символов

OCR, позволяет с высокой эффективностью извлекать буквенный текст из изображений с независимо от языка, размера или цвета текста. Эффективность OCR колеблется от 71% до 98%. Такая система способна достигать средних значений 85,1% для рукописного текста и 90,93% для печатного текста [24].

Была определена специальная функция для поиска всех скриншотов, сохраненных в файловой системе с заданной частью путевого имени. Затем был реализован цикл, в котором все найденные изображения анализировались в том порядке, в котором они были захвачены. После обработки нужного изображения результаты OCR сохраняются в текстовой переменной. Затем текст форматируется: все буквы заменяются на прописные, слова разделяются, устраняются пробелы и символы, не поддерживаемые в символьных строках в языке SQL, чтобы иметь возможность сопоставления с регулярными выражениями.

3.3 Регулярные выражения

После того, как информация о содержимом Web-страницы получена и преобразована в текстовые строки, производится доступ к локальному серверу для проверки содержимого Web-сайта с помощью базы данных, в трех таблицах которой хранится около 600 различных слов на основе регулярных выражений: i) слова, наиболее часто используемые в цифровом маркетинге; ii) торговые марки, учитывая их соответствующие суб-бренды в продуктах или услугах, предлагаемых компаниями; iii) тип продукта или сервиса, или из чего состоит продукт или услуга.

В табл. 1, 2 и 3 показаны фрагменты регулярных выражений, соответственно относящиеся к наиболее часто используемым словам в цифровом маркетинге, некоторым узнаваемым брендам и некоторым продуктам, наиболее часто упоминаемым в Мексике.

Табл. 1. Фрагмент наиболее часто используемых слов в цифровом маркетинге на испанском языке
Table 1. Fragment of the most used words in digital marketing in Spanish

Слово	Множественное число	Акцент	Символ
Ahorro	Ahorros	Null	Null
Bajo	Bajos	Null	Null
Comprar	Null	Null	Null
Cotiza	Null	Null	Cotizar
Descuento	Descuentos	Null	%
Dinero	Dineros	Null	Null

Especial	Especiales	Null	Null
Gratis	Gratis	Null	Gratis
Hasta	Null	Null	Null
Ilimitado	Ilimitados	Null	Null
Interes	Intereses	Interés	Null
Internet	Null	Null	Web
Oferta	Ofertas	Null	Null

Табл. 2. Фрагмент некоторых из самых продаваемых и самых дорогих брендов Мексики

Table 2. Fragment of some of the best-selling and highest-paid brands in Mexico

Бренд	Суб-бренд	Продукт	Акроним
Adidas	Boost	Boost	Null
Adidas	NMD	NMD	Null
Apple	iMac	iMac	Null
Apple	iPhone	iPhone	Null
Bancomer	BBVA	BBVA	BBVA
Banorte	Banorte	Banorte	Null
Levis	Trucker	Trucker	Null
Levis	Western	Western	Null
Mazda	Mazda2	Mazda2	Null
Mazda	Mazda3	Mazda3	Null
Microsoft	Azure	Azure	Null
Microsoft	Office	Office	Null
Nike	Jordan	Jordan	Null

Табл. 3. Фрагмент некоторой продукции, рекламируемой в Мексике

Table 3. Fragment of some products advertised in Mexico

Ключ	Тип	Категория
5G	Internet	Internet
Americano	Deportes	Entretenimiento
Basquetbol	Deportes	Entretenimiento
Béisbol	Deportes	Entretenimiento
Chamarra	Ropa	Vestimenta
Chico	Talla	Vestimenta
Compacto	Autos	Automóvil
Ellos	Género	Social
Familia	Género	Social
Grande	Talla	Vestimenta
Hatchback	Autos	Automóvil
Jeans	Ropa	Vestimenta
Laptops	Electrónicos	Electrónica
Licuada	Electrónicos	Electrodomésticos
Smartphones	Electrónicos	Electrónica

Эти слова и символы, составляющие регулярные выражения, обычно используются в рекламе в Интернете [25].

Кроме того, поскольку реклама играет не только с визуальными эффектами, но и с буквами, размерами и стилями, диапазон поиска был расширен, с тем чтобы охватить множественные числа слов, акценты и символы, относящиеся к некоторым ключевым словам.

Важно отметить, что эти регулярные выражения брендов и продуктов связаны с исследованиями брендов в Мексике и некоторыми исследованиями в Латинской Америке. Для Мексики статистические данные за 2019 год были найдены в базах данных Национального института статистики и географии (Instituto Nacional de Estadística y Geografía, INEGI) – мексиканской государственной организации, осуществляющей сбор и

систематизацию статистической, демографической, географической и экономической информации о стране. Другим источником данных была Экономическая комиссия для Латинской Америки и Карибского бассейна (Economic Commission for Latin America and the Caribbean, ECLAC/ЭКЛАК), которая является органом Организации Объединенных Наций, предоставляющим доступ к информации в некоторых странах Латинской Америки.

3.4 Псевдокод

Итак, основная идея алгоритма заключается в использовании автоматической прокрутки для захвата информации, содержащейся на Web-страницах. Затем эти захваченные изображения обрабатываются и преобразуются в текстовый формат, чтобы впоследствии идентифицировать существующие рекламные объявления на основе фильтров и сопоставлений с задаваемыми в виде регулярных выражений семействами слов, используемыми в цифровом маркетинге. Наконец, выявляются наиболее часто рекламируемые бренды, которые лидируют в поисковых запросах через Web-браузеры. Это полезно для их отслеживания. На рис. 3 представлен псевдокод разработанного алгоритма.

```
1 Open a browser with selenium
2 Enter the URL of the desired page
3 iteration = 1
4
5 while true:
6
7     Page Height = Height of the web page in pixels
        according to selenium function
8     Height = High computer screen size in pixels
9     Slip = Height * iteration
10    Capture with selenium tools
11    Save image with the capture number name
12    Slide page according to the number
        of pixels indicating that Slide
13    if Slip >= Page Height:
14        break
15        iteration +=1
16 Close selenium browser
17 Search file where captures were saved
18 Add in a list all the paths of the elements of this file
19 Connection is made to the database
20
21 for i in range (capture path list):
22     image = capture path list [i]
23     list = pytesseract.image_to_string (
24         img).upper().split ()
25         #Separate the text string returned by
        pytesseract into many smaller strings
26         #with atomic elements, that are words,
        because the segmentation process was
        carried out
27         #when a space is found. These words
        are now capitalized.
28     WordsFound = []
29     for j in range (list):
30         #We proceed to make queries to find each
        word in 'list' in the 3 tables of the base
31         Find = Result of queries looking
        for list [j]
32         if length (Find) != 0:
33             WordsFound.append (list [j])
34             Delete Capture already analyzed
35 print (Capture number analyzed)
print (Found Words without repeating, their number
of occurrences and precedence table)
```

Рис. 3. Псевдокод реализованного алгоритма
Fig. 3. Pseudocode of the implemented algorithm

В целом, как уже было описано в предыдущих пунктах, алгоритм состоит из трех основных этапов: i) необходимо получить URL-адрес страницы, а затем сделать снимки экрана с помощью автоматической прокрутки; ii) изображения скриншотов затем обрабатываются в том порядке, в котором они были сделаны, для преобразования их в текст с помощью OCR; и iii) текст анализируется, в результате получается список совпадений с регулярными выражениями, хранящимися в таблицах, характеризующий компании, которые размещают рекламу, продукты и стратегии.

Одним из ограничений работы является то, что в ней не используется Web-скрейпинг (Web Scraping), представляющий собой процесс автоматического сбора данных и информации из Интернета, которые затем анализируются для определенных нужд и целей [25]. Кроме того, не учитывались ни персональные расширения браузеров, ни связанные учетные записи для синхронизации с устройствами. Кроме того, не анализировалась реклама с боковым выдвиганием, потому что слайдинг в нашем случае вертикальный, сверху вниз.

4. Результаты

Для анализа веб-рекламы были рассмотрены три типа динамических Web-страниц, протестированных в трех разных браузерах: Google Chrome, Mozilla Firefox и Safari, разработанный компанией Apple. Были проанализированы следующие сайты:

- MSN: www.msn.com/es-mx;
- Sanborns: www.sanborns.com.mx;
- AhorraSeguros: <https://ahorraseguros.mx>.

В Табл. 4 обобщаются результаты, полученные для каждого URL-адреса в каждом Web-браузере, общее количество слов (регулярных выражений), которые появляются в рекламе на каждом оцениваемом Web-сайте, количество снимков экрана и время выполнения с момента открытия Web-браузера до завершения сравнения.

Табл. 4. Результаты выполнения программы

Табл. 4. Results of the algorithm evaluation

Web-сайт	Web-браузер	Реклама	Количество скриншотов	Время (сек)
URL 1, MSN	Chrome	106	9	11.636
	Firefox	103	9	12.144
	Safari	118	9	14.539
URL 2, Sanborns	Chrome	56	4	4.449
	Firefox	62	4	4.036
	Safari	68	4	4.209
URL 3, Ahorra Seguros	Chrome	133	9	13.547
	Firefox	149	9	14.547
	Safari	153	9	14.556

Как показывает Табл. 4, наибольшее количество рекламных объявлений алгоритм обнаружил при использовании браузера Safari. Такая лучшая идентификация рекламы обусловлена тем, что в этом случае алгоритм лучше выравнивает содержимое Web-страницы, и на обработку уходит больше времени. В Chrome и Firefox также обнаружено значительное количество рекламных объявлений, но некоторая информация была потеряна при прокрутке Web-страницы.

Для определения эффективности алгоритма был выполнен визуальный анализ информации, содержащейся на оцениваемых Web-страницах. Этот анализ заключался в подсчете RegEx в Web-рекламе, результаты которого должны соответствовать общему количеству слов,

обнаруженных алгоритмом. В табл. 5 приведены результаты, полученные вручную и автоматически при выполнении нашего алгоритма.

Табл. 5. Слова, идентифицированные алгоритмом по отношению к общему количеству слов с рекламным содержанием

Table 5. Words identified by the algorithm with respect to the total of words with advertising content

Web-сайт	Раздел сайта	Визуальный подсчет	Chrome	Firefox	Safari
URL 1, MSN	Microsoft	22	22	14	22
	Новости	9	9	9	9
	IOS	10	10	0	10
	Android	10	10	0	10
	MSN	12	3	10	2
	Rebaja	15	13	12	14
	Bcero	124	106	103	118
URL 2, Sanborns	\$	27	18	23	24
	Libros	3	2	1	3
	Perfumes	2	2	2	2
	Tecnología	3	3	3	3
	Videojuegos	2	2	2	2
	Bcero	75	56	62	68
URL 3, Ahorra Seguros	Seguros	57	48	54	56
	Seguro	22	22	22	21
	Beneficios	7	5	7	7
	Precios	5	4	5	4
	Servicios	5	2	1	4
	Bcero	172	133	149	153

В случае www.msn.com/es-mx, URL 1 была достигнута отличная производительность алгоритма обнаружения Web-рекламы через браузер Safari с достоверностью 95,16%. В то же время браузеры Chrome и Firefox также достигли значительной достоверности в 85,48 и 83,06% соответственно.

Для URL 2, www.sanborns.com.mx при работе через Safari получена достоверность 90,66%, для Chrome – 74,66% и для Firefox – 82,66%. Эти результаты обусловлены удивительным визуальным дизайном сайта, который понятен людям, но сложен для анализа из-за наложения некоторых слов, а также наличия логотипов, а также слов разного размера и с разными шрифтами, связанных друг с другом.

Для случая <https://ahorraseguros.mx>, URL 3 также были достигнуты значительные результаты с достоверностью 88,95% в Safari, 86,62% в Firefox и 77,32% в Chrome. Более высокий уровень достоверности не был достигнут из-за большого количества логотипов разных брендов с различными шрифтами и фоном. Это явилось основной проблемой при оптическом распознавании символов.

Специфика результатов тестов обусловлена обрезкой и выравниванием экрана при автоматическом скольжении, то есть в каждом Web-браузере конфигурация варьируется, изменяя способ организации информации на Web-сайте, что и вызывает потери содержимого.

6. Заключение

Достижения современных технологий имеют также и отрицательные стороны для пользователя, такие как наплыв рекламы в Интернете – персонализированные рекламные объявления, в которых учитываются потребности и интересы пользователей.

Интернет-реклама состоит не только из слов или предложений, привлекающих внимание пользователя к рекламным акциям или предложениям. Для рекламодателей важно, чтобы пользователь знал, кто занимается продвижением товаров и услуг, независимо от того,

действительно ли клиент планирует купить рекламируемый продукт. Для рекламодателей самое главное – привлечь внимание пользователя, чтобы он запомнил бренд для будущих покупок.

Использование регулярных выражений в нашей работе было полезным. Кроме того, реализация базы данных облегчила организацию обнаружения Web-рекламы, позволяя охватить больше случаев использования рекламных объявлений.

Для проведенных тестов были получены ожидаемые результаты с достоверностью от приемлемого до высокого, в диапазоне от 74,66% до 95,16%, а самый высокий уровень достоверности был получен при использовании MSN благодаря простому дизайну, использованию распространенных шрифтов постоянного размера, отсутствию составных слов, логотипов и т.д.

В будущем предполагается включить в базу данных больше регулярных выражений и произвести расширение алгоритма, то есть включить алгоритмы искусственного интеллекта, способные распознавать рекламу на основе цветовых узоров, размера и расположения баннеров и других возможностей современных рекламодателей.

Список литературы / References

- [1]. Marketing Digital. What is digital marketing? URL: <https://www.mdmarketingdigital.com/en/what-is-digital-marketing>, accessed 12/01/2020.
- [2]. Redes Semánticas. URL: <http://tesis.uson.mx/digital/tesis/docs/9049/Capitulo1.pdf>, accessed 12/01/2020 (in Spanish).
- [3]. Marketing Online: Potencial y Estrategias, 2019. URL: https://www.cecarm.com/Guia_Marketing_Online_Potencial_y_Estrategias_-_CECARM.pdf-6120, accessed 12/01/2020 (in Spanish).
- [4]. Pomol R., González C., González S. Una herramienta didáctica para el aprendizaje interactivo de expresiones regulares. 2013. URL: <http://repositorio.uigv.edu.pe/handle/20.500.11818/804>, accessed 12/01/2020 (in Spanish).
- [5]. Beltrán R. El uso de expresiones regulares en la detección de errores escritos: implicaciones para el diseño de un corrector gramatical, 2008. URL: <https://dialnet.unirioja.es/servlet/articulo?codigo=4007478>, accessed 12/01/2020 (in Spanish).
- [6]. Gallego A. La jerarquía de Chomsky y la facultad del lenguaje: consecuencias para la variación y la evolución. *Teorema: Revista internacional de filosofía*, vol. 27, no. 2, 2008, pp. 47-60 (in Spanish).
- [7]. García Campos I. Herramienta para la corrección automática de autómatas finitos, 2017. URL: <https://riull.ull.es/xmlui/handle/915/5846>, accessed 12/01/2020 (in Spanish).
- [8]. Sánchez J., López L., Martínez J. Solución para garantizar la privacidad en el Internet de las Cosas. *El profesional de la información*, vol. 24, no. 1, 2015, pp. 62-70 (in Spanish).
- [9]. Ortiz M., Aguilar L., Marín L. Los desafíos del marketing en la era del big data. *e-Ciencias de la Información*, vol. 6, no. 1, 2016, pp. 1-30 (in Spanish).
- [10]. Riaño D., Molero-Castillo G., Velázquez-Mena A., Bárcenas E. Expresiones regulares para el tratamiento de privacidad de navegadores Web. *Abstraction and Application*, vol. 25, 2019, pp.121-130 (in Spanish).
- [11]. Cerezo, P., Ad blocking: el modelo publicitario digital, a revisión, *Cuadernos de periodistas: revista de la Asociación de la Prensa de Madrid*, 2016, pp. 81-89 (in Spanish).
- [12]. Londaitz A., Publicidad en los celulares: Publicidad invasiva vs. derecho a la privacidad. Thesis, Universidad del Salvador, 2011.
- [13]. Bienvenido a Google, la mejor empresa para trabajar, 2013. URL: www.expansion.com/2013/08/23/directivos/1377273795.html, accessed 12/01/2020 (in Spanish).
- [14]. Jarvis J. Y Google, ¿cómo lo haría?, 2000. URL: <https://narrativabreve.com/2013/10/libro-google-jeff-harvis.html>, accessed 12/01/2020 (in Spanish).
- [15]. Leotta M., Clerissi D., Ricca F., Spadaro C. Comparing the maintainability of selenium webdriver test suites employing different locators: A case study. In *Proc. of 1st International Workshop on Joining AcadeMiA and Industry Contributions to Testing Automation*, 2013, pp. 53-58
- [16]. Gojare S., Joshi R., Gaigaware D., Analysis and Design of Selenium WebDriver Automation Testing Framework, *Procedia Computer Science*, vol. 50, 2015, pp. 341-346.
- [17]. Selenium Webdriver, 2017. URL: www.tutorialspoint.com/selenium/pdf/selenium_webdriver.pdf, accessed 12/01/2020.

- [18]. Yih W., Goodman J., Carvalho V. Finding Advertising Keywords on Web Pages, In Proceedings of the 15th International Conference on World Wide Web, 2006, pp. 213-222.
- [19]. Mei T., Li L., Tian X., Tao D., Ngo C. PageSense: Toward Stylewise Contextual Advertising via Visual Analysis of Web Pages. *IEEE Transactions on Circuits and Systems for Video Technology*, vol. 28, no. 1, 2018, pp. 254-266.
- [20]. Sánchez D., Viejo A. Privacy-preserving and advertising-friendly web surfing. *Computer Communications*, vol. 130, 2018, pp. 113-123.
- [21]. Krammer V. An Effective Defense against Intrusive Web Advertising. In Proc. of the Sixth Annual Conference on Privacy, Security and Trust, 2008, pp. 3-14.
- [22]. Sajjad K. Automatic license plate recognition using Python and Opencv. College of Engineering, 2010. URL: <https://pdfs.semanticscholar.org/bddf/1200eb17f239e4dce2a9cec938eb8cf305f5.pdf>, accessed 12/01/2020.
- [23]. Patel C., Patel A., Patel D. Optical Character Recognition by Open source OCR Tool Tesseract: A Case Study. *International Journal of Computer Applications*, vol. 55, no. 10, 2012, pp. 50-56.
- [24]. Vallez M. Keyword Research: métodos y herramientas para identificar palabras clave. *BiD: textos universitarios de biblioteconomía i documentació*, vol. 27, 2011, pp. 1-14 (in Spanish).
- [25]. Slamet C., Andrian R., Maylawati D. et al. Web Scraping and Naïve Bayes Classification for Job Search Engine. In Proc. of the 2nd Annual Applied Science and Engineering Conference, 2018, pp. 1-7.

Информация об авторах / Information about authors

Донован РИАНЬО-ЭНРИКЕС, студент.

Donovan RIAÑO-ENRIQUEZ, Student.

Родриго ПИНОН-АЯЛА, студент.

Rodrigo PINON-AYALA, Student.

Гильермо МОЛЕРО-КАСТИЛЬО, кандидат наук, доцент кафедры вычислительной техники. Область научных интересов: искусственный интеллект, наука о данных, интеллектуальный анализ данных, машинное обучение.

Guillermo MOLERO-CASTILLO, Ph.D. in Information Technologies, Associate Professor, Computing Engineering Department. Research interests: Artificial Intelligence, Data Science, Data Mining, Machine Learning.

Эверардо БАРСЕНАС, кандидат наук, доцент. Область научных интересов: модальная логика, теория доказательств, автоматизированные рассуждения, логика описания, model checking, представление знаний, планирование, компьютерное зрение.

Everardo BARCENAS, Ph.D., Assistant Professor. Research interests: modal logics, proof theory, automated reasoning, description logics, model checking, knowledge representation, planning, computer vision.

Алехандро БЕЛАСКЕС-МЕНА, магистр, доцент. Область научных интересов: сетевой компьютер, туманные вычисления, распределенные вычисления.

Alejandro VELAZQUEZ-MENA, Master of Science, Assistant Professor. Research interests: Network Computer, Fog Computing, Distributed computing.



Оценка пользовательских историй на основе декомпозиции сложности с использованием байесовских сетей

¹ М. Дуран, ORCID: 0000-0002-6866-2647 <mayra.duran@uabc.edu.mx>

¹ Р. Хуарес-Рамирес, ORCID: 0000-0002-5825-2433 <reyesjua@uabc.edu.mx>

² С. Хименес, ORCID: 0000-0003-0938-7291 <samantha.jimenez@tectijuana.edu.mx>

¹ К. Тона, ORCID: 0000-0003-4492-3432 <tona.caudia@uabc.edu.mx>

¹ Автономный университет Нижней Калифорнии (UABC),
Мексика, 21100, Нижняя Калифорния, Энсенада

² Тихуанский технологический институт,
Мексика, 22414, Нижняя Калифорния, Тихуана

Аннотация. На сегодняшний день в методологии разработки программного обеспечения Scrum предлагаются разные методы оценки трудозатрат и сложности пользовательских историй. В большинстве существующих методов факторы анализируются на уровне мелких структурных единиц, и эти методы не всегда точны. Хотя чаще всего для оценки пользовательских историй в Scrum используется покер планирования, он эффективен в основном для опытных команд, поскольку оценка строится на основе наблюдений экспертов, но у неопытных команд применение этого метода вызывает трудности. В данной статье мы предлагаем метод декомпозиции сложности на крупные блоки, позволяющий учитывать важные для оценки факторы. Для представления факторов и связей между ними используется байесовская сеть. Ребра взвешиваются на основе профессиональной оценки важности рассматриваемых факторов. Узлы сети представляют факторы. На этапе оценки Scrum-команда присваивает каждому фактору значение, что позволяет сети сгенерировать значения для сложности пользовательской истории с последующей трансформацией в номер карты покера планирования, которой представляет относительную оценку сложности пользовательской истории. Цель данного исследования состоит в том, чтобы предоставить командам разработчиков без опыта или без имеющихся статистических данных метод, позволяющий существенно повысить точность оценки сложности пользовательских историй с помощью модели, ориентированной на человеческий фактор разработчиков.

Ключевые слова: пользовательские истории; мелкоструктурная сложность; оценка; байесовская сеть; экспертное мнение.

Для цитирования: Дуран М., Хуарес-Рамирес Р., Хименес С., Тона К. Оценка пользовательских историй на основе декомпозиции сложности с использованием байесовских сетей. Труды ИСП РАН, том 33, вып. 2, 2021 г., стр. 77-92. DOI: 10.15514/ISPRAS-2021-33(2)-4.

Благодарности. Авторы выражают благодарность студентам программы по разработке программного обеспечения, обучавшимся в 2019 году в Автономном университете Нижней Калифорнии (UABC) и принявшим участие в Scrum-проектах, а также профессионалов этой области, представляющих различные компании региона. Благодаря вашему вкладу нам удалось собрать необходимые данные для успешного исследования. Мы также выражаем особую благодарность Национальному совету по науке и технологиям (CONACYT) за предоставленную стипендию для обучения на магистерской программе в области точных наук.

User Story Estimation based on the Complexity Decomposition using Bayesian Networks

¹ M. Durán, ORCID: 0000-0002-6866-2647 <mayra.duran@uabc.edu.mx>

¹ R. Juárez-Ramírez, ORCID: 0000-0002-5825-2433 <reyesjua@uabc.edu.mx>

² S. Jiménez, ORCID: 0000-0003-0938-7291 <samantha.jimenez@tectijuana.edu.mx>

¹ C. Tona, ORCID: 0000-0003-4492-3432 <tona.caudia@uabc.edu.mx>

¹ Universidad Autónoma de Baja California,
Ensenada, Mexico, 21100

² Instituto Tecnológico de Tijuana, México,
Tijuana, Mexico, 22424

Abstract. Currently, in Scrum, there are different methods to estimate user stories in terms of effort or complexity. Most of the existing techniques consider factors in a fine grain level; these techniques are not always accurate. Although Planning Poker is the most used method in Scrum to estimate user stories, it is primarily effective in experienced teams since the estimation mostly depends on the observation of experts, but it is difficult when is used by inexperienced teams. In this paper, we present a proposal for complexity decomposition in a coarse grain level, in order to consider important factors for complexity estimation. We use a Bayesian network to represent those factors and their relations. The edges of the network are weighted with the judge of professional practitioners about the importance of the factors. The nodes of the network represent the factors. During the user estimation phase, the Scrum team members introduce the values for each factor; in this way, the network generates a value for the complexity of a User story, which is transformed in a Planning Poker card number, which represents the story points. The purpose of this research is to provide to development teams without experience or without historical data, a method to estimate the complexity of user stories through a model focused on the human aspects of developers.

Keywords: user stories; fine-grained complexity; estimation; Bayesian network; expert judge

For citation: Durán M., Juárez-Ramírez R., Jiménez S., Tona C. User Story Estimation based on the Complexity Decomposition using Bayesian Networks. *Trudy ISP RAN/Proc. ISP RAS*, vol. 33, issue 2, 2021, pp. 77-92 (in Russian). DOI: 10.15514/ISPRAS-2021-33(2)-4.

Acknowledgments. The authors would like to thank the students of Software Engineering course enrolled in 2019-1 and 2019-2 at Universidad Autónoma de Baja California who participated in the Scrum projects and the professionals in the area who participated from various companies in the region. Thanks to your participation we managed to collect data and information for the formation of our study. We also have special thanks to Consejo Nacional de Ciencia y Tecnología (CONACYT) for the scholarship to realize the master in science studies.

1. Введение

Scrum – это Agile-метод разработки, который определяется как итеративный, инкрементальный и эмпирический процесс, позволяющий управлять разработкой и осуществлять контроль над ней [1-3]. Это фреймворк, который представляет собой набор практических приемов для поддержки наглядного представления, экспертизы и адаптации проектов по разработке программного обеспечения [4-5]. Ключевым элементом Scrum являются пользовательские истории (user stories). Первая версия пользовательских историй – это краткое (одно или два предложения), простое и конкретное описание взаимодействия пользователя и разрабатываемого продукта [6]. Проект по разработке программного обеспечения включает в себя оценку стоимости, трудозатрат, размера и продолжительности [7].

При подсчете очков пользовательских историй (story points) следует учитывать множество факторов [8-11]. На качество и точность оценки напрямую влияют опыт и знания разработчиков, которые помогают определить и учесть все возможные факторы. По этой причине процесс оценки может быть проблематичным для начинающих разработчиков, так

как у них недостаточно опыта [12]. На оценку также негативно влияет неопределенность в требованиях [13]. В Scrum для оценки пользовательских историй можно использовать несколько методов, таких как покер планирования (Planning Poker) [8], Wideband Delphi [14], Fist of Five, Affinity Estimation [5] и T-Shirt Sizes [15-16].

Покер планирования – это наиболее часто используемый в Scrum метод для выполнения оценки пользовательских историй. Покер планирования – это метод, основанный на экспертном суждении, который заключается в присвоении очков каждой пользовательской истории для определения ее сложности [14]. Каждая пользовательская история оценивается в очках, которые отражают сложность ее реализации в рамках проекта. При применении покера планирования основной проблемой является высокая субъективность в оценке пользовательских историй. Поэтому покер планирования считается сложным и ненадежным методом, особенно для членов команды без опыта. Кроме того, оценка каждого отдельного члена команды является расплывчатой, поскольку основана на обобщенном представлении о сложности, в то время как действительности представление о сложности связано с рядом атрибутов: опытом работы с использованием данного языка программирования [17-18], опытом работы в предыдущих проектах [19-20], знанием текущего проекта [21], опытом использования технологий [17, 22], размером команды [23-24] и т.д. [14, 25-26].

Основная проблема существующих методов заключается в том, что в них сложность пользовательских историй рассматривается как единое целое [14, 25], однако некоторые авторы предлагают разделять сложность на несколько атрибутов [8, 13, 26-27] и учитывать каждый из них при оценке. В данной статье мы предлагаем рассмотреть набор атрибутов, влияющих на оценку сложности пользовательских историй. Часть атрибутов была сформулирована на основе опубликованного нами ранее систематического обзора литературы [28]. Окончательный набор атрибутов прошел экспертную проверку. Мы предлагаем использовать байесовскую сеть (БС), в который узлы представлены атрибутами, со взвешиванием соответствующих ребер на основе мнения практикующих профессионалов. Мы использовали следующие критерии: опыт разработчика (опыт работы с технологиями, опыт работы с языком и опыт работы в предыдущих проектах), навыки разработчика, знание темы проекта и техническая сложность (функциональная зависимость, размер пользовательской истории (функции соединения, функции обработки данных)). Предложенная нами байесовская модель существенно облегчает процесс оценки. Во время оценки разработчики задают значения атрибутов, а байесовская модель вычисляет значение карты в очках.

Цель данного исследования состоит в том, чтобы предоставить командам разработчиков без опыта или накопленных данных метод, позволяющий существенно повысить точность оценки. Статья организована следующим образом. Разд. 2 содержит основные понятия. В разд. 3 представлены родственные работы. В разд. 4 описана методология исследования. В разд. 5 разбираются атрибуты, используемые в предлагаемой модели. Разд. 6 содержит информацию о процессе построения БС. В разд. 7 описано тестирование БС. В разд. 8 представлены результаты, в разд. 9 – выводы.

2. Основные понятия

2.1 Scrum

Scrum – самая популярная Agile-методология в индустрии программного обеспечения. Используя методы Scrum, многие компании улучшили качество результатов и продуктивность. Scrum – это Agile-метод разработки, который определяется как итеративный, инкрементный и эмпирический процесс управления проектом и контроля его выполнения [29-30]. В Scrum имеются три основные роли: ответственный за разработку продукта (Product Owner, PO), Scrum-мастер (Scrum Master, SM) и Scrum-команда (Scrum

Team, ST). РО является представителем заказчика. В Scrum имеются различные методы оценки, однако наиболее распространенным из них [8] является покер планирования.

2.2 Покер планирования

Покер планирования – это метод Scrum, использующийся для оценки командой разработчиков пользовательских историй с точки зрения трудозатрат. В игре фигурируют игральные карты с цифрами, которые основаны на модифицированной последовательности Фибоначчи (1, 2, 3, 5, 8, 13, 20, 40, 100). РО и все разработчики используют карты, чтобы достичь консенсуса в оценке требований в журнале пожеланий продукта (Product Backlog) [7]. Цель покера планирования заключается в измерении сложности пользовательских историй в очках [14].

2.3 Пользовательские истории

Пользовательская история представляет собой короткое текстовое описание, отражающее только основные элементы требования: для кого предназначен продукт, что ожидается от системы и, возможно, почему это важно [6]. Пользовательская история оценивается в очках, которые отражают сложность ее реализации [14][31].

2.4 Сложность

Когда мы говорим о сложности в контексте разработки программного обеспечения, речь может идти о сложности проблемы, на решение которой направлено разрабатываемое приложение, о структуре кода или об связях между элементами данных, которые будут обрабатываться в приложении. В случае пользовательских историй сложность кроется в различных элементах требуемой функциональности и аспектах их реализации. Определение сложности при оценке пользовательской истории – непростая задача, так как она включает в себя множество аспектов, эта проблема рассматривалась в нескольких исследованиях об использовании БС, вероятностный характер которых позволяет нам работать с неопределенностью [10].

2.5 Байесовская сеть

БС относятся к семейству вероятностных графовых моделей. Эти графовые структуры используются для представления знаний о неопределенной области. Каждый узел в графе представляет собой случайную переменную, а ребра – вероятностные зависимости между соответствующими случайными переменными.

3. Родственные исследования

Для повышения качества оценок использовались различные методы. В этом разделе мы описываем работы, связанные с оценкой сложности и имеющие отношение к нашему предложению.

В своем исследовании С. Драгичевич (Srdjana Dragicevic) [10] и др. представляют модель БС, пригодную для прогнозирования трудозатрат. Она включает шесть атрибутов: тип задачи, сложность требований (сложность форм, сложность функций и сложность отчетов), качество спецификаций и навыки разработчика. Такой подход принес хорошие результаты. Однако используемые атрибуты сосредоточены на технических аспектах, которые может быть сложно распознать начинающим разработчикам.

Модель Д. Алостад (Jasem M. Alostad) и др. [11] основана на нечеткой логике, которая может улучшить оценку трудозатрат в Scrum. Авторы рассматривают четыре атрибута: опытность команды разработчиков, сложность задачи, размер задачи и точность оценки. В результате уровень точности удалось поднять до 60%. Однако в этой модели унализировались только

четыре элемента, тогда как для определения сложности пользовательских историй требуется больше атрибутов.

В работк Д. Мартинес (Janeth López-Martínez) и др. [8] предлагается модель БС для оценки пользовательских историй, основанная на экспертных знаниях. В этом исследовании использовались пять атрибутов – время, трудозатраты, опыт, приоритетность и значимость пользовательской истории, каждый из которых представляет узел в сети. Этот подход позволил получить хорошую корреляцию между традиционными оценками и оценками, полученными с помощью предложенной модели. Основной недостаток этой модели заключается в то, что авторы рассматривали опыт как единый фактор.

Как видно, в большинстве родственных работ учтены важные аспекты оценки, но при этом они обладают двумя существенными недостатками: 1) нет единого мнения о значимости каждого фактора и уровне детализации композиции; 2) никто детально не фокусировался на индивидуальных особенностях разработчиков на мелкоструктурном уровне.

4. Методология

Шаг 1. Выбор атрибутов модели: атрибуты были идентифицированы, проанализированы и отобраны для интеграции в предложенную модель. В разд. 5 показан процесс выбора атрибутов.

Шаг 2. Построение БС: БС была построена с использованием выбранного набора атрибутов, составляющих количественные и качественные компоненты. В разд. 6 подробно описывается, как была построена сеть.

Шаг 3. Тестирование модели: для тестирования модели были разработаны и проведены два различных эксперимента: один для студентов и один для профессионалов. В разд. 7 описана процедура, использованная в обоих экспериментах.

Шаг 4. Результаты и их валидация: в разд. 8 приведены полученные результаты. Полученные результаты были проанализированы и сопоставлены с результатами других исследований.

5. Предлагаемая модель байесовской сети

В этом разделе объясняется выбор атрибутов, предлагаемых для использования в модели. Был проведен анализ атрибутов и классификаций из систематического обзора литературы [28], в котором собрано более 50 атрибутов различных категорий, способных повлиять на оценку пользовательских историй. Основываясь на мнениях профессионалов, мы отобрали наиболее релевантные из них: опыт работы с технологиями, опыт работы с используемым языком программирования, опыт работы в предыдущих проектах, навыки разработчика, знание темы проекта.

Хотя основное внимание в нашей модели уделяется индивидуальным атрибутам, мы добавили несколько атрибутов, связанных с технической сложностью пользовательских историй. Речь идет о таких важных и часто используемых в оценке атрибутах, указывающих на техническую сложность пользовательской истории, как размер и зависимость. Выбранные атрибуты первого уровня описаны в табл. 1. Атрибуты были сгруппированы в узлы в зависимости от их природы. Окончательная модель показана на рис. 1.

Табл. 1. Выбранные атрибуты

Table 1. Selected attributes

Атрибут	Определение
Опыт работы с технологиями	Разработчик знаком с задействованными инструментами, методами и технологиями и уверенно их использует.
Опыт работы в предыдущих проектах	Знания и навыки разработчика, накопленные благодаря участию в аналогичных проектах.

Опыт работы с языком	Разработчик знает и умеет эффективно использовать ресурсы и ограничения языка программирования.
Навыки разработчика	Знания и навыки, которыми обладает разработчик (разрешение проблем, доступность, производительность и т.д.)
Знание темы проекта	Разработчик понимает суть проекта и знаком с концепциями, составляющими тему проекта.
Зависимость	Степени важности пользовательской истории с точки зрения ее влияния на другие истории.
Функции соединения	Это требования заказчика к обработке – входные данные, выходные данные, запросы и т.д.
Функции обработки данных	Это требования заказчика к хранению (внутренние логические файлы и внешние интерфейсные файлы).

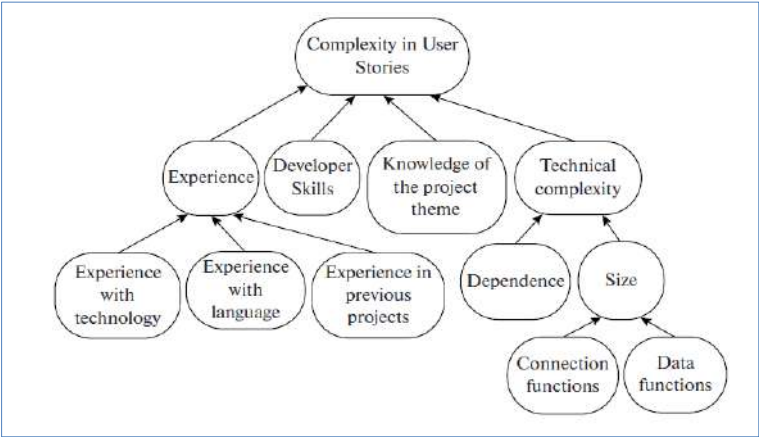


Рис. 1 Выбранная модель
Fig. 2.Selected model

6. Построение байесовской сети

6.1 Качественный компонент

Качественный компонент состоит из узлов и связей между ними, как показано на рис. 1. БС является многоуровневой иерархической структурой. Атрибуты первого уровня являются узлами, в которых собирается информация для оценки сложности; атрибуты второго уровня помогают организовать узлы иерархически; узел сложности пользовательских историй – это атрибут/узел, который показывает окончательное значение оценки сложности.

6.2 Количественный компонент

Чтобы сформировать количественный компонент, необходимо охватить три важных аспекта: значения связей между атрибутами, шкалы значений атрибутов первого уровня и построение таблицы условных вероятностей (Conditional Probability Table, CPT).

6.2.1 Значения отношений между атрибутами

Чтобы вывести значения связей между атрибутами, мы провели опрос среди профессионалов.

Опрос для подтверждения важности и значимости атрибутов модели

Цель опроса заключалась в валиации важности и значимости выбранных для модели атрибутов. Опрос проводился в цифровом формате. В опросе принял участие 21

профессионал и 19 студентов бакалавриата. При создании модели мы учитывали только ответы профессионалов. Мы подготовили вопросы для выяснения степени одобрения выбранных атрибутов, а также вопросы для сбора информации об опыте респондентов. Использовалась 5-балльная шкала Ликерта (Rensis Likert) от «полностью не согласен» до «полностью согласен».

Мы получили ответы 19 студентов, имеющих в среднем 15 месяцев опыта разработки программного обеспечения и 6 месяцев опыта использования Scrum. Мы также получили ответы от 21 профессионала, имеющих в среднем 57 месяцев опыта разработки программного обеспечения и 32 месяцев опыта использования Scrum. Обе группы использовали покер планирования для оценки пользовательских историй. Мы считаем, что участие 21 профессионала в опросе позволило нам собрать солидный объем данных в сравнении с исследованием [8], где было задействовано 16 профессионалов.

Применение уравнений

Для вычисления веса атрибутов мы использовали уравнения и определения, использованные в [8-9].

Определение 1. Пусть $V = \{v_1, \dots, v_n\}$ обозначает набор ответов профессионалов на один вопрос. Элемент $v_i = (weight, frequency)$ представляет вес одного варианта ответа по шкале Ликерта и частотность этого варианта среди всех экспертов. Формула (1) позволяет получить ненормализованное значение веса одного атрибута.

$$P_{p \in P} = \frac{\sum_{v \in V} (v_{weight} \times v_{frequency})}{\sum_{v \in V} v_{frequency}}. \quad (1)$$

Определение 2. Пусть $P = \{p_1, \dots, p_m\}$ представляет набор весов родительских узлов, которые относятся к одному и тому же дочернему узлу. Формула (2) обеспечивает нормализованные значения весов этих родительских атрибутов.

$$P_{norm_{p \in P}} = \frac{P}{\sum_{p \in P} P}. \quad (2)$$

Определение 3. Пусть $A = \{a_1, \dots, a_q\}$ – это набор атрибутов второго уровня. У них имеются два состояния – *low* и *high*. Их значения вычисляются с помощью формулы (3).

$$a_{a \in A} = \frac{\sum_{p \in P} P}{|P|}. \quad (3)$$

В табл. 2 приведены частоты и веса значений атрибутов в соответствии со шкалой Ликерта. В последней строке показаны нормализованные веса атрибутов.

Табл. 2. Веса и частота значений атрибутов

Table 2. Weights and frequencies of attributes

Веса	Шкала Ликерта	Частота			
		Опыт работы с технологиями	Опыт работы с языком	Опыт работы в предыдущих проектах	Навыки разработчика
0	1	0	0	0	0
0,25	2	0	0	0	0
0,5	3	1	1	4	2
0,75	4	8	9	10	13
1	5	12	11	7	6
Вес		0,347	0,343	0,310	0,253
Веса	Шкала Ликерта	Частота			
		Знание темы проекта	Зависимость	Функции соединения	Функции обработки данных
0	1	0	0	0	0

0,25	2	0	1	1	1
0,5	3	3	3	3	5
0,75	4	16	12	9	6
1	5	2	5	8	9
Bec		0,234	0,490	0,504	0,496

Рассмотрим атрибуты опыта, чтобы объяснить процесс взвешивания. Применим формулу (1) для получения ненормализованного веса каждого атрибута опыта первого уровня:
 $Exp_with_tech = \frac{0 \times 0 + 0.25 \times 0 + 0.5 \times 1 + 0.75 \times 8 + 1 \times 12}{21} = 0.881$, $Exp_with_lang = 0.869$,
 $Exp_in_prev_proj = 0.786$. Для получения нормализованного веса воспользуемся формулой (2). Необходимо рассмотреть узлы, являющиеся родителями одного и того же дочернего узла, в данном случае: $Exp_with_tech_norm = \frac{0.881}{0.881 + 0.869 + 0.786} = 0.347$,
 $Exp_lang_norm = 0.343$, $Exp_prev_norm = 0.310$. Чтобы получить вес узла опыта и других узлов второго уровня, применим формулу (3): $Experience = \frac{0.881 + 0.869 + 0.786}{3} = 0.845$.
Для получения нормализованного веса опыта необходимо учесть других родителей, которые влияют на один и тот же атрибут и его ненормализованный вес, в данном случае это навыки разработчика (0.798), знание темы проекта (0.738) и техническая сложность (0.770). Тогда $Experience_norm = \frac{0.845}{0.845 + 0.798 + 0.738 + 0.770} = 0.268$. На рис. 2 показаны результаты взвешивания узлов сети.

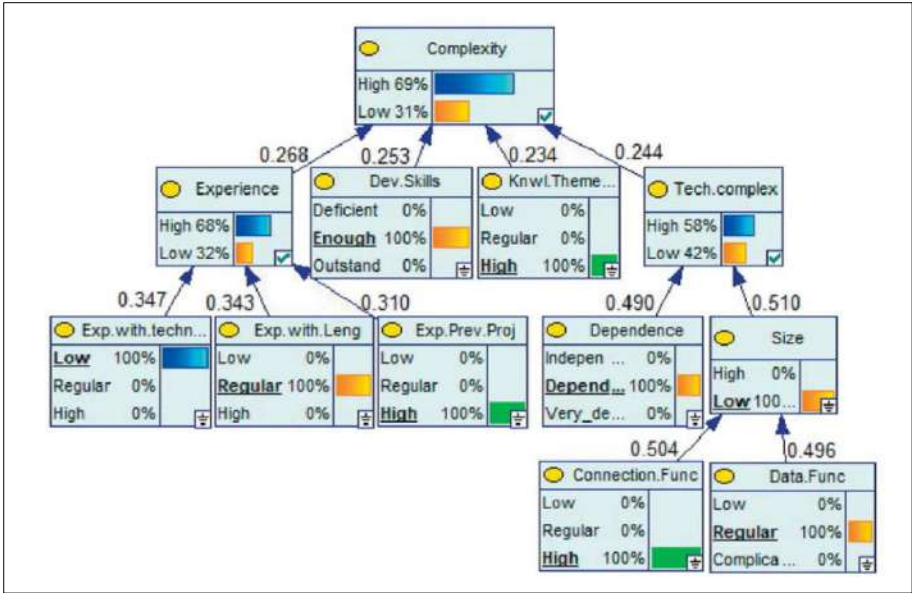


Рис. 3. Результирующие веса в байесовской сети
Fig. 4. Resulting weights in Bayesian network

Табл. 3. Вопросы для атрибутов первого уровня
Table 3. Questions for first level attributes

Опыт работы с технологиями	Как вы оцениваете свой опыт работы с технологиями, использующимися при создании пользовательской истории?	Незначительный	Я никогда не использовал технологии, с помощью которых будет разработана пользовательская история.
		Средний	Я знаком с технологиями, с помощью которых будет разработана пользовательская история, но не могу назвать себя уверенным пользователем.

		Значительный	Я уверенный пользователь технологий, с помощью которых будет разработана пользовательская история.
Опыт работы с языком	Как вы оцениваете свой опыт работы с языком программирования, использующимся при создании пользовательской истории?	Незначительный	Я никогда не имел дела с языком программирования, который будет использоваться при разработке пользовательской истории.
		Средний	Я знаком с языком программирования, который будет использоваться при разработке пользовательской истории, но не владею им уверенно.
		Значительный	Я уверенно владею языком программирования, который будет использоваться при разработке пользовательской истории.
Опыт работы в предыдущих проектах	Какой у вас опыт работы в проектах по разработке подобных пользовательских историй?	Незначительный	Я никогда не разрабатывал подобную пользовательскую историю.
		Средний	Я когда-то уже разрабатывал подобную пользовательскую историю.
		Значительный	Я разработал несколько подобных пользовательских историй.
Навыки разработчика	Как вы оцениваете свои навыки как разработчика в реализации пользовательских историй?	Недостаточные	Я не могу определить свои навыки, решение проблем дается мне с трудом.
		Достаточные	У меня есть определенные навыки, которые иногда приносятся мне в решении проблем.
		Выдающиеся	Я хорошо владею навыками решения проблем и умею своевременно их применять.
Знание темы проекта	Как вы оцениваете свою осведомленность о теме проекта, в рамках которого разрабатывается пользовательская история?	Незначительно	Я не знаю, с чем будет связана пользовательская история
		Средне	Я знаю основы темы проекта, в рамках которого разрабатывается пользовательская история.
		Значительно	Я хорошо ориентируюсь в теме проекта, в рамках которого разрабатывается пользовательская история.
Зависимость	Насколько разрабатываемая пользовательская история зависит от других?	Не зависит	Пользовательская история не зависит от других пользовательских историй.
		Зависит	Пользовательская история немного зависит от других пользовательских историй.
		Очень зависит	Пользовательская история является основой для разработки целого ряда пользовательских историй
Функции соединения	Насколько функции соединения добавляют сложности для разработки пользовательской истории?	Незначительно	Пользовательская история содержит простые функции input, output и query.
		Средне	Пользовательская история содержит функции input, output и query средней сложности.
		Значительно	Пользовательская история содержит сложные функции input, output и query
Функции обработки	Насколько функции	Незначительно	Пользовательская история содержит простые функции внутренней и

данных	обработки данных добавляют сложности для разработки пользовательской истории?		внешней логики.
		Средне	Пользовательская история содержит функции внутренней и внешней логики средней сложности.
		Значительно	Пользовательская история содержит сложные функции внутренней и внешней логики.

6.2.2 Априорные значения атрибутов первого уровня

Для каждого атрибута первого уровня был задан вопрос с тремя возможными ответами. В табл. 3 приведены вопросы, сгруппированные по атрибутам, и интерпретация ответов. На основе этих ответов для каждого атрибута первого уровня можно вычислить взвешенные показатели достоверности/важности количественных характеристик атрибутов. В этом исследовании мы не собирали соответствующую статистику и считали, что для всех узлов первого уровня эти показатели s одни и те же: $high\ s_1 = 0.500, regular\ s_2 = 0.330$, и $low\ s_3 = 0.170$.

6.2.3 Составление таблиц условных вероятностей

CPT показывают вероятности событий на основе комбинаций весов атрибутов и их показателей достоверности/важности [9]. CPT составляется для каждого атрибута/узла второго уровня сети. Для получения CPT применяются определения и формулы, используемые в [8-9].

Определение 4. Пусть $S = \{s_1, s_2, \dots, s_{|S|}\}$ – набор показателей достоверности/важности. Значение каждого атрибута первого уровня умножается на значение каждого показателя достоверности/важности: $p_1s_1, p \in P, s \in S$. Затем результаты первого шага для атрибутов первого уровня с общим узлом второго уровня объединяются: $p_1s_1 + p_2s_1 + \dots + p_ns_1$ и т.д. В общем виде элемент таблицы строится как

$$\sum_{1 \leq i_j \leq |S|, 1 \leq i \leq |P|} p_i s_{i_j}. \tag{4}$$

Для получения окончательных значений каждый элемента таблицы делится на максимальное значение элементов этой таблицы: $W = W/max$. Процесс повторяется для каждого атрибута второго уровня.

