

Министерство образования Республики Беларусь
Учреждение образования
«Белорусский государственный университет
информатики и радиоэлектроники»

КОДИРОВАНИЕ И ЦИФРОВАЯ ОБРАБОТКА СИГНАЛОВ В ИНФОКОММУНИКАЦИЯХ

МАТЕРИАЛЫ МЕЖДУНАРОДНОЙ НАУЧНО-ПРАКТИЧЕСКОЙ КОНФЕРЕНЦИИ
(Минск, 19 апреля 2021 г.)

Минск, 2021

УДК 621.391
ББК 32.811.4
К57

Редакционная коллегия:
В. К. Конопелько (руководитель конференции),
В. Ю. Цветков, Л. А. Шичко

Кодирование и цифровая обработка сигналов в инфо-
коммуникациях: материалы междунар. науч.-практ. конф. (Республика
Беларусь, Минск, 19 апреля 2021 года) / редкол.: В. К. Конопелько,
В. Ю. Цветков, Л. А. Шичко – Минск : БГУИР, 2021. – 80 с.: ил.
ISBN 978-985-543-616-5.

Сборник содержит статьи, тематика которых посвящена научно-теоретическим и практическим разработкам в области цифровой обработки сигналов, теории кодирования, обработки и передачи изображений, защиты инфокоммуникационных систем и сетей.

Предназначен для научных сотрудников в области телекоммуникаций, преподавателей, аспирантов, магистрантов и студентов технических вузов.

Научное издание

Корректор *В. В. Чепикова*
Ответственный за выпуск *В. Ю. Цветков*
Компьютерный дизайн и верстка *Е. Г. Макейчик*
Подписано в печать 13.04.2021. Формат 60×84 1/8. Бумага офсетная. Гарнитура «Таймс».
Отпечатано на ризографе. Усл. печ. л. 9,53. Уч.-изд. л. 7,6. Тираж 56 экз. Заказ 43.

Издатель и полиграфическое исполнение: учреждение образования
«Белорусский государственный университет информатики и радиоэлектроники»
Свидетельство о государственной регистрации издателя, изготовителя,
распространителя печатных изданий №1/238 от 24.03.2014,
№2/113 от 07.04.2014. №3/615 от 07.04.2014.
220013, Минск, П. Бровки, 6

ISBN 978-985-543-616-5

© УО «Белорусский государственный университет информатики
и радиоэлектроники», 2021

СОДЕРЖАНИЕ

Ma Jun, V.Yu. Tsviatkou, V.K. Konopelko Review on mainstream method of skeleton pruning	5
П.С. Маринич, А.А. Борискевич Выбор архитектуры нейронной сети для задачи распознавания физической активности.....	9
Е.К. Вашкевич, И.А. Борискевич Сравнительный анализ алгоритмов автоматического обобщения текста.....	13
Д.А. Качан, В.А. Вишняков Поддержка информационного управления в образовании с использованием блокчейн	19
С.Б. Саломатин, А.Э. Алексеенко, А.П. Турлай Сетевое кодирование кодом Рида-Соломона с учетом лидеров стохастической сети	23
А.В. Захаренко, О.Г. Шевчук, В.Ю. Цветков Оценка искажений изображений при субпиксельном сдвиге камеры	28
О.А. Азарчик, Е.А. Машков, Н.С. Давыдова Алгоритмы обработки изображений в комбинированных оптико-электронных системах защиты объектов от беспилотных летательных аппаратов.....	35
А.И. Митюхин Защита информации на основе спектрально-пространственного кодирования	40
Ren Xunhuan, V.K. Konopelko, V.Yu. Tsviatkou Method for generating two-dimensional dependent errors.....	44
Э.Б. Липкович, Е.А. Белоконь Помехоустойчивость систем связи с каскадным кодированием и многопозиционной модуляцией	47
М.А. Алисеенко, С.Б. Саломатин Оценка сложности алгоритма обучения с ошибками алгебраических решетчатых кодов	52
В.А. Аксенов, С.В. Смоляк Прогнозирование экстремального всплеска в сигнале с OFDM.....	55
Е.А. Машков, О.А. Азарчик, А.В. Тютрюмова Кодирование цифрового видеопотока в системах видеонаблюдения.....	60
U.A. Vishniakou, A.H. Al-Masri., S.Kh. Al-Hajj Modeling a network of things on the basis of a cloud platform	65
А.Ю. Ключкий, С.К. Лазарук Исследование оптических излучателей на основе наноструктурированного кремния	69
Н.Н. Сергеев Межканальные перекрестные помехи в восходящем направлении в сетях следующего поколения NG-PON2	73