Рассмотрим атрибут второго уровня *size*, чтобы объяснить, как составлялась CPT для этого узла. Веса атрибутов первого уровня функции *connection functions* и *data functions* составили 0,504 и 0,496 соответственно (табл. 2). Для двух атрибутов CPT может быть представлена матрицей (5), для получения которой сначала необходимо умножить вес атрибута (P) на значение каждого показателя достоверности/важности (S) и рассмотреть все возможные последовательности атрибутов, образующих узел размера. Результат складывается по каждой последовательности и делится на максимальное значение матрицы для получения конечных значений:

$$\begin{bmatrix} 1.00 & 0.83 & 0.67 \\ 0.83 & 0.66 & 0.50 \\ 0.67 & 0.50 & 0.34 \end{bmatrix}. \tag{5}$$

7. Тестирование байесовской сети

7.1 Эксперименты

7.1.1 Выборка студентов

Цель этого эксперимента состояла в том, чтобы получить оценки индивидуальных пользовательских историй с помощью покера планирования и построенной сети. Была проведена выборка из 19 студентов (в 4 группах) факультета компьютерной инженерии Автономного университета Нижней Калифорнии (UABC) и 112 пользовательских историй. Студенты имели базовые знания о Scrum и покере планирования. Эксперимент состоял из 3 частей. Часть 1: Каждый член Scrum-команды оценивает каждую историю пользователя индивидуально при помощи покера планирования. Часть 2: Студенты отвечают на вопросы из табл. 3 для каждой пользовательской истории. Часть 3: Студенты оценивают истории традиционным способом. Эта часть включала повторную оценку пользовательской истории после ее выполнения.

7.1.2 Выборка профессионалов

Цель этого эксперимента состояла в том, чтобы получить оценки индивидуальных пользовательских историй с помощью покера планирования и построенной сети. Была проведена выборка из 6 профессионалов, работающих с методом Scrum и 10 пользовательских историй. Эксперимент проводился в форме цифрового опроса. Эксперимент состоял из 3 частей. Часть 1: Был выбран проект и 10 пользовательских историй, написанных студентами. Часть 2: Была предоставлена важная информация о проекте (описание проекта, язык программирования и т. д.), профессионалы должны были представить, что такая история будет разработана и оценить ее с помощью покера планирования с учетом предоставленной информации и на основе своего опыта. Часть 3: Специалистам было предложено ответить на вопросы для каждой пользовательской истории с учетом предоставленной информации.

8. Результаты и обсуждение

Краткое описание результатов эксперимента со студентами представлена в табл. 4, эксперимента с профессионалами – в табл. 5. Столбец с очками (первичная оценка) содержит результаты покера планирования до реализации пользовательской истории. Столбец с очками (финальная оценка) содержит результаты после реализации пользовательской истории. Столбец сложности представляет значение, полученное с помощью БС. Чтобы сравнить оценки в очках, выставленные с помощью покера планирования, с оценками, полученными с использованием БС, необходимо было преобразовать полученные показатели сложности в карту покера планирования, чтобы она была эквивалентна очками (SP). Были выбраны девять наиболее часто используемых карт в покере планирования: карта 1 = 1SP, карта 2 = 2SP, карта 3 = 3SP, карта 4 = 5SP, карта 5 = 8SP, карта 6 = 13SP, карта 7 = 20SP, карта 8 = 40SP и карта 9 = 100SP. Полученные показатели сложности были преобразованы в значения карт покера планирования с помощью формулы (6):

$$Y = Y_1 + \left[\left(\frac{X - X_1}{X_2 - X_1} \right) (Y_2 - Y_1) \right], \quad (6)$$

где Y – карта покера планирования, X – показатель сложности на основе БС, X_1 – наименьшее значение, которое может получиться при использовании БС, X_2 – наибольшее значение, которое может получиться при использовании БС, Y_1 – самое низкое значение карты, Y_2 – самое высокое значение карты. Столбец очков (БС) представляет эквивалентное значение в соответствии с картой, полученной после применения формулы (5). Проверка полученных результатов проводилась с помощью теста ранговой корреляции Спирмена [33]. В случае

студентов были проанализированы следующие атрибуты: *estimationIni*, *estimationFin* и *estimationBN*.

Табл. 4. Результаты эксперимента со студентами
Table 4. Students results

Член команды	Сложность	Карта (покер планирования)	Очки (первичная оценка)	Очки (финальная оценка)	Очки (БС)
Участник 1	75	4	5	5	5
Участник 2	67	3	3	5	3
Участник 3	75	4	8	5	5
Участник 4	82	5	5	5	8
Участник 5	75	4	5	5	5

Табл. 5. Результаты эксперимента с профессионалами
Table 5. Professionals results

Член команды	Сложность	Карта (покер планирования)	Очки (первичная оценка)	Очки (БС)
Участник 1	77	4	5	5
Участник 2	67	3	3	3
Участник 3	77	4	5	5
Участник 4	56	1	1	1
Участник 5	75	4	8	5
Участник 6	81	5	8	8

Тесты применялись с учетом отношений: *estimationIni-estimationBN* и *estimationFin-estimationBN*. По результатам теста Спирмена для отношения *estimationIni-estimationBN* была получена корреляция 0,659 и р-значение 0,000, а для отношения *estimationFin-estimationBN* – корреляция 0,799 и р-значение 0,000. Как видно, корреляция между финальной оценкой и оценкой, полученной с использованием байесовской сети больше, то есть оценки БС сильно приближены к оценкам студентов, уже закончивших работу над пользовательской историей. В случае профессионалов были проанализированы переменные: *estimationIni* и *estimationBN*. Тесты применялись для отношения *estimationIni-estimationBN*. По результатам теста Спирмена для отношения *estimationIni-estimationBN* была получена корреляция 0,887 и р-значение 0,000. Если сравнить результаты студентов и профессионалов, то можно заметить, что уровень корреляции выше у вторых за счет их опыта.

В нашем исследовании предлагается модель для оценки сложности пользовательских историй. Совпадение результатов оценки, полученных с помощью предложенной модели, с результатами оценки профессиональными разработчиками – 88%. Точность результатов модели выше в случае использования опытными профессионалами, то есть модели можно доверять, как профессиональному разработчику – такой результат вполне жизнеспособен.

Если сравнить корреляцию, полученную в эксперименте с профессионалами (88%), то предложенная нами модель более точна, чем модели, предложенные в исследованиях [8] (87%) и [11] (60%). Наша модель имеет больше сходства с моделью [8], но в ней опыт разработчика разбит на три подкатегории: опыт работы с технологиями, опыт работы с языком и с предыдущими проектами. На основе этих факторов можно провести более глубокий анализ опыта как одного из наиважнейших факторов в процессе оценки. В свою очередь, в исследовании [10] точность оценок достигает 90%, но сравнивать нашу модель с этим результатом нельзя, поскольку там декомпозиция сложности производилась с учетом технических подкатегорий. Несмотря на разный характер атрибутов, достигнутая нами точность была приближена к результатам [10].

9. Заключение

В настоящем исследовании был предложен метод для оценки пользовательских историй с помощью модели, состоящей из атрибутов с упором на индивидуальные качества разработчиков. С помощью этой модели мы стремились повысить точность оценки пользовательских историй. Мы взяли за основу ранее опубликованную нами статью [28] и проанализировали упомянутые в ней атрибуты, чтобы сформировать предложенную модель. С помощью анализа были отобраны наиболее релевантные из них: опыт работы с технологиями, опыт работы с языком и опыт работы в предыдущих проектах, навыки разработчика, знание темы проекта. Мы добавили несколько важных атрибутов, связанных с технической сложностью пользовательских историй: зависимость и размер.

Для тестирования нашей модели мы использовали две выборки респондентов – студентов и профессионалов. Для проверки полученных результатов мы применили корреляционные тесты к полученным оценкам. В случае со студентами результаты показывают, что данные, полученные с помощью сети, имеют схожесть с оценками, сформулированными после выполнения пользовательских историй. В эксперименте с профессионалами показатели корреляции выше, чем у студентов, что связано с опытом первых. В результате оценки, полученные с помощью нашей модели, на 88% совпадают с оценками профессиональных разработчиков, что говорит об успешности и достоверности исследования.

Сравнение результатов нашего исследования с похожими работами – исследованиями [8] с аналогичной методологией и [11] с методологией, отличной от нашей, но схожим подходом – показало, что наша модель дает более точные результаты, возможно, за счет того, что оно основано на индивидуальных возможностях разработчиков, имеющих большое влияние на процесс оценки. В нашей модели мы рассмотрели опыт, но не как единый фактор. Разбив на составляющие такой важный атрибут, как опыт, мы смогли повысить точность наших оценок. Результаты, полученные с помощью байесовской сети, обладают высокой достоверностью, поскольку они схожи с результатами оценки профессиональными разработчиками; успех обусловлен тем, что вес атрибутов в сети основан на мнении профессионалов, это помогает справляться с ситуацией неопределенности и неопытности разработчиков и повысить точность оценки. Метод будет особенно полезен в работе с неопытными разработчиками и поможет обрести уверенность в процессе оценки. В будущем планируется расширение модели для повышения точности оценок, а также увеличение выборки команд для получения еще более проверенных данных.

Список литературы / References

- [1] K.S. and J. Sutherland. The Scrum Guide. The Definitive Guide to Scrum: the Rules of the Game. Scrum.Org and Scrum Inc, 2017, 19 p.
- [2] A. Srivastava, S. Bhardwaj, and S. Saraswat. SCRUM model for agile methodology. In Proc. of the International Conference on Computing, Communication and Automation, 2017, pp. 864-869.
- [3] K.S. Rubin. Essential Scrum: A Practical Guide to the Most Popular Agile Process. Addison-Wesley Professional, 2012, 496 p.
- [4] D. Pauly, B. Michalik, and D. Basten. Do Daily Scrums Have to Take Place Each Day? A Case Study of Customized Scrum Principles at an E-commerce Company. In Proc. of the 48th Hawaii International Conference on System Sciences, Kauai, 2015, pp. 5074-5083.
- [5] A Guide to the Scrum Body of Knowledge (SBOK Guide). VMEdu, Inc, 2016. 140 p.
- [6] G. Lucassen, F. Dalpiaz, J.M.E.M. van der Werf, and S. Brinkkemper. Improving agile requirements: the Quality User Story framework and tool. Requirements Engineering, vol. 21, no. 3, 2016, pp. 383-403.
- [7] J. Lopez-Martinez, A. Ramirez-Noriega, R. Juarez-Ramirez et al. Analysis of Planning Poker Factors between University and Enterprise. In Proc. of the 5th International Conference in Software Engineering Research and Innovation, 2017, pp. 54–60.
- [8] J. López-Martínez, A. Ramírez-Noriega, R. Juárez-Ramírez et al. User stories complexity estimation using Bayesian networks for inexperienced developers. Cluster Computing, vol. 21, no. 2, 2018, pp. 715-728.

- [9] J. López-Martínez, R. Juárez-Ramírez, A. Ramírez-Noriega et al. Estimating User Stories' Complexity and Importance in Scrum with Bayesian Networks. *Advances in Intelligent Systems and Computing*, vol. 569, 2017, 205-214.
- [10] S. Dragicevic, S. Celar, and M. Turic. Bayesian network model for task effort estimation in agile software development. *Journal of Systems and Software*, vol. 127, 2017, pp. 109-119.
- [11] J.M. Alostad, L.R.A. Abdullah, and L.S. Aali. A Fuzzy based Model for Effort Estimation in Scrum Projects. *International Journal of Advanced Computer Science and Applications*, vol. 8, no. 9, 2017, pp. 270-277.
- [12] E. Scott and D. Pfahl. Using developers' features to estimate story points. In *Proc. of the 2018 International Conference on Software and System Process*, 2018, pp. 106-110.
- [13] V. Lenarduzzi. Could social factors influence the effort software estimation? In *Proc. of the 7th International Workshop on Social Software Engineering*, 2015, pp. 21-24.
- [14] V. Mahnič and T. Hovelja. On using planning poker for estimating user stories. *Journal of Systems and Software*, vol. 85, no. 9, 2012, pp. 2086-2095.
- [15] S.Z. Ziauddin, S.K. Tipu, and S. Zia. An Effort Estimation Model for Agile Software Development. *Advances in Computer Science and Its Applications*, vol. 2, no. 1, 2012, pp. 2166-2924.
- [16] B. Peischl, M. Zanker, M. Nica, and W. Schmid. Constraint-based recommendation for software project effort estimation. *Journal of Emerging Technologies in Web Intelligence*, vol. 2, no. 4, 2010, pp. 282-290.
- [17] F. Zare, H. Khademi Zare, and M.S. Fallahnezhad. Software effort estimation based on the optimal Bayesian belief network. *Applied Soft Computing*, vol. 49, 2016, pp. 968-980.
- [18] S.K. Rath, B.P. Acharya, and S.M. Satapathy. Early stage software effort estimation using random forest technique based on use case points. *IET Software*, vol. 10, no. 1, 2016, pp. 10-17.
- [19] S. C. Hsu, K.W. Weng, Q. Cui, and W. Rand. Understanding the complexity of project team member selection through agent-based modeling. *International Journal of Project Management*, vol. 34, no. 1, 2016, pp. 82-93.
- [20] O. Malgonde and K. Chari. An ensemble-based model for predicting agile software development effort. *Empirical Software Engineering*, vol. 24, issue 2, 2018, pp. 1017-1055.
- [21] R. Sholiq, S. Dewi, and A.P. Subriadi. A Comparative Study of Software Development Size Estimation Method: UCPabc vs Function Points. *Procedia Computer Science*, vol. 124, 2017, pp. 470-477.
- [22] A. Qazi, J. Quigley, A. Dickson, and K. Kirytopoulos. Project Complexity and Risk Management (ProCRiM): Towards modelling project complexity driven risk paths in construction projects *International Journal of Project Management*, vol. 34, no. 7, 2016, pp. 1183-1198.
- [23] C. Drahý and O. Pastor. Relationship Between Project Complexity and Risk Kinds Identified on Projects. *EMI – Economics, Management, Innovation*, vol. 8, no. 1, 2016, pp. 42-50.
- [24] S.M. Abdou, K. Yong, and M. Othman. Project Complexity Influence on Project management performance – The Malaysian perspective. In *Proc. of the 4th International Building Control Conference*, 2016, article no. 00065.
- [25] H. Zahraoui and M.A. Janati Idrissi. Adjusting story points calculation in scrum effort & time estimation. In *Proc. of the 10th International Conference on Intelligent Systems: Theories and Applications*, 2015, pp. 1-8.
- [26] E. Mendes. Knowledge Representation using Bayesian Networks A Case Study in Web Effort Estimation. In *Proc. of the World Congress on Information and Communication Technologies*, 2011, pp. 612-617.
- [27] Р. Сильхавы, З. Попова, П. Сильхавы. Алгоритмический метод оптимизации оценки трудозатрат. Программирование, том 42, no. 3, 2016 г., стр. 64-71 / R. Silhavy, Z. Prokopova, and P. Silhavy. Algorithmic optimization method for effort estimation. *Programming and Computer Software*, vol. 42, no. 3, 2016, pp. 161-166.
- [28] M. Durán, R. Juárez-Ramírez, S. Jiménez, and C. Tona. Taxonomy for Complexity Estimation in Agile Methodologies: A Systematic Literature Review. In *Proc. of the 7th International Conference in Software Engineering Research and Innovation*, 2020, pp. 87-96.
- [29] M. Usman, E. Mendes, F. Weidt, and R. Britto. Effort estimation in Agile Software Development: A systematic literature review. In *Proc. of the 10th International Conference on Predictive Models in Software Engineering*, pp. 82-91, 2014.
- [30] A. Georgsson. Introducing Story Points and User Stories to Perform Estimations in a Software Development Organisation – A case study at Swedbank IT. Master's Thesis, UMEÅ University, 2011, 52 p.
- [31] N. R. Jennings and M. Wooldridge. Agent-Oriented Software Engineering. *Artificial Intelligence*, vol. 117, 2000, pp. 277-296.

[32] F.V. Jensen. Bayesian networks. *WIREs Computational Statistics*, vol. 1, no. 3, 2009, pp. 307–315.

[33] L. Myers, M.J. Sirois. Spearman Correlation Coefficients, Differences between. In *Encyclopedia of Statistical Sciences*. John Wiley & Sons, 2006.

Информация об авторах / Information about authors

Майра ДУРАН – магистр, профессор. Область научных интересов: гибкие методологии, оценка программных проектов.

Mayra DURÁN, Master of Science, Professor. Research interests include agile methodologies, estimation of software projects.

Рейес ХУАРЕС-РАМИРЕС, кандидат компьютерных наук, профессор. Область научных интересов: программная инженерия, оценка неопределенности программного обеспечения и взаимодействие человека с компьютером.

Reyes JUÁREZ-RAMÍREZ, Doctor of Computer Science, Full Professor. Research interests include software Engineering, software uncertainty estimation, and human-computer interaction.

Саманта ХИМЕНЕС, доктор наук, профессор. Область научных интересов: программная инженерия, удобство использования, образовательные технологии, взаимодействие человека и компьютера.

Samantha JIMÉNEZ, Doctor of Science, Full Professor. Research interests include Software Engineering, Usability, Educational Technology, Human-Computer Interaction.

Клаудия ТОНА, магистр, профессор. Сфера научных интересов: программная инженерия, гибкие методологии, Scrum, гибкие фреймворки.

Claudia TONA, Master of Science, Professor. Research interests include Software Engineering, Agile Methodologies, Scrum, Agile Frameworks.



Новая интеллектуальная система для обнаружения сахарного диабета 2-го типа с модифицированной функцией потерь и регуляризацией

¹ М. Г. Ч. <Mallikagc9@gmail.com>

^{1,2,3,4} А. Альсадун, ORCID: 0000-0002-2309-3540 <AAlsadoon@studygroup.com>

⁵ Д.Т.Х. Фам, ORCID: 0000-0002-0813-827X <hangpdt@ued.udn.vn>

⁶ С.Х. Абдулла, ORCID: 0000-0001-7972-210X <120015@uotechnology.edu.iq>

⁵ Х.Т. Май, ORCID: 0000-0002-0813-827X <mhthi@ued.udn.vn>

¹ П.В. Ч. Прасад, ORCID: 0000-0002-3007-687X <CWithana@studygroup.com>

⁵ Ч.К.В. Нгуен, ORCID: 0000-0003-2281-0429 <ntquocvinh@ued.udn.vn>

¹ Университет Чарльза Стерта,
Австралия, NSW 2650, Вагга-Вагга, Бурома

² Университет Западного Сиднея,
Австралия, NSW 2000, Сидней, ул. Элизабет, 255

³ Университет Саутерн Кросс,
Австралия, NSW 2010, Сидней, Сарри Хиллс, ул. Мэри, 84-86

⁴ Азиатско-Тихоокеанский международный колледж,
Австралия, NSW 2150, Сидней, Парраматта, ул. Фицвильям, 1-3

⁵ Данангский университет – Университет науки и образования,
Вьетнам, 550000, Дананг, ул. Тон Дык Тханг, 459

⁶ Иракский технический университет,
Ирак, 10066, Багдад, ул. Аль Синаа

Аннотация. Сахарный диабет 2-го типа (СД2) составляет около 90% случаев диабета, и одним из ключевых аспектов СД2 являются жесткие требования к постоянному мониторингу и выявлению. Это исследование направлено на разработку ансамбля из нескольких моделей машинного и глубокого обучения для раннего обнаружения СД2 с высокой точностью. При большом разнообразии моделей ансамбль обеспечивает больше возможностей, чем отдельные модели. Предлагаемый ансамбль моделей основан на использовании известных моделей логистической регрессии, случайного леса, опорных векторов и глубокой нейронной сети. Выходные данные каждой модели в модифицированном ансамбле используются для определения окончательных выходных данных системы. Датасеты, используемые для этих моделей, включают Practice Fusion HER, Pima Indians diabetic's data, UCI AIM94 Dataset и CA Diabetes Prevalence 2014. По сравнению с ранее разработанными решениями, наше решение на основе ансамблевой модели демонстрирует высокие показатели точности, чувствительности и специфичности. В среднем обеспечиваются точность 87,5% от 83,51%, чувствительность 35,8% от 29,59% и специфичность 98,9% от 96,27%. Время работы предлагаемого решения составляет 96,6 мс, в то время как у наиболее по архитектуре известной системы – 97,5 мс. Предлагаемая усовершенствованная система улучшает возможности прогнозирования СД2 на основе использования ансамбля из нескольких моделей машинного и глубокого обучения. Для получения окончательного точного прогноза с использованием результатов отдельных моделей применяется схема мажоритарного голосования. В работе также изменена функция регуляризации, чтобы учесть регуляризацию всех

моделей в ансамбле, что помогает предотвратить переобучение и поддержать возможность обобщений в предлагаемой системе.

Ключевые слова: прогнозирование диабета 2-го типа; машинное обучение; ансамбль моделей; глубокие нейронные сети; метод опорных векторов; логистическая регрессия; случайный лес

Для цитирования: Г.Ч. М., Альсадун А., Фам Т.Х.Ф., Абдулла С.Х., Май Х.Е., Прасад П.В.Ч., Нгуен Ч.К.В. Новая интеллектуальная система для обнаружения сахарного диабета 2-го типа с модифицированной функцией потерь и регуляризацией. Труды ИСП РАН, том 33, вып. 2, 2021 г., стр. 93-114. DOI: 10.15514/ISPRAS-2021-33(2)-5

Благодарности. Это исследование частично поддержано Университетом Дананга – Университетом науки и образования, Вьетнам, в рамках гранта «T2020-TD-03-BS».

A Novel Intelligent System for Detection of Type 2 Diabetes with Modified Loss Function and Regularization

¹ M. G.C. <Mallikagc9@gmail.com>

^{1,2,3,4} A. Alsadoon, ORCID: 0000-0002-2309-3540 <AAlsadoon@studygroup.com>

⁵ D.T.H. Pham, ORCID: 0000-0002-0813-827X <hangpdt@ued.udn.vn>

⁶ S. Abdullah, ORCID: 0000-0001-7972-210X <120015@uotechnology.edu.iq>

⁵ H.T. Mai, ORCID: 0000-0002-0813-827X <mhthi@ued.udn.vn>

¹ P.W.C. Prasad, ORCID: 0000-0002-3007-687X <CWithana@studygroup.com>

⁵ T.Q.V. Nguen, ORCID: 0000-0003-2281-0429 <ntquocvinh@ued.udn.vn>

¹ Charles Sturt University,

Boorooma Street, Wagga Wagga, NSW 2650, Australia

² University of Western Sydney,

255 Elizabeth St, Sydney, NSW 2000, Australia

³ Southern Cross University

84-86 Mary St., Surry Hills, Sydney, NSW 2010, Australia,

⁴ Asia Pacific International College,

1-3 Fitzwilliam Street, Parramatta NSW 2150

⁵ The University of Da Nang – University of Science and Education,

459, Ton Duc Thang St., Danang City, 550000, Vietnam

⁶ University of Technology, Iraq,

Al Sinaa Street, Baghdad, 10066, Iraq

Abstract. Type 2 Diabetes (T2DM) makes up about 90% of diabetes cases, as well as tough restriction on continuous monitoring and detecting become one of key aspects in T2DM. This research aims to develop an ensemble of several machine learning and deep learning models for early detection of T2DM with high accuracy. With high diversity of models, the ensemble will provide more excessive performance than single models. *Methodology:* The proposed system is modified enhanced ensemble of machine learning models for T2DM prediction. It is composed of Logistic Regression, Random Forest, SVM and Deep Neural Network models to generate a modified ensemble model. *Results:* The output of each model in the modified ensemble is used to figure out the final output of the system. The datasets being used for these models include Practice Fusion HER, Pima Indians diabetic's data, UCI AIM94 Dataset and CA Diabetes Prevalence 2014. In comparison to the previous solutions, the proposed ensemble model solution exposes the effectiveness of accuracy, sensitivity, and specificity. It provides an accuracy of 87.5% from 83.51% in average, sensitivity of 35.8% from 29.59% as well as specificity of 98.9% from 96.27%. The processing time of the proposed model solution with 96.6ms is faster than the state-of-the-art with 97.5ms. *Conclusion:* The proposed modified enhanced system in this work improves the overall prediction capability of T2DM using an ensemble of several machine learning and deep learning models. A majority voting scheme utilizes the output from several models to make the final accurate prediction. Regularization function in this work is modified in order to include the regularization of all the models in ensemble, that helps prevent the overfitting and encourages the generalization capacity of the proposed system.

Keywords: T2DM Prediction, Machine Learning, Ensemble, Deep Neural Networks, SVM, Logistic Regression, Random Forests

For citation: G.C. M., Alsadoon A., Pham D.T.H., Abdullah S., Mai H.T., Prasad P.W.C., Nguen T.Q.V. A Novel Intelligent System for Detection of Type 2 Diabetes with Modified Loss Function and Regularization. *Trudy ISP RAN/Proc. ISP RAS*, vol. 33, issue 2, 2021, pp. 93-114 (in Russian). DOI: 10.15514/ISPRAS-2021-33(2)-5.

Acknowledgement. This research is partially supported by The University of Da Nang – University of Science and Education, Vietnam under grant “T2020-TD-03-BS”.

1. Введение

Модели машинного обучения на основе случайного леса (Random Forest) и опорных векторов (Support Vector Machine, SVM) нашли широкое применение в медицинских исследованиях, в то время как нейронные сети также являются очень мощными моделями, способными изучать сильные нелинейные отношения в данных. Эти модели могут быть объединены в ансамбль, в котором применяется каждый из алгоритмов для построения окончательного результата. Предоставление большой размерности и сложной взаимосвязи между признаками позволяет получить лучшие результаты, чем отдельная модель. Эта работа направлена на использование ансамбля моделей с использованием не только одного алгоритма, но и ансамбля алгоритмов для достижения улучшенных характеристик системы [9].

Среди нескольких работ, использующих машинное обучение для прогнозирования сахарного диабета, лучшая эффективность была продемонстрирована при применении подхода с использованием ансамбля различных методов [5]. В этой работе использовался ансамбль глубокой и широкой нейронной сети. Широкая часть сети обрабатывает статические признаки, в глубокая сеть – регулируемые признаки. Демонстрируется, что в ансамбле можно успешно использовать разные модели, способные обрабатывать различные типы признаков. Основным ограничением использования единого алгоритма является то, что модели стремятся к получению очень похожих решений [11]. Когда используются совсем разные модели, их разнообразные решения могут дать лучшие результаты.

В этой статье мы предлагаем систему для прогнозирования СД2, опирающуюся на преимущества разнообразия нескольких моделей машинного обучения, включая случайный лес, логистическую регрессию, опорные вектора и глубокую нейронную сеть. Для всех моделей использовалась регуляризация L2. Для получения окончательного результата прогноза использовался мажоритарный алгоритм голосования Бойера-Мура (Boyer–Moore majority vote algorithm). Эксперименты проводились на 4 датасетах – Practice Fusion EHR Dataset, Pima Indians Diabetics Dataset, UCI AIM94 Dataset и CA Diabetes Prevalence 2014. Средняя точность классификации составила 87,5% (83,21% минимум и 92,98% максимум), а среднее время обработки – 96,6 мс. Эти показатели лучше тех, которые демонстрируют аналогичные известные нам системы.

2. Обзор литературы

В этом разделе представлен обзор различных методов и подходов, применяемых в этой области. Нгуен (Binh P. Nguyen) и др. [5] разработали систему для диагностики СД2, основанную на глубоком и широком обучении с использованием электронной истории болезни людей. Глубокая и широкая модель – это гибридная модель линейной и глубокой нейронной сети, которой свойственны преимущества как линейных моделей, так и нейронных сетей. Авторы достигают современных результатов с точностью 84,28% и AUC 84,13%, что значительно выше, чем у других моделей. Несмотря на то, что модель очень

хорошо работает с ансамблем данных, ограничение этой работы состоит в том, что обычно неясно, какие функции подходят для глубокой и широкой части модели.

В аналогичной работе [12] авторы разработали модель, которая могла бы использовать данные о пациентах с состоянием, похожим на состояние пациента, данные о котором отсутствуют. Этот остроумный метод помогает в практических ситуациях, когда данные не всегда полностью и/или постоянно доступны. Неполные данные и периоды отсутствия данных ежедневного измерения уровня глюкозы в крови приводят к ограничению точности. Авторы предложили метод прогнозирования HbA1c, в котором сверточная нейронная сеть (Convolutional Neural Network, CNN) используется для извлечения признаков из данных временных рядов; затем эти признаки объединяются со статическими признаками пациентов, такими как возраст, пол и т.д., и все это передается в полносвязную нейронную сеть. Авторы сообщают об улучшениях по сравнению с обычными моделями CNN со средней абсолютной ошибкой 4,8. Однако основным ограничением этой работы является то, что ансамбль данных, возможно, очень мал для использования больших глубинных моделей.

Диабетическая ретинопатия – одна из ведущих причин слепоты в мире. Раннее выявление может обеспечить лечение и предотвратить слепоту, но очень сложно диагностировать заболевание на ранних стадиях. Авторы [13] реализовали сиамскую модель CNN, каждая из частей которой обрабатывает данные сканирования левого/правого глаза, чтобы обеспечить предсказания для левого и правого глаза индивидуально. Сиамская сеть обеспечивает признаки для обоих глаз, которые затем передаются в полноценную CNN. Это решение предоставило новый подход, который был побужден тем, как медицинский персонал проводит диагностику ретинопатии у диабетиков с помощью сканирований глаз. Сообщается о получении впечатляющего показателя AUC в 0,95. Однако ограничение предлагаемой модели состоит в том, что она использует очень крупную структуру модели (Inception v3) для каждого глаза, что значительно снижает скорость обучения модели и усложняет обучение.

Контроль уровня глюкозы в крови необходим для лечения диабета. В статье [4] авторы предложили модель, в которой CNN получает многомерные входные данные и преобразует их в данные временных рядов, полезные для прогнозирования на основе рекуррентных нейронных сетей (Recurrent Neural Network, RNN). Это является новым способом совместного использования статических и динамических признаков. Используемые входные данные представляют собой временные ряды уровней глюкозы каждые 30 минут. В статье представлены значения средней квадратичной ошибки (Root Mean Square Error, RMS Error, RMSE) и среднего абсолютного относительного отклонения (Mean Absolute Relative Difference, MARD), которые являются более хорошими, чем у базовых моделей. Однако для модели требуется большой объем данных через равные промежутки времени, что трудно получить в реальном мире.

Авторы работы [2] занимаются прогнозированием послеродового гестационного сахарного диабета (ГСД) с помощью тестов GCT и OGTT. В этой работе использовался алгоритм XGBoost, который представляет собой усиленную модель случайного леса. Авторы сообщают о точности 91% при специфичности и чувствительности 74%. Результаты были выдающимися, потому что данные подходили для этого алгоритма, а сам алгоритм XGBoost прекрасно с ними справлялся и производил результаты на уровне последних мировых достижений.

С увеличением доступности носимых устройств стал значительно более доступным непрерывный мониторинг физиологических и поведенческих особенностей людей. Работа [8] интересна еще и тем, что в модели использовались данные временных рядов из непрерывных измерений глюкозы и активности, а затем они объединялись с независимыми от времени демографическими данными аналогично тому, как это делается в передовой системе [5]. Однако в этой работе использовалась долгая краткосрочная память (Long Short-Term Memory; LSTM), обеспечивающая хорошую производительность при работе с данными временных рядов и значительную общую точность. Это решение отличается еще и тем, что в 96

нем объединяются глубокие и широкие аспекты машинного обучения для использования как временных рядов, так и статических данных.

Авторы работы [3] применили на практике линейный дискриминантный анализ (Linear Discriminant Analysis, LDA), чтобы определить, ел ли пациент с СД1. Для этого используются данные автоматического мониторинга уровня глюкозы. Однако LDA – это относительно простая и линейная модель машинного обучения [10]. Важными особенностями LDA являются линейность метода и потребность в выборе правильного порогового значения классификатора. Для определения правильного порога в [3] использовалось много специфических знаний предметной области.

В работе [14] нечеткие правила выводятся из нечетких деревьев решений, построенных с использованием алгоритма муравейника. Классификатор достиг точности 87,7% и чувствительности 92,2% на Pima Indian Diabetes.

Благодаря использованию ансамбля моделей в [5] были получены многие результаты современного мирового уровня. В своей работе авторы использовали десять различных глубоких и широких моделей и усредняли результаты каждой модели на десяти прогонах для получения окончательного результата. Значительно помогает то, что каждая модель обучается независимо от других моделей, а окончательная ансамблевая модель обобщает знания отдельных моделей.

В работе [15] авторы также предпочли использовать ансамблевый метод. Они занялись проблемой раннего выявления СД2 и гипертонии. Их цель заключалась в развитии модели прогнозирования заболеваний для обеспечения прогнозирования СД2 и гипертонии на основе данных о личных факторах риска. Авторы использовали комбинацию мощных методов и сообщают о впечатляющих результатах. Кроме того, их решение также показало, как можно эффективно обучать ансамбли с помощью проверки K-Cross.

В аналогичной работе [6] авторы продемонстрировали эффективную предварительную обработку данных наряду с использованием ансамбля сильных моделей, который обеспечивает результаты, лучшие, чем у отдельных моделей. Они указали, что использование в моделях прогнозирования СД2 нерегулярно дискретизированных данных приводит к снижению производительности по причинам, связанным с аппроксимацией. В реальном мире для предсказания СД2 очень сложно получить периодические данные для каждого пациента [1].

2.1 Наиболее совершенная система

В работе [5] представлен расширенный ансамбль глубоких и широких нейронных сетей для обнаружения ранних симптомов T2DM. Авторы используют данные электронных медицинских карт (ЭМК) пациентов и применяют современную глубокую и широкую архитектуру для повышения производительности.

На блок-схеме на рис. 1 представлены функции (в черных прямоугольниках) и ограничение (красный прямоугольник) системы прогнозирования СД2 [5], а также разработанная авторами гибридная модель линейной и глубокой нейронной сети (глубокой и широкой нейронной сети). Авторы утверждают, что глубокие и широкие нейронные сети подходят для использования фиксированных и настраиваемых признаков, которые распространены в датасетах о здоровье. Для обучения ансамбля моделей используется перекрестная проверка. Это решение показывает, как модели глубокого обучения могут использоваться для прогнозирования состояния здоровья и как эти модели могут применяться в ансамбле глубоких и широких сетей для получения самых современных результатов для прогнозирования СД2. Представленная модель дает точность 83% в среднем для 10 моделей и 84,28% для ансамбля из 10 моделей. Работа системы состоит из четырех этапов:

предварительная обработка, извлечение признаков, выбор и обучение моделей-членов ансамбля.

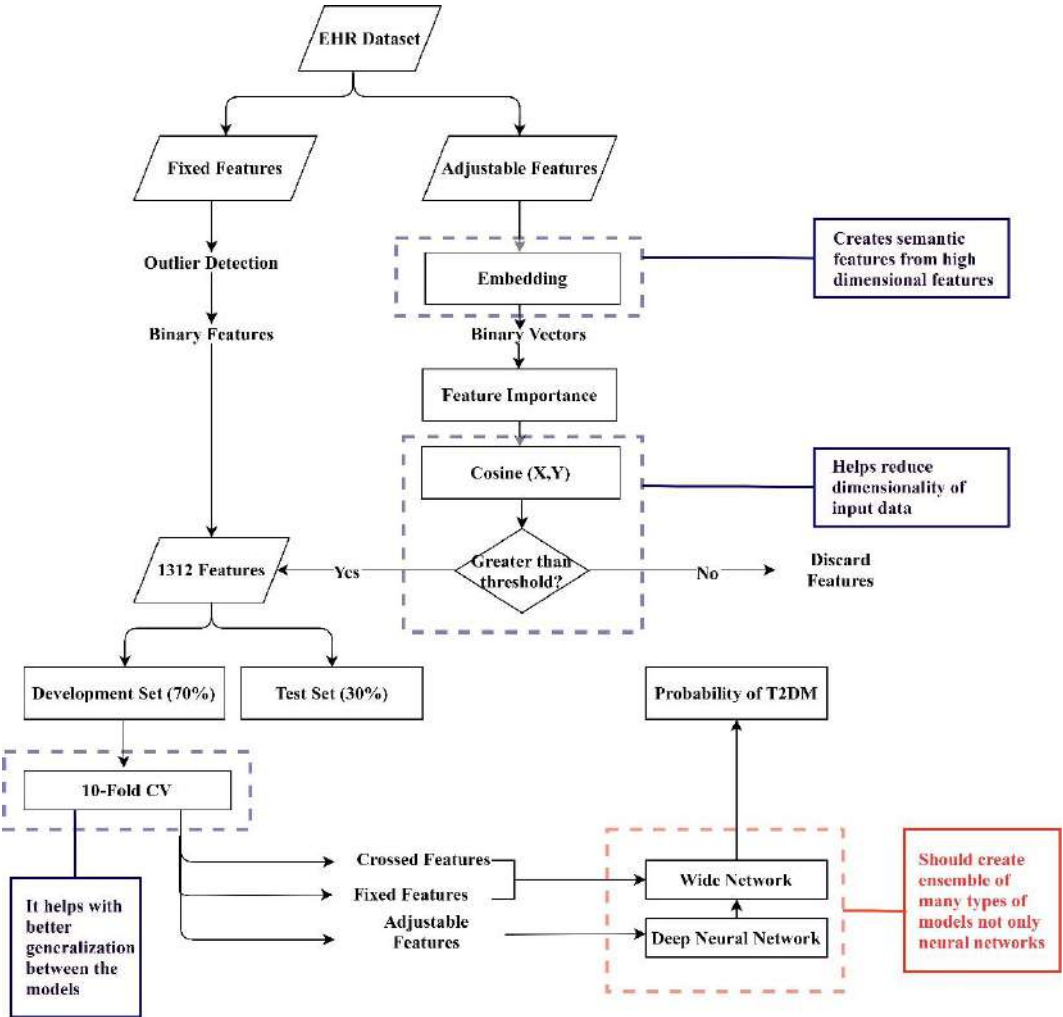


Рис. 1. Блок-схема системы [5]
Fig. 1. Block diagram of the state-of-the-art system [5]

Этап предварительной обработки: предварительная обработка данных, полученных из ЭМК пациентов. Датасет разделяется на фиксированные и настраиваемые признаки. К фиксированным признакам относятся статические характеристики пациентов – индекс массы тела (ИМТ), пол и т. д. Настраиваемые признаки собираются с датчиков и образцов. Фиксированные функции обрабатываются путем обнаружения и удаления выбросов. После этого они, где возможно, преобразуются в двоичные признаки.

Этап извлечения признаков: на этом этапе настраиваемые признаки проходят через слой вложений для построения семантических векторов. Затем их важность оценивается путем ранжирования в порядке косинусного расстояния от меток. Выбираются только наиболее важные признаки, отделяемые пороговыми значениями.

Этап выбора элементов ансамбля: на этом этапе датасет делится на поднаборы данных для разработки и тестирования сети с разделением 70/30. Настраиваемые признаки вводятся в нейронную сеть, состоящую из нескольких полностью связанных слоев с функциями активации ReLU. На заключительном выходном слое объединяется фиксированные и

перекрестные признаки. Объединенные признаки передаются на уровень логистической регрессии для получения окончательного результата в виде распределения вероятностей.

Этап обучения: на этом этапе строится ансамбль моделей. 10-кратная перекрестная проверка используется для создания 10 глубоких и широких моделей. Результаты этих моделей используются в схеме голосования для получения окончательного результата. Результат, который создается большинством моделей, выбирается как окончательный результат ансамбля.

Функция потерь, используемая в данной модели, представляет собой стандартную функцию кросс-энтропии (1):

$$L = \sum_N (-y \log(p)), \quad (1)$$

где L – функция потерь модели, N – количество образцов, y – истинная вероятность, p – прогнозируемая вероятность

Потеря регуляризации состоит в снижении веса, полученного с помощью параметра регуляризации, который является гиперпараметром модели. Это предотвращает переобучение моделей и способствует возможности обобщения моделей за счет ограничения их весов, как показано в формуле (2).

$$L_R = \lambda \times \sum_n |W|^2, \quad (2)$$

где L_R – потеря регуляризации, λ – параметр регуляризации, n – количество итераций, W – вес модели.

Табл. 1. Ансамбль моделей машинного обучения для прогнозирования начала СД2

Table1: Ensemble of machine learning models to predict onset of T2DM

Алгоритм: ансамбль моделей машинного обучения для прогнозирования начала СД2
Входные данные: 2D-матрица точек данных, где каждый столбец – это признак, а каждая строка – пациент
Выходные данные: вероятность обнаружения СД2 в диапазоне от 0,0 до 1,0
BEGIN
Шаг 1. Обнаружение выбросов: удалить данные для пациентов выше 95-го перцентиля.
Шаг 2. Баланс классов: определить вес каждого класса $w_i = \text{количество образцов класса} / \text{общее количество образцов}$.
Шаг 3. Нормализация: нормализовать каждый признак, чтобы он попал в диапазон (0,1)
Шаг 4. Обучить ансамбль моделей, минимизируя функцию потерь (1) и потери регуляризации (2).
Шаг 5. Использовать схему мажоритарного голосования для получения окончательного результата на основе результатов отдельных моделей.
END

3. Предлагаемая система

В нашем исследовании были проанализированы многие методы прогнозирования СД2. Основные проблемы, которые выделяются в статьях о прогнозировании СД2, – это предварительная обработка данных, выбор признаков и выбор модели.

Среди всех работ лучшим решением можно считать [5]. В ансамблевом методе для прогнозирования используется не одна модель, а коллекция моделей. Для учета всех различных получаемых результатов используется механизм голосования или усреднения. Достижению удовлетворительной производительности в целом помогает то, что разные алгоритмы подходит для выявления разных факторов и закономерностей в данных. Как сообщается в [5], эта система обеспечивают наилучшие точность и AUC ROC при решении задачи прогнозирования СД2.

Использованию ансамблевой модели для выявления сахарного диабета посвящена и впечатляющая работа [8], которую можно считать развитием [5] из-за разнообразия используемых моделей. Однако по своим результатам [8] не опережает [5]. Одна из основных

причин заключается в том, что авторы [8] не регуляризуют свою модель с помощью таких методов, как регуляризация L1/L2 и перекрестная проверка. В нашей работе эти ограничения преодолеваются путем применения соответствующих методов регуляризации. Предлагаемое нами решение состоит из четырех основных компонентов (рис. 2).

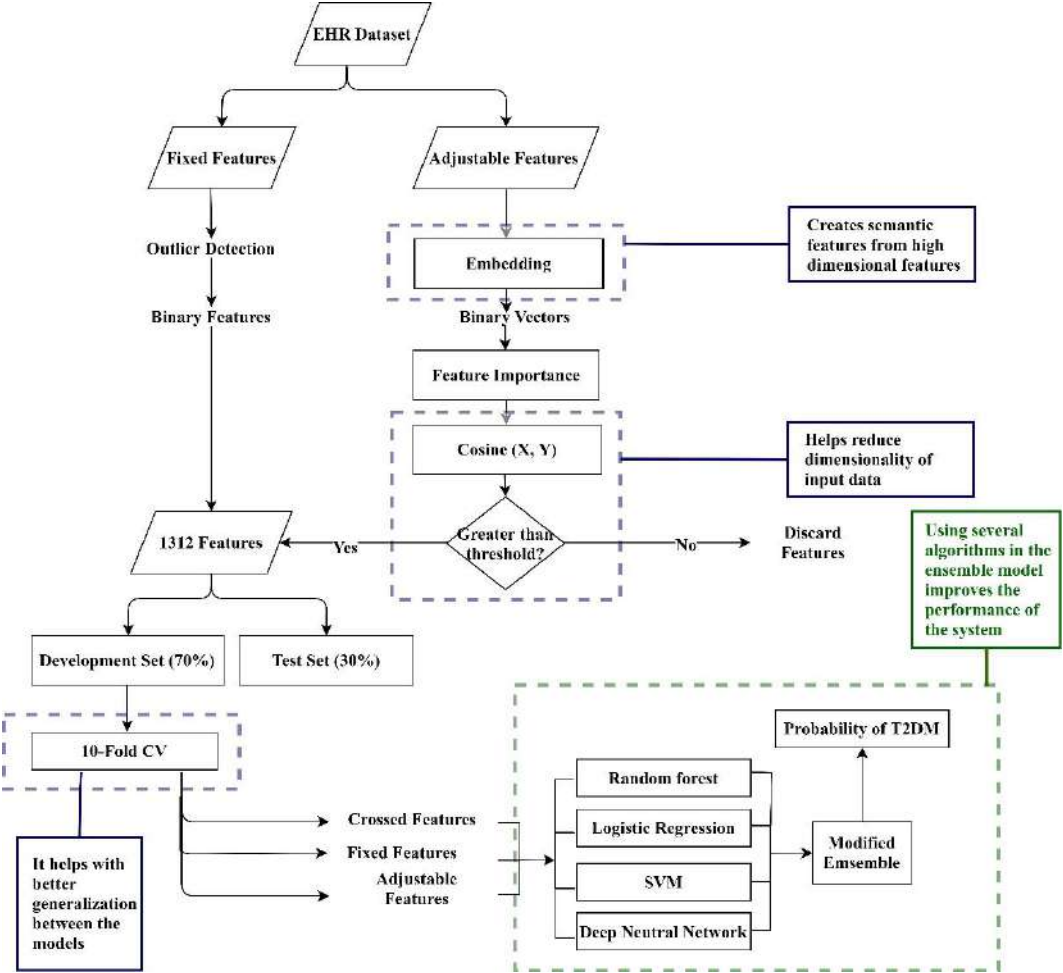


Рис. 2. Блок-схема предлагаемой системы
Зеленая рамка обозначает новые части в предлагаемой нами системе
Fig. 2. Block diagram of the purposed system
The green border refers to the new parts in our purposed system

Этап предварительной обработки: данные получаются из записей ЭМК пациентов, опубликованных компаниями Practice Fusion. Набор данных разбит на фиксированные и настраиваемые признаки, как в работе [5]. Фиксированные признаки включают статические характеристики пациента, такие как ИМТ, пол и т.д. Регулируемые признаки – это признаки, собираемые с датчиков и из образцов. Они обрабатываются путем обнаружения и удаления выбросов, а затем преобразуются в двоичные признаки, где это возможно.

Этап извлечения признаков: настраиваемые объекты проходят через слой встраивания для построения сематических векторов. Важность этих признаков оценивается путем ранжирования их в порядке косинусного расстояния от меток аналогично работе [5]. Выбираются только наиболее важные характеристики, отделенные пороговым значением. Однако в предлагаемом решении важность функции рассчитывается с использованием

алгоритма XGBoost, который опирается на использование коэффициента Джини расслоения деревьев решений в модели.

Этап выбора элементов ансамбля: датасет разделяется на наборы для разработки и тестирования с процентным соотношением 70/30, настраиваемые признаки передаются в нейронную сеть, состоящую из нескольких полностью связанных слоев с функциями активации ReLU. В заключительном выходном слое фиксированные признаки объединяются с перекрестными признаками. Затем объединенные признаки загружаются в слой логистической регрессии для получения окончательного результата в виде распределения вероятностей. Однако в предлагаемом решении все функции признаки во все модели ансамбле, который теперь включает не только глубокий и широкий компоненты.

Этап обучения: в работе [5] для построения ансамбля моделей используется 10-кратная перекрестная проверка для создания 10 глубоких и широких моделей. В предлагаемом нами решении с использованием 10-кратной перекрестной проверки создаются N моделей, и все они обучаются с использованием всех частей датасета. Затем результаты этих моделей используются в схеме голосования для получения окончательного результата – результата, полученного большинством моделей.

3.1 Предлагаемые формулы

Предлагаемые формулы ориентированы на ансамбль различных моделей машинного обучения, а не только на один тип модели [5]. Это увеличивает разнообразие решений, достигаемых каждой из разных моделей, повышая качество результатов ансамблевой модели. Формула (3) представляет собой функцию ансамблевых потерь. Она объединяет потерю регуляризации и потерю отдельной модели и является эффективной функцией потерь ансамбля. Это улучшает формулу (1): комбинирование потерь регуляризации отдельных моделей способствует оптимизации ансамбля. Наличие L_M уменьшает обобщения отдельных моделей:

$$L_E = \frac{1}{M} \sum_M (L_R + L_M), \quad (3)$$

где L_E – потеря ансамбля, M – количество моделей, L_R – потеря регуляризации каждой модели, L_M – потеря каждой модели.

Следующая формула (4) представляет собой функцию потери регуляризации всего ансамбля. Она помогает повысить степень обобщения моделей на протяжении всего процесса оптимизации:

$$ML_E = \frac{1}{M} \sum_M L_R, \quad (4)$$

где ML_E – измененная потеря регуляризации ансамбля, M – количество моделей, L_R – потеря регуляризации каждой модели.

Потеря модели ансамбля представляет собой функцию кросс-энтропии, которая эквивалентна сумме энтропии набора данных и расстояния Кўльбака-Лёйблера (Kullback-Leibler divergence) между истинными метками и прогнозами модели, как показывает формула (5):

$$ML = ML_E * \sum_N (-y \log(p)), \quad (5)$$

где ML – модифицированная функция потерь модели, ML_E – измененная потеря регуляризации ансамбля, N – количество образцов, y – истинная вероятность, p – прогнозируемая вероятность.

Следующая формула (6) заменяет потерю регуляризации $L1$ на потерю регуляризации $L2$, которая, как было эмпирически показано, обеспечивает большее обобщение:

$$R = \frac{\lambda}{2n} \sum_n |W|, \quad (6)$$

где R – регуляризация, n – количество итераций, W – веса нейронов, λ – параметр регуляризации, $\frac{\lambda}{2n}$ – масштабный коэффициент.

Потеря регуляризации масштабируется по следующей формуле. Параметр регуляризации – это гиперпараметр, который оптимизируется во время обучения, как показано в (7).

$$MR = \frac{\lambda}{2n}, \quad (7)$$

где λ – параметр регуляризации, $\frac{\lambda}{2n}$ – масштабный коэффициент.

В (8) формула (2) была изменено с использованием (6) и (7): потеря регуляризации определяется измененной регуляризацией и количеством итераций, производимых для взвешенной модели, чтобы увеличить разнообразие и обобщение модели.

$$ML_R = MR * \sum_n |W|^2, \quad (8)$$

где ML_R – потеря регуляризации, MR – модифицированная регуляризация, n количество итераций, W – вес модели.

Увеличение потерь в ансамбле за счет объединения модифицированной функции потерь модели и модифицированной потери регуляризации служит для увеличения разнообразия модели. Предлагаемая формула (9) представлена ниже:

$$EEL = ML + ML_R, \quad (9)$$

где EEL – усовершенствованная потеря ансамбля, ML – модифицированная функция потерь модели, ML_R – модифицированная потеря регуляризации.

3.2 Суть предлагаемого подхода

В предлагаемом решении, во-первых, используется ансамбль различных моделей машинного обучения вместо одного типа модели [5]. Это увеличивает разнообразие решений, достигаемых каждой из различных моделей, увеличивая качество за счет объединения результатов в ансамблевой модели.

Во-вторых, для всех моделей ансамбля используется регуляризация $L2$, чтобы обеспечить обобщение и предотвратить переобучение. Гарантируется, что ансамбль во время обучения приобретает разнообразные способности к обобщению.

Наконец, наряду с $L2$ -регуляризацией, используемой для моделей, модель нейронной сети в ансамбле также регуляризована путем своевременного прекращения обучения, чтобы избежать переобучения. В результате модели обучаются обобщению за счет дополнительной параметризации без потребности в изучении статистических нюансов набора данных. Улучшается способность к обобщению всего ансамбля.

В нашей системе использовалось несколько моделей вместо одной модели нейронной сети, чтобы снять ограничения на использование системы за счет улучшения результатов. В системе проводится голосование за различные результаты, и выбирается лучший из них по точности и AUC ROC. Предлагаемая система может диагностировать СД2, сводя к минимуму ограничения и обеспечивая дает высокую точность, чувствительность и специфичность.

На этапе обзора литературы мы поняли, что для диагностики СД2 лучше всего подходит ансамблевый метод. Однако в имеющихся решениях для снятия ограничения используется только модель нейронной сети. В своей работе мы пытаемся дополнительно уменьшить ошибку, чтобы повысить точность и получить лучшие результаты. Качественные результаты

можно получать путем обработки результатов многих моделей, таких как Random Forest, логистическая регрессия, SVM, глубокая нейронная сеть, заставляя ансамблевую модель голосовать за лучший результат, который будет выбран в качестве окончательного (табл. 2).

Табл. 2. Набор моделей машинного обучения для прогнозирования начала СД2

Table 2: Ensemble of machine learning models to predict onset of T2DM

Алгоритм: ансамбль моделей машинного обучения для прогнозирования начала СД2. Вход: данные ЭМК пациентов. Выход: Выходные данные классификации: 1 для пациентов с СД2 и 0 для пациентов без СД2.
BEGIN ШАГ 1. Обнаружение выбросов: удалить данные для пациентов выше 95 перцентиля. ШАГ 2. Баланс классов: определите вес каждого класса $w_t = \text{количество образцов класса} / \text{общее количество образцов}$. ШАГ 3. Нормализация: нормализовать каждый признак: $x = x - \text{mean}(x) / \text{std}(x)$. ШАГ 4. Обучить сумку из N моделей. Модель логистической регрессии с функцией потерь: $L_{LR} = -(y \log(p) + (1 - y) \log(1 - p))$ Модель SVM с функцией потерь: $L_{SVM} = \sum_{j \neq y_i} \max(0, s_j - s_{y_i} + \Delta)$ Модель Random forest с функцией потерь: $L_{RF} = \sum p(i) * (1 - p(i))$ Нейронная сеть с функцией кросс-энтропии: $L_{NN} = \sum_i (y_i \log(p_i))$ ШАГ 5. Окончательный вывод ансамблевой модели с использованием алгоритма мажоритарного голосования Бойера-Мура. END

4. Результаты и обсуждение

Среда, использованная для проведения экспериментов в данной работе, основана на языке программирования Python. Python обеспечивает экосистему для работы с машинным обучением и анализом данных. В нашем решении использовались четыре разных датасета, как показано в табл. 3.

Табл. 3. Наборы данных

Table 3. Datasets

Название датасета	Общее количество образцов	Количество положительных образцов	Количество отрицательных образцов
Practice Fusion EHR	9948	1904	8044
Pima Indians Diabetics	2500	1500	1000
UCI AIM94	2000	1000	1000
CA Diabetes Prevalence 2014	6000	1000	5000

Табл. 4. Средние значения показателей

Table 4. Average performance

	Чувствительность (%)	Специфичность (%)	Точность (%)	Время обработки (мс)
SOTA	29.59	96.27	83.51	97.5
Предлагаемое решение	35.8	98.9	87.5	96.6

Сначала сравним показатели результатов предлагаемого решения и наиболее совершенной известной нам системы (state-of-the-art, SOTA). Наша модель достигла точности 87,5% по сравнению с 83,51% у системы SOTA, демонстрируя эффективность предлагаемого решения. Среднее время обработки составило 96,6 мс, как показано в табл. 4.

Образцы создавались для разных возрастных групп мужчин и женщин, страдающих диабетом 2 типа. Распределение по возрасту и положительные/отрицательные случаи показаны на рис. 3. Распределение гендерного признака в датасете показано на рис. 4.

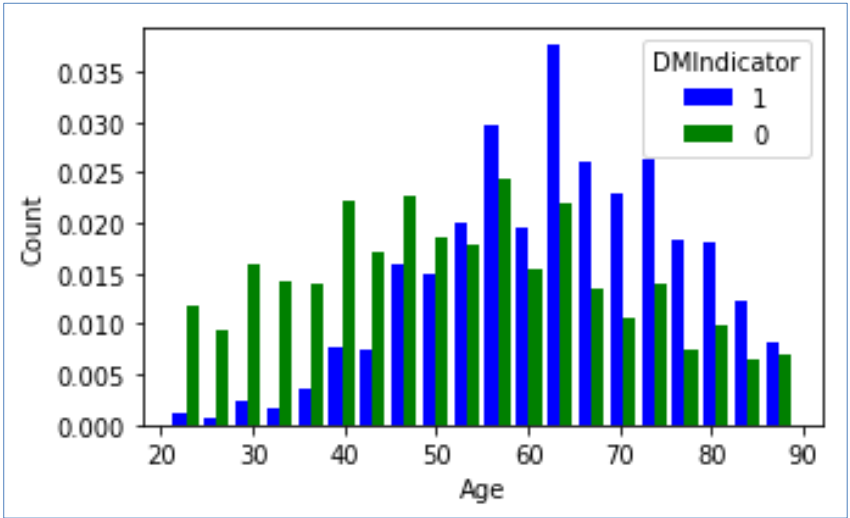


Рис. 3: Распределение по возрасту и положительные/отрицательные случаи
Fig. 3: Distribution of age and the positive/negative cases

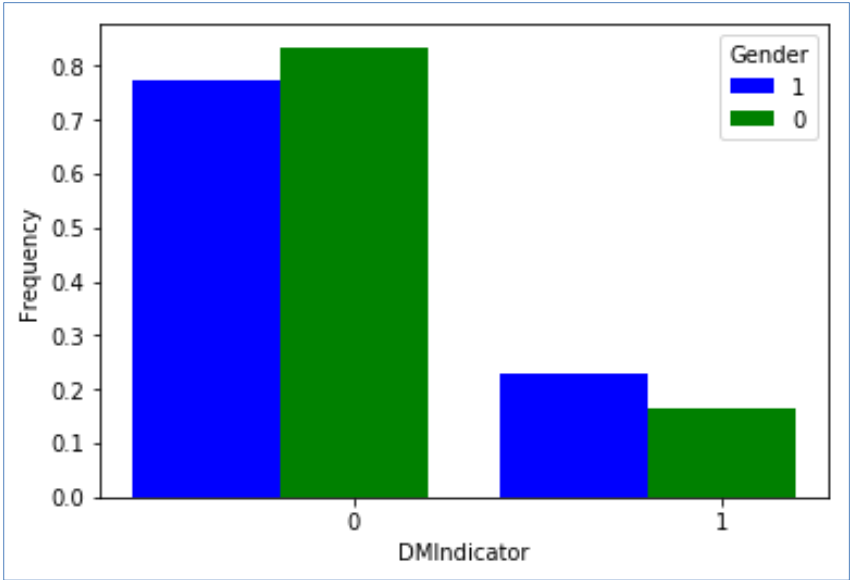


Рис. 4: Гендерный признак в датасете
Fig. 4: Gender feature in the dataset

Образцы данных использовались для обучения с использованием языка программирования Python и соответствующих инструментов анализа данных, таких как Scikit-learn, Pandas и Matplotlib. Сравнивались результаты (точность, чувствительность и специфичность) предложенного решения на основе ансамблевой модели и глубоким и широким решением SOTA. В табл. 5-12 приводятся средние результаты обеих систем, полученные при тестировании для 10 различных групп.

Табл. 5. Сравнение точности, чувствительности и специфичности SOTA и предлагаемого решения с использованием датасета Practice Fusion ЭМК (группа выборки 1: женщины в возрасте от 40 до 70 лет)

Table 5. Accuracy, Sensitivity and Specificity results comparison of SOTA (state-of-the-art) vs Proposed Solution using Practice Fusion EHR Dataset (Sample group 1: women of age 40 to 70)

№ выборки	Решение SOTA			
	Точность (%)	Чувствительность (%)	Специфичность (%)	Время обработки (мс)
1	82.98	27.22	97.12	97.5
2	83.01	27.23	97.23	97.3
3	83.4	26.56	97.04	97.5
4	82.93	27.23	97.33	97.4
5	83.24	27.04	97.27	97.3
6	83.21	27.52	97.1	97.3
7	83.41	26.8	97.22	97.4
8	82.99	27.1	97.02	97.2
9	83.14	27.23	97.3	97.3
10	83.12	26.82	97.12	97.4
№ выборки	Предлагаемое решение			
	Точность (%)	Чувствительность (%)	Специфичность (%)	Время обработки (мс)
1	84.72	28.03	97.36	96.7
2	85.23	27.64	97.22	96.5
3	84.16	27.14	97.26	96.7
4	84.52	28.09	97.15	96.6
5	84.11	27.15	98.15	96.6
6	85.25	27.83	97.54	96.5
7	84.13	27.94	97.61	96.7
8	84.81	27.42	97.42	96.5
9	84.33	27.25	97.73	96.6
10	84.14	27.73	97.59	96.6

Табл. 6. Сравнение точности, чувствительности и специфичности SOTA и предлагаемого решения с использованием датасета Practice Fusion ЭМК (группа выборки 2: мужчины в возрасте от 30 до 60 лет)

Table 6. Accuracy, Sensitivity and Specificity results comparison of SOTA vs Proposed Solution using Practice Fusion EHR Dataset (Sample group 2: men of age 30 to 60)

№ выборки	Решение SOTA			
	Точность (%)	Чувствительность (%)	Специфичность (%)	Время обработки (мс)
1	82.71	26.71	97.07	97.5
2	82.99	26.61	97.14	97.2
3	82.80	27.22	97.13	97.6
4	83.44	26.77	97.35	97.5
5	82.85	27.08	97.40	97.3
6	83.30	27.05	97.38	97.5
7	82.70	27.04	97.12	97.4
8	83.45	26.61	97.36	97.4
9	83.48	26.84	96.71	97.4
10	83.27	26.82	96.78	97.5

№ выборки	Предлагаемое решение			
	Точность (%)	Чувствительность (%)	Специфичность (%)	Время обработки (мс)
1	84.91	28.16	97.81	96.7
2	85.07	28.09	98.20	96.5
3	84.66	27.54	97.88	96.7
4	84.90	27.39	98.27	96.6
5	85.14	27.11	98.23	96.6
6	85.11	27.28	97.91	96.5
7	83.95	27.95	98.05	96.7
8	84.62	27.51	97.87	96.5
9	85.09	27.52	98.30	96.6
10	84.13	28.17	97.85	96.6

Табл. 7. Сравнение точности, чувствительности и специфичности SOTA и предлагаемого решения с использованием датасета Prima Indians (группа выборки 1: женщины в возрасте от 40 до 70 лет)
Table 7. Accuracy, Sensitivity and Specificity results comparison of SOTA vs Proposed Solution using Prima Indians Dataset (Sample group 1: Women of age 40 to 70)

№ выборки	Решение SOTA			
	Точность (%)	Чувствительность (%)	Специфичность (%)	Время обработки (мс)
1	82.30	27.11	96.67	43.1
2	83.26	26.72	97.13	43.0
3	82.59	26.77	97.38	43.2
4	83.13	27.16	97.03	43.3
5	83.05	26.81	97.10	43.0
6	82.74	26.72	97.07	43.3
7	81.68	27.11	97.16	43.1
8	82.17	26.72	97.25	43.2
9	81.24	26.71	96.79	43.3
10	82.82	27.27	96.79	43.4
№ выборки	Предлагаемое решение			
	Точность (%)	Чувствительность (%)	Специфичность (%)	Время обработки (мс)
1	83.83	27.37	98.28	42.9
2	83.75	27.78	97.83	42.5
3	84.11	27.29	98.10	42.9
4	84.19	27.14	97.93	42.5
5	84.10	27.37	98.03	42.9
6	84.31	27.17	97.89	42.6
7	84.29	27.71	98.22	42.6
8	83.23	27.33	98.18	42.9
9	84.36	27.17	98.05	42.6
10	83.64	27.48	97.85	43.0

Табл. 8. Сравнение точности, чувствительности и специфичности SOTA и предлагаемого решения с использованием датасета Prima Indians (группа выборки 2: мужчины в возрасте от 30 до 60 лет)
Table 8. Accuracy, Sensitivity and Specificity results comparison of SOTA vs Proposed Solution using Prima Indians Dataset (Sample group 2: Men of age 30 to 60)

№ выборки	Решение SOTA			
	Точность (%)	Чувствительность (%)	Специфичность (%)	Время обработки (мс)
1	82.99	27.26	97.16	43.1
2	82.68	26.87	97.06	43.0
3	82.98	27.15	97.09	43.2
4	81.24	27.05	96.90	43.3
5	82.43	27.25	96.85	43.0
6	81.34	26.80	96.74	43.3
7	81.69	27.08	97.27	43.1
8	82.94	27.12	96.78	43.2
9	81.26	27.10	96.88	43.3
10	83.03	26.97	97.18	43.4
№ выборки	Предлагаемое решение			
	Точность (%)	Чувствительность (%)	Специфичность (%)	Время обработки (мс)
1	83.65	27.34	98.18	42.9
2	84.32	27.56	97.97	42.5
3	83.84	27.82	98.04	42.9
4	83.84	28.15	97.84	42.5
5	83.37	27.77	98.07	42.9
6	84.28	27.12	97.80	42.6
7	84.28	27.42	98.00	42.6
8	83.48	27.52	98.19	42.9
9	83.21	28.17	98.09	42.6
10	83.29	27.67	97.90	43.0

Табл. 9. Сравнение точности, чувствительности и специфичности SOTA и предлагаемого решения с использованием датасета UCI AIM94 (группа выборки 1: женщины в возрасте от 40 до 70 лет)
Table 9. Accuracy, Sensitivity and Specificity results comparison of SOTA vs Proposed Solution using UCI AIM94 Dataset (Sample group 1: women of age 40 to 70)

№ выборки	Решение SOTA			
	Точность (%)	Чувствительность (%)	Специфичность (%)	Время обработки (мс)
1	87.14	26.90	97.31	42.8
2	88.06	26.74	96.77	42.8
3	87.87	27.22	96.69	43.0
4	87.12	27.18	97.04	42.7
5	88.19	26.73	97.29	43.0
6	87.24	26.74	97.11	42.6
7	87.97	27.29	97.22	42.9
8	87.53	26.80	97.34	42.8

9	87.50	27.21	97.08	42.6
10	88.24	27.27	96.96	42.9
№ выборки	Предлагаемое решение			
	Точность (%)	Чувствительность (%)	Специфичность (%)	Время обработки (мс)
1	87.95	28.03	97.96	42.4
2	88.79	27.15	97.93	42.1
3	88.22	28.01	97.87	42.3
4	88.14	27.16	97.82	42.3
5	88.88	28.14	98.19	42.4
6	88.37	28.11	98.10	42.2
7	89.14	27.83	98.24	42.1
8	89.12	27.93	98.30	42.3
9	89.04	27.55	98.13	42.4
10	88.77	28.15	97.99	42.1

Табл. 10. Сравнение точности, чувствительности и специфичности SOTA и предлагаемого решения с использованием датасета UCI AIM94 (группа выборки 2: мужчины в возрасте от 30 до 60 лет)
Table 10. Accuracy, Sensitivity and Specificity results comparison of SOTA vs Proposed Solution using UCI AIM94 Dataset (Sample group 2: men of age 30 to 60)

№ выборки	Решение SOTA			
	Точность (%)	Чувствительность (%)	Специфичность (%)	Время обработки (мс)
1	87.14	27.16	96.87	43.0
2	88.06	27.15	97.28	42.8
3	87.87	27.15	96.66	42.6
4	87.12	27.26	97.09	42.7
5	88.19	26.99	96.60	42.8
6	87.24	27.11	96.62	42.5
7	87.97	27.10	97.19	43.0
8	87.53	26.98	97.31	42.8
9	87.50	26.82	97.36	42.7
10	88.24	26.96	97.11	42.6
№ выборки	Предлагаемое решение			
	Точность (%)	Чувствительность (%)	Специфичность (%)	Время обработки (мс)
1	87.99	28.08	98.06	42.4
2	89.39	27.79	98.12	42.3
3	88.64	28.11	97.86	42.2
4	89.00	27.20	97.91	42.4
5	88.34	28.17	98.12	42.3
6	88.77	27.46	97.88	42.4
7	89.32	27.65	98.27	42.1
8	88.55	27.16	97.84	42.2
9	88.01	27.12	97.87	42.5
10	88.49	27.54	98.05	42.4

Табл. 11. Сравнение точности, чувствительности и специфичности SOTA и предлагаемого решения с использованием датасета CA Diabetes Prevalence 2014 (группа выборки 1: женщины в возрасте от 40 до 70 лет)

Table 11. Accuracy, Sensitivity and Specificity results comparison of SOTA vs Proposed Solution using CA Diabetes Prevalence 2014 Dataset (Sample group 1: Women of age 40 to 70)

№ выборки	Решение SOTA			
	Точность (%)	Чувствительность (%)	Специфичность (%)	Время обработки (мс)
1	90.96	50.71	97.05	69.0
2	91.30	51.13	96.62	69.3
3	90.55	50.58	97.01	69.1
4	90.26	50.67	96.76	69.0
5	90.12	51.32	96.60	69.0
6	90.65	51.88	96.66	69.2
7	90.03	52.76	96.75	69.1
8	91.30	50.94	97.10	69.0
9	90.50	51.09	97.08	69.2
10	90.38	51.13	97.12	69.1
№ выборки	Предлагаемое решение			
	Точность (%)	Чувствительность (%)	Специфичность (%)	Время обработки (мс)
1	91.78	27.47	97.99	69.1
2	92.87	27.47	97.85	68.7
3	92.31	27.72	98.25	68.6
4	91.04	27.49	98.25	68.9
5	92.26	27.97	97.96	68.6
6	91.62	27.25	97.82	68.5
7	92.18	28.00	98.01	68.7
8	92.92	27.94	97.85	68.7
9	92.75	27.37	97.91	68.9
10	91.43	27.92	98.15	68.8

Табл. 12. Сравнение точности, чувствительности и специфичности SOTA и предлагаемого решения с использованием датасета CA Diabetes Prevalence 2014 (группа выборки 2: мужчины в возрасте от 30 до 60 лет)

Table 12. Accuracy, Sensitivity and Specificity results comparison of SOTA vs Proposed Solution using CA Diabetes Prevalence 2014 (Sample group 2: men of age 30 to 60)

№ выборки	Решение SOTA			
	Точность (%)	Чувствительность (%)	Специфичность (%)	Время обработки (мс)
1	91.30	52.17	97.14	69.1
2	91.25	52.40	96.91	69.2
3	91.22	51.69	96.69	69.1
4	90.11	52.26	96.87	69.0

5	90.41	52.65	96.64	69.1
6	90.74	52.52	96.65	69.3
7	91.33	52.31	97.05	69.3
8	90.73	51.39	97.33	69.0
9	91.10	51.18	96.70	69.0
10	90.10	51.41	96.95	69.2
№ выборки	Предлагаемое решение			
	Точность (%)	Чувствительность (%)	Специфичность (%)	Время обработки (мс)
1	91.64	54.40	98.10	69.1
2	91.82	53.30	98.21	68.7
3	92.98	54.07	97.93	68.6
4	91.64	54.13	97.90	68.9
5	92.55	53.45	97.99	68.6
6	91.67	53.72	97.98	68.5
7	91.39	53.43	98.29	68.7
8	92.59	53.07	98.13	68.7
9	91.47	53.00	97.99	68.9
10	91.05	54.27	98.03	68.8

Используемая в предлагаемом решении ансамблевая модель обеспечивает улучшенные точность, чувствительность и специфичность по сравнению с решением SOTA. Точность в среднем составляет 87,5%, чувствительность – 35,8%, специфичность – 98,9%, время обработки – 96,6 мс.

Эксперименты подтвердили правильность нашего теоретического подхода с использованием нескольких различных моделей. Вместо использования одной модели или набора однотипных моделей, как в решении SOTA, мы используем набор моделей машинного обучения: логистическая регрессия, SVM, Random Forest и нейронные сети. Комбинируя прогнозы этих моделей, мы достигаем более качественных результатов, чем SOTA. Из-за индивидуальных особенностей разные модели изучают разные аспекты обучающих данных. Это приводит к получению более качественного решения, чем при использовании одной модели. Наше решение также снижает проблему переобучения за счет комбинирования регуляризации отдельных моделей. Точность, чувствительность и специфичность в среднем улучшились на 3-4%. Немного сократилось время обработки

Прогнозирование СД2 – активная область науки о данных и медицинских исследований. В решении SOTA используется глубокое и широкое обучение с впечатляющим улучшением точности, чувствительности и специфичности. В нашей исследовательской работе были рассмотрены ограничения модели SOTA. Нам удалось добиться повышения точности до 87,5% по сравнению с 84,28% решения SOTA. Это связано с расширением модели путем формирования ансамбля из четырех разных алгоритмов машинного обучения, что, в свою очередь, улучшает обобщение и общее качество решения. Предлагаемое решение неизменно превосходит современную модель в нескольких экспериментах, описанных в этой работе (см. табл. 13).

Табл. 13. Сравнение решения SOTS и предлагаемого решения
Table 13. Comparison table between state of art and proposed solutions

	Предложенное решение
Суть подхода	Обнаружение T2DM с использованием ансамбля моделей машинного обучения.
Точность	За счет использования ансамблевого подхода точность увеличена до 85,2%.

Чувствительность	Чувствительность увеличена с 29,59% до 32,1%.
Специфичность	Специфичность увеличена с 96,27% до 98,3%.
Предлагаемая формула	Потери ансамбля и регуляризации для увеличения разнообразия моделей ансамбля можно представить как $EEL = ML + ML_R$
Вклад 1	Предлагаемая модель использует преимущества различных моделей машинного обучения для обеспечения прогнозирования, включая глубокое обучение. Различные алгоритмы достигают разнообразных решений, позволяя ансамблю извлечь выгоду из их разнообразия.
Вклад 2	Предлагаемая модель использует алгоритм большинства голосов Бойера-Мура для получения окончательного результата прогноза. Это позволяет различным моделям голосовать за окончательный прогноз. Отдельные прогнозы можно сравнить с прогнозом большинства, чтобы определить, какие модели работают лучше в определенных ситуациях. Это делает результаты немного более интерпретируемыми и надежными.
Вклад 3	Предлагаемое решение использует регуляризацию $L2$ для всех моделей для предотвращения переобучения. Наряду с этим модели нейронных сетей обучаются со своевременным прекращением обучения, чтобы обеспечить правильное обобщение.
Решение SOTA	
Суть подхода	Обнаружение T2DM с использованием ансамбля глубоких и широких сетей.
Точность	Обеспечиваемая точность составляет 83,51%.
Чувствительность	Обеспечиваемая чувствительность составляет 29,59%.
Специфичность	Обеспечиваемая специфичность составляет 96,27%.
Предлагаемая формула	Потери ансамбля и регуляризации для увеличения разнообразия моделей ансамбля могут быть представлены как $EL = L + L_R$
Вклад 1	Решение SOTA использует ансамбль одной и той же модели, что обесценивает ансамбль, поскольку в решениях, достигаемых одной и той же моделью, очень мало или совсем нет разнообразия.
Вклад 2	Решение SOTA усредняет выходную вероятность 10 моделей, чтобы дать окончательную вероятность для двоичной классификации.
Вклад 3	Решение SOTA использует регуляризацию $L1$, что приводит к разреженности. Но более важно избежать переобучения, чем стимулировать разреженность.

6. Заключение и будущие исследования

Обнаружение СД2 является важной проблемой медицинского направления науки о данных. Если на предсказание или идентификацию болезни потребуется много времени, пациент будет позже предупрежден и позже начнет лечиться. Это можно считать основной причиной тысяч смертей ежегодно.

Предлагаемая в этой работе система улучшает общие возможности прогнозирования СД2 с использованием набора из нескольких моделей машинного обучения и глубокого обучения. В настоящее время большинство исследований в этой области сосредоточено на использовании только одной модели для прогнозирования диабета. Однако известно, что ансамбли превосходят отдельные модели. Когда используются разные типы моделей, разнообразие помогает делать более точные прогнозы. Схема мажоритарного голосования использует результаты нескольких моделей для окончательного точного прогноза. Функция регуляризации в нашей работе изменена, чтобы включить регуляризацию всех моделей в ансамбле, что помогает предотвратить переобучение и способствует возможности обобщения предлагаемой системы. Предлагаемая система была реализована на языке Python и при тестировании показала более высокую точность, чувствительность и специфичность.

Имеется много возможностей для совершенствования описанного решения. Для повышения качества ансамблевого решения может быть реализована обучаемая система взвешенного голосования. Это будет способствовать тому, чтобы ансамбль использовал лучшие стороны отдельных моделей. Можно ввести другие типы моделей машинного обучения, а также сложные модели глубокого обучения, чтобы еще больше улучшить показатели системы.

Список литературы / References

- [1] Bernardini M., Romeo L., Misericordia P., and Frontoni E. Discovering the Type 2 Diabetes in Electronic Health Records Using the Sparse Balanced Support Vector Machine. *IEEE Journal of Biomedical and Health Informatics*, vol. 24, no. 1, 2020, pp. 235-246.
- [2] Hourii. O., Gil. Y., Berezowsky A., Wiznitzer A. et al. 339: Future Type-2 diabetes prediction following pregnancy – using a novel machine learning, *American Journal of Obstetrics and Gynecology*, vol. 222, issue 1, Supplement, 2020, p. S228.
- [3] Kölle Konstanze, Biester Torben, Christiansen Sverre et al. Pattern Recognition Reveals Characteristic Postprandial Glucose Changes: Non-Individualized Meal Detection in Diabetes Mellitus Type 1. *IEEE Journal of Biomedical and Health Informatics*, vol. 24, no. 2, 2020, pp. 594-602.
- [4] Kezhi Li, John Daniels, Chengyuan Liu et al. Convolutional Recurrent Neural Networks for Glucose Prediction. *IEEE Journal of Biomedical and Health Informatics* vol. 24, no. 2, 2020, pp. 603-613.
- [5] Binh P. Nguyen, Hung N Pham, Hop Tran et al. Predicting the onset of type 2 diabetes using wide and deep learning with electronic health records. *Computer Methods and Programs in Biomedicine*, vol. 182, 2019, article id 105055.
- [6] Perveen S., Shahbaz M., Saba T. et al. Handling irregularly sampled longitudinal data and prognostic modeling of diabetes using machine learning technique. *IEEE Access*, vol. 8, 2000, pp. 21875-21885.
- [7] Vivek Rai, Daniel X. Quang, Michael R. Erdos et al. Single-cell ATAC-Seq in human pancreatic islets and deep learning upscaling of rare cells reveals cell-specific type 2 diabetes regulatory signatures. *Molecular Metabolism*, vol. 32, 2020, pp. 109-121.
- [8] Ramazi Ramin, Perndorfer Christine, Soriano C.E. et al. Multi-modal Predictive Models of Diabetes Progression. In *Proc. of the 10th ACM International Conference on Bioinformatics, Computational Biology and Health Informatics*, 2019, pp. 253-258.
- [9] Sierra-Sosa D., Garcia-Zapirain B., Castillo C. et al. Scalable Healthcare Assessment for Diabetic Patients Using Deep Learning on Multiple GPUs. *IEEE Transactions on Industrial Informatics*, vol. 15, no. 10, 2019, pp. 5682 – 5689.
- [10] Agata Wesolowska-Andersen, Grace Zhuo Yu, Vibe Nylander et al. Deep learning models predict regulatory variants in pancreatic islets and refine type 2 diabetes association signals. *eLife* 2020;9:e51503, 2020.
- [11] Tomohide Yamada, Kosuke Iwasaki, Shotaro Maedera et al. Myocardial infarction in type 2 diabetes using sodium–glucose co-transporter-2 inhibitors, dipeptidyl peptidase-4 inhibitors or glucagon-like peptide-1 receptor agonists: proportional hazards analysis by deep neural network based machine learning. *Current Media Research and Option*, vol. 36, no. 3, 2000, pp. 404-410.
- [12] Zaitcev A., Eissa R.M., Hui Z. et al. A Deep Neural Network Application for Improved Prediction of HbA1c in Type 1 Diabetes, *IEEE Journal of Biomedical and Health Informatics*, vol. 24, no. 10, 2020, pp. 2932-2941.
- [13] Xianglong Zeng, Haiquan Chen, Yuan Luo, and Wenbin Ye. Automated Diabetic Retinopathy Detection Based on Binocular Siamese-Like Convolutional Neural Network, *IEEE Access*, vol. 7, 2019, pp. 30744-30753.
- [14] Anuradha, Akansha Singh, and Gaurav Gupta. ANT FDCSM: A novel fuzzy rule miner derived from ant colony meta-heuristic for diagnosis of diabetic patients. *Journal of Intelligent & Fuzzy Systems*, vol. 36, no. 1, 2019, pp. 747-760.
- [15] Muhammad Fazal Ijaz, Ganjar Alfian, Muhammad Syafrudin, and Jongtae Rhee. Hybrid Prediction Model for Type 2 Diabetes and Hypertension Using DBSCAN-Based Outlier Detection, Synthetic Minority Over Sampling Technique (SMOTE), and Random Forest, *Applied Sciences*, vol. 8, no. 8, 2018, pp. 1-22.

Информация об авторах / Information about authors

Маллика Г.Ч., магистр, разработчик программного обеспечения. Область научных интересов: обработка изображений.

Mallika G.C., Master, Software Developer. Research interests: image processing.

Абир АЛЬСАДУН, кандидат наук, доцент. Область научных интересов: обработка изображений.

Abeer ALSADOON, Ph.D., Associate Professor. Research interests: image processing.

Дуонг Тху Ханг ФАМ, магистр, преподаватель. Область научных интересов: базы данных, интеллектуальный анализ данных, педагогическая методология.

Duong Thu Hang PHAM, Master, Lecturer. Research interests: databases, data mining, pedagogical methodology.

Salma Hameedi ABDULLAH, Ph.D, Lecturer. Research interests: machine learning, computer vision, deep learning.

Сальма Хамиди АБДУЛЛА, Ph.D, Lecturer. Область научных интересов: машинное обучение, компьютерное зрение, глубокое обучение.

Ха Тхи МАЙ, магистр, преподаватель. Область научных интересов: интеллектуальный анализ данных, программная инженерия.

Ha Thi MAI, Master, Lecturer. Research interests: data mining, software engineering.

П.В. Чандана ПРАСАД, доктор философии, доцент. Область научных интересов: обработка изображений.

P.W. Chandana PRASAD, Ph.D., Associate Professor. Research interests: image processing.

Чан Куок Винь НГУЕН, кандидат наук, преподаватель. Область научных интересов: базы данных, интеллектуальный анализ данных, программная инженерия.

Tran Quoc Vinh NGUYEN, PhD, Lecturer. Research interests: databases, data mining, software engineering.

DOI: 10.15514/ISPRAS-2020-33(2)-6



Классификация депрессивных эпизодов на основе ночных измерений: многомерный и одномерный анализ данных

*Х.Г. Родригес-Руис, ORCID: 0000-0002-8690-8049 <jr.ruiz68@uaz.edu.mx>
К.Э. Гальван-Техада, ORCID: 0000-0002-7635-4687 <ericgalvan@uaz.edu.mx>
С. Васкес-Рейес, ORCID: 0000-0002-9435-5499 <vazquez@uaz.edu.mx>
Х.И. Гальван-Техада, ORCID: 0000-0002-7555-5655 <gatejo@uaz.edu.mx>
Х. Гамбоа-Росалес, ORCID: 0000-0002-9498-6602 <hamurabigr@uaz.edu.mx>*

*Автономный университет Сакатекаса,
Мексика, 98000, Сакатекас*

Аннотация. Психические расстройства типа депрессии составляют 28% от общего числа заболеваний, связанных с инвалидностью, и около 7,5% инвалидов в мире имеют подобные расстройства. Депрессия – это распространенное расстройство, которое влияет на душевное состояние, ежедневную деятельность, эмоции, а также вызывает нарушения сна. По некоторым оценкам, примерно 50% пациентов с депрессией страдают от нарушений сна. В данной работе процесс извлечения данных для классификации депрессивных и недепрессивных эпизодов в ночное время осуществляется на основе формального метода обнаружения знаний в базах данных (Knowledge Discovery in Databases, KDD). При использовании KDD процесс интеллектуального анализа данных имеет следующие четко выраженные этапы: предварительное обнаружение знаний в базах данных, отбор, предварительная обработка, преобразование, собственно анализ, оценка и последующее обнаружение знаний в базах данных. Для классификации был использован набор данных DEPESION, в котором содержится информация о двигательной активности 23 пациентов с монополярным и биполярным расстройством и 32 здоровых людей. Для классификации депрессивных и недепрессивных эпизодов были использованы два различных метода – многомерный и одномерный анализ данных. Для многомерного анализа применяется алгоритм случайного леса с моделью, основанной на 8 признаках. Результаты классификации имеют специфичность 0,9927 и чувствительность 0,9991. Одномерный анализ показывает, что наиболее информативной характеристикой модели являются пики активности, обеспечивающие точность 0,908 при классификации депрессивных эпизодов.

Ключевые слова: депрессия; нарушение сна; классификация; добыча данных; обнаружение знаний в базах данных; случайный лес

Для цитирования: Родригес-Руис Х.Г., Гальван-Техада К.Э., Васкес-Рейес С., Гальван-Техада Х.И., Гамбоа-Росалес Х. Классификация депрессивных эпизодов на основе ночных измерений: многомерный и одномерный анализ данных. Труды ИСП РАН, том 33, вып. 2, 2021 г., стр. 115-124. DOI: 10.15514/ISPRAS-2021-33(2)-6

Classification of Depressive Episodes Using Nighttime Data: Multivariate and Univariate Analysis

J.G. Rodríguez-Ruiz, ORCID: 0000-0002-8690-8049 <jr.ruiz68@uaz.edu.mx>

C.E. Galván-Tejada, ORCID: 0000-0002-7635-4687 <ericgalvan@uaz.edu.mx>

S. Vázquez-Reyes, ORCID: 0000-0002-9435-5499 <vazquezs@uaz.edu.mx>

J.I. Galván-Tejada, ORCID: 0000-0002-7555-5655 <gatejo@uaz.edu.mx>

H. Gamboa-Rosales, ORCID: 0000-0002-9498-6602 <hamurabigr@uaz.edu.mx>

Autonomous University of Zacatecas,
Juarez 147, Zacatecas, 98000, Mexico

Abstract. Mental disorders like depression represent 28% of global disability, it affects around 7.5% percent of global disability. Depression is a common disorder that affects the state of mind, normal activities, emotions, and produces sleep disorders. It is estimated that approximately 50% of depressive patients suffering from sleep disturbances. In this paper, a data mining process to classify depressive and not depressive episodes during nighttime is carried out based on a formal method of data mining called Knowledge Discovery in Databases (KDD). KDD guides the process of data mining with stages well established: Pre-KDD, Selection, Pre-processing, Transformation, Data Mining, Evaluation, and Post-KDD. The dataset used for the classification is the DEPRESJON dataset, which contains the motor activity of 23 unipolar and bipolar depressed patients and 32 healthy controls. The classification is carried out with two different approaches; a multivariate and univariate analysis to classify depressive and non-depressive episodes. For the multivariate analysis, the Random Forest algorithm is implemented with a model construct of 8 features, the results of the classification are specificity equal to 0.9927 and sensitivity equal to 0.9991. The univariate analysis shows that the maximum of the activity is the most descriptive characteristic of the model with 0.908 in accuracy for the classification of depressive episodes.

Keywords: depression; sleep disorders; classification; data mining; kdd; random forest

For citation: Rodríguez-Ruiz J.G., Galván-Tejada C.E., Vázquez-Reyes S., Galván-Tejada J.I., Gamboa-Rosales H. Classification of Depressive Episodes Using Nighttime Data: Multivariate and Univariate Analysis. *Trudy ISP RAN/Proc. ISP RAS*, vol. 33, issue 2, 2021, pp. 115-124 (in Russian). DOI: 10.15514/ISPRAS-2021-33(2)-6

1. Введение

По данным Всемирной организации здравоохранения (ВОЗ), психоневрологические расстройства составляют в общей сложности 28% от общего числа заболеваний в мире, и более трети из них вызваны депрессивными расстройствами, от которых страдают около 322 миллионов человек во всем мире. Большое депрессивное расстройство (major depressive disorder, MDD) относится к числу наиболее распространенных причин инвалидности, встречающегося у около 7,5% населения мира, а также к числу основных причин самоубийств, число которых приближается к 800 000 в год [1]. MDD – это распространенное психическое расстройство, которое влияет на душевное состояние и наряду с аффективным, когнитивным, волевым и соматическим нарушением вызывает у пациентов отсутствие желания жить, а также сказывается на повседневной деятельности и работоспособности [2]. Диагностику MDD лучше всего проводить во время депрессивных эпизодов. При биполярном расстройстве они выражаются в виде резкой потери самообладания, а при монополярном – в виде ухудшения настроения. MDD можно диагностировать на основании шкалы Монтгомери-Асберга для оценки депрессии (Montgomery-Asberg Depression Rating Scale MADRS), которая является диагностическим опросом, используемым психиатрами для оценки тяжести протекающего MDD с помощью ряда вопросов, выявляющих факторы риска MDD [3].

В работе Р. Армитаж (Roseanne Armitage) [4] об электроэнцефалографии во время сна делается вывод о том, что сон является подходящим периодом для диагностики MDD. Нарушения сна влияют, по крайней мере, на четверть населения мира, и они напрямую связаны с MDD. Даже когда поставлен диагноз MDD и пациент начал лечение, нарушения сна, как правило, могут продолжаться в основном из-за побочных действий антидепрессантов [5], [6]. Бессонница, гиперсомния, гипосомния и заторможенность – это лишь некоторые виды нарушения сна, выявленные как при монополярной, так и биполярной депрессии [7].

Сегодня машинное обучение широко используется для прогнозирования, диагностики и выявления психических расстройств. Существует три типа обучения: с учителем, без учителя и с частичным привлечением учителя. Данные методы обычно используются для прогнозирования развития болезни или для выявления чрезвычайных ситуаций у пациентов с психическими заболеваниями, такими как болезнь Альцгеймера, депрессия и шизофрения. И среди них депрессия является заболеванием, наиболее изученным с помощью машинного обучения [8].

Й.И. Фрогнер (Joakim Ihle Frogner) и др. с использованием набора данных DEPRESJON определили оптимальный временной интервал для классификации пациентов с депрессией из контрольной группы. Они брали 48-часовые образцы и тренировали одномерную сверточную нейронную сеть для достижения 0,70 F-меры с трехкратной перекрестной проверкой [9]. Э. Гарсиа-Сеха (Enrique Garcia-Ceja) и др. получили 0,78 взвешенной полноты с использованием алгоритма случайного леса SMOTE для классификации пациентов с депрессией и здоровых людей, входящих в контрольную группу [10]. Другой подход с использованием набора данных DEPRESJON представлен в работе Л.А. Занела-Кальсада (Laura A. Zanella-Calzada) и др., в которой с помощью алгоритма случайных лесов классифицируются депрессивные эпизоды со статистическими характеристиками, извлекаемыми раз в час. В этом подходе используются полные данные из набора данных и достигается точность 0,893 [11]. Эти и другие ранние исследования подтверждают эффективность применения машинного обучения для классификации депрессии.

Известно, что по меньшей мере 50% пациентов с MDD не могут выполнить минимальный рекомендуемый объем физической активности [12]. По этой причине для распознавания закономерностей или аномалий поведения двигательная активность хорошо воспринимается через умную одежду [13].

Изменения поведения также происходят в ночное время суток. Так, можно увидеть, что у большинства пациентов с депрессией имеются нарушения сна и фазы сна изменены. Решение данных проблем способствует улучшению состояния пациентов [14]. Учитывая тот факт, что нарушения сна являются серьезной проблемой и имеют связь с психическими расстройствами, особенно с депрессией, являясь ее причиной или в большинстве случаев симптомом [15], в данной работе основное внимание уделяется данным, собранным в ночное время, с помощью которых можно найти закономерности, позволяющих выявить депрессивные эпизоды.

Структурно каждый раздел данной статьи подробно описывает определенный этап анализа данных, проводимого по методологии обнаружения знаний в базах данных (KDD) [16], и, наконец, завершается статья разделом «Заключение».

2. Датасет

Датасет Depresjon содержит информацию о двигательной активности 23 пациентов с MDD (монополярным и биполярным расстройством) и 32 участников контрольной группы, не имеющих депрессию. Все они проходили исследование актиграфом (Actiwatch, Cambridge Neurotechnology Ltd, Англия, Модель AW4) с частотой дискретизации 32 Гц и регистрацией движений свыше 0,05 г в минуту. Интенсивность движения прямо пропорциональна подсчету

активности в минуту. Актиграфы – это устройства, используемые для неинвазивного мониторинга деятельности человека и представляющие собой браслет с акселерометром внутри [17].

Датасет можно загрузить по ссылке с [18] и там же ознакомиться с его характеристиками и описанием.

3. Предварительная обработка

Первая задача на этапе предварительной обработки – это стандартизировать данные, сигнализирующие об активности, так, чтобы среднее значение было равно 0, а среднеквадратическое отклонение – 1 для масштабирования данных. Для этого используется z-оценка, задаваемая формулой (1):

$$z_i = \frac{x_i - \bar{x}}{s}, \quad (1)$$

где x – точка сигнала активности во времени, $\bar{x}$ – среднее значение активности, а s – среднеквадратическое отклонение активности. В процессе стандартизации учитывается общая выборка из датасета после очистки данных от записей с пропущенными значениями. В данной работе целевыми данными являются сигналы, которые были собраны с помощью актиграфа в ночное время (с 9 часов вечера до 7 часов утра). Выбрано именно 10 часов, поскольку необходимо учитывать время, предшествующее сну, и все часы сна от заката до рассвета.

Для процесса стандартизации все данные от каждого пациента и участников контрольной группы были объединены в один датасет. В результате этого процесса получается набор данных из 716700 наблюдений со столбцами, содержащими информацию о времени, дате, активности и лице, к которому относится наблюдение.

4. Преобразование

Выделение признаков выполняется для каждого наблюдения во временном и частотном доменах.

Сигнал активности изначально находится во временном домене, для преобразования его в частотный домен используется быстрое преобразование Фурье (БПФ). Таким образом, за часовой промежуток некоторые признаки извлекаются во временном домене, а другие – в частотном домене.

БПФ определяется формулой (2):

$$x(k) = \sum_{n=0}^{N-1} x(n) * e^{-j2\pi(x\frac{n}{N})}, \quad (2)$$

где $x(n)$ – выборки во временном домене, то есть подсчет активности в минуту в течение одного часа, N – общий размер выборки (60), k – текущее значение частоты диапазоне от 0 до $N-1$, и $x(k)$ – спектральные компоненты выборки. После применения БПФ необходимо вычислить спектральную плотность сигнала с помощью формулы (3) для устранения мнимой части БПФ и нормализации спектра по длине временного ряда:

$$P = \frac{1}{T} \int_0^T |x(k)|^2 dt. \quad (3)$$

Извлеченные из данных признаки представлены в табл. 1.

Табл. 1. Признаки, извлеченные во временной и частотной областях
Table 1. Features extracted on time and frequency domain

Признак	Временная область	Частотная область
Коэффициент эксцесса	*	*
Среднее значение	*	*
Медиана	*	*
Среднеквадратическое отклонение	*	*
Дисперсия случайной величины	*	*
Коэффициент вариации	*	*
Межквартильный размах	*	*
Минимальное значение	*	*
Максимальное значение	*	*
Усеченное среднее	*	*
Спектральная плотность		*
Энтропия		*
Коэффициент асимметрии		*
Спектральная плоскостность		*

Отбор признаков осуществляется с помощью рекурсивного исключения признаков с перекрестной проверкой (Recursive Feature Elimination with Cross-Validation, RFECV) с использованием алгоритма случайного леса (Random Forest, RF).

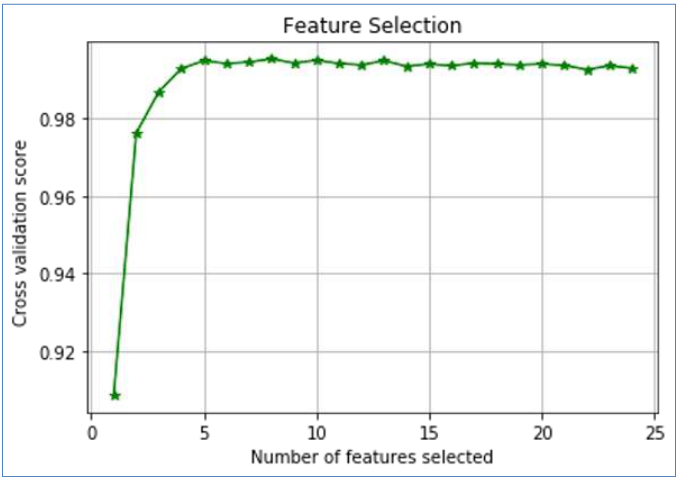


Рис. 1. График отбора признаков с использованием RFECV
Fig. 2. Feature selection plot using RFECV

На рис. 1 изображен график отбора признаков, на котором видно, что оптимальное количество признаков в модели для этого датасета равно 8. При оценивании значимости признаков с помощью RFECV было получено 8 «лучших» признаков из «лучшей» модели, перечисленные ниже.

Из временного домена:

- коэффициент эксцесса;
- медиана;
- среднеквадратическое отклонение;
- дисперсия случайной величины;
- максимальное значение.

Из частотной области:

- среднее значение;
- дисперсия случайной величины;
- коэффициент вариации.

5. Интеллектуальный анализ данных

Алгоритм случайного леса (RF) доказал свою эффективность в классификации подобных депрессивных эпизодов [11]. В нем объединены схожие функциональные возможности других алгоритмов машинного обучения для создания нового, более мощного алгоритма [19]. Именно RF был выбран для работы с моделью, имеющей 8 признаков, с использованием сбалансированной выборки для достижения заданной цели выявления депрессивных эпизодов и их классификации. Кроме того, с применением алгоритма случайного леса проводится одномерный анализ данных для повышения качества анализа отобранных признаков и их классификации.

Для реализации алгоритма случайного леса используется Python 3.6, применяется случайное разбиение данных на обучающие и тестовые, и соотношение обучающих и тестовых данных – 70:30 (8361 и 3584 наблюдения). Обучающий набор содержит признаки модели с известным результатом для того, чтобы классификатор учился на основе этих данных. Тестовый набор состоит из 2346:0 и 1238:1. Как в обучающем, так и тестовом наборе наблюдается проблема несбалансированных данных. Для ее решения в классификатор алгоритма случайного леса включается атрибут веса класса.

Сбалансированная подвыборка веса классов (bs) используется для корректировки весов классов на основе бутстрэп-выборки для каждого дерева, построенное с помощью формулы (4):

$$bs = \frac{nsamples}{nclases * np.bincount(y)}, \quad (4)$$

где $nsamples$ – общее число единиц выборки в наборе данных, $nclases$ – общее число классов в датасете, в данном случае – 2, а $np.bincount(y)$ – частоты классов во входных данных. Алгоритм обучается на обучающих данных, а после классификации проверяется на тестовых данных.

Данный процесс реализуется на полной модели с 8 признаками в одномерном режиме, т.е. алгоритм обучается и тестируется по каждому признаку в отдельности.

6. Оценка

6.1. Многомерный анализ данных

По итогам двоичной классификации методом RF получается матрица ошибок, изображенная на рис. 2. Значения в матрице ошибок соответствуют TN (true negative), TP (true positive), FN (false negative) и FP (false positive). На основе этих значений можно вычислить следующие показатели:

- чувствительность: 0.999147;
- специфичность: 0.992730;
- точность: 0.9969.