CONTENTS

Ma Jun, V.Yu. Tsviatkou, V.K. Konopelko	
Review on mainstream method of skeleton pruning.....	5
P.S. Marynich, A.A. Boriskevich	
Selecting neural network architecture for task of physical activity recognizing	9
E.K. Vashkevich, I.A. Baryskievic	
Comparative analysis of text summarization algorithms in English language	13
D.A. Kachan, U.A. Vishnyakou	
The support of information control in education using blockchain.....	19
S.B. Salomatin, A.E. Alekseenko, A.P. Turlay	
Network coding by Reed-Solomon code with stochastic network leading	23
A.V. Zakharenko, O.G. Shevchuk, V.Yu. Tsviatkou	
Estimation of image distortions with sub-pixel shift of the camera	28
O.A. Azarchik, E.A. Mashkov, N.S. Davydova	
Image processing algorithms in combined optoelectronic systems for protecting objects from unmanned aerial vehicles.....	35
A.I. Mitsiukhin	
Protection of information based on spectral-spatial coding.....	40
Ren Xunhuan, V.K. Konopelko, V.Yu. Tsviatkou	
Method for generating two-dimensional dependent errors.....	44
E.B. Lipkovich, E.A. Belakon	
Immunity of communication systems with cascade coding and multi-position modulation.....	47
M.A. Aliseyenka, S.B. Salomatin	
Complexity estimation of the learning with errors algorithm for algebraic lattice codes.....	52
V.A. Aksyonov, S.V. Smolyak	
Extremal pulse forecasting in the signal with OFDM	55
E.A. Mashkov, O.A. Azarchik, A.V. Tyutryumova	
Digital video stream encoding in video surveillance systems	60
U.A. Vishniakou, A.H. Al-Masri., S.Kh. Al-Hajj	
Modeling a network of things on the basis of a cloud platform	65
A.Yu. Kliutski, S.K. Lazarouk	
Study of optical light emitters based on porous silicon	69
N.N. Sergeev	
Uplink cross-channel crosstalk in next-generation networks NG-PON2	73

REVIEW ON MAINSTREAM METHOD OF SKELETON PRUNING

MA JUN, V.Yu. TSVIATKOU, V.K. KONOPELKO

*Belarusian State University of Informatics and Radioelectronics, Republic of Belarus**Submitted 9 March 2021*

Abstract. The technique of skeleton pruning was generally deployed after the skeletonization of the binary image for filtering unwanted branches caused by the boundary noise, which was a vital pre-processing method before analysis and recognition of the skeleton. It was still a challenging task since there were not standard measurements for distinct the noise branches and original structural branches. In the past decade, there were many approaches based on different perspective have emerged for trying tackle that problem. The review of the most 6 cited paper was conducted in this paper.

Keyword: 2D skeleton, skeleton branch, skeleton pruning.

Introduction

The skeleton, sometimes also named as the medial axis in some literatures, is the result of the skeletonization of the binary image. The foundation of the skeletonization is established by the Blum [1] in 1967. The skeleton is very useful in the field of recognition and object representation because it is a compact abstraction of the visual shape. An ideal skeleton of an object is expected to be have the following properties: it should be a thin subset of the object, consisting of the union of curves; it should be symmetrically placed within the object, it should maintain the same topological (connectivity) [2] as the original shape. However, skeleton has an inherent defect that it is very sensitive to the boundary noise. Slight noise or small perturbations in the boundary could generate redundant skeleton branches that may dramatically influence the topology of the skeleton graph (Fig. 1), which may increase the difficulty of the processing and analyzing of skeletons. Therefore, it is necessary to remove those unwanted branches by using skeleton pruning method.