Процедура классификации сопровождается «слепым» тестом с использованием 30 процентов случайно отобранных наблюдений, т.е. признаков, выявленных в двигательной активности за часовой период. Матрица ошибок на рис. 2 отражает общее количество условий и контрольных использований для тестирования модели с использованием алгоритма случайного леса.

6.2. Одномерный анализ

Из результатов одномерного анализа важно выделить то, что максимальный признак, извлеченный во временной области, имеет точность 0,908.

Максимальный признак представляет собой максимальную точку сигнала деятельности пациента или участника контрольной группы в течение одного часа. Напомним, что одно из значений чувствительности – это подсчет движений, достигающих 0,05 g в минуту, поэтому максимальное значение равно самому интенсивному движению в течение часа. Полученные результаты по признакам приведены в табл. 2.

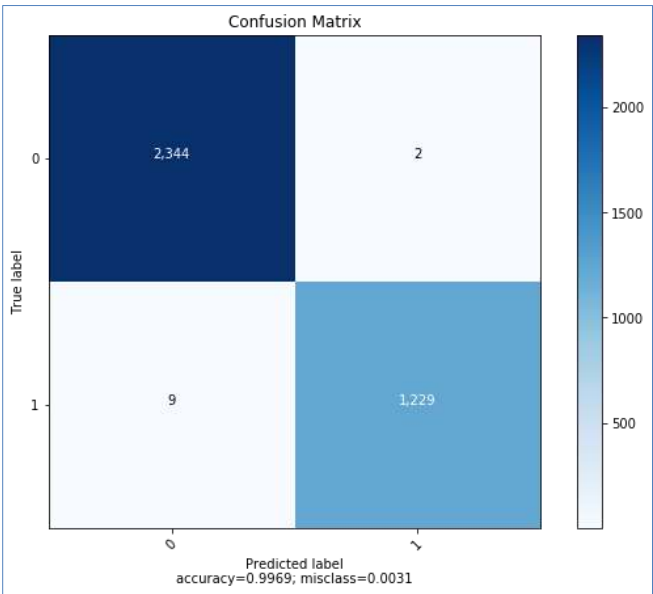


Рис. 2. Матрица ошибок бинарной классификации с использованием алгоритме случайного леса депрессивных эпизодов

Fig. 2. Confusion Matrix of the binary classification with Random Forest of depressive episodes

Табл. 2. Результаты одномерного анализа данных
Table 2. Results for univariate analysis

Признаки	TP	TN	FP	FN	Чувстви- тельность	Специ- фичность	Точность
Временной коэффициент экссесса	540	1578	702	764	0,401	0,677	0,583
Среднее значение времени	753	1945	522	364	0,553	0,876	0,702
Временное среднеквадратическое отклонение	612	1684	646	642	0,494	0,723	0,642
Временная дисперсия случайной величины	616	1685	657	626	0,540	0,721	0,647

Максимальное время	1074	2193	137	180	0,854	0,940	0,908
Среднее значение частоты	716	1761	547	716	0,540	0,765	0,689
Частотная дисперсия случайной величины	310	1666	619	689	0,497	0,730	0,647
Частотный коэффициент вариации	518	157	758	734	0,431	0,694	0,603

7. Обсуждение и выводы

Эффективность анализа депрессивных и недепрессивных эпизодов с использованием только данных о двигательной активности в ночное время суток наглядно демонстрируется в данной работе на примере добычи данных, проведенной с помощью метода обнаружения знаний в базах данных (KDD). При данной классификации у модели достигается точность 99%. Для оценки классификации алгоритма существуют четыре важных термина: истинно положительный (true positive – TP), истинно отрицательный (true negative – TN), ложно отрицательный (false negative – FN) и ложно положительный (false positive – FP).

Эти результаты могут быть лучше интерпретированы с помощью одномерного анализа. Максимальный сигнал двигательной активности за каждый часовой промежуток является наиболее описательным признаком явления. Даже если сигнал состояния представляет собой меньшее движение по сравнению с контрольным сигналом движения, закономерности между обоими сигналами отличаются до такой степени, что при классификации депрессивных эпизодов сама по себе достигается 90% точность.

Несмотря на первый аргумент об использовании только данных, собранных в ночное время, поскольку 50% пациентов с депрессией имеют расстройства сна, для диагностирования депрессии или выявления нарушений сна не может применяться исключительно этот тип классификации. Он может быть использован для выявления депрессивных эпизодов у ранее диагностированного пациента, для контроля лечения и вспомогательного участия в диагностике или оценке пациента.

Список литературы / References

- [1] Depression and other common mental disorders: global health estimates. World Health Organization, 2017, 24 p.
- [2] Espinosa-Aguilar J. Caraveo-Anduaga M., Zamora-Olvera A. et al. Guía de práctica clínica para el diagnóstico y tratamiento de depresión en los adultos mayores. *Salud Mental*, vol. 30, no. 6, 2007, pp. 69-80 (in Spanish).
- [3] S.A. Montgomery and M. Asberg. A new depression scale designed to be sensitive to change. *The British Journal of Psychiatry*, vol. 134, no. 4, 1979, pp. 382-389.
- [4] Armitage R. Sleep and circadian rhythms in mood disorders. *Acta Psychiatrica Scandinavica*, vol. 115, issue s403, 2007, pp. 104-115.
- [5] Koffel E., Polusny M.A., Arbisi P.A., and Erbes C.R. Pre-deployment day time and nighttime sleep complaints as predictors of post-deployment PTSD and depression in national guard troops. *Journal of Anxiety Disorders*, vol. 27, no. 5, 2013, pp. 512-519.
- [6] Wichniak A. Wierzbicka A., Walecka M., and Jernajczyk W. Effects of antidepressants on sleep. *Current Psychiatry Reports*, vol. 19, no. 9, 2017, article no. 63.
- [7] Kuhs H. and Reschke D.. Psychomotor activity in unipolar and bipolar depressive patients. *Psychopathology*, vol. 25, no. 2, 1992, pp. 109-116.
- [8] Shatte A.B., Hutchinson D.M., & Teague S.J. Machine learning in mental health: a scoping review of methods and applications. *Psychological Medicine*, vol. 49, no. 9, 2019, pp. 1426-1448.
- [9] Frogner J.I., Noori F.M., Halvorsen P. et al. One-Dimensional Convolutional Neural Networks on Motor Activity Measurements in Detection of Depression. In *Proc. of the 4th International Workshop on Multimedia for Personal Health & Health Care*, 2019, pp. 9-15.

- [10] García-Ceja E., Riegler M., Jakobsen P. et al. Depresjon: a motor activity database of depression episodes in unipolar and bipolar patients. In Proc. of the 9th ACM Multimedia Systems Conference, 2018, pp.472-477.
- [11] Zanella-Calzada L.A., Galván-Tejada C.E., Chávez-Lamas N.M. et al. Feature extraction in motor activity signal: Towards a depression episodes detection in unipolar and bipolar patients. *Diagnostics*, vol. 9, no. 1, 2019, pp. 1-13.
- [12] Schuch F.B., Vancampfort D., Firth J. et al. Physical activity and incident depression: a meta-analysis of prospective cohort studies. *American Journal of Psychiatry*, vol. 175, no. 7, 2018, pp. 631-648.
- [13] Gruenerbl A., Osmani V., Bahle G. et al. Using smartphone mobility traces for the diagnosis of depressive and manic episodes in bipolar patients. In Proc. of the 5th Augmented Human International Conference, 2014, p. 1-8.
- [14] Murphy M.J. and Peterson M.J. Sleep disturbances in depression. *Sleep Medicine Clinics*, vol. 10, no. 1, 2015, pp. 17-23.
- [15] Guglielmi O., Magnavita N., and Garbarino S. Sleep quality, obstructive sleep apnea, and psychological distress in truck drivers: A cross-sectional study. *Social Psychiatry and Psychiatric Epidemiology*, vol. 53, no. 5, 2018, pp. 531–536.
- [16] Dåderman A. and Rosander S. Evaluating Frameworks for Implementing Machine Learning in Signal Processing: A Comparative Study of CRISP-DM, SEMMA and KDD. Bachelor's Thesis. KTH, School of Electrical Engineering and Computer Science, 2018, 43 p.
- [17] Srinivasan R., Chen C., and Cook D. Activity recognition using actigraph sensor. In Proc. of the Fourth Int. Workshop on Knowledge Discovery from Sensor Data, 2010, pp. 25-28.
- [18] Garcia-Ceja E., Riegler M., Jakobsen P. et al. Depresjon: A Motor Activity Database of Depression Episodes in Unipolar and Bipolar Patients. In Proc. of the 9th ACM Multimedia Systems Conference, 2018, pp. 472-477.
- [19] Vijayashree J. & Sultana H. P. A machine learning framework for feature selection in heart disease classification using improved particle swarm optimization with support vector machine classifier. *Programming and Computer Software*, vol. 44, no. 6, 2018, pp. 388-397.

Информация об авторах / Information about authors

Джульета Г. РОДРИГЕС-РУИЗ, аспирант. Область научных интересов: программная инженерия, анализ биомедицинских данных, машинное обучение, интеллектуальный анализ данных.

Julieta G. RODRÍGUEZ-RUIZ, PhD Student. Research interests: Software Engineering, Biomedical Data Analysis, Machine Learning, Data Mining.

Карлос Эрик ГАЛЬВАН-ТЕХАДА, кандидат наук, профессор-исследователь. Научные интересы: искусственный интеллект, контекстные вычисления, повсеместные вычисления.

Carlos Eric GALVÁN-TEJADA, Ph.D., Professor Researcher. Research interests: Artificial Intelligence, Context Computing, Ubiquitous Computing.

Содель ВАСКЕС-РЕЙЕС, кандидат наук, разработчик. Область научных интересов: программная инженерия, интеллектуальный анализ текста, обработка естественного языка.

Sodel VÁZQUEZ-REYES, Ph.D., Developer. Research interests: Software Engineering, Text Mining, Natural Language Processing.

Хорхе Исак ГАЛЬВАН-ТЕХАДА, кандидат наук, профессор-исследователь. Область научных интересов: биоинформатика, цифровая обработка сигналов при остеоартрозе.

Jorge Issac GALVÁN-TEJADA, Ph.D., Professor Researcher. Research interests: Bioinformatics, Osteoarthritis Digital Signal Processing.

Хамурапи ГАМБОА-РОСАЛЕС, кандидат наук, доцент. Область научных интересов: обработка голоса, телекоммуникации, программирование на C/C ++.

Hamurabi GAMBOA-ROSALES, Ph.D., Associate Professor. Research interests: Voice Processing, Telecommunications, C / C ++ Programming.

DOI: 10.15514/ISPRAS-2020-33(2)-7



Стратегии управления спросом в умных сетях электроснабжения для центров данных

¹ Х. Муранья, ORCID: 0000-0002-9328-2320 <jmurana@fing.edu.uy>

^{1,2} С. Несмачнов, ORCID: 0000-0002-8146-4012 <sergion@fing.edu.uy>

¹ С. Итурриага, ORCID: 0000-0002-0212-7916 <siturria@fing.edu.uy>

¹ С. Монтес де Ока, ORCID: 0000-0002-0212-837X <smontes@fing.edu.uy>

¹ Г. Белькреди, ORCID: 0000-0002-3475-2062 <gbelcredi@fing.edu.uy>

¹ П. Монсон, ORCID: 0000-0001-7924-681X <monzon@fing.edu.uy>

² В. Шепелев, ORCID: 0000-0002-1143-2031 <v.shepelevvd@susu.ru>

^{2,3,4} А. Черных, ORCID: 0000-0001-5029-5212 <chernykh@cicese.mx>

¹ Республиканский университет,

Уругвай, 11200, Монтевидео, ул. 18 июля 1824-1850

² Южно-Уральский государственный Университет,

Россия, 454080, г. Челябинск, проспект Ленина, д. 76.

³ Центр научных исследований и высшего образования,

Мексика, 22860, Нижняя Калифорния, Энсенада, ш. Тихуана-Энсенада, 3918

⁴ Институт системного программирования им. В.П. Иванникова РАН,

109004, г. Москва, ул. А. Солженицына, дом 25

Аннотация. В данной статье представлены методы управления спросом для участия центров обработки данных в умных электроэнергетических рынках в рамках парадигмы умных сетей электроснабжения. Для определения энергопотребления задач с интенсивным использованием процессора и памяти используется модель центра данных, основанная на эмпирической информации. Предлагаются переговорный подход между центром обработки данных и клиентами и эвристический метод планирования для оптимизации энергосбережения. Экспериментальная оценка осуществляется на примере реальных проблем, моделирующих разные типы клиентов. Результаты исследования свидетельствуют о том, что предложенный подход эффективен для обеспечения соответствующих мер по управлению спросом на электроэнергию в соответствии с денежными стимулами.

Ключевые слова: умная сеть электроснабжения; вычислительный интеллект; управление спросом на электроэнергию

Для цитирования: Муранья Х., Несмачнов С., Итурриага С., Монтес де Ока С., Белькреди Г., Монсон П., Шепелев В., Черных А. Стратегии управления спросом в умных сетях электроснабжения для центров данных. Труды ИСП РАН, том 33, вып. 2, 2021 г., стр. 125-136. DOI: 10.15514/ISPRAS-2021-33(2)-7

Благодарности. Частичную поддержку данного исследования оказало Национальное агентство научных исследований и инноваций (ANII) (FSE_2017_1_144789). Исследование С. Несмачнова и С. Итурриага было проведено при частичной финансовой поддержки Национального агентства научных исследований и инноваций и Программы развития фундаментальных наук Уругвая (PEDECIBA). Авторы данного исследования также выражают благодарность Национальному суперкомпьютерному центру (Cluster.Uy).

Smart grid demand response strategies for datacenters

¹ J. Muraña, ORCID: 0000-0002-9328-2320 <jmurana@fing.edu.uy>

^{1,2} S. Nesmachnow, ORCID: 0000-0002-8146-4012 <sergion@fing.edu.uy>

¹ S. Iturriaga, ORCID: 0000-0002-0212-7916 <siturria@fing.edu.uy>

¹ S. Montes de Oca, ORCID: 0000-0002-0212-837X <smontes@fing.edu.uy>

¹ G. Belcredi, ORCID: 0000-0002-3475-2062 <gbelcredi@fing.edu.uy>

¹ P. Monzón, ORCID: 0000-0001-7924-681X <monzon@fing.edu.uy>

² V. Shepelev, ORCID: 0000-0002-1143-2031 <v.shepelevvd@susu.ru>

^{2,3,4} A. Tchernykh, ORCID: 0000-0001-5029-5212 <chernykh@cicese.mx>

¹ Universidad de la Republica

Av. 18 de Julio 1824-1850, Montevideo, Departamento de Montevideo, Uruguay, 11200

² South Ural State University,

454080, Russia, Chelyabinsk, Lenin Avenue, 76

³ Centro de Investigación Científica y de Educación Superior,

3918, Ensenada-Tijuana Highway, Ensenada, 22860, Mexico

⁴ Ivannikov Institute for System Programming of the Russian Academy of Sciences,
25, Alexander Solzhenitsyn st., Moscow, 109004, Russia

Abstract. This article presents demand response techniques for the participation of datacenters in smart electricity markets under the smart grid paradigm. The proposed approach includes a datacenter model based on empirical information to determine the power consumption of CPU-intensive and memory-intensive tasks. A negotiation approach between the datacenter and clients and a heuristic planning method for energy reduction optimization are proposed. The experimental evaluation is performed over realistic problem instances modeling different types of clients. Results indicate that the proposed approach is effective to provide appropriate demand response actions according to monetary incentives.

Keywords: smart grid; computational intelligence; demand response

For citation: Muraña J., Nesmachnow S., Iturriaga S., Montes de Oca S., Belcredi G., Monzón P., Shepelev V., Tchernykh A. Smart grid demand response strategies for datacenters. *Trudy ISP RAN/Proc. ISP RAS*, vol. 33, issue 2, 2021, pp. 125-136 (in Russian). DOI: 10.15514/ISPRAS-2021-33(2)-7.

Acknowledgments. The work is partially supported by Agencia Nacional de Investigación e Innovación (FSE_2017_1_144789). The work of S. Nesmachnow and S. Iturriaga has been partly funded by ANII and PEDECIBA, Uruguay. The authors also want to thank the Centro Nacional de Supercomputación (Cluster.Uy).

1. Введение

Умные сети электроснабжения (smart grid) – это новейшая технология, используемая в электрических сетях. Они включают в себя функции эксплуатации и управления, которые позволяют улучшить контроль над производством и распределением энергии [1].

В соответствии с парадигмой умных сетей электроснабжения на рынке электроэнергии могут присутствовать крупные потребители с гибким расходом электроэнергии. Это одна из основных идей, лежащих в основе стратегии использования современных интеллектуальных электрических сети: потребителям присваиваются роли как активных клиентов, так и участников рынка [1]. Как активный клиент, потребитель может корректировать свой спрос на электроэнергию в часы пиковых нагрузок, например, за счет снижения потребления электроэнергии, тем самым сглаживая кривую спроса всей электросистемы. Как участник рынка, потребитель может получать доход путем предоставления различных услуг (например, путем заключения двусторонних договоров с производителями электроэнергии или участия в периодических аукционах по управлению умными сетями).

В умных сетях электроснабжения стратегии планирования ответных мер на изменение спроса на электроэнергию необходимы для регулирования энергопотребления и для того, чтобы

иметь возможность участвовать на рынке в различных ролях. Для планирования энергозатратных операций необходимо использовать специальные приемы, например, приоритизация или отсрочка выполнения данных операций. Кроме того, необходимо проанализировать влияние на глобальную энергоэффективность возможного ухудшения качества обслуживания, предлагаемого пользователям. Эти меры планирования необходимы для обеспечения правильного использования энергоресурсов и гарантирования эффективного использования энергии крупными потребителями с гибким спросом на энергию.

В данной статье описывается предложение по развитию стратегий управления спросом на электроэнергию применительно к центрам обработки данных, что позволит им быть участниками рынка электроэнергии и предоставлять вспомогательные услуги. Центры обработки данных могут регулировать энергопотребление, тем самым способствуя выполнению особых функций электросети. Они могут расходовать доступный избыток энергии, выполняя сложные задачи, требующие больших временных затрат, или могут откладывать выполнение задач в периоды, когда энергия стоит дороже и/или производится ниже нормы. Кроме того, их инфраструктура теплоснабжения и охлаждения требует значительного энергопотребления, обеспечивая большую инерцию сети. Таким образом, центры данных могут использоваться для работы с умными сетями электроснабжения.

Исследование, представленное в данной статье, основано на процессе согласования с использованием механизма ценообразования, который обеспечивает возможность сброса нагрузки с арендаторов операторами центров обработки данных. Предлагаемая стратегия управления спросом на электроэнергию позволяет внедрить умное управление электросетями, добиться рационального использования источников энергии, а также правильного использования информационных технологий для совершенствования процессов принятия решений в современных умных сетях электроснабжения.

2. Проблема планирования нагрузки центра обработки данных

При наличии запроса на энергосбережение со стороны рынка проблема оптимизации заключается в минимизации общего денежного вознаграждения, предоставляемого арендаторам, и стоимости использования локальных генераторов электроэнергии с целью достижения целевых показателей энергосбережения. Режим пониженного энергопотребления должен поддерживаться в течение некоторого промежутка времени T . Проблему можно формализовать следующим образом.

- Промежуток времени T является множеством дискретных тактов t .
- β – коэффициент сокращения энергопотребления β , требуемого рынком.
- $C = \{c_1, \dots, c_{|C|}\}$ – множество арендаторов центра обработки данных.
- $W_j = \{w_j^1, \dots, w_j^{|W_j|}\}$ – рабочая нагрузка арендатора c_i .
- $DF_j^i = 1$, если задачу w_j^i можно отложить, и $DF_j^i = 0$, если ее отложить нельзя.
- DD_j^i – установленный срок выполнения задачи w_j^i .
- MP_j^i – денежный штраф, налагаемый на арендатора c_j , если установленный срок выполнения задачи w_j^i нарушается.
- Пусть RI – денежное поощрение каждого арендатора за каждую сэкономленную единицу электроэнергии.
- $WS_j = \{ws_1, \dots, ws_{|C|}\}$ – график выполнения рабочей нагрузки для каждого арендатора c_j без денежных поощрений (т.е. $RI = 0$),
- Пусть DP_j^t – требуемый объем энергопотребления на каждый такт t каждого графика рабочей нагрузки WS_j .
- FT_j^i — это время окончания выполнения задачи w_j^i для графика WS_j .

- $VD_j^i = 0$, если $FT_j^i \leq DD_j^i$ для графика WS_j , в противном случае $VD_j^i = 1$.
- Определим общий денежный штраф за отклонение от графика WS_j каждого арендатора c_j как $Y_j = \sum_{i=1...|W|} VD_j^i \times P_j^i$.
- Пусть функция γ_j определяет новый график $\overline{W'S_j}$ с требованием энергопотребления $\overline{D'P_j^t}$ для арендатора c_j с учетом поощрения RI : $\gamma_j(RI) = \overline{W'S_j}$.
- Определим функцию сокращения потребления электроэнергии при использовании графика $\overline{W'S_j}$ вместо графика WS_j следующим образом: $\delta(\overline{ws_j}) = \min \{\overline{DP_j^t} - DP_j^t, t \in T\}$.
- Пусть GP^t – объем электроэнергии, произведенной за каждый квант t с использованием локального генератора.
- Пусть α – стоимость в денежном выражении единицы электроэнергии, выработанной с использованием локального генератора.

Целевая функция представлена в виде формулы 2.1.

$$\min z = \sum_{j=1}^{|C|} \delta(\gamma_j(RI)) \times RI + \sum_{t=1}^T GP^t \times \alpha \quad (2.1a)$$

при

$$\beta \leq \delta(\gamma_j(RI)) + GP^t, \quad (2.1b)$$

$$z \leq \sum_{t=1}^T \overline{DP_j^t} \times \alpha \quad (2.1c)$$

Целью (2.1a) является минимизация затрат оператора центра обработки данных (т.е. сокращение средств, выплачиваемых арендаторам, и затрат на использование локального генератора) для достижения целевых показателей сокращения потребления электроэнергии. В условии (2.1b) указано, что общее экономия энергии должно быть не менее β за каждый такт. Наконец, согласно условию (2.1c), общая денежная стоимость должна быть меньше, чем стоимость питания всего центра обработки данных при помощи лишь одного локального генератора.

3. Подход к управлению спросом на электроэнергию

3.1 Работа центра обработки данных

Оператор энергосети предлагает администратору центра данных денежное вознаграждение для снижения определенного объема потребляемой электроэнергии в течение определенного времени. Для достижения необходимого уровня снижения потребления энергии был обеспечен следующий рыночный механизм.

В предлагаемом рыночном механизме оператор может побудить клиента к снижению энергопотребления и тем самым уменьшить потребность в использовании загрязняющих среду источников энергии, используя параметризованную функцию предложения, представленную формулой 3.1, где r_i – объем энергии, сэкономленной i -ым клиентом, D – значение целевого показателя сокращения энергопотребления центром обработки данных, b_i – предложение клиента о вознаграждении при снижении энергопотребления на r_i , и p – это равновесная рыночная цена, определяемая оператором [2]:

$$r_i(b_i, p) = D - \frac{b_i}{p}. \quad (3.1)$$

Рыночный механизм сокращения объема D потребляемой энергии реализуется в четыре этапа итеративным образом.

- (а) Центр обработки данных рассылает клиентам функцию предложения $r_i(b_i, p)$.
- (б) Каждый клиент i подает заявку на получение вознаграждения b_i за сокращение энергопотребления r_i в целях максимизации его полезности.
- (с) Центр обработки данных определяет равновесную рыночную цену p и объем энергии y , вырабатываемый локальным генератором (со стоимостью производства α), минимизируя общие затраты.

$$p(b_i, y) = \frac{\sum_i b_i}{(N-1)D + y}, \quad (3.2)$$

$$y = \operatorname{argmin} (D - y)p + \alpha y, \quad 0 \leq y \leq D. \quad (3.3)$$

Условие оптимальности первого порядка для формулы 4.3 задает значение для y :

$$y = \sqrt{\frac{(\sum_{i=1}^N b_i)ND}{\alpha}} - (N-1)D. \quad (3.4)$$

- (д) Если p и y сходятся, то принимаются последние предложения и клиенты планируют объем сокращения энергопотребления, в противном случае оператор рассылает новую функцию предложения с обновленным значением p .

Стратегия, используемая центром обработки данных, решает проблему размещения методом проксимального градиента [3]. Для каждого участника создается распределенное решение. В Алгоритме 1 D – значение целевого показателя сокращения энергопотребления, $price$ – равновесная рыночная цена за ватт, N – количество арендаторов и j – идентификатор арендатора. Функция $client_evaluation(price, j)$ соответствует оценке предложения арендатора j с учетом его SLA (Service Level Agreement, соглашение об уровне обслуживания). Эта функция выдает значение, характеризующее сокращение энергопотребления арендатора ($reduction[j]$) в соответствии с ценой, $bid[j]$ – это предложение арендатора j по сокращению потребления электроэнергии, y_k – итерационная переменная, которая в конце согласования соответствует объему электроэнергии, вырабатываемой локальным генератором. Стоимость производства одного ватта с помощью генератора обозначается α . Параметр ϵ является мерой достижения целевого показателя сокращения энергопотребления.

INPUT: D (целевой показатель энергосбережения), $price_0$

OUTPUT: $price$, $on - site generation$

```

1:  $k \leftarrow 0$                                 *шаг итерации
2:  $price_k \leftarrow price_0$ 
3: while  $\epsilon \geq \epsilon_{min}$  do
4:   for  $j = 1$  to  $N$  do
5:      $reduction[j] \leftarrow client\_evaluation(price_k, j)$ 
6:      $bid[j] \leftarrow (D - reduction[j]) \times price_k$ 
7:   end for
8:    $y_k \leftarrow \max\left(\sqrt{(\sum_j bid) \cdot \frac{ND}{\alpha}} - (N-1)D, 0\right)$ 
9:    $price_k \leftarrow \sum_j bid[j] / ((N-1)D + y_k)$ 
10:   $\epsilon \leftarrow \|(y_k + \sum_j reduction[j] - D) / D\|$ 
11:   $k \leftarrow k + 1$ 
12: end while
13:  $on - site generation \leftarrow y_k$ 

```

Алгоритм 1. Рыночный механизм центра обработки данных
Algorithm 1. Datacenter market mechanism

3.2 Оценка предложения клиентом

Для оценки денежного предложения администратора центра данных и определения размеров снижения энергопотребления клиенты моделируют выполнение своей рабочей нагрузки, применяя стратегию оптимизации энергопотребления. Денежное предложение администратора центра данных принимается, если чистая прибыль, получаемая в результате снижения энергопотребления, за вычетом убытков, которые клиент должен возместить в случае несоблюдения SLA со своими пользователями, больше нуля. В любом случае, в результате согласования разных денежных предложений достигаются различные компромиссы. Эти компромиссы могут быть приняты во внимание, если центр обработки данных не может достичь желаемого снижения энергопотребления, чтобы можно было учесть конфликтность поставленных целей (снижения энергопотребления и затрат).

Стратегия оптимизации энергопотребления, предлагаемая в данной статье, направлена на обеспечение максимальной прибыли клиентов за счет сокращения числа активных серверных вычислительных ядер, что приводит к снижению энергопотребления в соответствии с предложением, полученным от администрации центра обработки данных. Применяется эвристический метод сокращения количества активных ядер (Active Cores Reduction, ACR), основная идея которого заключается в выборе оптимального по рентабельности расписания и с учетом всех комбинаций активных ядер. Детали эвристики ACR представлены в Алгоритме 2, где *price* – ценовое предложение за сэкономленный ватт, а *reduction* – число ватт, на которое клиент готов сократить свое энергопотребление в соответствии с предложением. Функция *schedule* моделирует выполнение рабочей нагрузки при использовании *cores_number* активных ядер из общего числа *server_cores* ядер, доступных клиенту. В свою очередь, функция *eval* оценивает прибыль и размер снижения энергопотребления в результате планирования решения *sol*.

INPUT: *price*

OUTPUT: *reduction*

```
1: profit  $\leftarrow$  0
2: reduction  $\leftarrow$  0
3: cores  $\leftarrow$  server_cores  $\times$  serves_numbers
4: for cores_numbers in cores do
5:   sol  $\leftarrow$  schedule(cores_number)
6:   reduction_ax, profit_ax  $\leftarrow$  evaluate(price, sol)
7:   if profit_ax > profit then
8:     reduction  $\leftarrow$  reduction_ax
9:     profit  $\leftarrow$  profit_ax
10:  end if
11: end for
```

Алгоритм 2. Стратегия оптимизации энергопотребления

Algorithm 2. Energy optimization strategy

3.3 Моделирование планирования работы клиента

Клиенты являются поставщиками услуг в области высокопроизводительных вычислений (HPC) для отдельных пользователей. Пакетные задачи поступают в систему и находятся в очереди до тех пор, пока сервер не будет иметь возможность их выполнять (т.е. пока его возможности не соответствуют требованиям задачи, таким как доступность ядер, объем памяти и предполагаемое время выполнения).

Для определения стоимости реализации определенной стратегии оптимизации энергопотребления применяется подход, основанный на моделировании. В связи с ограниченными возможностями существующих центров обработки данных и облачных

симуляторов используется специализированный симулятор, обеспечивающий необходимую среду для реализации основных элементов предлагаемого подхода.

Период моделирования делится на интервалы равной продолжительности int_d . В каждом интервале планировщик распределяет поступившие задачи по серверам, учитывая текущую мощность каждого сервера и критерии стратегии планирования. Количество интервалов для выполнения задачи на сервере вычисляется следующим образом: $ct/int_d + 1$, где ct — время завершения задачи (в секундах), а int_d — продолжительность каждого интервала (в секундах). Время выполнения задачи определяется как ее размер в миллионах инструкций в секунду (MIPS), деленный на заданную скорость ядра (также в MIPS).

3.4 Модель энергопотребления

Для оценки расхода энергии при использовании некоторого плана использовалась новая квадратичная модель энергопотребления, разработанная с использованием реальных экспериментальных данных. [4].

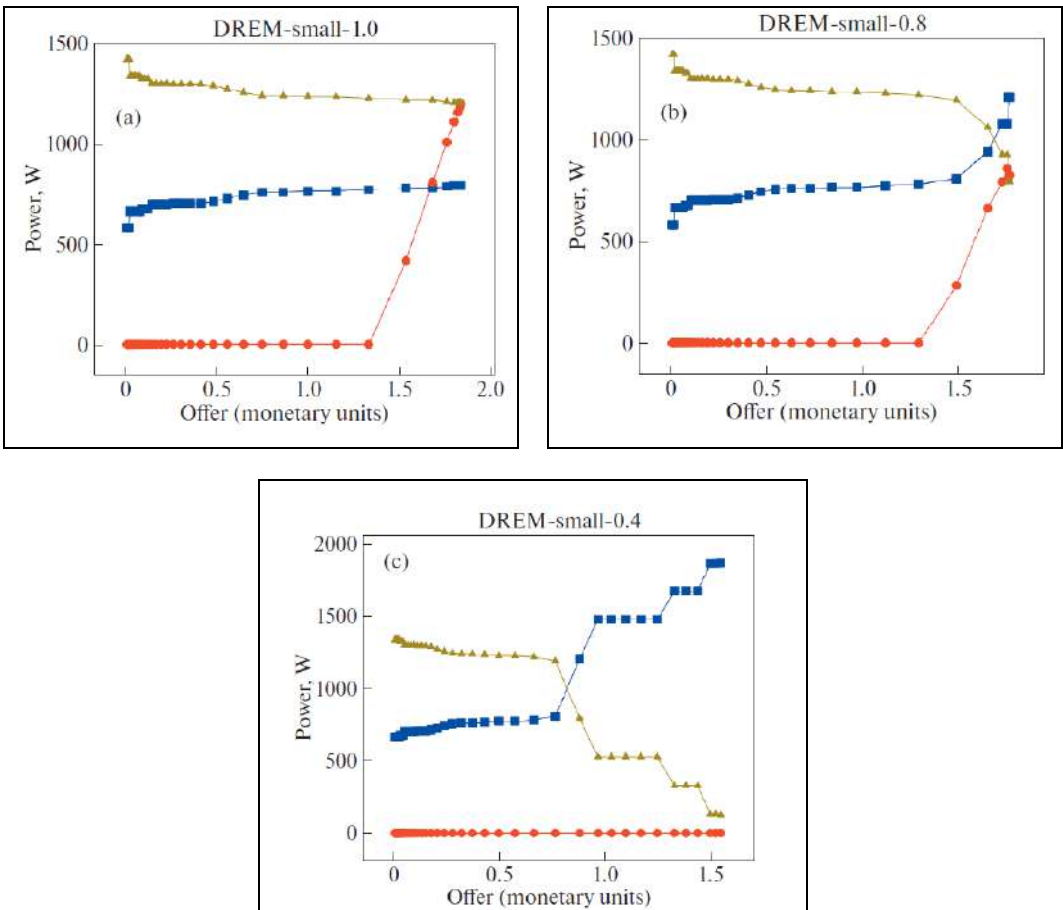


Рис. 1. Согласование на небольших примерах: красные точки – y_k , синие квадраты – CR, красные треугольники – производство энергии от локального генератора

Fig. 1. Negotiation for small instances: red dots – y_k , blue squares – CR, red triangles – on-site generation

4. Экспериментальный анализ

Для оценки и валидации предложенной модели участия центра обработки данных на рынке электроэнергии были подготовлены конкретные примеры. Предварительные результаты сообщаются для примеров с пятью клиентами, пятью серверами и 1500 задачами у каждого клиента. Более подробная оценка эксперимента изложена в [2]. Что касается вычислительной инфраструктуры, то рассматриваются 24-ядерные серверные процессоры Intel Xeon с быстродействием 3000 MIPS. Эксперименты были разработаны на Java SE 1.8 и выполнены в Национальном суперкомпьютерном центре Уругвая [5].

На рис. 1 показан результат согласования по снижению энергопотребления между администратором центра данных и клиентами для небольших примеров. Предложение (на ватт) для клиентов находится на оси абсцисс. Синими квадратами отмечена динамика сокращения клиентом объема потребления электроэнергии (Clients Reduction, CR) по мере согласования, красные круги отображают значения y_k , а зеленые треугольники обозначают мощность, вырабатываемую локальным генератором, чтобы достичь целевой показатель сокращения энергопотребления ($OG = D - CR$). Для небольших задач в результате согласования энергопотребление быстро сокращается на 600 Вт при низких денежных затратах. Однако для более масштабного сокращения энергопотребления требуется больше итераций и более выгодное денежное предложение. Замедление темпов сокращения энергопотребления связано с тем, что клиенты не могут сократить энергопотребление без существенного влияния на пользователей.

Табл. 2. Сводка результатов процесса согласования

Table 2. Negotiation summary

DPEM - small - 1.0						
k	$price$	CR	y_k	ϵ	OG	$cost$
1	0,012	582	0	70,90	1418	2842
20	0,227	701	0	64,95	1299	2757
41	1,840	794	1200	0,03	1206	3872
DPEM - small - 0.8						
k	$price$	CR	y_k	ϵ	OG	$cost$
1	0,012	582	0	70,90	1418	2842
21	0,260	705	0	64,75	1295	2773
41	1,760	1206	802	0,04	794	3710
DPEM - small - 0.4						
k	$price$	CR	y_k	ϵ	OG	$cost$
1	0,012	665	0	66,75	1335	2677
21	0,244	728	0	63,60	1272	2721
42	1,576	1894	0	0,06	126	3205

В табл. 2 приводится сводка процесса согласования на небольших примерах. Согласование состоит из трех этапов: первое предложение, промежуточное предложение и последнее предложение (когда алгоритм согласования достигает конца согласно критериям его завершения). Столбец k – этап согласования, $price$ – предложение на данном этапе k , CR – размер сокращения потребления электроэнергии клиентом, OG – производство электроэнергии локальным генератором, ϵ оценивает уровень достижения требуемого уровня сокращения потребления электроэнергии, а $cost$ – сумма в денежном выражении, которую администратор центра данных должен вложить, чтобы достичь целевого показателя сокращения энергопотребления (формула 4.1).

$$cost = CR \times price + OG \times \alpha. \quad (4.1)$$

Сравнение примеров с разными значениями допуска показывает, что в случае с менее гибкими клиентами на последнем этапе согласования определяются низкие значения в столбце *CR* и высокие значения в столбце *OG*. Это наблюдение подтверждает очевидную идею о том, что в центрах обработки данных, в которых клиенты имеют менее гибкие SLA, производство электроэнергии на местах является основным способом достижения целевого показателя энергосбережения, установленного поставщиком электроэнергии. Кроме того, для менее гибких клиентов (*small - 1.0*) требуются более крупные предложения.

Результаты исследования подтверждают, что предложенный подход согласования позволяет оптимально использовать возможность отложенного выполнения задач для достижения требуемого сокращения энергопотребления.

5. Заключение

В данной статье был рассмотрен подход к согласованию участия центров обработки данных и суперкомпьютерных центров в умных сетях электроснабжения, что является важной задачей в современных системах умных сетей электроснабжения.

Был изучен конкретный случай применения стратегии управления спросом на электроэнергию в центрах обработки данных, предоставляющих услуги аппаратного хостинга. Снижение энергопотребления достигается в соответствии с предложениями, направляемыми клиентам. При согласовании применялся децентрализованный подход, когда от клиентов не требуется предоставление стратегической информации администратору центра данных. Вместо этого каждый клиент согласовывает цену с учетом эвристики планирования и особенностей задач, которые требуется выполнить. Представлена модель, основанная на данных, полученных от реальных центров обработки данных, для определения энергопотребления задач с интенсивным использованием ресурсов процессоров и памяти.

Алгоритм согласования и эвристический метод планирования для оптимизации снижения энергопотребления были экспериментально проверены на девяти реалистичных примерах, моделирующих различные характеристики решаемых задач и гибкости клиентов центра данных.

Полученные результаты свидетельствуют о том, что предложенный подход эффективен для обеспечения надлежащих мер по управлению спросом на электроэнергию в соответствии с денежным стимулом. Система принесла экономические выгоды операторам центров обработки данных и для арендаторов (за счет предоставления вознаграждений за сокращение расходов), а также способствовала защите окружающей среды за счет сокращения использования дизельного топлива.

Подводя итог, можно сказать, что клиенты быстро добились соответствующих показателей сокращения энергопотребления, тем самым ограничив потребность центра обработки данных в использовании локального генератора. Результаты подтвердили, что задача по своей сути является многокритериальной. При формулировке должны учитываться как эксплуатационные расходы, так и предлагаемое пользователям качество обслуживания, а также должны быть изучены компромиссы между общей стоимостью и предложениями, вынесенными на согласование. Предлагаемый подход является практическим и эффективным; он может быть реализован в современных центрах обработки данных и суперкомпьютерных центрах.

Список литературы / References

- [1] J. Momoh. Smart Grid: Fundamentals of Design and Analysis. Wiley-IEEE Press, 2012, 232 p.
- [2] J. Muraña, S. Nesmachnow, S. Iturriaga, S. Montes de Oca, G. Belcredi, P. Monzón, V. Shepelev, A. Tchernykh. Negotiation approach for the participation of datacenters and supercomputing facilities in smart electricity markets. *Programming and Computer Software*, vol. 46, no. 8, 2020, pp. 636–651.

- [3] N. Parikh and S. Boyd. Proximal Algorithms. In: *Foundations and Trends in Optimization*, vol. 1, issue 3, 2014, pp. 127-239.
- [4] J. Muraña, S. Nesmachnow, F. Armenta, and A. Tchernykh. Characterization, modeling and scheduling of power consumption of scientific computing applications in multicores. *Cluster Computing*, vol. 22, no.3, 2019, pp. 839-859.
- [5] S. Nesmachnow and S. Iturriaga. Cluster-UY: Collaborative Scientific High Performance Computing in Uruguay. *Communications in Computer and Information Science*, vol. 1151. 2019, pp. 188-202.

Информация об авторах / Information about authors

Джонатан МУРАÑЯ-СИЛЬВЕРА – магистр компьютерных наук, младший научный сотрудник. Область научных интересов: метаэвристика, высокопроизводительные вычисления, экологически чистые вычисления.

Jonathan MURAÑA-SILVERA, M.Sc. in Computer Science, Researcher Assistant. Research interests include metaheuristics, high performance computing, green computing.

Серджио Энрике НЕСМАЧНОВ-КАНОВАС, кандидат наук, профессор, научный сотрудник. Область научных интересов: оптимизация, метаэвристика, высокопроизводительные вычисления, умные города.

Sergio Enrique NESMACHNOW-CÁNOVAS, Ph.D. in Computer Sciences, Full Professor and Researcher. Research interests: optimization, metaheuristics, high performance computing, smart cities.

Сантьяго Дамиан ИТУРРИАГА-ФАБРА, кандидат компьютерных наук, адъюнкт-профессор. Область научных интересов: высокопроизводительные вычисления, эволюционные алгоритмы, оптимизация.

Santiago Damián ITURRIAGA-FABRA, Ph.D. in Computer Sciences, Adjunct Professor. Research interests include HPC, evolutionary algorithms, optimization.

Себастьян МОНТЕС ДЕ ОКА, магистр электротехники, младший научный сотрудник. Область научных интересов: оптимизация, переходные энергетические системы, реагирование на спрос в умных сетях электроснабжения.

Sebastián MONTES DE OCA, M.Sc. in Electrical Engineering, Researcher Assistant. Research interests: optimization, transactive energy systems, demand response on smart grids.

Гонсало БЕЛКРЕДИ, инженер-электрик, ассистент. Область научных интересов: беспроводная связь, программно-конфигурируемая радиосвязь, реакция на спрос в умных сетях электроснабжения.

Gonzalo BELCREDI, Electrical Engineer, Assistant. Research interests: wireless communication, software defined radio, demand response on smart grids.

Пабло Ариэль МОНЗОН-РАНГЕЛОФФ, кандидат наук, начальник отдела систем и контроля. Область научных интересов: нелинейные системы управления, энергосистемы, реакция на спрос в умных сетях электроснабжения.

Pablo Ariel MONZÓN-RANGELOFF, Ph.D. in Electrical Engineering, Chief of the Systems and Control Department. Research interests include nonlinear control systems, power systems, demand response on smart grids.

Владимир Дмитриевич ШЕПЕЛЁВ – кандидат наук, доцент. Область научных интересов: интеллектуальный анализ данных, искусственный интеллект, искусственные нейронные сети, транспортная инженерия, умные города.

Vladimir Dmitrievitch SHEPELEV, PhD, Associate Professor. Research interests include Data Mining, Artificial Intelligence, Artificial Neural Network, Transportation Engineering, Smart city.

Андрей Николаевич ЧЕРНЫХ получил степень кандидата наук в Институте точной механики и вычислительной техники РАН. Он является профессором CICESE. В научном плане его интересуют многоцелевая оптимизация распределения ресурсов в облачной среде, проблемы безопасности, планирования, эвристики и метаэвристики, интернет вещей.

Andrei Nikolaevitch TCHERNYKH received his PhD degree at the Institute of Precision Mechanics and Computer Engineering of the Russian Academy of Sciences. He is holding a full professor position in computer science at CICESE. He is interesting in grid and cloud research addressing multiobjective resource optimization, both, theoretical and experimental, security, uncertainty, scheduling, heuristics and meta-heuristics, adaptive resource allocation, and Internet of Things.

DOI: 10.15514/ISPRAS-2020-33(2)-8



Классификатор изображений транспортных средств для определения корреляции их воздействия со смещением координат моста

*В. Флорес-Фуэнтес, ORCID: 0000-0002-1477-7449 <flores.wendy@uabc.edu.mx>
Автономный университет Нижней Калифорнии (UABC),
Мексика, 21100, Нижняя Калифорния, Энсенада*

Аннотация. Современные компьютерные технологии открывают возможности для инноваций в широком спектре приложений. На смену традиционным подходам, в основе которых лежат визуальные и ручные методы, все чаще приходят автоматизированные процессы за счет использования киберфизических систем. Настоящая работа является примером этой тенденции и посвящена использованию машинного зрения на основе глубокого обучения для классификации действующих нагрузок на мосты и поддержки систем оптического сканирования для мониторинга состояния строительных конструкций. Система оптического сканирования контролирует состояние строительных конструкций (здания, мосты, дамбы и т.д.) путем измерения возможного смещения координат характерных точек, что позволяет выявить аномальное поведение конструкции, возможно, связанного с ее повреждением. При анализе мостовых конструкций использование подобной оптической сканирующей системы несколько затруднено из-за проезда транспортных средств по мосту, что вызывает его колебания. Причем эти колебания моста и, соответственно, смещения координат не обязательно связаны с повреждением моста. Таким образом, требуется классификатор нагрузок на мост для определения корреляции измеряемых смещений координат характерных точек с колебаниями, вызванными взаимодействием с ним транспортных средств, чтобы отличить нормальное поведение конструкции от аномального состояния и выявить тенденции, указывающие на нежелательную деформацию моста, или спрогнозировать поведение моста во времени под действием нагрузки.

Ключевые слова: структурный мониторинг состояния; глубокое обучение; мост; машинное зрение; классификатор; киберфизические системы

Для цитирования: Флорес-Фуэнтес В. Классификатор изображений транспортных средств для определения корреляции их воздействия со смещением координат моста. Труды ИСП РАН, том 33, вып. 2, 2021 г., стр. 137-148. DOI: 10.15514/ISPRAS-2021-33(2)-8

Vehicle Image Classifier for Bridge Displacement Correlation

*W. Flores-Fuentes, ORCID: 0000-0002-1477-7449 <flores.wendy@uabc.edu.mx>
Autonomous University of Baja California (UABC),
Ensenada, Mexico, 21100*

Abstract. Advanced computing brings opportunities for innovation in a broad gamma of applications. Traditional practices based on visual and manual methods tend to be replaced by cyber-physical systems to automate processes. The present work introduces an example of this, a machine vision system research based on deep learning to classify bridge load, to give support to an optical scanning system for structural health monitoring tasks. The optical scanning system monitors the health of structures, such as buildings, warehouses, water dams, etc. by the measurement of their coordinates to identify if a coordinate displacement befalls that could indicate an anomaly in the structure that can be related to structural damage. The use of this optical scanning system to monitor the structural health of bridges is a little more complicated due to the vehicle's

transit over the bridge that causes a vehicle-bridge interaction which manifests as a bridge oscillation. Under this scheme, the bridge oscillation corresponds to their coordinate's displacement due to the vehicle-bridge interaction, but not necessarily due to bridge damage. So, a bridge load classifier is required to correlate the bridge coordinates measurements behavior with the bridge oscillation due to vehicle-bridge interaction to discriminate the normal behavior of the structure to abnormal behavior or identify tendencies that could indicate bridge deformation or discover if the bridge behavior due to loads is changing through the time.

Keywords: Structural Heal Monitoring; Deep Learning; Bridge; Machine Vision; Classifier; Cyber-physical

For citation: Flores-Fuentes W. Vehicle Image Classifier for Bridge Displacement Correlation. *Trudy ISP RAN/Proc. ISP RAS*, vol. 33, issue 2, 2021, pp. 137-148 (in Russian). DOI: 10.15514/ISPRAS-2021-33(2)-8.

1. Введение

Развитие городов [1, 2] сопровождается строительством новых транспортных сетей, в том числе автомобильных мостов.

Мостовые конструкции подвержены износу в результате строительных дефектов и различных внешних воздействий, в частности, а) использование неадекватных материалов и методов строительства, б) чрезмерный рост динамических нагрузок вследствие увеличения массы транспортного потока и его скоростей, с) перегрузка, d) естественные природные явления, такие как старения во времени, изменения температуры, ветровой нагрузки, влажности и сейсмических воздействий и даже e) столкновения автомобилей между собой и/или с элементами конструкции моста. Мониторинг целостности конструкций (Structural Health Monitoring, SHM) мостов является весьма актуальной практически важной задачей [3], он позволяет получить данные, помогающие принимать решения о профилактическом обслуживании, использовании и ремонте конструкций, а также позволяет отчетливо представлять требования к техническому обслуживанию моста и, следовательно, прогнозировать размер ресурсов, необходимых для его сохранения, и тем самым защищать инвестиции, вложенные в строительство моста.

Традиционная технология SHM мостов включает разрушающие испытания элементов мостовых конструкций, систематические визуальные осмотры и ручные оценки специализированным персоналом [4]. Это требует присутствия специалистов, что приводит к значительным финансовым затратам и длительному времени ожидания, прежде чем конструкция может быть снова использована, особенно в тех случаях, когда оценка состояния моста осложнена действием разрушающих факторов, например, после землетрясения, урагана или внезапного наводнения.

В то же время, системы SHM, основанные на нетрадиционных методах, обычно интегрируются с датчиками, подсистемами сбора, обработки и хранения данных, подсистемами связи, а также интеллектуальными системами для поведенческого моделирования конструкции и обнаружения повреждений на основе информации о поведении конструкции.