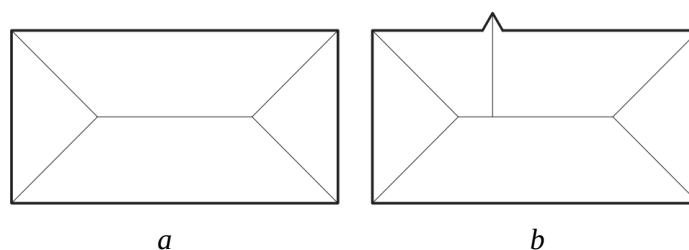


Fig. 1. An example of the instability of skeletons:

a – clean skeleton resulted from normal case;

b – skeleton resulted from the pattern with small perturbations

Mainstream method of skeleton pruning

1. Skeleton Pruning by contour partitioning with DCE.

Bai and his team [6] proposed a skeleton pruning algorithm combining one of the contour partitioning method that is Discrete Curve Evolution (DCE) [3–5]. They proved that the obtained skeletons generated by their method are in accord with human visual perception and stable, even in the presence of significant noise and shape variations, and have the same topology as the original skeletons from the perspective of the theoretical properties and the experiments. The main idea of them is to remove all skeleton points whose generation points all lie on the same contour segment. The generation points are referring to those boundary points that are tangential to the maximal circle of a certain skeleton point. A hierarchical skeleton structure obtained by their approach is illustrated in Fig. 2, where the (red)

bounding polygons represent the contours simplified by DCE. Because DCE can reduce the boundary noise without displacing the remaining boundary points, the accuracy of the skeleton position is guaranteed. The continuity, which implies stability in the presence of noise, of their pruning methods follows from the continuity of the DCE.

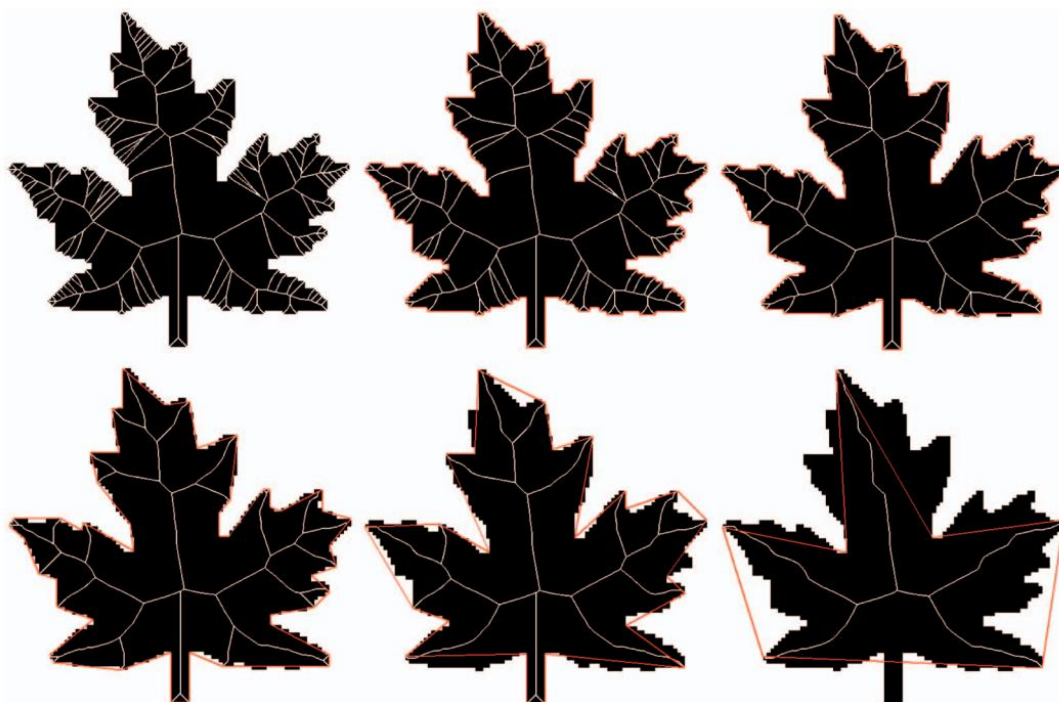


Fig. 2. Hierarchical skeleton of leaf obtained by pruning the input skeleton

The groupwise Medial Axis Transform for fuzzy skeletonization and pruning.

Aaron D. Ward proposed their pruning framework [7] that yields a fuzzy significance measure for each branch, derived from information provided by the group of shapes. They used geometric and topological branch feature information provided by the group as a whole. This groupwise approach represents a third paradigm in computation of branch pruning order. Their approach is more effective in terms of removal of noisy skeleton branches, classification of noisy artificial and real skeletons into different object classes and in producing skeletons that are similar for similar objects.

Skeleton growing and pruning with bending potential ratio.

Pruning skeleton based on bending potential ratio (BPR) is proposed by Wei Shen and his co-workers [8], in which the decision regarding whether a skeletal branch should be pruned or not is based on the context of the boundary segment that corresponds to the branch. The BPR is a measure of the significance of contour segments and depicts the bending potential of a contour segment. Unlike other significance measures that only contain local shape information, the BPR evaluates both local and global shape information. Thus, it is insensitive to local boundary deformation. The skeleton obtained by this method are medially placed, insensitive to boundary noise, multi-scale and provide intuitive ordering of skeleton branches in that negligible skeleton branches are pruned while significant branches remain.

A skeleton pruning algorithm based on information fusion.

HongZhi Liu and his group developed a pruning algorithm [10] based on their previous works [9]. In their opinion, the relative significance of the same branches will be different if see them from different perspectives with different objectives. Different objective measurements have their advantages and limitations. To integrate the advantages of different objective measurements, they consider skeleton pruning as a multi-objective decision-making problem and propose a skeleton pruning algorithm based on information fusion. Additionally, they divided the pruning procedure into coarse-pruning and fine-grain pruning procedure, in which they used combinatorial fusion analysis and the concept of cognitive diversity to fuse various measurements of branch significance including region reconstruction, contour reconstruction and visual contribution. A butterfly example produced by their algorithm is shown in Fig. 3. In general, their algorithm succeeds in avoiding removing any significant branches in advance. not

only efficiently delete those redundant branches due to noise, but also be able to generate multi-scale skeletons according with visual judgement