В научной литературе можно найти некоторые сведения об инновациях и разработке автономных систем, способных к непрерывному мониторингу, таких как оптические сканирующие системы (Optical Scanning System, OSS) [5, 6], которые позволяют непрерывно отслеживать контролируемые координаты конструкции для проверки, что она претерпевает смещения, которые могут указывать на то, что мост разрушается. Однако из-за взаимодействия транспортного средства с мостом (Vehicle-Bridge Interaction, VBI) мост колеблется с частотой приблизительно от 1 до 8 Гц [7], что приводит к тому, что контролируемая координата также имеет колебательное поведение. По этой причине для поддержки OSS требуется классификатор нагрузки моста (Bridge Load Classifier, BLC), чтобы можно было определить, соответствуют ли смещения отслеживаемых координат ухудшению состояния моста или же влиянию VBI.

Настоящее исследование посвящено разработке системы машинного зрения, основанной на машинном обучении и совмещенной с BLC, для обнаружения корреляции структурных смещений на мосту.

2. Глубокое обучение для классификации

В современном компьютерном мире возможно объединение сенсоров, информационных и коммуникационных технологий для разработки интеллектуальных систем [8]. Глубокое машинное обучение – это мощный аппарат, используемый для классификации. Он основан на искусственных нейронных сетях, состоящих из нескольких слоев.

Искусственный нейрон, также называемый перцептроном, имеет несколько двоичных входов $x_1, x_2, \dots, x_n$, значения которых умножаются на веса $w_1, w_2, \dots, w_n$, где w_i – значение, определяющее важность i -го входа для выходного значения нейрона. Если сумма входных значений, умноженных на их соответствующие веса, достигает определенного порога, перцептрон выводит двоичное значение. Когда выходы группы перцептронов питают другую группу перцептронов, каждая из этих групп называется слоем. Массив с несколькими слоями называется многослойным персептроном, в котором слои классифицируются на входные, скрытые и выходные. Без использования функции активации (f) нейроны являются двоичными, а веса и пороговые значения должны быть установлены программистом.

Когда нейрон для вычисления своего выходного значения использует функцию активации, говорят, что нейрон обучается в одиночку; нейроны получают определенное число входных значений и смещение (b), его входы могут быть реальными значениями, а смещение всегда представляет значение 1. Наиболее распространенными функциями активации являются линейные – двоичная ступенчатая и линейная, а также нелинейные – сигмоид, гиперболический тангенс, блок линейной ректификации (Rectified Linear Unit, ReLU), ReLU с утечкой (leaky ReLU), параметрический ReLU, softmax и переключатель. Выбор правильной функции активации для конкретного приложения определяет точность и вычислительную эффективность. Функции нелинейной активации позволяют создавать сложные сопоставления между входами и выходами сети для обучения и моделирования сложных данных. Современные нейронные сети основаны на нелинейных функциях активации. Наиболее популярным является ReLU, поскольку эта функция активации позволяет сети очень быстро сходиться и допускает обратное распространение ошибки. Новая функция активации, предложенная в [9], называется switch и обещает быть лучше, чем ReLU. Для выходного слоя полезна функция активации softmax, поскольку она позволяет классифицировать по нескольким классам.

Сети с разными соединениями слоев уступают место более сложным сетевым структурам, наиболее популярными из которых являются нейронная сеть прямого распространения (Feedforward Neural Network, FNN), рекуррентная нейронная сеть (Recurrent Neural Network, RNN) и сверточная нейронная сеть (Convolutional Neural Network, CNN).

Все эти конфигурации с несколькими слоями – глубокие нейронные сети (Deep Neural Network, DNN) определяются как модели глубокого обучения, задаваемые гиперпараметрами, количеством скрытых слоев, функцией активации и числом повторений (эпох) обучения.

2.1 Родственные работы

В последних приложениях обнаруживаются новшества: сочетание глубокого обучения с различными методами обнаружения объектов, отслеживания, классификации, аугментации, сжатия данных и прогнозирования.

Технология глубокого обучения в проектах структурного мониторинга использовалась для обнаружения трещин в бетоне [10, 11], для обнаружения структурных повреждений поверхности, таких как расслоение и обнажение арматуры моста [12, 13].

В [14] предлагается мониторинг смещений с использованием метода полного поля оптического потока, основанного на глубоком обучении. Реализованы алгоритмы оптического потока для визуального отслеживания и расчета полного поля структурного движения. В результате получается вектор движения в каждом пикселе изображения, а смещение получается путем преобразования смещения в пикселях в смещение в физических единицах после того, как движение, вызванное вибрацией камеры, смягчается путем вычитания движения статических частей в изображениях. Глубокое обучение также реализовано для классификации сигналов, например, в [15], где анализируются данные об перегрузках моста, получаемые с помощью акселерометров.

Для оценки транспортной нагрузки на мост в [16, 17] предлагается методология определения трафика, сочетающая методы компьютерного зрения и традиционные инструменты на основе измерения деформации для простых и сложных сценариев движения. Предлагается классификатор транспортных средств на основе глубокого обучения и сети тензометрических датчиков для сбора обширной информации (вес транспортного средства, скорость, количество, тип и траектория), наличие которой приводит к повышению точности результатов идентификации.

Было разработано несколько типов решений для классификации транспортных средств посредством глубокого обучения, а некоторые из них – для расчета нагрузки на мосты, но не было создан BLC, основанный на классификации транспортных средств с помощью глубокого обучения и случаев для корреляции структурных смещений.

Некоторые примеры обнаружения транспортных средств с помощью глубокого обучения для других приложений можно найти в [18-21].

3. Киберфизическая система для разработки SHM мостов

Смещение координат контролируемых точек – важный индикатор состояния моста. Киберфизическая система, предложенная для SHM, продемонстрировала свою эффективность при отслеживании смещений трехмерных координат [22]. Однако измерение смещения моста с помощью OSS по-прежнему является сложной задачей из-за колебаний моста, вызываемых нагрузкой от проходящего транспорта. Взаимодействие между автомобилем и мостом создает трехмерное координатное смещение, которое нужно учитывать особым образом. Таким образом, киберфизическая система SHM для моста состоит из а) системы оптического сканирования (OSS), описываемой в подразделе 3.1 и б) предлагаемого классификатора нагрузки на мост (BLC), описываемого в подразделе 3.2.

3.1 Система оптического сканирования (OSS)

Происхождение OSS связано с намерениями разработать систему технического зрения (Technical Vision System, TVS) для мобильных роботов, функционирующих в условиях сложной ландшафтной навигации [23, 24]. Система основывалась на системе лазерного позиционирования (Positioning Lazer, PL) и сканирующей апертуре (Scanning Aperture, SA), установленной на штативе; оба эти которые связаны по принципу триангуляции для определения координат объекта, обнаруженных TVS [25-27]. Когда система PL была заменена некогерентным источником света, который устанавливался над конструкцией в выбранной точке (x_i, y_i, z_i) , подлежащей мониторингу, эту новую конфигурацию стали называть OSS [28, 29]. На рис. 1 показана, как работает OSS.

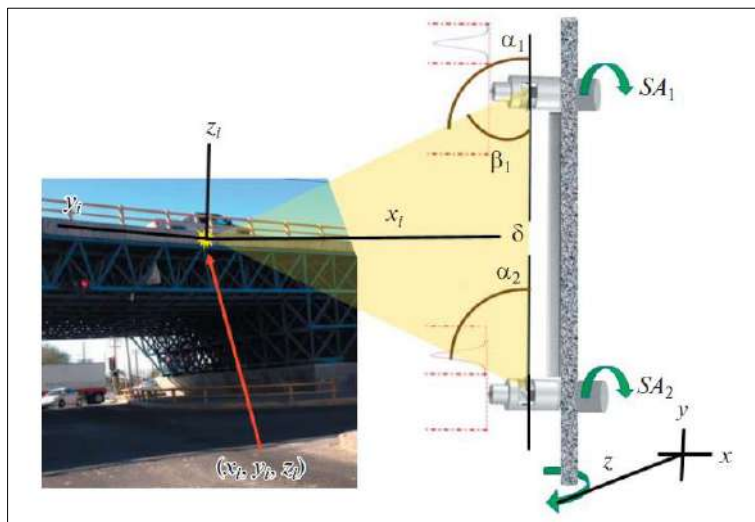


Рис. 1. OSS для измерения смещения моста
Fig. 1. OSS for bridge displacement measurement

TVS и OSS продолжают дорабатываться для использования в различных приложениях: автономная навигация [30], воздушная навигация [31], биометрические измерения, SHM различных конструкций и в различных средах [32]. В частном случае, эти системы могут быть использованы и для SHM мостовых конструкций. Вместе с тем, в последнем случае для системы SHM требуется оценка эффектов, связанных с воздействием транспортного средства на мост по трем координатам, измеряемых классификатором нагрузки моста. Эта оценка может быть реализована на основе глубокого обучения SHM для корреляции возникающих структурных смещений.

3.2. Классификатор нагрузки моста (BLC)

В этом подразделе описывается методология экспериментов и результаты исследований, выполненных при разработке классификатора нагрузки моста.

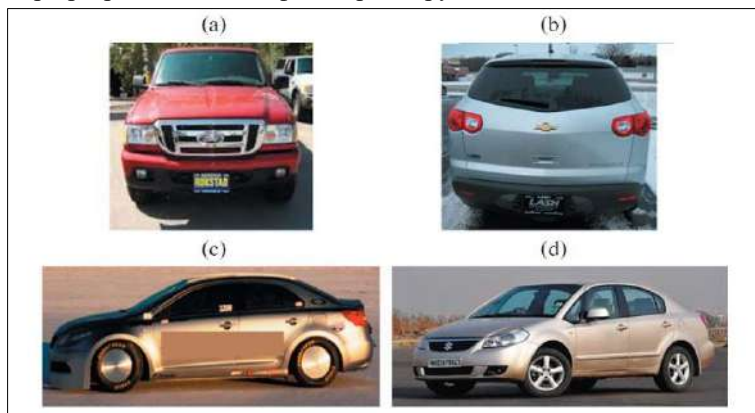


Рис. 2. Классы изображений автомобилей: а) вид спереди, б) вид сзади, в) вид сбоку, г) вид под углом
Fig 2. Vehicles classes: a) Front view, b) Rear view, c) Side view, d) Tilt view

Поскольку нагрузкой на мост являются проезжающие транспортные средства, то предлагается классификатор транспортных средств, основанный на анализе изображений посредством глубокого обучения изображений, полученных с уже установленных камер наблюдения транспортных средств. На этом этапе экспериментов желательно: определить

наилучшее положение и ориентацию камер для наблюдения за автомобилем. При грубой настройке изображения транспортных средств были классифицированы: вид спереди, вид сзади, вид сбоку и вид под углом, как показано на рис. 2.

Был использован существующий набор данных, который содержит 16185 изображений 196 классов транспортных средств. Полученные данные разделены на 8144 обучающих изображения и 8041 тестовое изображение, где каждый класс разделен примерно поровну. В каждом классе учитывались марка автомобиля, его модель, год выпуска, например, Tesla Model S 2012 года или купе BMW M3 2012 года [33]. Однако для этого эксперимента набор данных был заново разбит на классы в соответствии с табл. 1. Было сохранено 15935 изображений, а 250 изображений были отклонены из-за того, что они не подходят ни к одному классу.

Табл. 1. Заново разбитый на классы набор данных транспортных средств
Table.1. Vehicles dataset reclassification

Маркировка класса	Описание	Количество изображений
Класс 1	Вид спереди	733
Класс 2	Вид сзади	275
Класс 3	Вид сбоку	1020
Класс 4	Вид под углом	13907

Изображения в исходном наборе данных получались в разных условиях, из-за чего у них разные соотношения ширины и высоты. Поэтому изображения из новых четырех классов были приведены к размеру 100×100 пикселей с обычным режимом кадрирования (кадрирование по центру тяжести). Архитектура глубокого обучения была разработана на языке Python с использованием библиотеки Keras на платформе TensorFlow по общей схеме, показанной на рис. 3.

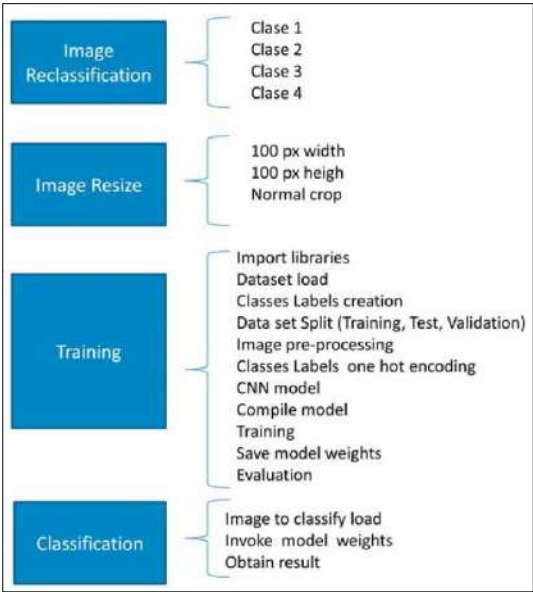


Рис. 3. Общая схема BLC
Fig 3. General flow diagram

Датасет был разделен следующим образом: 80% – для процедуры обучения и 20% – для тестирования. В свою очередь обучающий для обучения был разделен на 80% для обучения

и 20% для валидации между эпохами. Модель состояла из 9 слоев, различающихся формой вывода в соответствии с табл. 2, где параметр Cqty соответствует количеству классов, обученных в различных экспериментах.

Табл. 2. *Формы выхода модели CNN*
Table 2. *CNN model structure output shape*

Тип слоя	Форма вывода
Conv2D	(None, 100, 100, 32)
LeakyReLU	(None, 100, 100, 32)
MaxPooling	(None, 50, 50, 32)
Dropout	(None, 50, 50, 32)
Flatten	(None, 80000)
Dense	(None, 32)
LeakyReLU	(None, 32)
Dropout	(None, 32)
Dense	(None, Cqty)

В эксперименте 1 (E1) было использовано 15935 изображений, соответствующих четырем классам (см. табл. 1). Изображения были разделены следующим образом: 12748 – для обучения и 3187 – для тестирования. По результатам тестирования 2832 изображения были классифицированы правильно, а 355 изображений – неправильно. Полученные оценки точности обучения, валидации и тестирования представлены в строке 1 табл. 3. Результаты показателей классов приведены в табл. 4.

Точность (*precision*) модели P , т.е. доля правильных положительных прогнозов определяется формулой (1):

$$P = \frac{TP}{TP + FP}, \quad (1)$$

где TP – число истинно-положительных (*true positive*) предсказаний, FP – число ложно-положительных (*false positive*) предсказаний.

Полнота (*recall*) R , т.е. доля истинно-положительных результатов, которые получены для всех объектов, реально относящихся к тому же классу, вычисляется по формуле (2):

$$R = \frac{TP}{TP + FN}, \quad (2)$$

где FN – это число ложно-отрицательных (*false negative*) предсказаний.

F-мера (*F-Score*) – линейная комбинация точности и полноты, являющаяся единым показателем оценки качества методов классификации, определяется формулой (3).

$$F = 2 \times \frac{P * R}{P + R}. \quad (3)$$

В эксперименте 2 (E2) был исключен класс 2 (изображения с видом сзади) из-за неудовлетворенности результатами E1. В этом эксперименте использовалось 15660 изображений, соответствующих трем классам 1, 3 и 4 (табл. 1). Изображения были разделены следующим образом: 12528 – для обучения и 3132 – для тестирования. По результатам тестирования 2806 изображений были классифицированы правильно, а 326 изображений – неправильно. Точности обучения, валидации и тестирования представлены в строке 2 табл. 3. Результаты показателей классов приведены в табл. 5.

Во время эксперимента 3 (Е3) классы 1 и 2 (изображения с видами спереди и сзади) были исключены из-за неудовлетворенности результатами Е1 и Е2. В этом эксперименте было использовано 14927 изображений, соответствующих двум классам 3 и 4 (табл. 1). Изображения были разделены: 11941 – для обучения и 2986 – для тестирования. По результатам тестирования 2818 изображений были классифицированы правильно, а 168 изображений – неправильно. Точности обучения, валидации и тестирования представлены в строке 3 табл. 3. Результаты показателей классов приведены в табл. 6.

Табл. 3. Точность при выполнении экспериментов
Table 3. Experiments run accuracy

Эксперимент	Точность тренировки	Точность валидации	Точность тестирования
Е1	88.46%	89.57%	88.86%
Е2	89.58%	88.87%	89.56%
Е3	94.74%	94.35%	94.37%

Табл. 4. Показатели классов в эксперименте 1
Table 4. Experiment 1 Classes Metrics Results

Класс	Точность	Полнота	F-мера
1	0.88	0.15	0.26
2	0.00	0.00	0.00
3	0.86	0.20	0.32
4	0.89	1.00	0.94

Табл. 5. Показатели классов в эксперименте 2
Table 5. Experiment 2 Classes Metrics Results

Класс	Точность	Полнота	F-мера
1	0.82	0.06	0.11
3	1.00	0.02	0.04
4	0.90	1.00	0.94

Табл. 6. Показатели классов в эксперименте 3
Table 6. Experiment 3 Classes Metrics Results

Класс	Точность	Полнота	F-мера
3	0.80	0.25	0.39
4	0.95	1.00	0.97

4. Заключение

Структурному мониторингу состояния моста способствовали современные компьютерные технологии, которые позволили заменить традиционные методы, основанные на визуальных и ручных методах, киберфизическими системами для автоматизации процессов. Настоящая работа представляет киберфизическую систему для мониторинга состояния конструкций моста, состоящую из системы оптического сканирования и классификатора нагрузки на мост. Классификатор нагрузки моста – это система машинного зрения, основанная на глубоком обучении для классификации транспортных средств с целью поддержки измерений системы оптического сканирования. Данная система позволяет по измерениям смещения координат выявить аномалию в структуре моста и установить причину этих смещений, которая может

быть связана со структурными повреждениями или с временной нагрузкой от транспортного потока.

Предлагаемый классификатор транспортных средств основан на использовании изображений, полученных с установленных камер наблюдения, и их анализе посредством глубокого обучения. На этом этапе экспериментов были оценены наилучшее положение и ориентация камер для просмотра автомобиля. При грубом подходе к настройке изображения автомобилей были классифицированы: вид спереди, вид сзади, вид сбоку и вид под углом. Наилучшие результаты с точностью 0,95% были получены по изображениям автомобиля сбоку и под углом.

В будущей работе в центре внимания будет создание датасета изображений для проверки классификатора. В этот датасет будет включен класс изображений вида сверху. Кроме того, планируется внедрение инфраструктуры облачных вычислений для эффективной обработки больших объемов данных [34-37]. Мы также продолжим измерение координат и анализ корреляции данных нагрузки моста с помощью методов двумерного и многомерного анализа данных, основанных на коэффициентах корреляции Фишера и Пирсона, анализе главных компонент, регрессии методом дробных наименьших квадратов и каноническом корреляционном анализе с целью выявления наиболее эффективного метода для решения поставленной задачи.

Список литературы / References

- [1]. R. Massobrio, S. Nesmachnow, A. Tchernykh et al. Towards a cloud computing paradigm for big data analysis in smart cities. *Programming and Computer Software*, vol. 44, no. 3, 2018, pp. 181-189.
- [2]. Р. Массобрио, С. Несмачнов, А. Черных и др. Применение облачных вычислений для анализа данных большого объема в умных городах. *Труды ИСП РАН*, том 28, вып. 6, 2016 г., стр. 121-140 / R. Massobrio, S. Nesmachnow, A. Tchernykh et al. Towards a Cloud Computing Paradigm for Big Data Analysis in Smart Cities. *Trudy ISP RAN/Proc. ISP RAS*, vol. 28, issue 6, 2016. pp. 121-140 (in Russian). DOI: 10.15514/ISPRAS-2016-28(6)-9.
- [3]. A. Mufti, K. Helmi. A case for structural health monitoring (SHM) and civionics enhances the evaluation of the load carrying capacity of aging bridges. *Innovative Infrastructure Solutions*, vol. 4, no. 1, pp. 3, 2019.
- [4]. M. Barousse Moreno, A. Galindo Solozano. Sistema de Administracion de Puentes (SIAP). Instituto Mexicano del Transporte, Publicación Técnica, vol. 49, 1994, 88 p. (in Spanish).
- [5]. W. Flores-Fuentes, M. Rivas-Lopez, O. Sergiyenko et al. Energy center detection in light scanning sensors for structural health monitoring accuracy enhancement. *IEEE Sensors Journal*, vol. 14, no. 7, 2014, pp. 2355-2361.
- [6]. J. Rivera-Castillo, W. Flores-Fuentes, M. Rivas-López et al. Experimental image and range scanner datasets fusion in SHM for displacement detection. *Structural Control and Health Monitoring*, vol. 24, no. 10, 2017, article e1967.
- [7]. L. Ma, W. Zhang, W. S. Han, & J. X. Liu. Determining the dynamic amplification factor of multi-span continuous box girder bridges in highways using vehicle-bridge interaction analyses. *Engineering Structures*, vol. 181, 2019, pp. 47-59.
- [8]. C. Emmanouilidis, P. Pistofidis, L. Bertonecelj et al. Enabling the human in the loop: Linked data and knowledge in industrial cyber-physical systems. *Annual Reviews in Control*, vol. 47, 2019, pp. 249-265.
- [9]. P. Ramachandran, B. Zoph, & Q.V. Le. Swish: a self-gated activation function. *arXiv preprint arXiv:1710.05941*, 2017.
- [10]. Y. Ren, J. Huang, Z. Hong et al. Image-based concrete crack detection in tunnels using deep fully convolutional networks. *Construction and Building Materials*, vol. 234, 2020, article no. 117367.
- [11]. N. S. Gulgec, M. Takáč, & S. N. Pakzad. Experimental Study on Digital Image Correlation for Deep Learning-Based Damage Diagnostic. In *Dynamics of Civil Structures*, vol. 2, Springer, 2020, pp. 205-210.
- [12]. W. Deng, Y. Mou, T. Kashiwa et al. Vision based pixel-level bridge structural damage detection using a link ASPP network. *Automation in Construction*, vol. 110, 2020, article no. 102973.
- [13]. J. J. Rubio, T. Kashiwa, T. Laiteerapong et al. Multi-class structural damage segmentation using fully convolutional networks. *Computers in Industry*, vol. 112, 2019, article no. 103121.

- [14]. W. Deng, Y. Mou, T. Kashiwa et al. Vision based pixel-level bridge structural damage detection using a link ASPP network. *Automation in Construction*, vol. 110, 2020, article no. 102973.
- [15]. Y. Bao, Z. Tang, H. Li, & Y. Zhang. Computer vision and deep learning-based data anomaly detection method for structural health monitoring. *Structural Health Monitoring*, vol. 18, no. 2, 2019, pp. 401-421.
- [16]. Y. Xia, X. Jian, B. Yan, & D. Su. Infrastructure Safety Oriented Traffic Load Monitoring Using Multi-Sensor and Single Camera for Short and Medium Span Bridges. *Remote Sensing*, vol. 11, no. 22, 2019, article no. 2651.
- [17]. X. Jian, Y. Xia, J. A. Lozano-Galant, & L. Sun. Traffic Sensing Methodology Combining Influence Line Theory and Computer Vision Techniques for Girder Bridges. *Journal of Sensors*, Special Issue, 2019, Article ID 3409525.
- [18]. D.K. Amara, R. Karthika, & K.P. Soman. DeepTrackNet: Camera Based End to End Deep Learning Framework for Real Time Detection, Localization and Tracking for Autonomous Vehicles. *Advances in Intelligent Systems and Computing book series*, vol. 1039, 2019, pp. 299-307.
- [19]. C. Cardenas, & M. Gonzalez-Mendoza. Distributed System Based on Deep Learning for Vehicular Re-routing and Congestion Avoidance. *Advances in Intelligent Systems and Computing book series*, vol. 1071, 2019, pp. 159-172.
- [20]. S. Zhang, C. Wang, Z. He et al. Vehicle global 6-DoF pose estimation under traffic surveillance camera. *ISPRS Journal of Photogrammetry and Remote Sensing*, vol. 159, 2020, pp. 114-128.
- [21]. C.K. Ng, S.N. Cheong, & Y.L. Foo. Low Latency Deep Learning Based Parking Occupancy Detection By Exploiting Structural Similarity. *Lecture Notes in Electrical Engineering*, vol. 603, 2020, pp. 247-256.
- [22]. J. E. Miranda-Vega, W. Flores-Fuentes, O. Sergiyenko et al. Optical cyber-physical system embedded on an FPGA for 3D measurement in structural health monitoring tasks. *Microprocessors and Microsystems*, vol. 56, 2018, pp. 121-133.
- [23]. V. Tyrsa, O. Sergiyenko, L. Burtseva et al. Mobile transport object control by technical vision means. In *Proc. of the Electronics, Robotics and Automotive Mechanics Conference (CERMA'06)*, vol. 2, 2005, pp. 74-82.
- [24]. М.В. Иванов, О.Ю. Сергиенко, В.В. Тырса и др. Интеграция беспроводной связи для оптимизации распознавания окружения и расчёта траектории движения группы роботов. *Труды ИСПРАН*, том 31, вып. 2, 2019 г., стр. 67-82. DOI: 10.15514/ISPRAS-2019-31(2)-6 / M. Ivanov, O. Sergiyenko, V. Tyrsa et al. Software Advances using n-agents Wireless Communication Integration for Optimization of Surrounding Recognition and Robotic Group Dead Reckoning. *Programming and Computer Software*, vol. 45, no. 8, 2019, pp. 557-569.
- [25]. M. Rivas, O. Sergiyenko, M. Aguirre et al. Spatial data acquisition by laser scanning for robot or SHM task. In *Proc. of the 2008 IEEE International Symposium on Industrial Electronics*, 2008, pp. 1458-1462.
- [26]. M. R. López, O. Sergiyenko, & V. Tyrsa. Machine vision: approaches and limitations. In *Computer vision*. IntechOpen, 2008, pp. 395-428.
- [27]. O. Sergiyenko, V. Tyrsa, D. Hernandez-Balbuena et al. Precise optical scanning for practical multi-applications. In *Proc. of the 2008 34th Annual Conference of IEEE Industrial Electronics* (pp. 1656-1661). IEEE. 2008, November.
- [28]. W. Flores-Fuentes, M. Rivas-Lopez, O. Sergiyenko et al. Energy center detection in light scanning sensors for structural health monitoring accuracy enhancement. *IEEE Sensors Journal*, vol. 14, no. 7, 2014, pp. 2355-2361.
- [29]. W. Flores-Fuentes, M. Rivas-Lopez, O. Sergiyenko et al. Combined application of power spectrum centroid and support vector machines for measurement improvement in optical scanning systems. *Signal Processing*, vol. 98, 2014, pp. 37-51.
- [30]. M. Reyes-Garcia, C. Sepulveda-Valdez, O. Sergiyenko et al. Digital Control Theory Application and Signal Processing in a Laser Scanning System Applied for Mobile Robotics. In *Control and Signal Processing Applications for Mobile and Aerial Robotic Systems*, IGI Global, 2020, pp. 215-265.
- [31]. L. Lindner, O. Sergiyenko, M. Rivas-López et al. Machine vision system errors for unmanned aerial vehicle navigation. In *Proc. of the 2017 IEEE 26th International Symposium on Industrial Electronics (ISIE)*, 2017, pp. 1615-1620.
- [32]. W. Flores-Fuentes, J. Miranda-Vega, M. Rivas-López et al. Comparison between Different Types of Sensors Used in the Real Operational Environment Based on Optical Scanning System. *Sensors*, vol. 18, no. 6, 2018, article no. 1684.
- [33]. J. Krause, M. Stark, J. Deng, & L. Fei-Fei. 3d object representations for fine-grained categorization. In *Proc. of the IEEE International Conference on Computer Vision Workshops*, 2013, pp. 554-561.

- [34].Н.П. Варновский, С.А. Мартишин, М.В. Храпченко, А.В. Шокуров. Методы пороговой криптографии для защиты облачных вычислений. Труды ИСП РАН, том 26, вып. 2, 2015 г., стр. 269-274 / N.P. Varnovskiy, S.A. Martishin, M.V. Khrapchenko, & A.V. Shokurov (2015). Secure cloud computing based on threshold homomorphic encryption. *Programming and Computer Software*, vol. 41, no. 4, 2015, pp. 215-218.
- [35].Miranda-López, V., Tchernykh, A., Cortés-Mendoza et al. Experimental analysis of secret sharing schemes for cloud storage based on RNS. *Communications in Computer and Information Science (CCIS)*, vol. 79, 2017, pp. 370-383.
- [36].A. Tchernykh, M. Babenko, N. Chervyakov et al. Towards mitigating uncertainty of data security breaches and collusion in cloud computing. In *Proc. of the 28th International Workshop on Database and Expert Systems Applications (DEXA)*, 2017, pp. 137-141.
- [37].A. Tchernykh, M. Babenko, N Chervyakov et al. AC-RRNS: Anti-collusion secured data sharing scheme for cloud storage. *International Journal of Approximate Reasoning*, vol. 102, 2018, pp. 60-73.

Информация об авторе / Information about the author

Венди ФЛОРЕС-ФУАНТЕС, кандидат наук, профессор-исследователь инженерного факультета. Область научных интересов: оптоэлектроника, машинное зрение, искусственный интеллект.

Wendy FLORES-FUENTES, Doctor in Sciences, Research-Professor at Engineering Faculty. Research interests: Optoelectronics, Machine vision, Artificial Intelligence.

DOI: 10.15514/ISPRAS-2020-33(2)-9



Обнаружение объектов в аэронавигации с использованием вейвлет-преобразования и сверточных нейронных сетей: первый подход

¹ *Х.М. Фортуна-Сервантес, ORCID: 0000-0002-9229-3159
<juan.manuel.fortuna@hotmail.com>*

¹ *М.Т. Рамирес-Торрес, ORCID: 0000-0002-7457-7318 <tulio.torres@uaslp.mx>*

² *Х. Мартинес-Карранса, ORCID: 0000-0001-8123-0008 <carranza@inaoep.mx>*

¹ *Х.С. Мургуиа-Ибарра, ORCID: 0000-0001-7239-8968 <ondeleto@uaslp.mx>*

¹ *М. Мехиа-Карлос, ORCID: 0000-0003-2872-9461 <marce.mejia@uaslp.mx>*

¹ *Автономный университет Сан-Луис-Потоси,*

Мексика, 78000, Сан-Луис-Потоси, Альваро Обрегон, 64

² *Национальный институт оптической и электронной астрофизики,
Мексика, 72840, Пуэбла, Тонанцинтла, Луис Энрике Эрро, 1*

Аннотация. Предлагается первый подход, основанный на применении вейвлет-анализа при обработке изображений с целью обнаружения объектов с повторяющимися чертами и двоичной классификации в плоскости изображения, в частности, для навигации в симулируемых средах. На сегодняшний день стало привычным использовать алгоритмы на основе сверточных нейронных сетей (Convolutional Neural Networks, CNN) для обработки изображений, полученных с бортовых камер беспилотных летательных аппаратов (Unmanned Aerial Vehicles, UAV), в пространственной области, что способствует решению задач обнаружения и классификации. Архитектура CNN позволяет обучать сеть, используя в качестве входных данных изображения без предварительной обработки. Это позволяет извлекать характерные признаки изображения. Тем не менее, в этой работе мы утверждаем, что спектральные характеристики изображений на разных частотах, низких и высоких, также влияют на производительность CNN во время обучения. Мы предлагаем архитектуру CNN, дополненную двумерным дискретным вейвлет-преобразованием как методом выделения признаков. Такая информация улучшает способность сети к обучению, устраняет переобучение и обеспечивает более высокую эффективность при обнаружении цели.

Ключевые слова: сверточная нейронная сеть; вейвлет-анализ; обнаружение объектов; дрон; классификация объектов; среда симуляции Gazebo

Для цитирования: Фортуна-Сервантес Х.М., Рамирес Торрес М.Т., Мартинес-Карранса Х., Мургуиа-Ибарра Х.С., Мехиа Карлос М. Обнаружение объектов в аэронавигации с использованием вейвлет-преобразования и сверточных нейронных сетей: первый подход. Труды ИСП РАН, том 33, вып. 2, 2021 г., стр. 149-162. DOI: 10.15514/ISPRAS-2021-33(2)-9

Благодарности. Х.М. Фортуна-Сервантес – докторант CONACYT (Мексика) по программе «Ciencias Aplicadas» в ИСКО-UASLP. Мы благодарим INAOE за предоставление условий для проведения исследовательской стажировки, в рамках которой была выполнена часть этой работы.

Object Detection in Aerial Navigation using Wavelet Transform and Convolutional Neural Networks: A first Approach

¹ J.M. Fortuna-Cervantes, ORCID: 0000-0002-9229-3159 <juan.manuel.fortuna@hotmail.com>

¹ M.T. Ramírez-Torres, ORCID: 0000-0002-7457-7318 <tulio.torres@uaslp.mx>

² J. Martínez-Carranza, ORCID: 0000-0001-8123-0008 <carranza@inaoe.mx>

¹ J.S. Murguía-Ibarra, ORCID: 0000-0001-7239-8968 <ondeleto@uaslp.mx>

¹ M. Mejía-Carlos, ORCID: 0000-0003-2872-9461 <marce.mejia@uaslp.mx>

¹ Universidad Autónoma de San Luis Potosí,

Álvaro Obregón #64, Col. Centro, San Luis Potosí, C.P. 78000, México

² Instituto Nacional de Astrofísica Óptica y Electrónica,

Luis Enrique Erro # 1, Tonantzintla, Puebla, C.P. 72840, México

Abstract. This paper proposes a first approach based on wavelet analysis inside image processing for object detection with a repetitive pattern and binary classification in the image plane, in particular for navigation in simulated environments. To date, it has become common to use algorithms based on convolutional neural networks (CNNs) to process images obtained from the on-board camera of unmanned aerial vehicles (UAVs) in the spatial domain, being useful in detection and classification tasks. CNN architecture can receive images without pre-processing, as input in the training stage. This advantage allows us to extract the characteristic features of the image/ Nevertheless, in this work, we argue that characteristics at different frequencies, low and high, also affect the performance of CNN during training. Thus, we propose a CNN architecture complemented by the 2D discrete wavelet transform, which is a feature extraction method. The information improves the learning capacity, eliminates the overfitting, and achieves a better efficiency in the detection of a target.

Keywords: CNN; Wavelet Analysis; Object Detection, Drone, Object Classification, Gazebo Simulation Environment.

For citation: Fortuna-Cervantes J.M., Ramírez-Torres M.T., Martínez-Carranza J., Murguía-Ibarra J.S., Mejía-Carlos M. Object Detection in Aerial Navigation using Wavelet Transform and Convolutional Neural Networks: A first Approach. *Trudy ISP RAN/Proc. ISP RAS*, vol. 33, issue 2, 2021, pp. 149-162 (in Russian). DOI: 10.15514/ISPRAS-2021-33(2)–.

Acknowledgments. J.M. Fortuna-Cervantes is a doctoral fellow of CONACYT (México) in the program of «Ciencias Aplicadas» at IICO-UASLP. We thank INAOE for giving the facilities to carry out the research internship where part of this work was done.

1. Введение

В последние годы в различных областях робототехники, в частности для автономной навигации беспилотных летательных аппаратов (Unmanned Aerial Vehicles, UAV), или дронов, активно используются алгоритмы визуального сервоуправления [1, 2]. В некоторых случаях возникает необходимость обеспечивать более высокую надежность при решении задач обнаружения, поскольку это оказывает существенное влияние на автономность робота. Для визуального восприятия, подобно цвету и форме, важна текстура изображения, так как она обеспечивает информацию о структуре поверхности и объектов на изображении [3]. Однако иногда обнаружение и отслеживание объекта на сцене бывает осложнено изменениями в освещении, масштабе и перспективе камеры [4, 5]. Поэтому использование таких методов обработки изображений, как вейвлет-анализ, глубокие нейронные сети или сверточные нейронные сети (Convolutional Neural Networks, CNN), стало очевидной альтернативой традиционным подходам [6-8].

В этой работе мы предлагаем подход, основанный на вейвлет-анализе как методе извлечения спектральных характеристик в сочетании с архитектурой CNN. Архитектура CNN явно допускает, что входными данными являются изображения в вейвлет-доме. Использование изображений в вейвлет-доме повышает способность к обучению на этапе обучения по

сравнению со случаем, когда для обучения используются изображения только в пространственном домене. Кроме того, если обучающий датасет невелик, то это позволяет избежать переобучения при обобщении. В результате при использовании для навигации оказывается возможна валидация модели обучения, в которой дрон может распознавать объекты с повторяющимися чертами (например, текстурой) и обучаться на этом. Модель обнаружения изображения с бортовой камеры дрона кадр за кадром и классифицирует изображения; предопределены два выходных класса: объект с текстурой (Texture) или не объект с текстурой (NotTexture).

Статья организована следующим образом: в разд. 2 описываются работы по родственной тематике; в разд. 3 представлена предлагаемая методология; разд. 4 демонстрирует экспериментальную часть и результаты; в разд. 5 представлено заключение.

2. Работы по родственной тематике

В последние годы проблема обнаружения объектов в приложениях UAV стала объектом активного исследования [8, 9]. Техника визуальной обработки, которая в настоящее время дает отличные результаты, основывается на архитектуре глубокого обучения. В некоторых работах по автономной навигации предлагается использовать CNN в реальных и моделируемых средах; например, в работе [10] авторы предлагают методологию обнаружения и избегания препятствий. Они используют предварительно обученную сеть AlexNet [11], обучение которой производит нейронная сеть меньшего размера.

Метод, предложенный в [12], – это нейронная архитектура YOLO, которая хорошо зарекомендовала себя в области обнаружения объектов, обрабатывая изображения в реальном времени со скоростью 45 кадров в секунду. Помимо этого, для решения задач распознавания объектов авторы работы [13] предлагают подход, основанный на глубоком обучении, для надежного определения центра объекта. Генерация линии прямой видимости в качестве направляющей позволяет избежать столкновений с другими объектами, которые могут происходить вследствие изменения таких условий, как освещение, геометрия объекта и наложение в плоскости изображения.

С другой стороны, в ряде проектов в области визуальной обработки используются методы глубокого обучения и вейвлет-анализа. Например, применительно к классификации изображений, метод, предложенный в [14], преобразует изображения из базы данных CIFAR-10 и KDEF в вейвлет-домен, получая таким образом временные и частотные характеристики. Различные представления изображений используются в нескольких архитектурах CNN; эта комбинация информации в вейвлет-домене обеспечивает более высокую эффективность обнаружения и более короткое время выполнения по сравнению с использованием только пространственного домена.

Авторы [15] предлагают другую альтернативу – вейвлет-пулинг как слой в архитектуре CNN. Этот метод разделяет карты признаков на два поддиапазона, отбрасывая карты первого уровня, чтобы уменьшить размер карт признаков. Преобразование карты признаков в вейвлет-домен улучшает классификацию изображений базы данных MNIST. Кроме того, этот подход поддерживает структурное сжатие данных, уменьшает образование неровных краев и других дефектов в изображении.

Объединение в архитектуре CNN инфракрасных и видимых фотографий обеспечивает эффективный метод обнаружения. Объединение основывается на декомпозиции изображения на основе вейвлет-анализа, и реконструируемое изображение лучше воспринимается зрительной системой человека [16].

В работе [17] представлены два метода выделения границ изображений с целью их классификации. Первый метод декомпозирует изображения на основе вейвлет-преобразования, а затем реконструирует их ограниченным образом. Второй метод, который

создает улучшенные изображения для ввода в нейронную сеть с использованием модулей локальных максимумов вейвлет-коэффициентов. Оба метода применяются для предварительной обработки изображений.

Говоря о классификации текстур в приложениях обработки изображений, авторы [18] предлагают вейвлет-CNN для обеспечения возможности обобщения спектральной информации, которая теряется в обычных CNN. Эта информация полезна для классификации текстур, поскольку обычно содержит достаточно сведений о форме объекта. Модель позволяет иметь меньше параметров, чем в традиционных CNN, и поэтому проводить обучение с меньшим объемом памяти.

Таким образом, обзор современных публикаций показывает, что алгоритмы вычислительного интеллекта улучшают стратегии обнаружения в приложениях UAV. Достигается приспособляемость к изменениям окружающей среды, освещенности, масштаба и проч. В отличие от работ, упомянутых выше, в настоящей работе основное внимание уделяется архитектуре CNN в сочетании с вейвлет-анализом. В результате такого подхода дрон лучше обнаруживает объекты с повторяющимися чертами, например, текстурой. Кроме того, дрон использует спектральную информацию о форме объекта, что увеличивает способность к обучению. Также исключается переобучение на этапе обучения, в отличие от случая наличия только пространственных данных и использования традиционной архитектуры CNN.

3. Материалы и методы

3.1 Кратномасштабный анализ

Алгоритм кратномасштабного анализа Малла (Stéphane Georges Mallat, Multiresolution Analysis, MA) обеспечивает связь между вейвлетами и наборами фильтров [19-21]. Кратномасштабная декомпозиция двумерной функции или изображения представляется серией приближений и деталей во вспомогательных изображениях. На первом уровне декомпозиции применяются два фильтра соответственно – низкочастотный (h) и высокочастотный (g), за каждым из которых следует операция субдискретизации с коэффициентом 2, как показано на рис. 1.

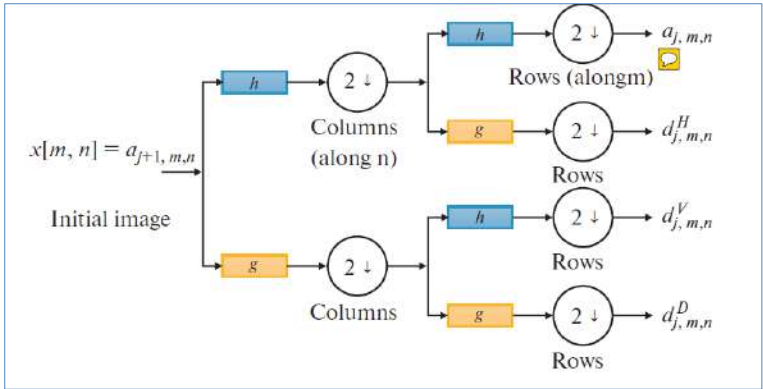


Рис 1. Первый уровень декомпозиции, примененный к изображению с использованием набора фильтров [22]

Fig 1. The first level of decomposition applied to an image using the filter bank [22]

Результат применения трех уровней вейвлет-декомпозиции к изображению ($x[m, n]$) размером $M \times N$ пикселей показан на рис. 2. После этого двумерный сигнал проходит через структуру набора фильтров, показанную на рис. 1. Получаются четыре вспомогательных

изображения с $M/2$ строками и $N/2$ столбцами; то есть каждое из четырех подизображений имеет четверть пикселей исходного изображения. Аппроксимация вспомогательного изображения достигается аппроксимационными вычислениями сначала по строкам, а затем по столбцам исходного изображения. Это подизображение представляет собой осредненную версию изображения ($x[m, n]$) с одной четвертой разрешения и аналогичными статистическими свойствами, аналогичными исходному сигналу [23]. Остальные подизображения показывают специфические характеристики исходного изображения в определенном направлении, то есть обеспечивают горизонтальный, вертикальный и диагональный коэффициенты детализации. Такое же преобразование формы сигнала применяется аппроксимированному подизображению для определения следующего уровня декомпозиции. Снова получаются четыре подизображения, но теперь с $M/2^2$ строками и $N/2^2$ столбцами. Эта итерация повторяется до достижения желаемого уровня разрешения или до уровня, допускаемого размерами изображения [22].

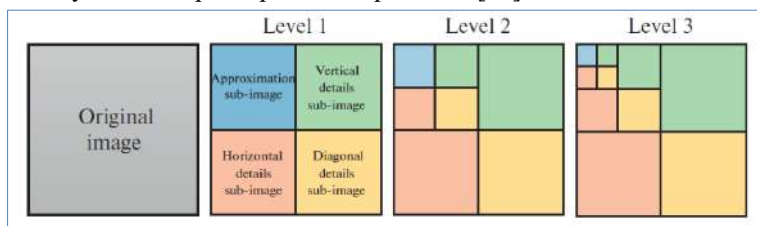


Рис. 2. Процесс декомпозиции с применением трех уровней набора фильтров, результатом которой является некоторая аппроксимация и детализация подизображений

Fig 2. Decomposition process applying three levels of the filter bank, which results are some approximation and detail sub-images

В общем случае кратномасштабное разложение двумерного сигнала выявляет различия в уровнях разрешения. Детали показываются в различных ориентациях, из чего следует, что метод двумерного дискретного вейвлет-преобразования (Two-Dimensional Discrete Wavelet Transform, 2D-DWT) хорошо подходит для обнаружения важной информации из исходного двумерного сигнала или изображения. Это существенно для таких задач обработки изображений, как обнаружение границ, распознавание изображений, классификация текстур и повышение качества изображений [3].

3.2 Сверточные нейронные сети

CNN широко используются при решении задач компьютерного зрения. CNN формируются из нейронов и обладают параметрами в виде весов и смещений, которые позволяют сети обучаться [24-26]. Эти сети состоят из входного и выходного слоев, а также нескольких скрытых слоев, некоторые из которых являются сверточными, откуда и происходит название этого вида нейронных сетей [24]. Операция свертки выполняется внутри сети на всех картах признаков сверточного слоя. Кроме того, при распространении весов и смещений эти операции применяются в направлении от первого входного слоя к последнему скрытому слою [14]. Формула (1) задает математическое представление распределения веса желаемого фильтра:

$$y_{ij} = \sigma \left(b + \sum_{l=0}^{n-1} \sum_{m=0}^{n-1} W_{l,m} a_{j+l,k+m} \right), \quad (1)$$

где $W_{l,m}$ представляет распределенные веса, b – смещение, $a_{j+l,k+m}$ – функция активации в заданном положении, n – размер окна фильтра.

Как следствие, использование сверточных слоев позволяет CNN обучаться разным уровням абстракции. Отличительной особенностью сверточной нейронной сети является наличие

явного предположения, что входными данными являются изображения; это позволяет закладывать в архитектуру специальные возможности распознавания конкретных элементов [27]. В общем случае сети с несколькими слоями могут выявлять во входных данных более сложные структуры. При наличии в CNN нескольких слоев увеличивается число параметров для обучения и поиска наилучшего решения. Имеются архитектуры глубоких нейронных сетей, подобные VGG16 [28], VGG 19 [28], AlexNet [11], ConvNet [29], SD [30], YOLO [12], с положительными результатами в областях классификации изображений и обнаружения объектов [31,32].

Тем не менее, в данном исследовании, в котором мы хотели обойтись небольшим датасетом для обучения модели, мы решили использовать архитектуру CNN, на которой выполнялась классификация изображений собак и кошек [33]. Использованный нами датасет не размещен в библиотеке Keras и поэтому нам пришлось создать наш собственный экспериментальный датасет. CNN представляет собой стек двумерных сверточных слоев с функцией активации блока линейной ректификации (Rectified Linear Unit, ReLU), чередующихся со слоями MaxPooling 2D. Кроме того, значение глубины скрытых слоев постепенно увеличивается с 32 до 128, в то время как размер карт признаков уменьшается с 62×62 до 2×2 , как показано на рис. 3. При использовании бинарной классификации сеть завершается одним блоком (сжатый слой размером 1) и сигмоидальной функцией активации.

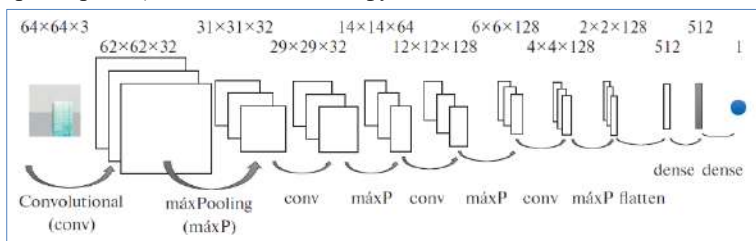


Рис 3. Архитектура глубокой нейронной сети (ConvNet) с двоичным выходом
Fig 3. Architecture of deep neural network (ConvNet) with binary output

4. Эксперименты и результаты

В этом разделе представлена экспериментальная система с двумя моделями обнаружения. В первой модели используются изображения в пространственном домене, то есть исходные изображения без предварительной обработки. Вторая модель использует изображения в вейвлет-доме, поэтому перед входом в нейронную сеть изображения предварительно обрабатываются на трех уровнях декомпозиции. Эти два датасета ранее были получены на этапе распознавания и навигации в среде моделирования Gazebo. Кроме того, необходимо отметить, что в обеих моделях была использована архитектура ConvNet, показанная на рис. 3. Нейронная сеть обучалась с использованием среды машинного обучения Keras, а в качестве бэкенда использовалась библиотека TensorFlow.

Как правило, оценка обучающей способности модели и эффективности обнаружения производится сравнением показателей точности и потерь. Это две статистики обычно собираются на этапе обучения, валидации и тестирования. Наконец, наша модель оценивалась в аэронавигационном приложении с использованием ROS и среды моделирования Gazebo.

4.1 Датасеты

Две модели обнаружения имеют бинарный выход, так что задача обнаружения сосредоточена на двух классах: первый класс обнаруживает присутствие текстурированного объекта в плоскости изображения, а во второй класс попадают изображения, для которых объект

находится вне сцены. Для каждого класса датасет содержит 700 изображений для этапа обучения, 150 для валидации и 105 для тестирования (955 изображений на класс). Архитектура ConvNet для двух моделей обнаружения параметризована для обучения и прогнозирования только изображений с размером 64×64 пикселя и тремя каналами RGB (красный, зеленый, синий). Поэтому для первой модели обнаружения размер исходных изображений (640×380 пикселей) изменяются до 64×64 пикселей. Во второй модели используются вейвлет-изображения; поэтому, во-первых, исходные изображения (640×380 пикселей) приводятся к размеру 512×512 пикселей для кратномасштабного анализа. В этом процессе генерируются одно аппроксимационное и три детализированных (горизонтальное, диагональное и вертикальное) вспомогательных изображения. Как показано на рис. 4, вейвлет-процесс применяется к двум классифицированным датасетам (классам). В этом случае предлагается использовать только аппроксимационные подмножества (изображения размером 64×64 пикселей), поскольку именно в этом месте можно сэкономить больше всего энергии.

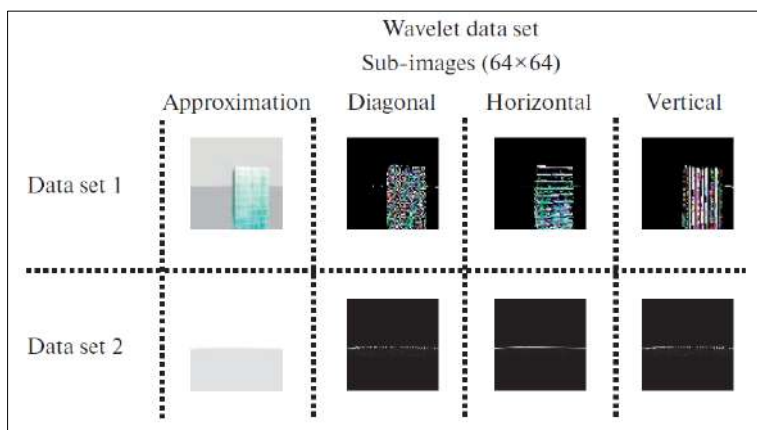


Рис. 4. Набор данных вейвлет-подизображения; в этом случае мы сосредотачиваемся только на аппроксимированных подизображениях

Fig 4. Wavelet sub-image dataset for the second detection model; in this case, we only focus on the approximation sub-images

4.2 Оценка модели

Эти две модели обнаружения используются для демонстрации вклада вейвлет-анализа в сочетании с архитектурой CNN. Кроме того, экспериментальная разработка позволяет наблюдать за поведением обеих моделей при обучении, поэтому показатели точности и потерь выбираются за десять эпох (количество итераций, в которых должно производиться обучение на основе датасета) на этапах обучения и валидации. Таким образом, обе модели обнаружения обучаются с использованием 504 001 параметра на компьютере с процессором Intel Core i5-2450M.

На рис. 5 приведены результаты, демонстрируемые во всех десяти эпохах первой моделью обучения, для которой характерно использование только исходных изображений и ConvNet для бинарной классификации [33]. Точность на этапе обучения (зеленая линия) начинается с 78%, затем достигает почти 100% во вторую эпоху; с этого момента поведение обучения является случайным между 98% и 100%. Что касается точности на этапе валидации (синяя линия), то обобщение обучения уменьшается при обучении на новых данных. Этот эффект вызван переобучением (после трех эпох); то есть сеть начинается обучаться паттернам, которые характерны для обучающих данных, но некорректны или неуместны по отношению к новым данным. Между тем, на рис. 6 показаны показатели потерь для оценки модели

обучения. Значение потерь, близкое к нулю, очень быстро достигается на этапе обучения (зеленая линия); в результате этого эффекта сеть становится более восприимчивой к переобучению. Таким образом, в случае новых данных потери на этапе валидации (синяя линия) уменьшаются и значительно увеличиваются (три и семь эпох), сохраняя эффект переобучения.

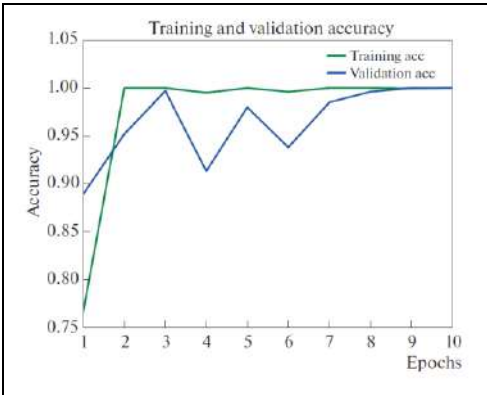


Рис. 5. Точность обучения и валидации модели без предварительной обработки входных изображений в Convnet
Fig. 5. Capacity of the model in the accuracy of training and validation, without pre-processing of the input images to the ConvNet

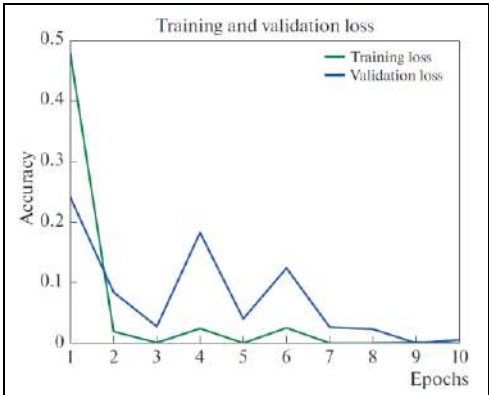


Рис. 6. Потери обучения и валидации модели без предварительной обработки входных изображений в ConvNet
Fig 6. Capacity of the model in the loss of training and validation, without pre-processing of the input images to the ConvNet

Вторая модель обучения находится в тех же условиях архитектуры ConvNet, но с добавлением изображений в вейвлет-доме (только аппроксимационные подизображения). Ниже представлены результаты нашего подхода с использованием вейвлет-анализа за десять эпох обучения и валидации. Как показано на рис. 7, модель имеет лучшую производительность при обобщении обучения, а также позволяют избежать эффекта переобучения по отношению к новым изображениям датасета валидации. На этапе обучения сначала достигается 68% точности (зеленая линия), но на четвертой эпохе достигается почти 100%.

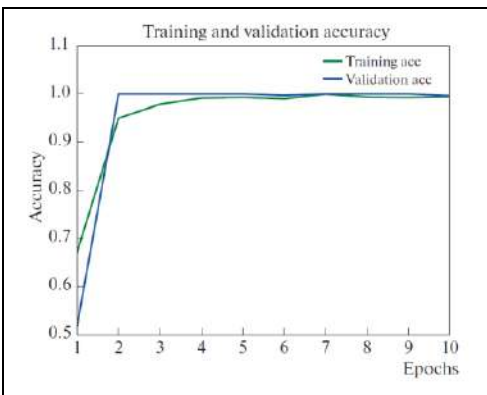


Рис. 7. Точность обучения и валидации модели при использовании вейвлет-датасета
Fig. 7. Capacity of the model in the accuracy of training and validation, with the wavelet dataset

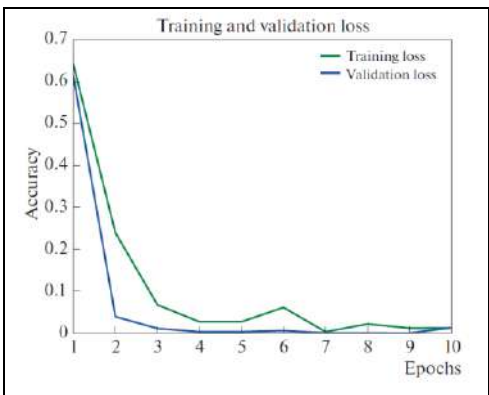


Рис. 8. Потери обучения и валидации модели при использовании вейвлет-датасета
Fig 8. Capacity of the model in the loss of training and validation, with the wavelet dataset

Показатель точности работает лучше на этапе валидации (синяя линия), поскольку обобщение обучения выше, чем обучающие данные, что позволяет избежать переобучения для новых данных и достичь почти 100% на второй эпохе. Более того, рис. 8 показывает, что модель достигает нулевых потерь на этапе обучения (зеленая линия), очень медленно приближаясь к нулю на седьмой эпохе; этот эффект позволяет нам иметь модель, которая не поддается переобучению для новых данных. На стадии валидации значение потерь, близкое к нулю, получается очень быстро (синяя линия); это связано с обучаемостью и качеством изображений в вейвлет-домене.

Таким образом, мы имеем модель с более высокой производительностью обнаружения и классификации по сравнению с первой моделью обучения, которая использует исходный набор данных. Обобщение знаний позволяет нам адекватно обучать новой информации; в результате мы получаем меньшую разницу между потерями на этапе обучения и на этапе валидации. Наша работа позволяет оптимизировать обучение для обнаружения объектов в аэронавигационных приложениях. Кроме того, некоторые преимущества преобразования данных в вейвлет-домен заключаются в возможности обучения физических характеристикам объектов, ориентированным на детали текстуры, и устранении переобучения при наличии небольшого обучающего набора данных. В табл. 1 приведены показатели результативности модели при использовании вейвлет-датасета для валидации. Показатели точности и потерь валидируют модель обучения, полученную слиянием методов CNN и вейвлет-анализа.

Табл.1. Результаты валидации модели обучения, созданной с применением вейвлет-анализа и глубокого обучения.

Показатель	Значение [%]
Точность тестирования	100
Потери тестирования	0.61

4.3 Эксперименты с симуляцией

В нашей работе использовался симулятор Gazebo, который применяется для проектирования UAF типа Parrot AR.Drone 2.0 и разработки реалистичных 3D-сценариев для симуляции [34]. Симулятор, в частности, обеспечивает быстрое выполнение алгоритмов, предоставляет пользовательский интерфейс и гибко контролирует навигацию дрона. Кроме того, Gazebo обеспечивает такое же управление, как и в реальных мобильных аппаратах.

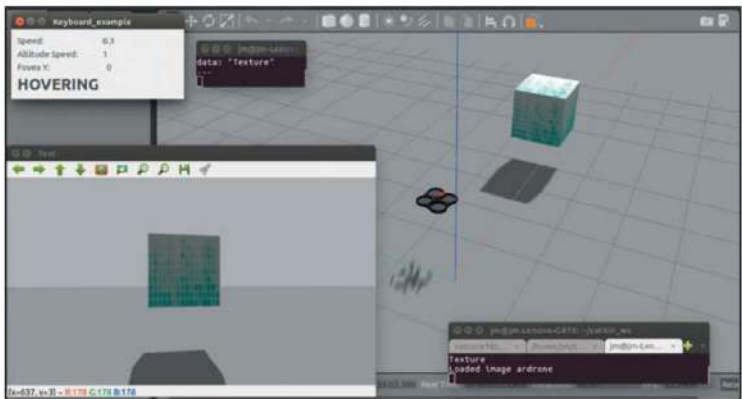


Рис. 9. Тест 1: Обнаружение объектов на сцене предложенным методом с применением подхода вейвлет-анализа и глубокого обучения

Fig 9. Test 1: Object detection in scene applying the proposed method with the approach of wavelet analysis and deep learning

При разработке приложений симулятор позволяет подключать операционную систему ROS с открытым исходным кодом, созданную под лицензией Berkeley Software Distribution (BSD).

Обеспечивается функционирование операционной системы в гетерогенном компьютерном кластере [35]. Узлы обмениваются сообщениями, что позволяет программировать их на любом языке, имеющем клиентские библиотеки для ROS (например, C, C ++, Python, Java, Matlab) [36].

Что касается второй модели обнаружения, симулятор обеспечивает поддержку оценки и выполнения в реальном времени. Важно подчеркнуть, что изображения на этапе оценки – это изображения с камеры на борту AR.Drone (вид спереди), как если бы это была настоящая камера. Изображения в системе изначально имеют размер 640×380 пикселей, поэтому они преобразуются к размеру 512×512 пикселей. Потом эти изображения преобразуются в вейвлет-домен посредством кратномасштабного анализа на уровне масштабирования, равном трем. В результате мы имеем четыре подизображения с разрешением 64×64 пикселя (значение, допустимое для модели обнаружения). При исполнении в реальном времени мы показываем разные ракурсы с бортовой камеры. Показаны два класса предсказания модели обнаружения: первый класс, в котором объект с текстурой полностью появляется на сцене, как показано на рис. 9, и второй класс, в котором он не появляется в плоскости изображения, см. рис. 10.

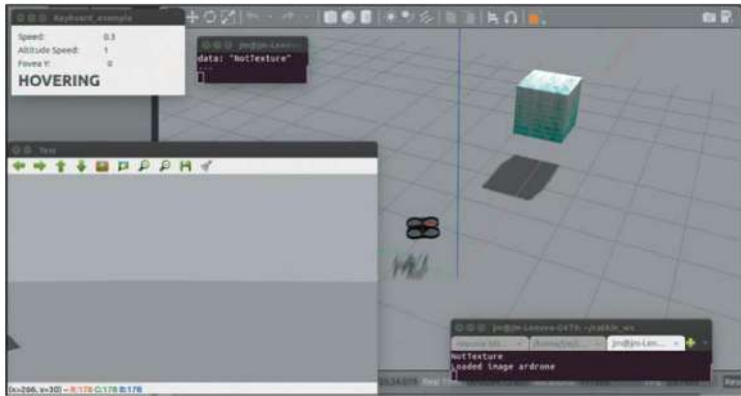


Рис. 10. Тест 2: Объект не обнаруживается на сцене, так как он находится вне поля зрения бортовой камеры дрона. Видео по работе доступно на <https://youtu.be/MOSrJyf14T8>
Fig 10. Test 2: Object is not detected on the scene, as it is out of view from the on-board camera of the drone. A video of this work for review purposes is available at <https://youtu.be/MOSrJyf14T8>

Средняя частота обнаружения достигает 98% в различных ракурсах (вариации масштаба сцены и освещения), как показано на рис. 11. Модель обнаружения продемонстрировала отличные значения времени прогнозирования. Это время значительно меньше того, через которое становятся видимыми исходные изображения, как показано в табл. 2. Время обнаружения зависит от сети, глубины, размера изображения, количества нейронов в последнем слое и возможностей используемой аппаратуры.

Табл. 2. Экспериментальные результаты выполнения нашего приложения в операционной системе ROS и симуляторе Gazebo

Table 2. Experimental results of running our application on the ROS operating system and Gazebo simulator

Время	Значение [с]
Функция предсказания	0.01608
Видимое исходное изображение	0.02587

Мы показали, что взаимодействие нейронной сети ConvNet с вейвлет-набором данных в роботизированных приложениях может дать многообещающие результаты на этапе обучения. Среда моделирования обеспечивает полный анализ предложенного метода, следовательно, он может быть адаптирован к реальным условиям для классификации и

обнаружения объектов с текстурами. Однако следует заметить, что при этом целесообразно использовать компьютер с высокой вычислительной мощностью, поскольку для моделирования требуется много ресурсов.

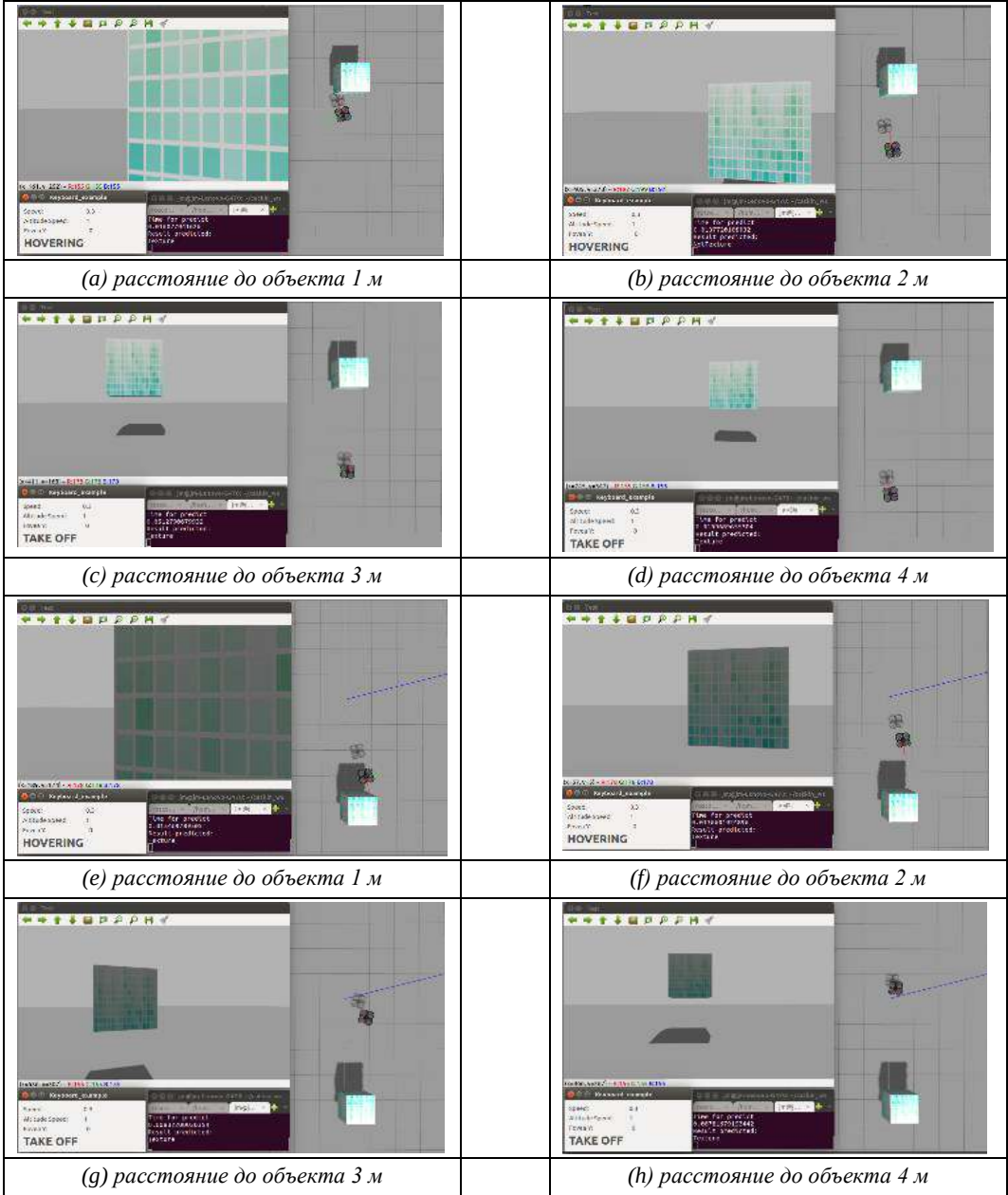


Рис. 11. Результаты обнаружения на четырех различных расстояниях от цели до камеры дрона и при разном освещении; (а)-(d) основаны на снимках в одном ракурсе, а (е)-(h) – на снимках в противоположном ракурсе (изображение цели содержит тень)

Fig 11. Detection results at four different distances from the target to the drone camera, and different scales and illumination (a)-(d) are results based on a perspective in the environment, while (e)-(f) are results based on an opposite perspective in the environment (the image of the target contains shadow)

5. Заключение

Мы представили новый метод обнаружения объектов с текстурами. Наше предложение основано на использовании методов вейвлет-анализа и глубоких нейронных сетей, которые были адаптированы к области аэронавигации. Для этого мы кадр за кадром классифицировали изображения, полученные с бортовой камеры дрона. В сочетании с системой ROS и моделью обучения полученная информация преобразуется в вейвлет-домен, в котором это трехуровневое преобразование проявляет важные характеристики аппроксимации и детализации. В получаемом аппроксимационном датасете сохраняется большинство свойств исходного изображения. Результаты оценки показывают, что наличие изображений в вейвлет-домене и с более низким разрешением значительно повышает эффективность обнаружения по сравнению с использованием изображений в пространственном домене в качестве входных данных архитектуры ConvNet. Кроме того, вейвлет-датасет был полезен для смягчения эффекта переобучения при обобщении обучения на этапе обучения.

В заключение отметим, что приложение показывает высокую точность в случае, когда в результате прогноза на сцене обнаруживается объект с текстурой (Texture) или не объект с текстурой (NotTexture). Существенно и то, что система предсказывает почти вдвое быстрее, чем передаются кадры камеры. Мы считаем, что полученные результаты перспективны, поэтому было бы целесообразно экспериментировать с новыми семействами вейвлетов для расширения возможностей обнаружения объекта, улучшения характеристики текстур, особенно использования нашего подхода в более сложной автономной навигационной системе.

Список литературы / References

- [1]. L.O. Rojas-Pérez. Autonomous Navigation System for Micro Aerial Vehicles. Dissertation, Instituto Tecnológico Superior de Atlixco, Puebla, México, 2018.
- [2]. J. Martínez-Carranza, L. Valentin, F. Márquez-Aquino et al. Obstacle Detection during Autonomous Flight of Drones Using Monocular SLAM. *Research in Computing Science*, vol. 114, 2016, pp. 111-124.
- [3]. Н.С. Васильева. Методы поиска изображений по содержанию. Программирование, том 35, no. 3, 2009 г., стр. 51-80 / N.S. Vassilieva. Content-based image retrieval methods. *Programming and Computer Software*, vol. 35, no. 3, 2009, pp. 158-180.
- [4]. А.Д. Жданов, Д.Д. Жданов, Н.Н. Богданов и др. Проблемы дискомфорта зрительного восприятия в системах виртуальной и смешанной реальностей. Программирование, том 45, no. 4, 2019 г., стр. 9-18 / A.D. Zhdanov, D.D. Zhdanov, N.N. Bogdanov et al. Discomfort of Visual Perception in Virtual and Mixed Reality Systems. *Programming and Computer Software*, vol. 45, no. 4, 2019, pp. 147-155.
- [5]. В.В. Санжаров, В.А. Фролов. Уровень детализации для предрасчитанных процедурных текстур. Программирование, том 45, no. 4, 2019 г., стр. 54-63 / V.V. Sanzharov, V.A. Frolov. Level of Detail for Precomputed Procedural Textures. *Programming and Computer Software*, vol. 45, no. 4, 2019, pp.187-195.
- [6]. C. Vargas-Olmos. Procesamiento de imágenes con métodos de ondeleta. Dissertation, Facultad de Ciencias, UASLP, San Luis Potosí, México, 2010 (in Spanish).
- [7]. J. D. Cárdenas-Amaya. Extracción y análisis de características con la Transformada Wavelet para el reconocimiento de imágenes. Dissertation, Instituto de Investigación en Comunicación Óptica, UASLP, San Luis Potosí, México, 2018 (in Spanish).
- [8]. Y. Bazi and F. Melgani, Convolutional SVM Networks for Object Detection in UAV Imagery. *IEEE Transactions on Geoscience and Remote Sensing*, vol. 56, no. 6, 2018, pp. 3107-3118.
- [9]. Y. Pi, N.D. Nath, A.H. Behzadan, Convolutional neural networks for object detection in aerial imagery for disaster response and recovery. *Advanced Engineering Informatics*, vol. 43, 2020, Article 101009.
- [10]. S. Dionisio-Ortega, L.O. Rojas-Perez, J. Martinez-Carranza, I. Cruz-Vega. A deep learning approach towards autonomous flight in forest environments. In *Proc. of the International Conference on Electronics, Communications and Computers (CONIELECOMP)*, 2018, pp. 139-144.

- [11]. A. Krizhevsky, I. Sutskever, and G.E. Hinton. Imagenet classification with deep convolutional neural networks. In Proc. of the 26th Annual Conference on Neural Information Processing Systems (NIPS), 2012, pp. 1097-1105.
- [12]. J. Redmon, S. Divvala, R. Girshick, A. Farhadi. You Only Look Once: Unified, Real-Time Object Detection. In Proc. of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR), 2016, pp. 779-788.
- [13]. S. Jung, S. Hwang, H. Shin, D. H. Shim. Perception, Guidance, and Navigation for Indoor Autonomous Drone Racing Using Deep Learning. IEEE Robotics and Automation Letters, vol..3, issue 3, 2018, pp. 2539-2544.
- [14]. T. Williams, R. Li. An Ensemble of Convolutional Neural Networks Using Wavelets for Image Classification. Journal of Software Engineering and Applications, vol. 11, no 2, 2018, pp. 69-88.
- [15]. T. Williams, R. Li. Wavelet Pooling for Convolutional Neural Networks. International Conference on Learning Representations, 2018, 12 p.
- [16]. J. Piao, Y. Chen, H. Shin. A New Deep Learning Based Multi-Spectral Image Fusion Method. entropy, vol. 21, issue 6, 2019, 16 p.
- [17]. D. N. De Silva, S. Fernando, I.T.S. Piyatilake, A.V.S. Karunaratne. Wavelet based edge feature enhancement for convolutional neural networks. In Proc. of the Eleventh International Conference on Machine Vision (ICMV 2018), 2018.
- [18]. S. Fujieda, K. Takayama, and T. Hachisuka. Wavelet convolutional neural networks for texture classification. arXiv preprint arXiv:1707.07394, 2017.
- [19]. C.S. Burrus, R.A. Gopinath and H. Guo. Introduction to Wavelets and Wavelet Transforms: A Primer. Pearson, 1998, 288 p.
- [20]. S. Mallat. A Wavelet Tour of Signal Processing: The Sparse Way. 3rd ed. Academic Press, 2009, 832 p.
- [21]. G. Strang and T. Nguyen. Wavelets and Filter Banks. 2nd revised ed. Wellesley-Cambridge Press, 1996 500 p.
- [22]. C. Vargas-Olmos. Procesamiento de señales y solución de problemas con la transformada wavelet. Ph.D. dissertation. Instituto de Investigación en Comunicación Óptica, UASLP, San Luis Potosí, México, 2017 (in Spanish).
- [23]. J.S. Walker. A Primer on Wavelets and Their Scientific Applications. 2nd ed. Chapman & Hall/CRC, 2008, 320 p.
- [24]. I. Goodfellow, Y. Bengio, and A. Courville. Deep Learning. MIT Press, 2016, 800 p.
- [25]. И.В. Степанян. Методология и инструментальные средства проектирования бинарных нейронных сетей. Программирование, том 46, no. 1, 2020 г., стр. 54-62 / I.V. Stepanyan, Methodology and Tools for Designing Binary Neural Networks. Programming and Computer Software, vol. 46, no. 1, 2020, pp. 49-56.
- [26]. Ю.Г. Сметанин, Л.Е. Карпов, Ю.Л. Карпов. Адаптация общих концепций тестирования программного обеспечения к нейронным сетям. Программирование, том 44, no. 5, 2020 г., стр. 43-56 / Y.L. Karpov, L.E. Karpov, and Y.G. Smetanin. Adaptation of General Concepts of Software Testing to Neural Networks. Programming and Computer Software, vol. 44, no. 5, 2018, pp. 324-334.
- [27]. J. Torres. Convolutional Neural Networks for Beginners. Practical Guide with Python and Keras, 2018. URL: <https://towardsdatascience.com/convolutional-neural-networks-for-beginners-practical-guide-with-python-and-keras-dc688ea90dca>, accessed 27 April 2019.
- [28]. K. Simonyan and A. Zisserman. Very deep convolutional networks for large-scale image recognition. arXiv preprint arXiv:1409.1556, 2014.
- [29]. A. Yayık, K. Yakup, and G. Altan. Deep Learning with ConvNET Predicts Imagery Tasks Through EEG. arXiv preprint arXiv:1907.05674, 2019.
- [30]. W. Liu, D. Anguelov, D. Erhan et al. SSD: Single shot multibox detector. Lecture Notes in Computer Science, vol. 9905, 2016, pp. 21-37.
- [31]. H. Kaiming, G. Gkioxari, P. Dollár, and R. Girshick. Mask R-CNN. In Proc. of the IEEE international conference on computer vision, 2017, pp. 2961-2969.
- [32]. R. Shaoqing, K. He, R. Girshick, and J. Sun. Faster R-CNN: Towards real-time object detection with region proposal networks. In Proc. of the 28th Annual Conference on Neural Information Processing Systems (NIPS), 2015, pp. 91-99.
- [33]. F. Chollet. Deep Learning with Python. Manning Publications, 2018, 384 p.

- [34]. N. Koenig, A. Howard. Design and use paradigms for Gazebo, an open-source multi-robot simulator, IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS), vol. 3, 2004, pp. 2149-2154.
- [35]. E. Rodríguez-Martín. Sistema de posicionamiento para un drone. Dissertation, Universidad de La Laguna, España, 2015 (in Spanish).
- [36]. L. Joseph. Robot Operating System for Absolute Beginners: Robotics Programming Made Easy. Apress, 2018, 295 p

Информация об авторах / Information about authors

Хуан Мануэль ФОРТУНА-СЕРВАНТЕС – аспирант. Область научных интересов: компьютерное зрение, цифровая обработка изображений, глубокое обучение, беспилотные летательные аппараты, робототехника.

Juan Manuel FORTUNA-CERVANTES, PhD Student. Research interests include Computer Vision, Digital Image Processing, Deep Learnig, Unmanned Aerial Vehicles, Robotics.

Марко Тулио РАМИРЕС-ТОРРЕС – кандидат прикладных наук, профессор. Область научных интересов: криптография, вейвлет-анализ, виртуальные приборы.

Marco Tulio RAMÍREZ-TORRES, Ph.D. in Applied Sciences, Full time Professor. Research interests include Cryptography, Wavelet Analysis, Virtual Instrumentation.

Хуан МАРТИНЕС-КАРРАНСА – кандидат прикладных наук, профессор. Область научных интересов: обработка сигналов, вейвлет-анализ, системы шифрования и прикладная математика.

Juan MARTINEZ-CARRANZA, Ph.D. in Applied Sciences, Full Professor. Research interests include signal processing, wavelet and scaling analysis, encryption systems and applied mathematics.

Хосе Саломе МУРГУИЯ-ИБАРРА, кандидат наук, профессор-исследователь кафедры электроники факультета естественных наук. Область научных интересов: обработка сигналов, нелинейные динамические системы.

José Salomé MURGUÍA-IBARRA, Ph.D., Professor Researcher of the Electronic Department of the Sciences Faculty. Research interests include Signal Processing, Nonlinear Dynamical Systems.

Марсела МЕХИЯ-КАРЛОС, кандидат наук, профессор-исследователь. Область научных интересов: коммуникации и обработка сигналов.

Marcela MEJÍA-CARLOS, Ph.D., Professor Researcher. Research interests: Communication and Signal Processing.

DOI: 10.15514/ISPRAS-2020-33(2)-10



Решение проблемы обеспечения качества дерева многоадресной рассылки услуг

*К. Риссо, ORCID: 0000-0003-0580-3083 <crisso@fing.edu.uy>
Ф. Робледо, ORCID: 0000-0003-4235-4221 <frobledo@fing.edu.uy>
С. Несмачнов, ORCID: 0000-0002-8146-4012 <sergion@fing.edu.uy>*

*Республиканский университет,
Уругвай, 11200, Монтевидео, ул. 18 июля 1824-1850*

Аннотация. В данной статье представлена основанная на потоках формулировка проблемы обеспечения качества дерева многоадресной рассылки услуг в терминах смешанного целочисленного программирования. Это актуальная проблема, связанная с современными телекоммуникационными сетями, обеспечивающие распространение мультимедийного контента через облачные Internet-системы. Насколько нам известно, для проблемы обеспечения качества дерева многоадресной рассылки услуг формулировка в терминах смешанного целочисленного программирования ранее не предлагалась. Экспериментальная оценка выполняется на наборе реалистичных примеров из SteinLib, чтобы показать применимость стандартных точных решателей для нахождения решений реальных задач. Точный метод применяется для бенчмаркинга предлагаемых формулировок, а также для поиска оптимальных или близких к оптимальным решений за приемлемое время исполнения.

Ключевые слова: многоадресная рассылка; качество; дерево многоадресной рассылки услуг; целочисленное программирование

Для цитирования: Риссо К., Робледо Ф., Несмачнов С. Решение проблемы обеспечения качества дерева многоадресной рассылки услуг. Труды ИСП РАН, том 33, вып. 2, 2021 г., стр. 163-172. DOI: 10.15514/ISPRAS-2021-33(2)-10

Solving the Quality of Service Multicast Tree Problem

*C. Risso, ORCID: 0000-0003-0580-3083 <crisso@fing.edu.uy>
F. Robledo, ORCID: 0000-0003-4235-4221 <frobledo@fing.edu.uy>
S. Nesmachnow, ORCID: 0000-0002-8146-4012 <sergion@fing.edu.uy>*

*Universidad de la Republica Uruguay,
Av. 18 de Julio 1824-1850, Montevideo, 11200, Uruguay*

Abstract. This article presents a flow-based mixed integer programming formulation for the Quality of Service Multicast Tree problem. This is a relevant problem related to nowadays telecommunication networks to distribute multimedia over cloud-based Internet systems. To the best of our knowledge, no previous mixed integer programming formulation was proposed for Quality of Service Multicast Tree Problem. Experimental evaluation is performed over a set of realistic problem instances from SteinLib, to prove that standard exact solvers can find solutions to real-world size instances. Exact method is applied for benchmarking the proposed formulations, finding optimal solutions and low feasible-to-optimal gaps in reasonable execution times.

Keywords: multicasting; Quality; Service Multicast Tree; Integer Programming

For citation: Risso C., Robledo F., Nesmachnow S. Solving the Quality of Service Multicast Tree Problem. Trudy ISP RAN/Proc. ISP RAS, vol. 33, issue 2, 2021, pp. 163-172 (in Russian). DOI: 10.15514/ISPRAS-2021-33(2)-10.

1. Введение

В этой статье исследуется актуальная проблема проектирования сети: проблема обеспечения качества дерева многоадресной рассылки услуг (Quality of Service Multicast Tree Problem, QoSMTTP). QoSMTTP относится к построению оптимального дерева для многоадресной передачи и доставки мультимедийного контента с использованием минимального ожидаемого количества ресурсов полосы пропускания, например, программно-определяемых сетей (Software Defined Network, SDN).

Основным вкладом исследования, изложенного в данной статье, является новая формулировка QoSMTTP в терминах смешанного целочисленного программирования (Mixed-Integer Programming, MIP), основанная на потоковой формулировке Гёманса (Michel Goemans) и Мюна (Young-Soo Myung) [1] проблемы дерева Штейнера (Steiner Tree Problem – STP). При экспериментальной оценке применяется точный метод для решения экземпляров из широко известной библиотеки SteinLib, модифицированных соответствующим образом, с целью доказать, что стандартные точные решатели могут находить решения реальных задач. Результаты экспериментов указывают на то, что оптимальные и близкие к оптимальным решения вычисляются за приемлемое время. Наиболее эффективные существующие решения для задачи QoSMTTP вычисляются с использованием приближенных алгоритмов, коэффициенты которых выше $3/2$ [2]; таким образом, предлагаемый подход и результаты являются непосредственным вкладом в данное направление исследований.

Статья имеет следующую структуру. В разд. 2 приводится описание QoSMTTP. Описание предлагаемой модели для задачи QoSMTTP представлено в разд. 3. Результаты экспериментальной оценки на основе стандартных экземпляров из репозитория SteinLib изложены в разд. 4. В заключение, в разд. 5 представлены выводы и дальнейшие перспективы исследования.

2. Проблема обеспечения качества дерева многоадресной рассылки услуг

QoSMTTP формулируется следующим образом. Пусть $G = (V, E, l, r)$ – неориентированный взвешенный граф, где V – множество узлов, а E – множество ребер. Также рассматривается подмножество $T \subseteq V$ выделенных терминальных вершин, одна из которых является корнем ($R \in T$). Цель состоит в том, чтобы найти подграф остовного дерева с минимальной стоимостью, который должен содержать все терминальные узлы. Основное отличие от STP в данном случае заключается в функции стоимости, так как оцениваются и узлы, и ребра. Каждый терминальный узел имеет вес от 0 до r_N , определяющий значимость данного терминального узла в сети. Функция $l: E \rightarrow R^+$ представляет длину каждого ребра, а функция $r: V \rightarrow R_0^+$ выдает вес каждого ребра. В этой проблеме корневой узел R представляет собой центральный узел проектируемой сети (например, узел, в котором находится весь доступный контент, который планируется совместно использовать / рассылать по многим адресам).

QoSMTTP предполагает нахождение поддерева с минимальной стоимостью $F = (V_F, E_F)$ графа G , растущего от заданного корня R и покрывающее множество терминальных узлов. Узлы, не принадлежащие T , как и STP, называются узлами Штейнера; они представляют собой узлы с нулевым весом, для которых не требуется соединение с корневым узлом, но которые можно использовать для уменьшения стоимости решения.

Что касается функции стоимости, то основная идея заключается в том, что чем выше стоимость, тем больше число различных копий информационных единиц, которые должны быть переданы в данный узел от корня. Таким образом, стоимость непосредственного соединения на пути к корню многоадресного дерева должна быть умножена на некоторое число, в котором учитываются эти веса (r_i).

Известным свойством деревьев является то, если имеются два дерева $T_1 = (V_1, E_1)$, $T_2 = (V_2, E_2)$ и две произвольные вершины $n_1 \in V_1$, $n_2 \in V_2$, то результатом соединения этих графов ребром (n_1, n_2) также является дерево. Обратное также верно, так как в результате удаления любого ребра из любого нетривиального дерева (т.е. дерева, содержащего не менее одного ребра), образуются два дерева, которые, тем самым, являются двумя независимыми и связанными подкомпонентами.

Стоимость решения F частного случая (экземпляра) G QoSMTIP определяется путем сложения стоимости $c(e)$ каждого ребра e в дереве F , где $c(e) = l(e) \times \bar{r}_e$. Выражения $l(e)$ являясь частью самого примера, в то время как значения $\bar{r}_e$ зависят от конкретного решения.

Вычисления $\bar{r}_e$ производятся следующим образом: для ребра $e = (v_1, v_2) \in F$ без потери общности предположим, что v_2 находится в пути от v_1 к корню R или v_2 является самим корнем. После удаления e из F получаются два поддерева F_1 и F_2 , при этом v_1 находится в F_1 , а v_2 и корень (если они разные) – в F_2 . Тогда значение $\bar{r}_e$ является максимальным среди всех r_i , $i \in F_1$.

Компоненты целевой функции QoSMTIP зависят от самого решения, что значительно усложняет эту проблему по сравнению с STP. Насколько нам известно, до сих пор QoSMTIP не формулировалась в терминах MIP, а существующие решения основаны на аппроксимационных алгоритмах, коэффициенты которых выше 3/2 [3]. Формально, как STP, так и QoSMTIP являются NP-трудными задачами. STP является классической NP-трудной задачей [4], сводимой к QoSMTIP за полиномиальное время (STP является частным случаем QoSMTIP, когда веса всех узлов равны 1). Таким образом, в вычислительном отношении задача QoSMTIP является не менее трудной, чем STP.

3. Предлагаемая MIP-модуль для QoSMTIP

Предлагаемая модель основана на потоковой формулировке Гёманса и Мюна [1] для решения STP. Основная идея предлагаемой модели заключается в следующем. Для удовлетворения требования балансировки потока для каждого узла в графе $G = (V, E)$, за исключением одного (т.е. корня R), при вводе единицы трафика в каждую терминальную вершину T , за исключением корня, должен существовать путь от каждого терминального узла k к R для слива этой единицы потока; таким образом, результат должен быть связным. Предположим, что найдено оптимальное решение G_T . Реализуемость при соблюдении ограничения потока гарантируют наличие связности. Кроме того, поскольку ребра имеют положительные веса, а множество ребер, используемое для маршрутизации этих потоков, должно иметь минимальную стоимость, решение будет минимально связным; следовательно, получаемый граф является деревом.

Известным фактом теории потоков в сетях является то, что если параметры являются целыми числами, то крайние точки многогранника, определяемого уравнениями баланса и ограничений потока, также являются целыми числами [5]. Этот результат позволяет представлять целочисленные потоковые проблемы как задачи линейного программирования, что ограничивает оптимум замкнутой выпуклой оболочкой линейной целевой функции. Существование алгоритмов с полиномиальным временем для решения задач линейного программирования [6] уменьшает сложность чисто целочисленных потоковых проблем до сложности класса P.

Для правильного учета целевой функции STP необходимо также принимать во внимание некоторые целочисленные переменные и ограничения. Наличие этих переменных подрывает предыдущие рассуждения о сложности, перемещая проблему в категорию NP-трудных проблем. Однако целостность потоковых переменных можно, по меньшей мере, ослабить. Это не меняет теоретическую сложность проблемы, но облегчает ее решение за счет уменьшения количества целочисленных переменных и полиномиальной сложности

проблемы после присвоения целочисленным переменным возможных значений (как это делается, например, в методах ветвей и границ, используемых точными решателями).

Любой экземпляр STP определяется неориентированным взвешенным графом $G = (V, E, c)$, где V – множество вершин, E – множество потенциальных ребер (т.е. тех, которые допустимы в решении), а функция $c: E \rightarrow R^+$ определяет стоимость каждого ребра. Кроме того, для обеспечения замкнутости экземпляра проблемы требуется подмножество $T \subseteq V$ терминальных вершин. Без потери общности в модели предполагается, что одной из терминальных вершин $R \in T$ должен быть корень создаваемого дерева. Этот выбор является произвольным, но он упрощает формулировку STP и позволяет легко построить формулировку QoSMTSP, так как для полного определения экземпляра второй проблемы требуется наличие фиксированного корня.

Для $G = (V, E, c)$ возьмем ориентированный граф $G' = (V, E', c')$, который получается в результате дублирования каждого ребра E в G с использованием обоих направлений, за исключением только тех ребер, вторая узлом которых является R . Это означает, что если $\{ij\} \in E$ и $i, j \neq R$, то $(ij) \in E'$ и $(ji) \in E'$. Кроме того, если $\{iR\} \in E$, то только (iR) находится в E' . Перечисленные ребра – это единственные ребра G' . Наконец, $c': E' \rightarrow R^+$ определяется как $c'(\{ij\}) = c(\{ij\})$. Соотношения (3.1)–(3.3) задают формулировку потоковой подпроблемы:

$$\begin{cases} \sum_{(ij) \in E'} x_{ij} - \sum_{(ki) \in E'} x_{ki} = 1 & \forall i \in T \setminus \{R\}, \\ \sum_{(ij) \in E'} x_{ij} - \sum_{(ki) \in E'} x_{ki} = 0 & \forall i \in S, \\ 0 \leq x_{ij} & \forall (ij) \in E'. \end{cases} \quad (3.1)$$

$$\sum_{(ij) \in E'} x_{ij} - \sum_{(ki) \in E'} x_{ki} = 0 \quad \forall i \in S, \quad (3.2)$$

$$0 \leq x_{ij} \quad \forall (ij) \in E'. \quad (3.3)$$

В соотношениях (3.1)–(3.3) переменные $x_{ij} (ij \in E')$ учитывают поток от i по направлению к j . Таким образом, соотношение (3.1) просто выражают тот факт, что баланс сетевого потока равен 1 для каждого терминального узла i , кроме корня, что эквивалентно вводу в эти вершины единицы трафика. Множество S содержит все узлы Штейнера: $S = V \setminus T$. Соотношение (3.2) устанавливают баланс потока, равным 0, для узлов Штейнера, так как разница между суммарными исходящими и входящими потоками должна быть равна 0. Наконец, соотношение (3.3) устанавливает нижние пределы значений потока. Кроме того, потоки не могут быть больше или равны $|T|$, поскольку суммарный поток, введенный в граф, совпадает с $|T| - 1$ (количество терминальных вершин, за исключением корня), а значит, это является объемом суммарного потока, протекающего через R . Это ограничение задается отдельно в соотношениях (3.4) и (3.5):

$$\begin{cases} (|T| - 1) y_{ij} \geq x_{ij} & \forall (ij) \in E' \\ y_{ij} \in \{0, 1\} & \forall \{ij\} \in E \end{cases} \quad (3.4)$$

$$y_{ij} \in \{0, 1\} \quad \forall \{ij\} \in E \quad (3.5)$$

В соотношениях (3.4) и (3.5) вводятся булевские переменные y_{ij} , чтобы показать, входит ли в решение ребро $(ij) \in E'$. Соотношения (3.1)–(3.3) и (3.4) задают классический набор ограничений для потоковой проблемы, область допустимых решений которой – это выпуклая оболочка, что упрощает нахождение значений x_{ij} . Число переменных x_{ij} совпадает с числом переменных y_{ij} , поэтому объем вычислений уменьшается. Полная MIP-формулировка для STP объединяет уравнения (3.1)–(3.5) с целевой функцией $\min \sum_{(ij) \in E} c_{ij} y_{ij}$, которая обеспечивает суммарную стоимость решения в соответствии со значениями y_{ij} .

Значения y_{ij} должны обеспечить достижимость решения потоковой подпроблемы (3.1)–(3.4). И, наоборот, для любого потока с $x_{ij} > 0$ должно выполняться равенство $y_{ij} = 1$, чтобы удовлетворить ограничения (3.4)–(3.5). Поскольку оптимизация будет искать минимальное значение $\sum_{(ij) \in E} c_{ij} y_{ij}$, она не будет выбирать какое-либо излишнее значение $y_{ij} = 1$, поэтому

оптимум минимален, а результат должен быть древовидным графом. Следовательно, это и есть оптимальное решение STP, как и задумывалось.

Экземпляр QoSMTSP определяется неориентированным взвешенным графом $G = (V, E, l, r)$, множеством терминальных вершин $T \subseteq V$ и корневым узлом R , который для этой проблемы не является произвольным. Функция $l: E \rightarrow R^+$ определяет длину ребер, а функция $r: V \rightarrow R_0^+$ выдает значение весов узлов. Узлы Штейнера имеют нулевой вес, т.е. $r(s) = 0$ для всех $s \in S$. Поскольку $r(R)$ в QoSMTSP является бессмысленным, в модели предполагается, что $r(R) = 0$. Наконец, $r(t) > 0$ для каждого терминального узла за исключением корня ($t \in T \setminus \{R\}$). Если $r(t) = 1$ для каждого $t \in T \setminus \{R\}$, то экземпляр QoSMTSP эквивалентен экземпляру STP с тем же корнем, так как максимальный вес r_t в любом поддереве, содержащем некоторые терминальные узлы, кроме корня, равен 1. В отличие от STP, узлы Штейнера в экземпляре QoSMTSP могут входить в оптимальное решение, так как они имеют нулевой вес, и поэтому длины исходящих из них ребер также равны нулю.

Аналогичная процедура применяется к формулировке QoSMTSP. Создается ориентированная версия G' для экземпляра G QoSMTSP. Поточковые переменные x_{ij} и соответствующие им y_{ij} такие же, как в STP, то же относится и к соотношениям (3.1)–(3.5). Такая формулировка исключает возможность потоковых циклов, т.е. позволяет иметь либо $y_{ij} = 1$, либо $y_{ji} = 1$, но не одновременно. Этот факт важен для согласования блоков соотношений, используемых для определения весов поддеревьев.

Чтобы обеспечить веса $\bar{r}_e$ (максимум среди весов узлов поддеревьев), используются дополнительные переменные z_u , $u \in V \setminus \{R\}$. Для заданного экземпляра QoSMTSP введем константу $C = \max\{r_v, v \in T \setminus \{R\}\}$ и рассмотрим соотношения (3.7)–(3.8). Основная идея введения дополнительных переменных заключается в том, что если значения переменных y_{uv} определяют поддерево в G , то соотношения (3.6)–(3.7) заставляют переменные z_v быть большими или равными $\bar{r}_e$, где e является звеном в единственном пути, который соединяет v с корнем внутри этого дерева.

$$\begin{cases} z_v \geq r_v & \forall v \in V \setminus \{R\}, \end{cases} \quad (3.6)$$

$$\begin{cases} z_v \geq z_u + C(y_{uv} - 1) & \forall v \in V \setminus \{R\}, (uv) \in E'. \end{cases} \quad (3.7)$$

Таким образом, если переменные y_{uv} определяют дерево, ориентированное к корню, и процесс оптимизации ищет минимальное значение функции, определенное в уравнении (3.9), то процесс оптимизации будет опускать значения как можно ниже значения z_v , достигая при этом значений $\bar{r}_e$ для этого дерева, определенного переменными y .

$$\sum_{(uv) \in E'} l(uv) \times z_u \times y_{uv}. \quad (3.8)$$

Соотношение (3.6) определяют нижние границы для переменной z_v . В тех случаях, когда $y_{uv} = 0$, соотношение (3.7) принимает вид $z_v \geq z_u - C$ и деактивируется, так как C является максимумом среди разрядов, а $z_v \geq 0$ является не менее ограничивающим. И наоборот, для тех ребер, где $y_{uv} = 1$, соотношение (3.7) преобразуются в $z_v \geq z_u$, указывая на то, что значение z_v должно быть не больше максимума, определяемого активными соотношениями. Соотношение (3.7) действует в поддеревьях, содержащих v , в число которых не входит корень. Последние замечания верны при отсутствии ребер в E' , выходящих из R , что изолирует значения между поддеревьями, и в направлении путей дерева от терминальных узлов к корню, что добавляет ограничения нижних границ в пути от каждого узла к R . Результатом оптимизации должно быть дерево, так как при любом возможном решении добавление ребра увеличило бы число активных нижних границ для z_v , тем самым увеличивая целевое значение, за исключением узлов-заглушек Штейнера.

Целевая функция в (3.8) квадратичная, что значительно затрудняет решение задачи. Для устранения этого недостатка в предлагаемой модели вводится вспомогательное множество

переменных $\eta_{ij}, (ij) \in E'$ и переформулируется целевая функция, определяемая теперь соотношениями (3.9)–(3.11):

$$\begin{cases} \min \sum_{(ij) \in E'} l(ij) \times \eta_{ij} & (3.9) \\ \eta_{ij} \geq z_i + C(y_{ij} - 1) & \forall (ij) \in E' & (3.10) \\ \eta_{ij} \geq 0 & \forall (ij) \in E' & (3.11) \end{cases}$$

После объединения уравнений (3.1)–(3.7), (3.9)–(3.11) переменные η_{ij} захватывают значения $z_i \times y_{ij}$ в оптимальных решениях, так что новая целевая функция соответствует (3.9) для целей оптимизации. Соотношение (3.10) деактивируется всякий раз, когда $y_{ij} = 0$, оставляя соотношение (3.11) единственной нижней границей для η_{ij} . Оптимизация приводит к $\eta_{ij} = 0$, таким образом, $\eta_{ij} = z_i \times y_{ij}$. Если же $y_{ij} = 1$, то соотношение (3.10) преобразуется в $\eta_{ij} \geq z_i$, а предыдущее ограничение должно быть активным при оптимуме.

4. Экспериментальный анализ

В данном разделе представлена экспериментальная оценка предложенной модели для QoSMTp на конкретном наборе новых экземпляров. Топологическая структура новых экземпляров была выбрана из стандартных экземпляров классов B и I080 библиотеки SteinLib, общеизвестном репозитории экземпляров STP [7]. Предложенная модель для QoSMTp была реализована в оптимизационном решателе IBM ILOG CPLEX Interactive Optimizer 12.6.3. Экспериментальная оценка проводилась на сервере HP ProLiant DL385 G7 с 24 процессорами AMD Opteron 6172 и 64 Гб оперативной памяти.

Табл. 1. Экспериментальные результаты предложенной модели для QoSMTp на модифицированных экземплярах b01-b18

Table 1. Experimental results of the proposed model for the QoSMTp over modified instances b01 to b18

Экземпляр	V	T	E	Число переменных	Число ограничений
mdfb01	50	9	63	424	472
mdfb02	50	13	63	418	464
mdfb03	50	25	63	421	468
mdfb04	50	9	100	640	686
mdfb05	50	13	100	634	678
mdfb06	50	25	100	634	678
mdfb07	75	13	94	635	708
mdfb08	75	19	94	632	704
mdfb09	75	38	94	623	692
mdfb10	75	13	150	959	1028
mdfb11	75	19	150	965	1036
mdfb12	75	38	150	962	1032
mdfb13	100	17	125	840	936
mdfb14	100	25	125	843	940
mdfb15	100	50	125	837	932
mdfb16	100	17	200	1287	1382
mdfb17	100	25	200	1287	1382

mdfb18	100	50	200	1296	1394
Пример	Оптимум	Разрыв	TC(с)	TV(с)	
mdfb01	6080*	0.01%	0,59	0,59	
mdfb02	5770*	0.01%	0,5	0,5	
mdfb03	7925*	0.01%	0,71	0,71	
mdfb04	3883*	0.01%	13,39	145,86	
mdfb05	3146*	0.01%	10,06	41,84	
mdfb06	7820	5.96%	9476,51	timeout	
mdfb07	7616*	0.01%	10,11	10,11	
mdfb08	7364*	0.01%	11,12	29,35	
mdfb09	15877	0.01%	18,08	18,08	
mdfb10	5096*	0.01%	15,67	3667,31	
mdfb11	6718	11.78%	4600,61	timeout	
mdfb12	10716	0.01%	1988,57	3276,68	
mdfb13	12076	0.01%	32,19	797,12	
mdfb14	15159	3.61%	18,54	timeout	
mdfb15	20599	0.01%	33,56	118,26	
mdfb16	5288	11.12%	56,34	timeout	
mdfb17	6807	15.67%	19499,5	timeout	
mdfb18	9884*	0.01%	19513,19	21588,79	

В табл. 1 приведены результаты предложенной модели для QoSMTF, с указанием названия каждого примера, количества узлов, терминальных узлов и ребер в графе ($|V|$, $|T|$, и $|E|$). Оптимальные решения для экземпляров QoSMTF неизвестны, метод ветвей и границ и целочисленных релаксаций, применяемый в CPLEX, позволяет вычислить пороговые значения ошибок для вычисляемых решений (значение разрыва указано).

С учетом процедуры определения веса узлов (целочисленные значения, выбранные между 1 и 10 в соответствии с равномерным распределением), если абсолютная погрешность меньше одной единицы, то вычисленная оптимальная стоимость соответствует оптимальному значению (которое обозначается звездочкой *). Также сообщается количество переменных и ограничений, связанных с решением, время, требующееся решателю для расчета решения (ТС, в секундах), время, прошедшее до того момента, как решатель подтвердил/проверил, что решение найдено (TV, в секундах). Все результаты соответствуют выполнению предложенной модели, реализованной в CPLEX, с указанным выше временным ограничением в шесть часов.

Результаты, приведенные в табл. 1, показывают, что предлагаемая модель способна вычислить точные решения для всех экземпляров проблемы класса B, включая нахождение оптимального решения для девяти примеров. В 15 из 18 случаев разрыв составил менее 6%, а худший показатель был на уровне 15,67%. Тринадцать решений были рассчитаны эффективно (т.е. время исполнения составило менее минуты), но предложенной модели не удалось проверить пять решений в установленный срок – за шесть часов. Остальные результаты представлены в [8] для экземпляров i080* в библиотеке SteinLib.

5. Заключение и перспективы исследований

В данной статье представлена основанная на потоках формулировка QoSMTDP в терминах смешанного целочисленного программирования, расширяющая применение существующих потоковых моделей для STP, чтобы охватить дополнительную сложность QoSMTDP.

Поскольку для QoSMTDP не существует стандартных экземпляров проблемы, был создан новый набор путем модификации экземпляров STP класса B из библиотеки SteinLib, следуя рандомизированной процедуре для включения особенностей QoSMTDP. В результате, 13 примеров задачи были решены оптимально. В 15 из 18 случаев максимальный разрыв составил менее 6%, а худший показатель был на уровне 15,67%. Большинство решений были рассчитаны менее, чем за одну минуту. Предлагаемая модель является эффективным и гибким вариантом решения QoSMTDP, учитывая, что прежде для нее не предлагалось никаких формулировок смешанного целочисленного программирования или исчерпывающих экспериментальных оценок.

Основное направление будущих исследований связано с расширением экспериментальной оценки предлагаемой модели на большее количество экземпляров, тем самым совершенствуя вычислительные возможности предлагаемого подхода, например, путем применения методов предварительной обработки для повышения производительности решателей.

Список литературы / References

- [1] Michel Goemans and Young-Soo Myung. A catalog of Steiner tree formulations. *Networks*, vol. 23, no.1, 1993, pp. 19-28.
- [2] Mathias Hauptmann and Marek Karpinski. A Compendium on Steiner Tree Problems. Technical report. Department of Computer Science and Hausdorff Center for Mathematics, University of Bonn, 2014.
- [3] Marek Karpinski, Ion Mandoiu, Alexander Olshevsky, and Alexander Zelikovsky. Improved Approximation Algorithms for the Quality of Service Multicast Tree Problem. *Algorithmica*, vol. 42, no. 2, 2005, pp. 109-120.
- [4] Richard Karp. Reducibility among Combinatorial Problems. In *Complexity of Computer Computations*. The IBM Research Symposia Series. Springer, 1972, pp. 85-103.
- [5] Lester Ford and Delbert Fulkerson. *Flows in Networks*. Princeton University Press, 2010. 212 p.
- [6] Narendra Karmarkar. A new polynomial-time algorithm for linear programming. *Combinatorica*, vol. 4, no. 4, 1984, pp. 373-395.
- [7] Thorsten Koch, Alexander Martin, and Stefan Voß. SteinLib: An Updated Library on Steiner Tree Problems in Graphs. *Combinatorial Optimization*, vol. 11, 2001, pp. 285-325.
- [8] Claudio Risso, Sergio Nesmachnow, and Franco Robledo. Mixed Integer Programming Formulations for Steiner Tree and Quality of Service Multicast Tree problems. *Programming and Computer Software*, vol. 46, no. 8, 2020, pp. 661-678.

Информация об авторах / Information about authors

Клаудио Энрике РИССО-МОНТАЛЬДО, кандидат наук, научный сотрудник. Его исследовательские интересы включают оптимизацию сетей, многоуровневые сети, проектирование устойчивых сетей, комбинаторную оптимизацию, метаэвристику, теорию графов, оптические транспортные сети.

Claudio Enrique RISSO-MONTALDO, Ph.D., Researcher. His research interests include Network Optimization, Multilayer Networks, Design of Resilient Networks, Combinatorial Optimization, Metaheuristics, Graph Theory, Optical Transport Networks.

Франко Рафаэль РОБЛЕДО-АМОЗА, кандидат наук, профессор. Его исследовательские интересы включают исследования операций, модели надежности сетей, комбинаторную оптимизацию, моделирование методом Монте-Карло, дискретные события.

Franco Rafael ROBLED-AMOZA, Ph.D., Full Professor. His research interest include Operational Research, Network Reliability Models, Combinatorial Optimization, Monte Carlo Simulation, Discrete Events.

Серджио Энрике НЕСМАЧНОВ-КАНОВАС, кандидат наук, профессор, научный сотрудник. Область научных интересов: оптимизация, метаэвристика, высокопроизводительные вычисления, умные города.

Sergio Enrique NESMACHNOW-CÁNOVAS, Ph.D. in Computer Sciences, Full Professor and Researcher. Research interests: optimization, metaheuristics, high performance computing, smart cities.

DOI: 10.15514/ISPRAS-2020-33(2)-11



Исследование задачи обеспечения безопасности при хранении и обработке конфиденциальных данных

¹ С.А. Мартишин, ORCID: 0000-0001-5437-4049 <mart@ispras.ru>

¹ М.В. Храпченко, ORCID: 0000-0002-5147-5132 <khrapm@gmail.com>

^{1,2} А.В. Шокуров, ORCID: 0000-0002-6801-7728 <shok@ispras.ru>

¹ Институт системного программирования им. В.П. Иванникова РАН,
109004, Россия, г. Москва, ул. А. Солженицына, д. 25

² Московский физико-технический институт,
141700, Россия, Московская область, г. Долгопрудный, Институтский пер., 9

Аннотация. Приведен обзор современных подходов к работе с конфиденциальными данными в облачных вычислениях. Значительная часть хранилищ данных и систем их обработки используют облачные сервисы. Пользователи и организации рассматривают такие сервисы как поставщика услуг. В этом случае у них нет необходимости закупать, устанавливать и поддерживать дорогостоящее оборудование, они могут получить доступ к данным и результатам их обработки с любого устройства. Такое использование облачных сервисов несет в себе определенные риски, связанные с возможными угрозами раскрытия конфиденциальной информации, поскольку одним из участников протокола обеспечения безопасности доступа к облачным хранилищам данных может быть противник. Рассмотренные подходы предназначены для баз данных, в которых, с одной стороны, информация хранится в зашифрованном виде, а с другой, позволяют работать в привычной парадигме SQL-запросов. Такие подходы имеют, наряду с преимуществами, некоторые ограничения, в связи с необходимостью выбора метода шифрования, а также соблюдения баланса между его надежностью и необходимым пользователю набором запросов. Для случая, когда пользователь не ограничен рамками SQL-запросов, предложена организация облачных вычислений над конфиденциальными данными, основанная на использовании лямбда-архитектуры с ограничением на разрешенные дедуктивно безопасные запросы к базам данных и реализованная с использованием свободного программного обеспечения.

Ключевые слова: облачные вычисления; конфиденциальные данные; обеспечение безопасности; SQL; криптография; лямбда-архитектура; свободное программное обеспечение

Для цитирования: Мартишин С.А., Храпченко М.В., Шокуров А.В. Исследование задачи обеспечения безопасности при хранении и обработке конфиденциальных данных. Труды ИСП РАН, том 33, вып. 2, 2021 г., стр. 173-190. DOI: 10.15514/ISPRAS-2021-33(2)-11

Study of the problem of ensuring security in the storage and processing of confidential data

¹S.A. Martishin, ORCID: 0000-0001-5437-4049 <mart@ispras.ru>

¹M.V. Khrapchenko, ORCID: 0000-0002-5147-5132 <khrapm@gmail.com>

^{1,2}A.V. Shokurov, ORCID: 0000-0002-6801-7728 <shok@ispras.ru>

¹ *Ivannikov Institute for System Programming of the Russian Academy of Sciences,
25, Alexander Solzhenitsyn st., Moscow, 109004, Russia,*

² *Moscow Institute of Physics and Technology,
9, Institutskiy per., Dolgoprudniy, 141701, Russia*

Abstract. We introduce an overview of modern approaches to cloud confidential data processing. A significant part of data warehouse and data processing systems is based on cloud services. Users and organizations consider such services as a service provider. This approach allows users to take benefit from all of these technologies: they do not need to purchase, install and maintain expensive equipment, they can access the data and the calculation results from any device. Such data processing on cloud services carries certain risks because one of the participants of the protocol for securing access to cloud data storage may be an adversary. This leads to the threat of confidential information leakage. The above approaches are intended for databases in which information is stored in the encrypted form and they allow to work in the familiar paradigm of SQL queries. Despite the advantages such approach has some limitations. It is necessary to choose an encryption method and to maintain a balance between the reliability of encryption and the set of requests required by users. In the case if users are not limited by the framework of SQL queries, we propose another way of implementation of cloud computing over confidential data using free software. It is based on lambda architecture combined with certain restrictions on allowed deductively safe database queries.

Keywords: cloud computing; confidential data; security; SQL; cryptography; lambda architecture; free software

For citation: Martishin S.A., Khrapchenko M.V., Shokurov A.V. Study of the problem of ensuring security in the storage and processing of confidential data. *Trudy ISP RAN/Proc. ISP RAS*, vol. 33, issue 2, 2021, pp. 173-190 (in Russian). DOI: 10.15514/ISPRAS-2021-33(2)-11.

1. Введение

С развитием интернет-технологий возможность распределенного хранения и обработки данных, ранее обсуждавшаяся только теоретически, стала реальной, например, на облачных серверах. Основным компонентом такого облака является хранилище данных. В нем данные хранятся на многочисленных распределённых в сети серверах, предоставляемых в пользование клиента. Клиент рассматривает облако как большой виртуальный сервер, вне зависимости от того, как он реализован и где физически расположен.

Развитие аппаратных и программных средств привело к появлению облачных вычислений, что позволило рассматривать облако как поставщика различных IT-услуг: SaaS (программное обеспечение как услуга, в которой клиент использует программное обеспечение (ПО) облачного провайдера), PaaS (платформа как услуга, в которой клиент имеет возможность размещать принадлежащее ему ПО на облаке), IaaS (инфраструктура как услуга, в которой клиент использует только облачную инфраструктуру, а задачу управления вычислительными ресурсами, хранением и обработкой данных решает при помощи системного и прикладного ПО), DBaaS (база данных как услуга, предназначенная для управления и хранения структурированных или неструктурированных данных).

Преимущества использования облака для хранения данных и операциями над ними очевидны. К данным можно получить доступ с любого устройства, организовать совместную работу, нет необходимости закупать, устанавливать и поддерживать дорогостоящее оборудование. Однако при работе с конфиденциальными данными, хранение и обработка

которых происходит на облаке, следует учитывать, что одним из участников протоколов обеспечения безопасности доступа к облачным хранилищам данных может быть противник. В криптографии рассматриваются два типа противников.

- Активный противник, который может вмешиваться в ход выполнения криптографического протокола, как правило, может быть обнаружен путем полного анализа результатов однократного выполнения криптографического протокола.
- Пассивный противник – противник, который может получать некоторую информацию о процессе выполнения криптографического протокола, но не может в него вмешиваться. При этом полный анализ результатов неоднократного выполнения криптографического протокола не позволяет обнаружить присутствие пассивного противника.

Очевидно, что в подавляющем большинстве случаев наличие активного противника выявляется достаточно быстро. В то же время пассивный противник часто представляет даже большую опасность, поскольку он, собирая конфиденциальную информацию, может представлять значительную угрозу, оставаясь при этом необнаруженным.

Поэтому далее предполагается, что противник является пассивным, то есть может получать некоторую информацию о выполнении криптографического протокола, но не может в него вмешиваться.

Решению вопроса обеспечения безопасности хранения и обработки конфиденциальных данных посвящено большое количество работ, предложены различные подходы, основанные как на теоретических моделях, так и на готовых решениях, применяемых на практике. Ниже подробно рассмотрены примеры готовых решений, описаны их особенности, определенные недостатки, предложены пути использования теоретических разработок на практике.

2. Исследование задачи обеспечения безопасности при хранении и обработке конфиденциальных данных

Ниже приведена классификация типов облачных систем вычислений:

- частное облако (эксклюзивное использование одной организацией, состоящей из нескольких пользователей);
- облако сообщества пользователей (эксклюзивное использование определенным сообществом пользователей из нескольких организаций, которые решают задачи совместно);
- публичное облако (открытое использование широким кругом пользователей);
- гибридное облако (комбинация как минимум двух из трех, упомянутых выше).

Очевидно, что пользователи, хранящие данные на любом типе облака и использующие облачные сервисы для вычислений (хотя в первых двух случаях, возможно в меньшей степени) будут сталкиваться с проблемами, связанными с потерей данных, их искажением, утечкой результатов вычислений и пр.

Основными требованиями при хранении данных на облаке являются целостность данных, доступность и безопасность данных [1].

Под целостностью данных понимается свойство, при котором данные не могут быть изменены или уничтожены несанкционированным образом [2].

В облачном хранилище само облако рассматривается как противник, которому не доверяют. Данные могут быть неумышленно потеряны или умышленно стерты, или искажены. Поэтому необходимо иметь механизм проверки целостности данных, и регулярно проверять целостность данных на серверах облака. Таким образом проверка целостности данных направлена в первую очередь на защиту от активного противника, пытающегося внести нежелательные изменения в данные участников протокола.

Основными механизмами проверки целостности данных являются:

- POR (Proofs Of Retrievability, доказательство извлекаемости) и Provable Data Possession (PDP, доказательство владения данными) – схема подтверждает, что сохраненные данные не повреждены на сервере, в процессе хранения и извлечения клиентом;
- доказательство стираемости Proof of Erasability (POE), где облако обеспечивает полное уничтожение сохраненных данных в хранилище, когда клиент удаляет данные и разрывает соглашение с поставщиком услуг хранения.

Доступность означает, что данные хранятся в базе данных и могут использоваться по любому допустимому запросу авторизованных пользователей. Доступность может быть обеспечена при помощи различных методов резервного копирования данных.

Безопасность (конфиденциальность) данных означает, что информация, хранящаяся в базе данных и/или результаты вычислительного процесса, не предоставляется или не раскрывается неавторизованным пользователям. Основными способами защиты являются контроль доступа и шифрование данных (включая полностью гомоморфное шифрование).

Возможные атаки на облачное хранилище данных.

- Моментальный снимок (Snapshot). Можно получить моментальный снимок состояния атакованной системы. Угроза: атакующий может получить изображение виртуальной машины, выполняющей действия в системе управления базами данных (СУБД) или руткит (rootkit) операционной системы (ОС). Руткит – набор программных средств (например, исполняемых файлов, скриптов, конфигурационных файлов), которые злоумышленник устанавливает на взломанной им компьютерной системе, обеспечивающих в том числе сбор данных (параметров системы).
- Мониторинг (постоянный пассивный противник) – пассивно наблюдает за всеми операциями с базами данных включая выданные запросы, и также образом осуществляется механизм получения доступа к зашифрованным данным. Угроза, например, для зашифрованных баз – раскрытие принадлежности данных или схем шифрования, сохраняющих порядок зашифрованных данных.
- SQL-инъекция – инъекция произвольного SQL кода. Урон: возможна утечка данных за счет внедрения в SQL запрос вредоносного кода.
- Угроза всей системе в целом (нарушение данных) – рутинг (rooting, получение прав администратора) СУБД. Урон: получение полного доступа к ОС и БД.

Поскольку системы облачных вычислений получают все более широкое распространение, за последние несколько лет были предприняты усилия по разработке программных средств, которые предлагают различные подходы к проблеме безопасности облачных вычислений над конфиденциальными данными, сочетающие в себе определенные компромиссы между безопасностью, эффективностью и универсальностью.

Наиболее естественным подходом представляются теоретические схемы, основанные на современном криптографическом инструментарии:

- протокол конфиденциальных вычислений (также используют название *secure multi-party computation* – MPC);
- полностью гомоморфное шифрование (Fully Homomorphic Encryption, FHE);
- функциональное шифрование (Functional Encryption, FE).

Эти схемы могут обеспечить надежные гарантии безопасности, но их вычислительные и коммуникационные требования несовместимы с большинством практических приложений. Рассмотренные ниже системы управления базами данных (СУБД) и файловая система с возможностью хранения и обработки зашифрованных данных основаны на широко используемых СУБД MySQL, PostgreSQL, MongoDB и распределенной файловой системе HDFS.

Все перечисленные выше программные средства (ПС) являются свободным программным обеспечением (СПО), что подразумевает, в том числе, право на его изменение

(совершенствование), а также распространение копий и результатов изменения, что делает возможным адаптацию данных ПС для работы с зашифрованными данными. Таким образом, рассмотренные ниже продукты основаны на привычных и хорошо известных пользователям программных средствах. Кроме того, эти продукты практически полностью скрывают от пользователей особенности реализации шифрования данных и способы обработки зашифрованных данных за счет введения в архитектуру дополнительных модулей, что позволяет пользователям работать в привычной среде.

На практике предлагается использовать широко известные криптографические алгоритмы, если это возможно, модификацию этих алгоритмов с адаптацией к аппаратным особенностям и ограничениям или разработку новых специализированных решений.

Требования к такого рода алгоритмам хорошо известны: безопасность, стоимость и производительность. Под стоимостью понимается стоимость аппаратно-программной составляющей, необходимой для реализации алгоритмов шифрования. На практике легко оптимизировать любые две из трех целей проектирования: безопасность и стоимость, безопасность и производительность или затраты и производительность, но очень трудно оптимизировать все три цели проектирования одновременно.

3. ZeroDB

ZeroDB – СУБД со сквозным шифрованием, которая позволяет клиентам работать с зашифрованными данными (искать, сортировать, запрашивать и обмениваться) без предоставления ключей шифрования или открытого (незашифрованного) текста серверу базы данных. ZeroDB написана на Python (с некоторыми скомпилированными расширениями C) и предназначена для использования приложениями, написанными на Python. Проект больше не поддерживается (github.com/zerodb/zerodb).

ZeroDB [3] обеспечивает конфиденциальность данных, хранящихся на сервере базы данных. Предполагается, что клиент изначально владеет данными и ключами шифрования. Клиент работает с зашифрованными данными на ненадежном сервере, предполагая, что имеется пассивный (честный, но любопытный) противник, который только наблюдает, но не вмешивается в процесс хранения и обработки данных, и клиент получает правильные результаты.

Протокол запроса состоит в следующем. Клиент взаимодействует с сервером во время выполнения запроса в течение ряда циклов приема-передачи. Зашифрованный индекс хранится на сервере как B-Tree. BTree является обобщением бинарного дерева поиска; вместо того, чтобы хранить один ключ и иметь 2 дочерних узла, узлы B-Tree имеют n ключей и $(n + 1)$ дочерних узлов. Клиент обходит дерево индексов, чтобы получить необходимые зашифрованные записи. Этот обход происходит постепенно, причем каждое следующее обращение соответствует шагу вниз по дереву B-Tree индексов. Индекс состоит из блоков, которые зашифровываются перед загрузкой на сервер и расшифровываются только на стороне клиента. Таким образом, сервер не знает, как отдельные объекты организованы в B-Tree.

Поскольку протокол запросов ZeroDB включает в себя несколько раундов между клиентом и сервером, часто используемые запросы кэшируются на стороне клиента, чтобы избежать ненужных сетевых вызовов.

Клиент запрашивает у сервера зашифрованный узел, расшифровывает его, принимает решение, на какой узел пройти дальше, а затем просит сервер вернуть следующий зашифрованный узел. Клиент проходит B-Tree удаленно, пока не получит требуемую запись/записи или объект/объекты. Тем не менее, этот протокол позволяет всю обработку данных производить на стороне клиента, в открытом тексте.

Вот сводка основных особенностей ZeroDB:

- система позволяет клиентам выполнять зашифрованные операции без предоставления ключей шифрования или незашифрованных данных серверу базы данных;
- поддерживаются операции поиска, сортировки, запросов, обмена;
- ZeroDB основана на объектной ZODB, которая вместо реляционных соединений таблиц используются ссылки на объекты;
- сервер БД действует как система хранения с гарантией согласованности; шифрование, дешифрование и сжатие данных происходят на стороне клиента, следовательно, сервер не имеет представления о типах хранимых данных, что исключает возможность раскрытия данных в результате утечки данных на стороне сервера (в случае успешного проникновения будут доступны только зашифрованные данные);
- индексы в ZeroDB хранятся в B-Tree, дерево состоит из зашифрованных сегментов, каждый из которых может быть корневым, ветвлением или листовым узлом; листовые узлы дерева указывают на фактически сохраняемые объекты; таким образом, поиск в базе данных представляет собой простой обход дерева;
- узлы B-Tree – $n + 1$ дочерних узлов; поскольку высота дерева меньше, чем у двоичного дерева, поиск требует гораздо меньше доступа к памяти; временная сложность – $O(\log n)$ для поиска, вставки и удаления.
- сервер, на котором хранятся данные, никогда не знает используемый ключ шифрования; объекты, на которые ссылаются конечные узлы индексов B-Tree, также зашифрованы на стороне клиента; в результате сервер не знает, как отдельные объекты организованы в B-Tree и распространяется ли на них индексирование вообще.

4. CryptDB

CryptDB – проект для решения проблемы безопасного хранения данных в БД, обслуживаемых в облачных сервисах и других неподконтрольных системах [4]. В основе лежит использование возможностей MySQL (или PostgreSQL).

Основное преимущество CryptDB заключается в выполнении запросов над зашифрованными данными, что делает возможным применение CryptDB на практике – использование четко определенного набора SQL-операторов, каждый из которых может эффективно обрабатывать зашифрованные данные.

В CryptDB рассматриваются две основные угрозы.

- любопытный администратор базы данных (DBA) – пассивный противник, который пытается узнать конфиденциальные данные (путем отслеживания на сервере СУБД); здесь CryptDB не позволяет администратору базы данных это сделать;
- противник, который может получить полный контроль над приложениями и серверов СУБД; в этом случае, CryptDB не может предоставить никаких гарантий для пользователей, которые работали с приложением во время атаки, но все же может обеспечить конфиденциальность данных остальных пользователей.

Архитектура CryptDB (рис. 1) состоит из двух частей: прокси-сервера базы данных и немодифицированной СУБД.

CryptDB использует функции, определяемые пользователем (User Defined Functions, UDF) для выполнения криптографических операций в СУБД. Прямоугольники и закругленные прямоугольники представляют процессы и данные соответственно. Затенение указывает компоненты, добавленные CryptDB. Пунктирные линии обозначают разделение между компьютерами пользователей, сервером приложений и сервером, на котором работает база данных CryptDB.

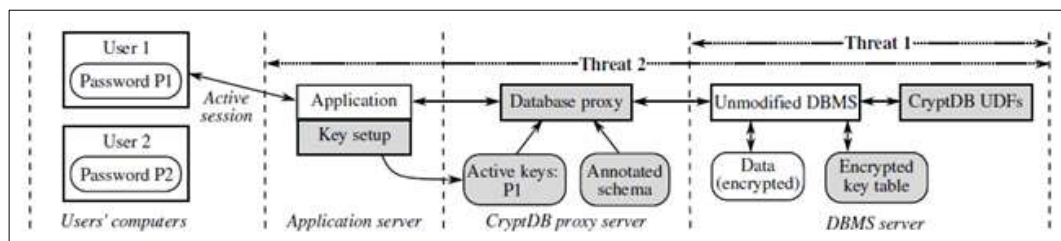


Рис.1. Архитектура CryptDB [4]

Fig. 1. CryptDB Architecture [4]

Прокси-сервер использует секретные ключи для шифрования всех вставленных данных. Основная идея вычислений над зашифрованными данными состоит в том, чтобы позволить серверу СУБД выполнять обработку запросов к зашифрованным данным, как это было бы с незашифрованной базой данных, то есть разрешить ему вычислять определенные функции над элементами данных на основе зашифрованных данных. Например, если СУБД необходимо выполнить команду группировки GROUP BY по некоторому столбцу, то сервер СУБД должен уметь определять, какие элементы в этом столбце равны друг другу, но не фактическое значение каждого.

Противник может получить доступ к ключам, используемым для шифрования всей базы данных. Решение состоит в том, чтобы зашифровать различные элементы данных (например, данные, принадлежащие разным пользователям) с разными ключами. Противник, который атакует сервер приложений или прокси-сервер, теперь может расшифровать только данные пользователей, вошедших в систему в данный момент (данные, которые хранятся на прокси-сервере). Данные неактивных в настоящее время пользователей зашифрованы ключами, которые недоступны злоумышленнику и останутся конфиденциальными.

Для успешной работы CryptDB использует три основных подхода.

- Система выполняет SQL-запросы к зашифрованным базам данных. Поскольку CryptDB выполняет SQL-запросы на зашифрованных наборах данных с четко определенными операторами (такими как проверка равенства, сравнения порядка, агрегации/суммирования и объединения), это делает возможным осуществление на практике обработку зашифрованных данных (в основном используется шифрование с симметричным ключом).
- Производится настраиваемое шифрование на основе запросов. Некоторые схемы шифрования менее устойчивы к раскрытию информации, но требуются для обработки определенных запросов. Чтобы избежать этого, CryptDB тщательно настраивает схему шифрования для любого заданного элемента данных, используя луковичное шифрование.
- Ключи шифрования связываются с паролями пользователей, чтобы каждый элемент данных в базе данных можно было бы расшифровать только через цепочку ключей. Таким образом, если пользователь не работает в данный момент с приложением и если противник не знает пароль, он не сможет расшифровать данные пользователя, даже если СУБД и сервер приложений атакованы.

Поэтому CryptDB имеет возможность переключаться на лету между различными криптографическими схемами в зависимости от типа выполняемой операции. Это реализовано за счёт «луковичного» многоступенчатого шифрования, когда данные зашифрованы в несколько слоёв разными алгоритмами. У каждого слоя – свой ключ и свой список поддерживаемых операций. На нижнем слое используются самые надёжные алгоритмы, а операции в верхних слоях возможны без расшифровки нижних слоёв.

Таким образом, CryptDB

- способна эффективно обслуживать SQL-запросы к БД – поиск, сортировку,

математические функции и др. без расшифровки записей базы;

- выполняет запросы SQL над зашифрованными данными, используя набор эффективных схем шифрования с учетом SQL;
- может связывать ключи шифрования с паролями пользователей, так что элемент данных может быть расшифрован только с помощью пароля одного из пользователей, имеющих доступ к этим данным; в результате администратор базы данных никогда не получает доступ к расшифрованным данным, даже в том случае, когда все серверы скомпрометированы.

Однако в CryptDB количество запросов ограничено из-за используемых схем шифрования.

На практике предлагается использовать StealthyCRM: приложение для безопасной облачной CRM-системы (Customer Relationship Management, управление взаимоотношениями с клиентами), поддерживающее полностью гомоморфное шифрование баз данных [5] поверх среды шифрования базы данных CryptDB. StealthyCRM использует реализацию FHE под названием HELib (библиотека гомоморфного шифрования с поддержкой гомоморфного сложения и умножения) с открытым исходным кодом для интеграции FHE с CryptDB [6].

Система StealthyCRM направляет запросы на сервер CryptDB в случае простых запросов и к зашифрованным обработчикам объектов StealthyCRM для более сложных запросов. StealthyCRM анализирует запрос и преобразует его в два новых запроса – один для CryptDB (который анализирует текстовые запросы) и другой для обработчика зашифрованных объектов StealthyCRM. Прокси-сервер решает, CryptDB или StealthyCRM необходимы для выполнения конкретного запроса и пересылает зашифрованные запросы на сторону сервера. Сервер отправляет зашифрованный результат приложению после выполнения. Модуль дешифрования StealthyCRM расшифровывает полученный результат и передает его обратно в приложение, генерирующее запрос.

CryptDB использует для различных запросов к БД следующие типы шифрования.

- **Случайный (RND).** RND обеспечивает максимальную безопасность в CryptDB: IND-CPA (indistinguishability under chosen plaintext attack), неразличимость шифротекста – криптосистема надёжна в смысле IND-CPA, если любой вероятный злоумышленник за полиномиальное время имеет лишь пренебрежимо малое «преимущество» в различении шифротекстов над случайным угадыванием. Схема является вероятностной, что означает, что два равные значения отображаются в разные шифротексты с вероятностью почти единица. С другой стороны, в случае RND не существует эффективного выполнения вычислений над зашифрованным текстом. Обычно используется блочный шифр AES или Blowfish вместе со случайным вектором инициализации.
- **Детерминированный (DET).** DET имеет немного более слабую, надежную безопасность: возможна утечка только тех зашифрованных значений, которые соответствуют одному и тому же значению данных, детерминировано генерируя один и тот же зашифрованный текст для одного и того же открытого текста. Это шифрование позволяет серверу выполнять проверки равенства, что означает, что он может выполнять выборки с предикатами равенства, соединениями равенства, GROUP BY, COUNT, DISTINCT и т. д. В криптографических терминах DET должен быть псевдослучайной перестановкой (pseudo-random permutation, PRP).
- **Шифрование с сохранением порядка (OPE).** OPE (Order-Preserving Encryption) позволяет установить отношения порядка между элементами данных на основе их зашифрованных значений, не раскрывая самих данных. Если $x < y$, то $OPE_K(x) < OPE_K(y)$ для любого секретного ключа K . Следовательно, если столбец зашифрован с помощью OPE, сервер может выполнять запросы диапазона, когда заданы зашифрованные константы $OPE_K(c_1)$ и $OPE_K(c_2)$, соответствующие диапазону $[c_1; c_2]$. Сервер также может выполнять SQL запросы с операторами ORDER BY, MIN, MAX, SORT и т.д. OPE – более слабая схема шифрования, чем DET, потому что она раскрывает

порядок.

- **Гомоморфное шифрование (НОМ).** НОМ – безопасная вероятностная схема шифрования (безопасность IND-CPA), позволяющая серверу выполнять вычисления с зашифрованными данными с расшифровкой окончательного результата на прокси-сервере. Хотя полностью гомоморфное шифрование очень медленное, гомоморфное шифрование для определенных операций является эффективным.

В частности, для поддержки суммирования был реализован метод Пэйе [7]. Это метод, например, для произведения двух зашифрованных чисел дает зашифрованную сумму значений: $НОМ_K(x)НОМ_K(y) = НОМ_K(x + y)$, где умножение выполняется по модулю некоторого значения открытого ключа.

- **Присоединение (JOIN и OPE-JOIN).** Отдельная схема шифрования, которая необходима для обеспечения проверки равенства между двумя столбцами, поскольку используются разные ключи для DET, чтобы предотвратить корреляцию между столбцами. JOIN также поддерживает все операции, разрешенные DET, а также позволяет серверу определить повторяющиеся значения между двумя столбцами. OPE-JOIN позволяет соединения по порядку отношений.
- **Поиск слова (ПОИСК).** ПОИСК используется для поиска соответствия зашифрованного текста, аналогично работе оператора MySQL LIKE.

На практике был реализовали криптографический протокол поиска по зашифрованным данным, разработанный в работе Суна (Dawn Xiaoding Song) и др. [8]. Для каждого столбца, требующего ПОИСКА, текст разбивается на ключевые слова с использованием стандартных разделителей. Далее убираются повторяющиеся слова, произвольно изменяются позиции слов, а затем каждое слово шифруется по Суну [8], так, что все слова имеют одинаковый размер. ПОИСК почти как же безопасен как RND: шифрование не раскрывает серверу СУБД, повторяется ли определенное слово в нескольких строках. При этом он может знать количество ключевых слов, зашифрованных с помощью ПОИСКА. Противник может уметь оценивать количество различных или повторяющихся слов (например, сравнивая размер шифротекстов SEARCH и RND для тех же данных).

В предлагаемой реализации прокси-сервер CryptDB состоит из библиотеки C++ и модуля Lua. Библиотека C++ состоит из парсера запросов; блока шифрования/перезаписи запросов, который шифрует поля или включает UDF в запрос; и модуль дешифрования.

5. MONOMI

MONOMI – система для безопасного выполнения операций над конфиденциальными данными на ненадежном сервере базы данных. MONOMI работает путем шифрования всей базы данных и выполнения запросов к зашифрованным данным. MONOMI была реализована с использованием PostgreSQL и предназначена для выполнения аналитических SQL-запросов [9]. По-прежнему сервер рассматривается как пассивный противник.

Существующие проблемы, для решения которых предназначена MONOMI:

- во-первых, запросы к большим наборам данных часто ограничены возможностями системы ввода-вывода, например, чтение данных с диска или потоковая передача через память; в результате схемы шифрования, которые значительно увеличивают размер данных, могут замедлить обработку запросов.
- во-вторых, аналитические запросы требуют сложных вычислений, которые могут быть неэффективными для выполнения над зашифрованными данными; вместо этого на практике должны использоваться эффективные схемы шифрования, которые могут выполнять только определенные вычисления (например, использование шифрования с сохранением порядка для сортировки и сравнения и пр.); задача состоит в том, чтобы разделить запрос на части, которые могут быть выполнены с использованием доступных схем шифрования на ненадежном сервере, и части, которые должны быть выполнены на

доверенном клиенте;

- в-третьих, некоторые методы обработки запросов по зашифрованным данным могут ускорить одни запросы, но замедлить другие, что требует тщательного планирования выполнения каждого запроса для рассматриваемой базы данных и комбинации запросов.

MONOMI решает эти проблемы тремя способами:

- во-первых, вводится разделенное клиент-серверное выполнение сложных запросов, при котором выполняется столько запросов, сколько возможно, по зашифрованным данным на сервере, а остальные компоненты запросов выполняются путем отправки зашифрованных данных доверенному клиенту, который расшифровывает данные и обрабатывает запросы в обычном режиме;
- во-вторых, вводится ряд методов, которые улучшают производительность для определенных типов запросов (но не обязательно для всех), включая предварительное вычисление каждой строки, эффективное шифрование, групповое гомоморфное сложение и предварительную фильтрацию;
- в-третьих, в архитектуру добавляются проектировщик для оптимизации физического размещения данных на сервере и планировщик для принятия решения о том, как разделить выполнение запроса между клиентом и сервером.

Поскольку разработка MONOMI основывалась на CryptDB, то она имеет аналогичные свойства, касающиеся безопасности. Хотя ненадежный сервер хранит только зашифрованные данные, он все же может получить информацию об исходных данных в виде открытого текста тремя способами:

- во-первых, некоторые схемы шифрования раскрывают информацию, необходимую для обработки запросов (например, детерминированное шифрование выявляет дубликаты для выполнения проверки равенства);
- во-вторых, некоторые схемы шифрования могут пропускать больше информации, чем необходимо; например, схема с сохранением порядка дает частичную утечку информации об открытом тексте;
- в-третьих, сервер узнает, какие строки соответствуют каждому предикату, вычисленному на сервере, например, строки, соответствующие столбцу LIKE '%keyword%'.

Комбинируя эти источники информации, сервер, являющийся противником, может получить дополнительную информацию о строках или запросах, используя статистические методы.

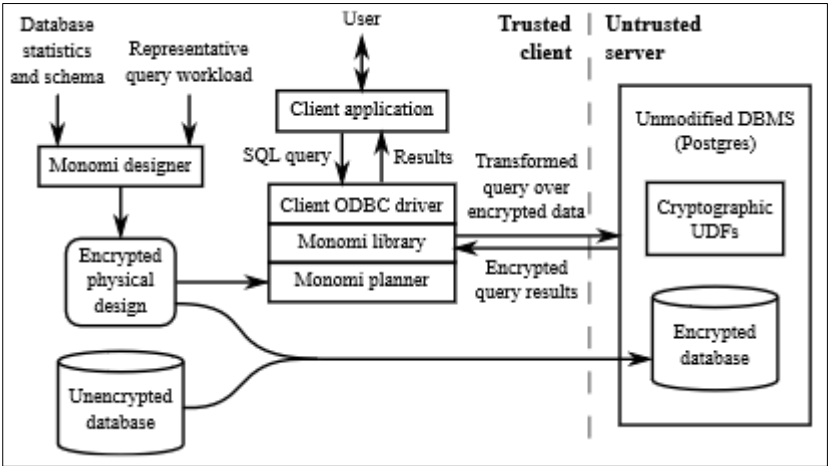


Рис. 2. Архитектура MONOMI [9]
Fig. 2. MONOMI Architecture [9]

MONOMI никогда не хранит текстовые данные на сервере и использует только схемы шифрования, необходимые для работы приложения. MONOMI позволяет администратору

182

дополнительно ограничивать схемы шифрования, используемые для особо чувствительных столбцов: например, требование шифрования с сохранением порядка (самая слабая схема MONOMI) не должна использоваться для столбцов, в которых хранятся номера кредитных карт или номеров социального страхования.

Архитектура MONOMI представлена на рис. 2.

- Во время настройки системы проектировщик MONOMI запускается на доверенной клиентской машине и определяет ее эффективную физическую конфигурацию для ненадежного сервера. Чтобы определить основные характеристики рабочей нагрузки для достижения хорошей производительности, проектировщик принимает в качестве входных данных репрезентативное подмножество запросов и статистику по данным, предоставленным пользователем. Пользователи не обязаны применять проектировщик, и вместо этого могут вручную вводить стратегию шифрования или изменять стратегию, им созданную.
- При нормальной работе приложения выдают немодифицированные запросы SQL с использованием библиотеки MONOMI ODBC, которая является единственным компонентом, имеющим доступ к ключам дешифрования. Библиотека ODBC использует планировщик, чтобы определить наилучший план выполнения запроса с разделением для приложений клиент/сервер.
- Учитывая план выполнения, библиотека выдает один или несколько запросов к зашифрованной базе данных, которая не имеет доступа к ключам дешифрования и может выполнять операции только с зашифрованными данными. База данных запускает немодифицированное программное обеспечение СУБД, такое как PostgreSQL, с несколькими определяемыми пользователем функции (UDF), предоставляемые MONOMI, которые реализуют операции с зашифрованными данными. MONOMI шифрует все данные, хранящиеся в базе данных, хотя на практике неконфиденциальные данные могут быть сохранены в виде открытого текста для повышения эффективности. После того, как клиентская библиотека получает промежуточные результаты из базы данных, расшифровывает их и выполняет любые оставшиеся операции, которые не могут быть эффективно выполнены на сервере, результаты отправляются в приложение, как если бы они выполнялись в стандартной базе данных SQL.

Вот сводка основных характеристик MONOMI:

- прототип MONOMI реализован поверх PostgreSQL;
- работает путем шифрования всей базы данных и выполнения запросов к зашифрованным данным;
- библиотека MONOMI ODBC является единственным компонентом, имеющим доступ к ключам расшифровки;
- MONOMI представляет новый подход, основанный на раздельном исполнении запроса на клиент-сервере; это позволяет выполнить часть запроса на ненадежном сервере поверх зашифрованных данных; для других частей запроса MONOMI загружает промежуточные результаты на клиента;
- библиотека MONOMI ODBC получает промежуточные результаты из базы данных, расшифровывает их и выполняет все оставшиеся операции, которые не могут быть эффективно выполнены на сервере;
- реализовано несколько методов, которые улучшают производительность: предварительное вычисление строки, пространственно-эффективное шифрование, групповое гомоморфное сложение и предварительная фильтрация;
- поскольку эти оптимизации хорошо работают для одних запросов и неэффективны для других, MONOMI имеет планировщик, чтобы определить лучший способ выполнения запроса.

6. Seabed

В Seabed используются HDFS и Spark. То есть, в отличие от CryptDB и Monomi, система основана не на реляционной базе данных, а на файловой системе, предназначенной для хранения неструктурированной информации. Прототип реализован на Apache Spark [10].

Цель – поддержка Business intelligence (BI) – компьютерных методов, которые обеспечивают перевод транзакционной деловой информации в удобную для восприятия человеком форму, пригодную для бизнес-анализа, а также средства для массовой работы с такой обработанной информацией.

Предполагается, что противник является пассивным (честным, но любопытным), то есть противник будет пытаться узнать конфиденциальные данные, но не будет их повреждать или иным образом вмешиваться в работу системы.

В плане шифрования Seabed

- использует новую, аддитивно-симметричную схему гомоморфного шифрования (Additively Symmetric Homomorphic Encryption Scheme, ASHE) для эффективного выполнения крупномасштабных агрегаций;
- представляет новую рандомизированную схему шифрования под названием Splayed ASHE, или SPLASHE.

Схема Seabed приведена ниже (рис. 3).

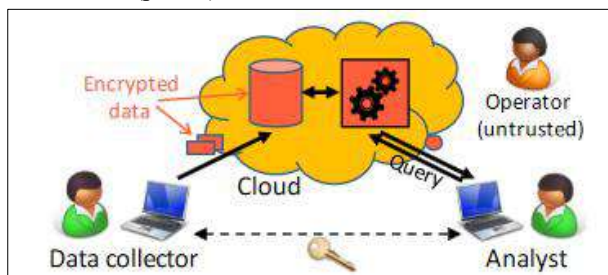


Рис. 3. Схема Seabed [10]

Fig. 3. Seabed outline [10]

Сборщик данных (Data collector) собирает большое количество данных, шифрует их и загружает на облако, которому не доверяет. Аналитик может генерировать запросы для обработчика запросов в облаке. Ответы будут зашифрованы, но аналитик может расшифровать их с помощью секретного ключа, который является общим со сборщиком данных. Рабочая нагрузка, которую предполагается поддерживать, состоит из запросов в стиле OLAP для больших наборов данных. OLAP (OnLine Analytical Processing, интерактивная аналитическая обработка) – технология обработки данных, заключающаяся в подготовке суммарной (агрегированной) информации на основе больших массивов данных, структурированных по многомерному принципу. Реализации технологии OLAP являются компонентами программных решений класса Business Intelligence.

Один из распространенных подходов к решению вышеуказанной проблемы – использовать гомоморфное шифрование. Например, есть криптосистемы с аддитивным гомоморфизмом, такие как Пэе (Pascal Paillier) [7], что позволяет облаку выполнять агрегирование непосредственно на зашифрованных данных.

Однако бывают ситуации, когда необходимо, чтобы облако имело доступ к какому-то свойству зашифрованных значений (шифрование с сохранением свойств). Например, чтобы вычислить соединение (JOIN), облако должно иметь возможность согласовывать зашифрованные значения. В этом случае можно использовать детерминированное шифрование, где каждое значение отображается ровно на один зашифрованный текст. Однако такие схемы подвержены частотным атакам, особенно если столбец может принять только небольшое количество значений, и облако знает, что какое-то значение будет

наиболее распространенным. Еще одним примером операции, доступной облаку в случае использования схемы шифрования с сохранением свойств, является выбор строки на основе диапазона значений в зашифрованном столбце. Здесь можно использовать шифрование OPE. Очевидно, что в этих схемах есть компромисс между конфиденциальностью, производительностью и функциональностью.

Примеры операций в Seabed, которые могут быть полностью выполнены на сервере, – вычисление суммы; среднего, минимума и т.д. Пример операции, которые могут быть выполнены при наличии предварительной обработки клиента, – квадратичные вычисления, необходимые для более сложной аналитики, такой как обнаружение аномалий, линейной регрессии в одном измерении, и деревьев решений.

ASHE предполагает, что открытые тексты взяты из аддитивной группы $Z_n = \{0, 1, \dots, n - 1\}$. Также предполагается, что отправитель и получатель, шифрующие входную и выходную информацию соответственно, имеют общий секретный ключ k , а также общую псевдослучайную функцию (PRF) $F_k: I \rightarrow Z_n$, преобразующая идентификатор из набора I в случайный номер из Z_n .

Один из возможных вариантов PRF – $F_k(i) := H(i || k) \bmod n$, где $i \in I$, H – криптографическая хеш-функция (моделируется как случайная функция), $||$ обозначает конкатенацию, размер диапазона H кратен n . Другим вариантом может быть использование AES (Advanced Encryption Standard – симметричного алгоритма блочного шифрования), когда он используется как псевдослучайная перестановка.

Ниже рассмотрена архитектура Seabed (рис. 4).

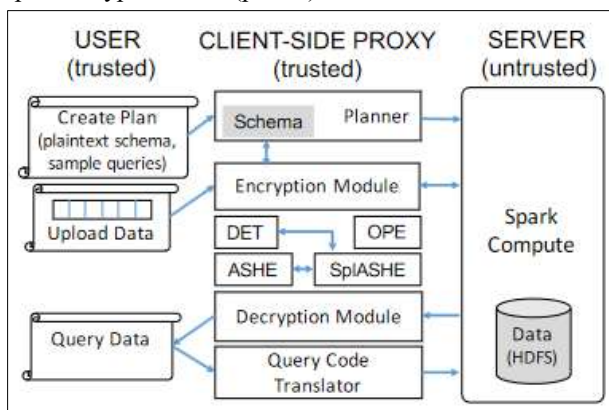


Рис. 4. Архитектура Seabed [10]

Fig. 4. Seabed architecture [10]

Пользователь взаимодействует с клиентским прокси Seabed, который является доверенным. Прокси, в свою очередь, взаимодействует с сервером Seabed, которому не доверяют.

Пользователь может отправлять запросы трех типов.

- **Создание плана:** сначала пользователь предоставляет схему в виде открытого текста и образец запроса, заданный для планировщика Seabed. Планировщик использует их и специальную процедуру для определения схем шифрования столбцов.
- **Загрузка данных.** Затем пользователь отправляет данные в виде открытого текста в Модуль шифрования Seabed. Данные шифруются с использованием необходимой схемы шифрования и записи добавляются в таблицу, хранящуюся в облаке. Это непрерывный процесс; вставки в базу данных обрабатываются таким же образом.
- **Запрос данных:** во время анализа пользователь отправляет скрипт запроса в транслятор запросов Seabed, который изменяет запросы для обработки зашифрованных данных перед их отправкой на сервер. Сервер выполняет запросы и выдает ответ модулю

дешифрования прокси. После расшифровки и дальнейшей обработки (если она требуется) результаты отправляются обратно пользователю.

Планировщик данных определяет, как зашифровать каждый столбец в схеме, учитывая список конфиденциальных столбцов. Пользователь также предоставляет образец набора запросов, который используется планировщиком для выбора алгоритмов шифрования.

Модуль шифрования шифрует открытый текст.

Транслятор запросов предназначен для переписывания запросов пользователей, чтобы адаптировать их к выбранному способу шифрования данных.

7. Arx

Система Arx реализована на базе MongoDB, СУБД категории NoSQL [11]. Декларация разработчиков:

- отказ от слабых схем шифрования;
- в основе схемы шифрования для Arx используется AES, (что связано с наличием соответствующей аппаратуры); возможно использование других методов шифрования;
- уровень безопасности данных IND-CPA.

Предполагается, что сервер БД может быть размещен в частном или публичном облаке. Модель угроз Arx предполагает, что противник не имеет доступа к данным на стороне клиента (пользователи, приложения и клиентский прокси Arx), а имеет доступ только к информации на стороне сервера (прокси-сервера Arx и серверы баз данных).

Предполагается, что

- противником, например, может быть любопытный администратор облачного провайдера;
- противник может видеть всю информацию на сервере – все содержимое базы данных, любые хранящиеся данные или ключи, в памяти, и любые сетевые сообщения; следовательно, если секретный ключ находится в основной памяти, злоумышленник может расшифровать базу данных;
- противник пассивен, то есть не изменяет содержимое базы данных или результаты запроса.

Arx не полагается на какое-либо доверенное оборудование на сервере.

Основная задача предложенной архитектуры – не менять имеющийся сервер БД и приложения. Ниже показана архитектура Arx (рис. 5).

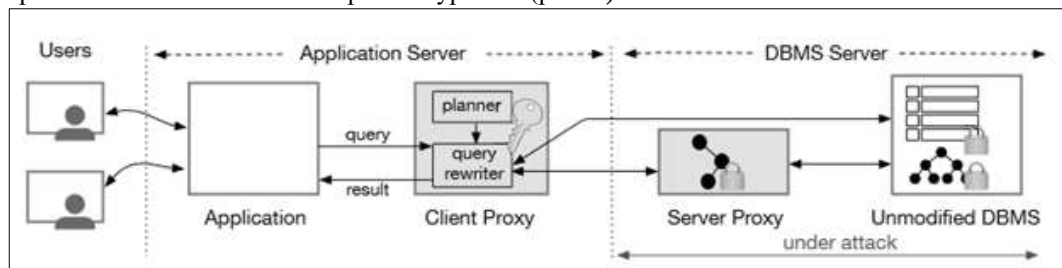


Рис. 5. Архитектура Arx [11]

Fig. 5. Arx's architecture [11]

- Доверенный клиентский прокси развернут на сервере приложений, ненадежный прокси-сервер развернут на сервере СУБД, сервер СУБД размещен на частном или публичном облаке.
- Клиентский прокси перехватывает запросы и шифрует конфиденциальную информацию. Прокси-сервер поддерживает индексы зашифрованных данных и выполняет входящие запросы.
- Серые прямоугольники изображают компоненты, представленные Arx, а остальные

прямоугольники представляют существующие компоненты.

- Ключ указывает, что конфиденциальные данные в компоненте всегда остаются строго зашифрованными

Заметим, что клиентский прокси экспортирует в приложение тот же API, что и сервер БД, поэтому приложение не нужно менять. Прокси-сервер взаимодействует с сервером БД, вызывая его неизменный API (например, выдавая запросы); другими словами, прокси-сервер ведет себя как постоянный клиент сервера БД.

- Клиентский прокси хранит ключ (master key). Переписывает запросы, шифрует конфиденциальные данные и пересылает зашифрованные запросы на прокси-сервер для выполнения
- Прокси-сервер помогает серверу выполнять зашифрованную обработку данных.
- Прокси клиента не хранит базы данных, а хранит метаданные (информацию о схеме) и небольшой дополнительный кеш. Почти во всех случаях клиентский прокси обрабатывает только результаты запросов (например, для их расшифровки).

Мастер-ключ – ключ шифрования для симметричного алгоритма шифрования (например, в Windows 10 доступен AES).

Arx представляет два новых индекса базы данных:

- Arx-RANGE для запросов из диапазона;
- Arx-EQ для запросов на равенство.

Оба индекса являются структурами данных, построенными на основе AES.

Индекс – бинарное дерево поиска. Arx-RANGE позволяет серверу самостоятельно обходить дерево при сохранении безопасности.

Для обеспечения сложных вычислений с высокой степенью безопасности, Arx-RANGE использует одноразовую обфускацию на каждом узле в дереве индексов для выполнения сравнения.

Планировщик Arx принимает в качестве входных данных:

- набор шаблонов запросов, специфичных для Arx;
- список регулярных индексов

и создает:

- план шифрования данных;
- список индексов Arx;
- план выполнения запроса для каждого шаблона запроса.

Чтобы использовать Arx, разработчик приложения должен указать:

- какие поля являются конфиденциальными и должны быть зашифрованы;
- операции, выполняемые на конфиденциальных полях.

Поддерживаемые операции:

- чтение:
 - ВЫБОРКА (сумма, сортировка, группировка, агрегирование, например, с минимальной постобработкой на клиентском прокси, среднее или стандартное отклонение);
- запись:
 - INSERT;
 - DELETE;
 - UPDATE.

8. Обсуждение

Напомним, что во всех случаях сервер рассматривается как пассивный противник.

Следует заметить, что все рассмотренные выше инструменты для работы с базами данных основываются на хорошо известных СУБД (MySQL, PostgreSQL, известных NoSQL-базах) или структурах (Btree), а также на распределенной файловой системе (HDFS).

Основная задача – сохранить конфиденциальность данных при работе на облаке, которому не доверяют, и оставить возможность использовать запросы SQL-типа на зашифрованных данных.

Для этого применяются различные методы шифрования. Так, например, в CryptDB используются различные схемы шифрования (детерминированное, гомоморфное, шифрование, сохраняющее порядок, шифрование, позволяющее выполнять операции присоединения). Следует заметить, что не всегда эти методы шифрования одинаково пригодны. Например, при использовании шифрования, сохраняющего порядок, пассивный противник путем наблюдения может собрать статистику. При этом он не должен определять точное значение данных p , соответствующее зашифрованному значению c . Нарушение конфиденциальности может произойти даже в том случае, если противник может c определенной вероятностью оценить, что значение данных p находится в интервале $[p_1; p_2]$. Более того, несмотря на гибкий выбор шифрования, количество типов SQL запросов остается ограниченным.

Кроме того, для осуществления шифрования и расшифровки данных, планирования порядка выполнения запросов, предварительной обработки данных в архитектуру вводятся дополнительные компоненты (доверенные прокси-сервера, планировщики и пр.).

Заметим, что все это замедляет выполнение запросов. При этом некоторые схемы шифрования, необходимые для выполнения определенных запросов, не обеспечивают достаточную конфиденциальность данных, то есть при длительном наблюдении пассивный противник может получить значительную часть информации о конфиденциальных данных.

Таким образом, все системы, описанные выше, позволяют пользователям работать в привычной парадигме SQL-запросов. Однако при этом ограничивается спектр возможных запросов, а для некоторых разрешенных запросов заведомо повышается возможность утечек конфиденциальных данных за счет выбора менее надежных схем шифрования. Кроме того, большая часть запросов выполняется не в один раунд, что сильно замедляет время его выполнения.

Все это заставляет исследовать альтернативные алгоритмы для работы с базами данных, особенно тогда, когда пользователь не ограничен рамками SQL-запросов, а имеет возможность задавать запрос на вычисление некоторых разрешенных для вычисления функций над зашифрованными данными.

В этом случае протокол и модели облачных вычислений на основе криптосерверов, предложенные в [12,13], являются более предпочтительным способом работы с конфиденциальными данными, поскольку для этой модели доказана ее стойкость [14] и дедуктивная безопасность запросов к базам данных [15].

Кроме того, в [16] были предложены методы реализации вычислений над конфиденциальными данными для предложенной модели при помощи пакетной обработки данных на основе СПО. Это позволяет использовать хорошо известное программное обеспечение (HDFS, PostgreSQL и т.д.) для работы с конфиденциальными данными, а также выполнять их быструю обработку.

Такой подход прекрасно работает, когда, например, необходимо отнести человека (пациента) к определенной группе лечения по некоторому заболеванию. При этом имеется значительное число показателей, которые должны быть учтены. Для этого потребуется вычислить только статистические функции, которые являются дедуктивно безопасными. Соответствующая выборка должна производиться на основе некоторого идентификатора личности, который необходимо хранить в зашифрованном виде. Идентификатором может быть, например, UUID (Universally Unique Identifier – универсальный уникальный идентификатор), который позволяет распределенным системам уникально идентифицировать информацию без обращения к единому центру координации. Во многих СУБД СПО, например, в MySQL, существует встроенная функция, генерирующая UUID и возвращающая значение, которое представляет собой 128-битное число, в виде строки формата utf8 из пяти

шестнадцатеричных чисел.

В этом случае в СУБД хранится таблица, аналогичная таблице паролей. Этот первичный ключ позволит сгруппировать записи, относящиеся к человеку, и собрать статистическую информацию, не раскрывая конфиденциальные данные. Далее производится обработка больших объемов информации при помощи упомянутого СПО (уровень пакетной обработки) и результат возвращается пользователю.

9. Заключение

Вопрос сохранения конфиденциальности данных на облаке является одним из наиважнейших, поскольку одним из участников протокола доступа к облачным хранилищам данных может быть противник.

В предложенном обзоре рассмотрены различные подходы к работе с конфиденциальными данными на облаке. Эти подходы базируются на широко используемых СУБД (MySQL, PostgreSQL, NoSQL базе данных MongoDB), структурах (Btree), а также на распределенной файловой системе (HDFS).

Программные средства ZeroDB, CryptDB, MONOMI, Arx и Seabed, разработанные на основе вышеупомянутых СУБД и файловой системы, позволяют работать в парадигме SQL-запросов. Конфиденциальность обеспечивается при помощи использования различных схем шифрования (в зависимости от разрешенных запросов) и введения в архитектуру дополнительных элементов, таких как доверенные прокси-сервера, планировщики и некоторые другие, которые позволяют организовать вычислительный процесс.

Из недостатков рассмотренных средств следует отметить, что в некоторых случаях при выполнении запросов приходится использовать схемы шифрования, которые не могут обеспечить полную конфиденциальность данных. Также работа планировщиков и других вспомогательных элементов архитектуры замедляет выполнение запросов.

Предложены методы работы с конфиденциальными данными для тех случаев, когда пользователь не ограничен парадигмой SQL-запросов.

Список литературы / References

- [1] Huang C-T, Huan L, Qin Z, Yuan H, Zhou L, Varadharajan V, Jay Kuo C.-C. Survey on securing data storage in the cloud. *APSIPA Transactions on Signal and Information Processing*, vol. 3, 2014, article e7.
- [2] X.800: Security architecture for Open Systems Interconnection for CCITT applications. URL: <https://www.itu.int/rec/T-REC-X.800-199103-I/en>, accessed 25.12.2020.
- [3] Egorov M., Wilkison M. ZeroDB white paper. arXiv:1602.07168, 2016, 11 p.
- [4] Popa A., Redfield C., Zeldovich N., and Balakrishnan H. CryptDB: Protecting Confidentiality with Encrypted Query Processing. In *Proc. of the Twenty-Third ACM Symposium on Operating Systems Principles*, 2011, pp 85-100.
- [5] Felipe M.R., Mi Aung K.M., Ye X. and Yonggang W. StealthyCRM: A Secure Cloud CRM System Application that Supports Fully Homomorphic Database Encryption. *International Conference on Cloud Computing Research and Innovation (ICCCRI)*, 2015, pp. 97-105.
- [6] Halevi S. HELib. URL: <https://github.com/shaih/HElib>, accessed 25.12.2020.
- [7] P. Paillier. Public-key cryptosystems based on composite degree residuosity classes. *Lecture Notes in Computer Science*, vol. 1592, 1999, pp. 223-238.
- [8] Song D.X., Wagner D., and Perrig A. Practical techniques for searches on encrypted data. In *Proc. of the 21st IEEE Symposium on Security and Privacy*, 2000, pp. 44-55.
- [9] Tu S., Kaashoek M. F., Madden S., Zeldovich N. Processing Analytical Queries over Encrypted Data, *Proceedings of the VLDB Endowment*, vol. 6, no. 5, 2013, pp. 289-300.
- [10] Papadimitriou A., Bhagwan R., Chandran N., Ramjee R., Singh H., Modi A. Big Data Analytics over Encrypted Datasets with Seabed. In *Proc. of the 12th USENIX conference on Operating Systems Design and Implementation*, 2016, pp. 587-602
- [11] Poddar R., Boelter T., Popa A. Arx: An Encrypted Database using Semantically Secure Encryption. *Proceedings of the VLDB Endowment*, vol. 12, no. 11, 2019, pp. 1664-1678.
- [12] Варновский Н.П., Мартиншин С.А., Храпченко М.В., Шокуров А.В. Методы пороговой

- криптографии для защиты облачных вычислений, Труды ИСП РАН, том 26, вып. 2, 2014 г., с. 269-274 / Varnovskiy N.P., Martishin S.A., Khrapchenko M.V., Shokurov A.V. A Threshold Cryptosystem in Secure Cloud Computations. *Trudy ISP RAN/Proc. ISP RAS*, vol. 26, issue 2, 2014, pp. 269-274 (in Russian). DOI: 10.15514/ISPRAS-2014-26(2)-12.
- [13] Варновский Н.П., Мартишин С.А., Храпченко М.В., Шокуров А.В. Пороговые системы гомоморфного шифрования и защита информации в облачных вычислениях. Программирование, том 41, no. 4, 2015 г., стр. 47-51 / Secure cloud computing based on threshold homomorphic encryption. Varnovskiy N.P., Martishin S.A., Khrapchenko M.V., Shokurov A.V. *Programming and Computer Software*, vol. 41, no. 4, 2015, pp. 215-218.
- [14] Варновский Н.П., Захаров В.А., Шокуров А.В. К вопросу о существовании доказуемо стойких систем облачных вычислений. Вестник Московского университета. Сер. 15. Вычислительная математика и кибернетика, no. 2, 2016 г., стр. 32-45. / Varnovsky N.P., Zakharov V.A., Shokurov A.V. On the existence of provably secure cloud computing systems. *Moscow University Computational Mathematics and Cybernetics*, vol. 40, no. 2, 2016, pp. 83-88
- [15] Варновский Н.П., Захаров В.А., Шокуров А.В. О дедуктивной безопасности запросов к базам конфиденциальных данных в системе облачных вычислений. Вестник Московского университета. Серия 15: Вычислительная математика и кибернетика, no. 1, 2017 г., стр. 38-44 / Varnovsky N.P., Zakharov V.A., Shokurov A.V. On the deductive security of queries to confidential databases in cloud computing systems. *Moscow University Computational Mathematics and Cybernetics*, vol. 41, no. 1, 2017, pp. 38-43.
- [16] Мартишин С.А., Храпченко М.В. Организация облачных вычислений над конфиденциальными данными на СПО. Труды 15-й конференции «Свободное программное обеспечение в высшей школе», 2020 г., стр. 171-174 / Martishin S.A., Khrapchenko M.V. Organization of cloud computing over confidential data on open source software. In *Proc/ of the 15th Conference on Free Software in Higher Education*, 2020, pp. 171-174 (in Russian).

Информация об авторах / Information about authors

Сергей Анатольевич МАРТИШИН, кандидат наук, научный сотрудник. Научные интересы: защита информации, облачные вычисления, вычисления над зашифрованными данными, гомоморфные вычисления, базы данных, СУБД.

Sergey Anatolyevich MARTISHIN, Candidate of Science, Research Fellow. Research interests: information security, cloud computing, computing over encrypted data, homomorphic computing, databases, DBMS.

Marina Valeryevna KHRAPCHENKO, Research Fellow. Research interests: cryptography, cloud computing, free software.

Марина Валерьевна ХРАПЧЕНКО, научный сотрудник. Научные интересы: криптография, облачные вычисления, свободное программное обеспечение.

Александр Владимирович ШОКУРОВ, кандидат наук, доцент, ведущий научный сотрудник ИСПРАН, доцент МФТИ. Сфера научных интересов: гомоморфное шифрование, облачные вычисления, криптография на решетках, алгебра.

Alexander Vladimirovich SHOKUROV, Candidate of Physics and Mathematics, Associate Professor, Leading Researcher, ISP RAS, Associate Professor at MIPT. Research interests: homomorphic encryption, cloud computing, lattice cryptography, algebra.

DOI: 10.15514/ISPRAS-2020-33(2)-12



Выполнимость мю-исчисления с арифметическими ограничениями

¹ Й. Лимон, ORCID: 0000-0002-3754-6926 <zs16017367@estudiantes.uv.mx>

² Э. Барсенас, ORCID: 0000-0002-1523-1579 <barcenas@unam.mx>

¹ Э. Бенитес-Герреро, ORCID: 0000-0001-5844-4198 <edbenitez@uv.mx>

² Г. Молеро-Кастильо, ORCID: 0000-0002-6330-6408 <gmolero@fi-b.unam.mx>

² А. Веласкес Мена, ORCID: 0000-0003-3509-6236 <mena@fi-b.unam.mx>

¹ Университет Веракрузана.

Мексика, 91000, Веракрус, Халапа

² Национальный автономный университет Мексики.

Мексика, 04510, Мехико, Мэрия Койоакана

Анотация. Пропозициональное модальное μ -исчисление является хорошо известным языком спецификаций систем помеченных переходов. В этой работе мы изучаем расширение этой логики с помощью обратных модальностей и арифметических ограничений Пресбургера, интерпретируемых на древовидных моделях. Мы описываем алгоритм выполнимости, основанный на построении по уровням моделей Фишера-Ларднера. Также сообщается о совместном выполнении нескольких экспериментов. Кроме того, мы описываем применение алгоритма для решения задач статического анализа полуструктурированных данных.

Ключевые слова: модальные логики; автоматизированное построение логического вывода; формальная верификация; пресбургерова арифметика; полуструктурированные данные

Для цитирования: Лимон Й., Барсенас Э., Бенитес-Герреро Э., Молеро-Кастильо Г., Веласкес Мена А. Выполнимость мю-исчисления с арифметическими ограничениями. Труды ИСП РАН, том 33, вып. 2, 2021 г., стр. 191-200. DOI: 10.15514/ISPRAS-2021-33(2)-12

Mu-Calculus Satisfiability with Arithmetic Constraints

¹ Y. Limón, ORCID: 0000-0002-3754-6926 <zs16017367@estudiantes.uv.mx>

² E. Bárcenas, ORCID: 0000-0002-1523-1579 <barcenas@unam.mx>

¹ E. Benítez-Guerrero, ORCID: 0000-0001-5844-4198 <edbenitez@uv.mx>

² G. Molero-Castillo, ORCID: 0000-0002-6330-6408 <gmolero@fi-b.unam.mx>

² A. Velázquez-Mena, ORCID: 0000-0003-3509-6236 <mena@fi-b.unam.mx>

¹ Universidad Veracruzana,

Xalapa, Veracruz, México, 91000

² National Autonomous University of Mexico,

Coyoacan Mayor's Office, Mexico City, CDMX, C.P. 04510, Mexico

Abstract. The propositional modal μ -calculus is a well-known specification language for labeled transition systems. In this work, we study an extension of this logic with converse modalities and Presburger arithmetic constraints, interpreted over tree models. We describe a satisfiability algorithm based on breadth-first construction of Fischer-Lardner models. An implementation together several experiments are also reported. Furthermore, we also describe an application of the algorithm to solve static analysis problems over semi-structured data.

Keywords: Modal Logics; Automated Reasoning; Formal Verification; Presburger Arithmetic; Semi-Structured Data

For citation: Limón Y., Bárcenas E., Benítez-Guerrero E., Molero-Castillo G., Velázquez-Mena A. Mu-Calculus Satisfiability with Arithmetic Constraints. *Trudy ISP RAN/Proc. ISP RAS*, vol. 33, issue 2, 2021, pp. 191-200 (in Russian). DOI: 10.15514/ISPRAS-2021-33(2)-12.

1. Введение

Модальные логики широко известны как языки спецификации для верификации аппаратных и программных систем. Эти логики получили известность в сообществе формальной верификации благодаря присущему им тонкому балансу между их выразительной силой и вычислительной сложностью связанных с ними алгоритмов формирования рассуждений.

Можно рассматривать μ -исчисление как расширение модальной логики высказываний K операторами неподвижной точки. Известно, что проблема разрешимости μ -исчисления, а существования или отсутствия алгоритма для проверки валидности любой формулы, является EXPTIME-полной [1]. Обратные модальности (converse modality) – это конструкторы, способные моделировать свойства в прошлом, если мы рассматриваем переходы в моделях как временные отношения. Номиналы (nominal) – это отдельные формулы, обозначающие одиночные состояния в моделях. Градуированные модальности (graded modality) ограничивают количество последующих состояний модели некоторым натуральным числом. Если расширить μ -исчисление чем-то одним – обратными модальностями, номиналами или градуированными модальностями – проблема разрешимости останется EXPTIME-полной. Расширение любыми двумя из этих трех конструкций также оставляет проблему разрешимости EXPTIME-полной. Однако полностью обогащенное μ -исчисление, включающее обратные модальности, номиналы и градуированные модальности и номиналами, разрешимым не является [1]. При интерпретации над древовидными структурами для полностью обогащенного μ -исчисления EXPTIME-полная разрешимость сохраняется [2]. Кроме того, показано, если полностью обогащенное μ -исчисление расширить еще и ограничениями арифметики Пресбургера (Mojżesz Presburger) – ограничениями последующих узлов арифметическими выражениями Пресбургера, – проблема разрешимости останется EXPTIME-полной [3]. Доказательство основано на алгоритме выполнимости. Однако о реализации не сообщается.

В настоящей работе мы описываем реализацию вместе с несколькими экспериментами. Начальные результаты этой работы были впервые представлены в [4].

Кроме расширения μ -исчисления градуированными модальностями, изучались и другие логики с числовыми ограничениями. В [5] показана PSPACE-полная разрешимость расширения модальной логики K ограничениями арифметики Пресбургера. В [6] показано, что EXPTIME-полной разрешимостью обладает логика с фиксированной точкой над деревьями с ограничениями Пресбургера. В [7] показана неразрешимость монадической логики второго порядка с арифметикой Пресбургера. Хотя в некоторых из этих работ явно или неявно сообщается о соответствующем алгоритме принятия решений, реализация не упоминается. Высокая вычислительная сложность этих алгоритмов исключает возможность непосредственной реализации. В настоящей работе мы предлагаем бинарно-векторное представление узлов в древовидных моделях. Это позволяет справиться с потенциально экспоненциальным размером моделей.

В литературе известно несколько арифметических решателей. Z_3 – это инструмент, разработанный над арифметическим расширением логики первого порядка [8]. Расширение Z_3 оператором неподвижной точки, называемое μZ , описано в [9]. Этот конструктор позволяет выражать рекурсивные свойства над реляционными структурами. Авторы работы [11] представили би-интуиционистский модальный алгоритм μ -исчисления, который переписывает логические формулы первого порядка в формулы модальной логики и решает

рекурсивные неравенства, а его модели являются древовидными [11]. Еще одним известным инструментом, способным производить логический вывод при наличии арифметических ограничений, является Мона [10]. Этот инструмент построен на монадической теории второго порядка. Следует отметить, что Мона может представлять или выражать то же самое, что и μ -исчисление. Хотя доказано, что вычислительная сложность этой теории не элементарна, обсуждаются несколько приложений с логическим выводом над полуструктурированными данными.

В отличие от этих инструментов, реализация алгоритма, описываемая в настоящей статье, строится на менее дорогостоящей теории.

Насколько нам известно, в современной литературе упоминаются два фреймворка для логического вывода над полуструктурированными данными [12-14]. Точнее, эти фреймворки решают несколько проблем вывода для разрешимых фрагментов запросов XPath и XML-схем, таких как вложенность результатов (query containment) запросов и проверка типов. Оба схем основаны на алгоритмах выполнимости μ -исчисления без арифметических ограничений. В этой статье мы описываем несколько экспериментов по логическому выводу над запросами XPath с арифметическими ограничениями.

2. μ -исчисление с арифметическими ограничениями

В этом разделе мы вводим синтаксис и семантику μ -исчисления с арифметическими ограничениями на древовидных моделях. Пусть заданы алфавиты $PROP$ и MOD . $PROP$ используется как множество высказываний, а MOD – как множество модальностей. Множеством модальностей, например, может быть $\{1, 2, 3, 4\}$. Тогда на интуитивном уровне в древовидных моделях 1 обозначает отношение непосредственных потомков, 2 обозначает правого брата, 3 – родителя и 4 – левого брата. Множество формул μ -исчисления определяется следующей грамматикой:

$$\begin{aligned}\varphi &::= p|X|\neg\varphi|\varphi \vee \varphi|(t)\varphi|tX.\varphi|\gamma > b, \\ \gamma &::= \alpha\varphi|\gamma + \gamma,\end{aligned}$$

где p – высказывание, t – модальность, X – переменная, $\alpha \in \mathbb{Z} \setminus \{0\}$, и $b \in \mathbb{N}$.

Теперь мы дадим интуитивное представление об интерпретации формул как подмножества узлов деревьев: высказывания интерпретируются как метки узлов; отрицание – как дополнение множества; дизъюнкция – как объединение множеств; модальные формулы $(t)\varphi$ – как узлы с t -доступностью для узла, верифицирующего φ ; формулы с неподвижной точкой – как рекурсия и арифметические формулы $\gamma > b$ – как узлы с дочерними элементами, удовлетворяющими арифметическому выражению. Другие конструкции могут быть определены обычным образом: $\varphi \wedge \psi := \neg(\neg\varphi \vee \neg\psi)$, $\gamma \leq b := \neg(\gamma > b)$. Например, формула $\varphi: \varphi = a \wedge (1)b$ обозначает узлы с именем a или с потомком a .

Для того чтобы описать точную семантику формулы, рассмотрим определение дерева в стиле структуры Крипке (Saul Aaron Kripke) (N, R^m, L) , представленное в [3]. Рассмотрим дерево T и валюацию $V: Var \rightarrow 2^N$, где Var – множество переменных. В [3] представлено формальное определение μ -исчисления со счетчиками. Касательно дерева T мы говорим, что формула φ выполнима тогда и только тогда, когда существует интерпретация φ над деревом T такая, что валюация V не пуста, то есть $\llbracket \varphi \rrbracket_V^T \neq \emptyset$. Формула φ валидна тогда и только тогда, когда она выполнима для каждого дерева.

3. Выполнимость

В этом разделе мы описываем алгоритм выполнимости, как он представлен в [3]. Алгоритм основан на восходящем построении деревьев Фишера-Ладнера (Michael John Fischer, Richard Emil Ladner) [15]. Сначала мы опишем это понятие. Для технического удобства и без потери

общности мы рассматриваем только бинарные деревья. Будем считать, что при интерпретации формул в бинарных деревьях модальности 1, 2, 3, 4 обозначают первого ребенка (слева направо), следующего родного брата, родителя и предыдущего родного братом соответственно. В этой работе мы для определения алгоритма выполнимости рассматриваем формулы только в нормальной форме отрицания (negation normal form) [3].

Для заданной входной формулы φ узлы в соответствующем дереве Фишера-Ладнера интуитивно определяются как наборы подформул φ , такие что каждая подформула выполняется в соответствующем узле. Поскольку арифметические формулы подсчитывают дочерние узлы, мы вводим в дочерние узлы новые высказывания, чтобы отслеживать число узлов. Мы называем эти предложения счетчиками, и они кодируются в двоичной системе. Это кодирование основано на булевой комбинации высказываний, встречающихся без отрицаний. Затем мы определяем $(C(\varphi = b) := \overline{\psi^i}$, где $i \in \{0, \dots, [\log(b)]\}$, $\overline{\psi^i}$ – последовательность высказываний, а b – целое число.

Определение 1 (дерево Фишера-Ладнера).

Дерево Фишера-Ладнера T определяется как

- пустое дерево или
- (n_r, T_1, T_2) , где n_r – корневой узел, а T_1 или T_2 являются первым потомком и последующими поддеревом соответственно.

Пример 1: Рассмотрим формулу: $\varphi = a \wedge (b - c > 0) \wedge \langle 1 \rangle \mu X. (a \vee \langle 2 \rangle X)$, и соответствующий набор узлов N^φ . Тогда $T = (n_1, (n_2, \emptyset, (n_3, (n_4, \emptyset, \emptyset)), \emptyset)$, где

- $n_1 = \{a, (b - c > 0), \langle 1 \rangle \mu X. (a \vee \langle 2 \rangle X), \langle 1 \rangle \top, \langle 1 \rangle \mu X. (b \vee \langle 2 \rangle X), \langle 1 \rangle \mu X. (c \vee \langle 2 \rangle X)\}$;
- $n_2 = \{b, C(b) = 2, C(c) = 1, \langle 2 \rangle \mu X. (a \vee \langle 2 \rangle X), \langle 2 \rangle \top, \langle 2 \rangle \mu X. (a \vee \langle 2 \rangle X)\}$;
- $n_3 = \{c, C(b) = 1, C(c) = 1, \langle 2 \rangle \mu X. (a \vee \langle 2 \rangle X), \langle 2 \rangle \top, \langle 2 \rangle \mu X. (b \vee \langle 2 \rangle X)\}$;
- $n_4 = \{a, b, C(b) = 1, C(c) = 0, \langle 4 \rangle \top\}$;
- $n_5 = \{a, C(b) = 0, C(c) = 0, \langle 4 \rangle \top\}$;
- $n_6 = \{p', C(b) = 0, C(c) = 0, \langle 4 \rangle \top\}$;
- $n_7 = \{c, C(b) = 0, C(c) = 1, \langle 4 \rangle \top\}$.

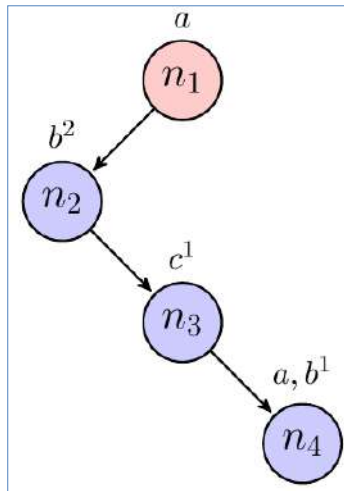


Рис. 1. Дерево Фишера-Ладнера для $a \wedge (b - c > 0) \wedge \langle 1 \rangle \mu X. (a \vee \langle 2 \rangle X)$
 Fig. 1. A Fischer-Ladner tree for $a \wedge (b - c > 0) \wedge \langle 1 \rangle \mu X. (a \vee \langle 2 \rangle X)$

Множество узлов N^Φ определяется исходя из множества подформул, называемого *lean* (φ) и содержащего высказывания, модальные формулы и счетчики. На рис. 1 изображено графическое представление соответствующего дерева Фишера-Ладнера. В корневом узле n_1 счетчики установлены в 0. Это связано с тем, что счетчики появляются только на узлах-братьях: n_2, n_3 и n_4 . Счетная формула выполняется в корневом узле n_1 потому что дочерние узлы n_1 удовлетворяют ограничениям Пресбургера.

Теперь определим алгоритм проверки выполнимости. Для входной формулы φ алгоритм строит восходящим образом дерево Фишера-Ладнера, выполняющее φ , если такое дерево существует. Алгоритм выполнимости определен на рис. 2. Функция *Init* вычисляет листья дерева. Между узлами Фишера-Ладнера и формулами можно определить отношение выполнимости. Функция *Update* последовательно добавляет родителей к ранее построенным деревьям.

```


$$Y \leftarrow N_i^\Phi$$


$$X \leftarrow \text{Init}(\Phi)$$


$$X' \leftarrow \emptyset$$

while  $N_i^\Phi \not\models \Phi$  and  $X = N^\Phi$  do
     $X' \leftarrow X$ 
     $X \leftarrow \text{Update}(X', Y)$ 
     $Y \leftarrow Y \setminus \text{root}(X)$ 
end while
if  $X \models \Phi$  then
    return  $\Phi$  is satisfiable
end if
return  $\Phi$  is not satisfiable

```

Рис. 2. Алгоритм определения выполнимости для μ -исчисления с арифметическими ограничениями

Fig. 2. Satisfiability algorithm for the μ -calculus with arithmetic constraints

Теорема 1 (Корректность [3])

Для заданной входной формулы φ алгоритм устанавливает выполнимость φ тогда и только тогда, когда существует древовидная структура T , такая что $\llbracket \varphi \rrbracket_V^T \neq \emptyset$, для любой валиуации V .

4. Реализация

Для реализации алгоритма мы используем битово-векторное представление узлов Фишера-Ладнера. Поскольку узлы определяются как подмножества *lean*, формулы в узлах в векторном представлении битовым единицам, а битовые нули соответствуют отсутствующим формулам. Например, предположим, что *lean* определяется следующим образом: $\text{lean}(\varphi) = \{\psi_1, \psi_2, \dots, \psi_n\}$. Если узел $n \subseteq \text{lean}(\varphi)$ содержит следующие формулы $n = \{\psi_{i_1}, \psi_{i_2}, \dots, \psi_{i_k}\}$, то битово-векторное представление n содержит единицы в битах $i_1, i_2, \dots, i_k$ и нули во всех остальных битах. Это битово-векторное кодирование узлов дерева Фишера-Ладнера подразумевает битово-векторное кодирование всех функций алгоритма.

Алгоритм выполнимости был реализован на языке Java с использованием технологии JavaService Faces. Для тестирования этой реализации мы использовали компьютер со характеристиками: операционная система Windows 10, процессор Intel i7-6700HQ 2,6 ГГц, 16 Гб оперативной памяти.

Таблица 1. Выполнимые формулы
Table 1. Satisfiable Formulas

Формула	N^φ	<i>Lean</i>	Время (мс)
a	36	6	22
$\neg a$	35	5	11
$a \vee b$	84	7	20
$a \wedge b$	84	7	124
$a \vee b \vee c \vee d \vee e \vee f$	756	10	30
$a \wedge b \wedge c \wedge d \wedge e \wedge f$	756	10	20
$\langle 1 \rangle a$	72	7	523
$\langle 1 \rangle a \wedge b$	168	8	2199
$\langle 1 \rangle \langle 1 \rangle > \top$	24	6	24814
$\langle 1 \rangle \langle 2 \rangle a$	144	8	100907
$1 * a > 0$	198	10	12333
$(1 * a > 0) \wedge c$	462	11	9189
$(a > 0) \wedge c \wedge \langle 1 \rangle b$	2070	13	9189
$1 * a > 1$	3258	10	179309
$1 * a + 1 * b > 0$	966	13	1560
$1 * a + 1 * b > 2$	966	19	20359493
$a \leq 0$	198	10	120.0

В табл. 1 и 2 приводится количество узлов N^φ , сгенерированных при выполнении алгоритма, размер *lean* и время выполнения алгоритма в миллисекундах. Табл. 1 демонстрирует лучшую производительность для булевых формул, например $(a \wedge b)$. Для модальных формул, например, $\langle 1 \rangle \varphi$ соответствующее время выполнения меньше 1 секунды, но по мере увеличения числа модальностей время выполнения также увеличивается (см., например, формулу $\langle 1 \rangle \langle 1 \rangle a$). В случае формул с арифметическими ограничениями эти ограничения проверяются в деревьях на более низком уровне, то есть ограничиваются дочерние узлы. Требуется построение более глубоких деревьев Фишера-Ладнера, что вызывает увеличение времени выполнения. Алгоритмы выполнимости, основанные на явном построении моделей, такие как деревья Фишера-Ладнера, обычно работают плохо, потому что при наличии небулевых противоречивых условий, как правило, требуется построение всех возможных древовидных моделей. Это видно из табл. 2.

Табл. 2. Невыполнимые формулы
Table 2. Unsatisfiable Formulas

Формула	N^{φ}	Lean	Время (мс)
$\neg$	12	5	3631
$\neg a \wedge a$	36	5	204811
$(\neg a \wedge a) \vee (\neg b \wedge b)$	36	7	94799
$(a \wedge \neg a) \vee (b \wedge \neg b) \vee (c \wedge \neg c)$	45	8	433768
$\langle 1 \rangle (\neg a \wedge a)$	72	7	1129042
$(1 * a > 0) \wedge (1 * a \leq 0)$	19810	10	12812479

5. Логические рассуждения о запросах

Язык путевых выражений XML (XML Path Language, XPath) – это язык запросов для полуструктурированных данных, стандартизованный W3C [16]. В этом разделе мы покажем представление навигационного ядра XPath, расширенного арифметическими ограничениями, в терминах формул μ -исчисления. Вот синтаксис XPath с арифметическими ограничениями.

$$\begin{aligned}
 P &::= \alpha : p \mid P/P \mid P[Q], \\
 Q &::= A > b \mid P \mid Q \vee Q \mid \neg Q, \\
 A &::= children : p \mid A + A,
 \end{aligned}$$

где α – направление навигации – непосредственный потомок, непосредственный предок, следующий брат, предыдущий брат, потомок или предок; b – положительное целое число. Формальная семантика этой версии XPath с арифметическими ограничениями представлена в работе [3].

Выражения XPath могут быть записаны в терминах формул μ -исчисления. Например $children : p_1[children : p_2 > 5]$ можно записать следующим образом: $(p \wedge \langle 3 \rangle >) \wedge (p_2 > 5)$. В работе [3] описаны правила преобразования выражений XPath в формулы μ -исчисления.

К числу общих проблем области логических рассуждений о запросах относятся проблема пустоты результатов запроса (emptiness), проблема включения результатов одного запроса во множество результатов другого запроса (containment) и проблема эквивалентности запросов (equivalence). Проблема пустоты заключается в том, что требуется определить, является ли данный запрос пустым в какой-либо базе данных (модели). Запрос P_1 включается в запрос P_2 тогда и только тогда, когда множество результатов P_1 содержится в множестве результатов P_2 в любой модели. Два запроса эквивалентны, если каждый из них включает другой запрос. Решение этих проблем для языка XPath на основе проверки выполнимости формул μ -исчисления обеспечивает следующая теорема.

Теорема 2 ([3]).

Для заданных выражений XPath P , P_1 и P_2

- P пусто тогда и только тогда, когда $F(P, T)$ выполнима и
- P_1 содержится в P_2 тогда и только тогда, когда $F(P_1, T) \wedge \neg F(P_2, T)$ невыполнимо.

Здесь $F(P, T)$ – формула μ -исчисления, представляющая запрос P .

В табл. 3 и 4 приведены результаты применения алгоритма выполнимости для обнаружения свойств запросов XPath.

Табл. 3. Непустые запросы XPath
Table 3. Non-empty XPath queries

Формула	N^{φ}	<i>Lean</i>	Время (мс)
a	8	144	6
$\top$	7	48	7
$\downarrow: a$	9	240	1667
$\uparrow: a$	10	406	1048
$\uparrow: *$	9	192	147
$\downarrow: *$	8	80	212
$\downarrow: *$	9	40	335
$\downarrow: a$	10	336	784
$\downarrow: a \cup \downarrow: b$	11	1008	8968
$a/\downarrow: b$	8	1008	7
$a[\downarrow: b]$	11	144	6158
$a[\downarrow: *]$	10	1344	1122
$a[\downarrow: b > 0]$	13	1932	490488
$a[\downarrow: b + \downarrow: c > 0]$	16	8460	77116123
$\downarrow a/\downarrow: a$	12	1232	13143385

Табл. 4. Пустые запросы XPath
Table 4. Empty XPath queries

Формула	N^{φ}	<i>Lean</i>	Время (мс)
$a[\neg a]$	8	144	343387
$a[\neg \top]$	8	144	390785
$a[\neg \top]$	9	240	2368813
$a[\neg a]$	9	204	1815198
$\top[\neg \downarrow: a] \cap \top[\downarrow: a]$	8	144	336591
$a[\downarrow: b \wedge \neg \downarrow: b]$	9	336	4206084

5. Выводы

В этой работе мы описали, как и в [3], алгоритм выполнимости для μ -исчисления, расширенного обратными модальностями, номиналами и арифметическими ограничениями Пресбургера. Алгоритм основан на построении моделей Фишера-Ладнера восходящим образом.

Известно, что вычислительная сложность μ -исчисления относится к категории EXPTIME, поэтому непосредственная реализация невозможна. Мы предложили битово-векторное представление узлов в деревьях Фишера-Ладнера. Это позволило справиться с экспоненциальным количеством узлов в памяти. Насколько нам известно, это первая реализация алгоритма принятия решений для μ -исчисления с арифметическими ограничениями.

Также мы сообщили о нескольких экспериментах. Удалось проверить некоторые формулы даже с арифметическими ограничениями. Однако сочетание неподвижных точек, формул с глубокой модальностью и арифметических ограничений по-прежнему выходит за рамки наших возможностей. В перспективе нас интересуют более эффективные конструкции деревьев Фишера-Ладнера [13].

В этой статье мы также говорим о представлении запросов XPath формулами μ -исчисления. Это позволяет решать проблемы логических рассуждений об XPath, такие как пустота, включение и эквивалентность, на основе алгоритма выполнимости μ -исчисления. В [3] мы показали, что фрагмент языка XPath с арифметическими ограничениями является EXPTIME-разрешимым, и нам неизвестен никакой другой фреймворк логических рассуждений для этого фрагмента языка XPath.

Мы также описали некоторые эксперименты с определением пустоты XPath-запросов.

Мы будем продолжать работать над формализацией XML и SXML-схем в терминах μ -исчисления, чтобы гарантировать согласованность этих спецификаций, поскольку из-за количества элементов в них может иметься некоторая несогласованность [17]. Кроме того, мы будем использовать инструменты машинного обучения в промежуточном процессе оптимизации в алгоритме, представленном в этой статье, чтобы использовать его для решения проблем на основе логического вывода в контекстно-зависимых системах [18]. Авторы [19] представили методы машинного обучения и продемонстрировали хорошие результаты. С помощью машинного обучения они проанализировали информацию в Интернете с целью выявления возможных терактов. Иные интересы исследователей касаются верификации контекстно-зависимых систем. В частности, сложной задачей оказались моделирование и верификация временных и количественных свойств этих систем. Мы считаем, что μ -исчисление с арифметическими ограничениями можно использовать в качестве языка спецификации. Соответствующий эффективный алгоритм выполнимости позволил бы осуществить практическую верификацию контекстно-зависимых систем.

Список литературы / References

- [1]. P. A. Bonatti, C. Lutz, A. Murano, and M. Y. Vardi. The complexity of enriched mu-calculi. *Logical Methods in Computer Science*, vol. 4, no. 3, 2008, pp. 1-27.
- [2]. E. Bárcenas and J. Lavalley. Global numerical constraints on trees. *Logical Methods in Computer Science*, vol. 10, no. 2, 2014, pp. 1-28.
- [3]. E. Bárcenas, E. Benítez-Guerrero, J. Lavalley, and G. Molero-Castillo. Presburger constraints on tree. *Computación y Sistemas*, vol. 24, no. 1, 2020, pp. 281-303.
- [4]. Y. Limón, E. Bárcenas, E. Benítez-Guerrero, G. Molero-Castillo, and A. Velázquez-Mena. A satisfiability algorithm for the mu-calculus for trees with presburger constraint. In *Proc. of the 7th International Conference in Software Engineering Research and Innovation (CONISOFT)*, 2019, pp. 72-79.
- [5]. S. Demri and D. Lugiez. Complexity of modal logics with Presburger constraints. *Journal of Applied Logic*, vol. 8, no. 3, 2010, pp. 233-252.
- [6]. H. Seidl, T. Schwentick, and A. Muscholl. Counting in trees. In *Logic and Automata: History and Perspectives [in Honor of Wolfgang Thomas]*, vol. 2. Amsterdam University Press, 2008, pp. 575-612.
- [7]. H. Seidl, T. Schwentick, and A. Muscholl. Numerical document queries. In *Proc. of the 22nd ACM SIGACT-SIGMOD-SIGART Symposium on Principles of Database Systems*, 2003, pp. 155-166.
- [8]. L. de Moura and N. Bjørner. Z3: An efficient SMT solver. *Lecture Notes in Computer Science*, vol. 4963, 2008, pp. 337-340.

- [9]. K. Hoder, N. Bjørner, and L. de Moura. μz – an efficient engine for fixed points with constraints. *Lecture Notes in Computer Science*, vol. 6806, 2011, pp. 457-462.
- [10]. W. Conradie, Y. Fomatati, A. Palmigiano, and S. Sourabh. Algorithmic correspondence for intuitionistic modal mu-calculus. *Theoretical Computer Science*, vol. 564, 2015, pp. 30-62.
- [11]. M. Biehl, N. Klarlund, and T. Rauhe. Mona: Decidable arithmetic in practice. *Lecture Notes in Computer Science*, vol. 1135, 1996, pp. 459-462.
- [12]. P. Genevès, N. Layaïda, and A. Schmitt. Efficient static analysis of XML paths and types. In *Proc. of the ACM SIGPLAN Conference on Programming Language Design and Implementation*, 2007, pp. 342-351.
- [13]. Y. Limón, E. Bárcenas, E. Benítez-Guerrero, and M. A. M. Nieto. Depth-first reasoning on trees. *Computación y Sistemas*, vol. 22, no. 1, 2018, pp. 189-201.
- [14]. M. Junedi, P. Genevès, and N. Layaïda. XML query-update independence analysis revisited. In *Proc. of the ACM Symposium on Document Engineering*, 2012, pp. 95-98.
- [15]. M. J. Fischer and R. E. Ladner. Propositional modal logic of programs. In *Proc. of the ACM Symposium on Theory of Computing*, 1977, pp. 286-294.
- [16]. The World Wide Web Consortium (W3C), XPath 2.0 W3C Recommendation, 2010.
- [17]. К.Ю. Лисовский. Разработка XML-приложений на языке Scheme. Программирование, том 28, no. 4, 2002 г., стр. 20-32 / K.Yu. Lisovsky. XML applications development in Scheme. *Programming and Computer Software*, vol. 28, no. 4, 2002, pp. 197-206.
- [18]. Y. Limón, E. Bárcenas, E. Benítez-Guerrero, and G. Molero. On the consistency of context-aware systems. *Journal of Intelligent and Fuzzy Systems*, vol. 34, no. 5, 2018, pp. 3373-3383.
- [19]. И.В. Машечкин, М.И. Петровский, Д.В. Царев, М.Н. Чикунов. Методы машинного обучения для задачи обнаружения и мониторинга экстремистской информации в сети Интернет. Программирование, том 45, no. 3, 2019 г., стр. 18-37 / I.V. Mashechkin, M.I. Petrovskiy, D.V. Tsarev, and M.N. Chikunov. Machine learning methods for detecting and monitoring extremist information on the Internet. *Programming and Computer Software*, vol. 45, no. 3, 2019, pp. 99-115.

Информация об авторах / Information about authors

Йенсен ЛИМОН-ПРИЕГО – кандидат наук. Научные интересы: модальные логики, логический вывод.

Yensen LIMÓN-PRIEGO, Ph. D. Research interests: Modal Logic, Automated Reasoning.

Исмаэль Эверардо БАРСЕНАС-ПАТИНЬО – кандидат наук, доцент. Научные интересы: искусственный интеллект, логический вывод, формальные методы.

Ismael Everardo BÁRCENAS-PATÍÑO, Ph. D. in Computer Science, Assistant Professor. Research interests: Artificial Intelligence, Automated Reasoning, Formal Methods.

Эдгард Иван БЕНЕТЕС-ГЕРРЕРО – кандидат наук, профессор. Область научных интересов: искусственный интеллект, взаимодействие человека с компьютером, коллективная работа с использованием компьютеров, системы баз данных.

Edgard Iván BENÍTEZ-GUERRERO, Ph. D. in Computer Science, Full-time Professor. Research interests: Artificial intelligence, human computer interaction, Computer-Supported Cooperative Work, database systems.

Гильермо Хильберто МОЛЕРО-КАСТИЛЬО – кандидат наук, доцент. Научные интересы: искусственный интеллект, наука о данных, интеллектуальный анализ данных, машинное обучение.

Guillermo Gilberto MOLERO-CASTILLO, Ph.D. in Information Technologies, Associate Professor. Research interests: Artificial Intelligence, Data Science, Data Mining, Machine Learning.

Алехандро БЕЛАСКЕС-МЕНА – магистр, доцент. Область научных интересов: сетевые вычисления, туманные вычисления, распределенные вычисления.

Alejandro VELAZQUEZ-MENA, Master of Science, Associate Professor. Research interests: Network Computing, Fog Computing, Distributed Computing.