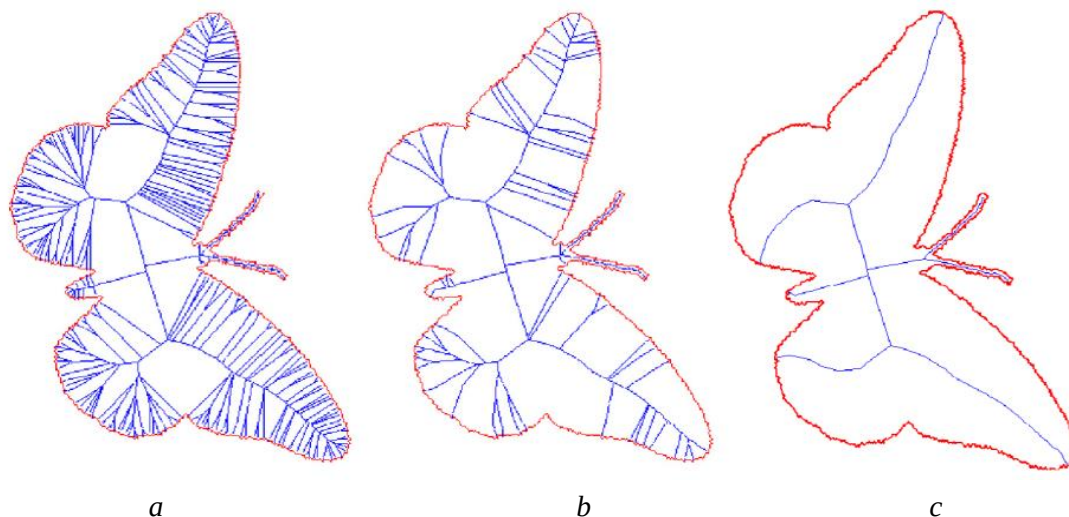


Fig. 3. A butterfly example: *a* – is the raw skeleton; *b* – is the skeleton after coarse-pruning; *c* – the result after fine-pruning

A new strategy for skeleton pruning.

Luca Serino [11] has proposed a new skeleton pruning algorithm. They building the skeleton hierarchy and remove the redundant branches by referring the significance measures that take into account the loss in object recovery caused by pruning. The topology of obtained skeleton yield by this method is maintained and the pruning does not dramatically decrease its ability in object recovery.

A skeleton pruning method based on saliency sorting.

To overcome the drawbacks of that the quality of pruning is hard to control because find a suitable threshold is a tough task and of that when the objects of skeletonization vary significantly, it is not necessarily the case that a threshold suitable for most objects even exists, Guo Siyu [12] and his co-workers have proposed a new skeleton pruning algorithm based on saliency sorting. Their method decomposes a skeleton into a number of skeletal components (SCs), and the terminal SCs are removed one by one according to a saliency measure. SCs may be merged into a new one after each removal. The removal process continues until the desirable number of terminals are achieved. Their approach can effectively prune the skeletons from the test image bases, and it has the advantage that the terminal number parameter is very intuitive and easy to set for the applications under concern.

Conclusion

We have reviewed a several popular methods of skeleton pruning. Almost of them are based on the salience measures of the significance of contour segment or skeleton branches to identify the noise branch and then the skeleton pruning can be seemed as the task of seeking a suitable threshold, which procedure need the interaction with the people for tuning. It will be more convenient if there is a skeleton pruning algorithm without any manual pruning. In the future, we will try to develop such an algorithm.

References

1. Blum H., Nagel R.N. Shape description using weighted symmetric axis features. // Pattern Recognition. 1978. 10. P. 167–180.
2. Saha P.K., Borgefors G., Sanniti di Baja G. A survey on skeletonization algorithms and their applications // Pattern Recognition Letters. 2016. 76. P. 3–12.
3. Latecki L.J., Lakämper R. Convexity Rule for Shape Decomposition Based on Discrete Contour Evolution // Computer Vision and Image Understanding. 1999. 73. P. 441–454.
4. Latecki L.J., Lakämper R. Polygon evolution by vertex deletion // Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics). 1999. 1682. P. 398–409.
5. Latecki L.J., Lakamper R. Shape Similarity Measure Based on Correspondence of Visual Parts // IEEE Transactions on Pattern Analysis and Machine Intelligence. 2000. 22. P. 1185–1190.

6. Bai X., Latecki L.J., Liu W.Y. Skeleton pruning by contour partitioning with discrete curve evolution // IEEE Transactions on Pattern Analysis and Machine Intelligence. 2007. 29. P. 449–462.
7. Ward A.D., Hamarneh G. The groupwise medial axis transform for fuzzy skeletonization and pruning // IEEE Transactions on Pattern Analysis and Machine Intelligence. 2010. 32. P. 1084–1096.
8. Shen W., Bai X., Hu R., [et. al] Skeleton growing and pruning with bending potential ratio // Pattern Recognition. 2011. 44. P. 196–209.
9. Liu H., Wu Z., Hsu D.F., [et. al] On the generation and pruning of skeletons using generalized Voronoi diagrams // Pattern Recognition Letters. 2012. 33. P. 2113–2119.
10. Liu H., Wu Z.H., Zhang X., [et. al] A skeleton pruning algorithm based on information fusion // Pattern Recognition Letters. 2013. 34. P. 1138–1145.
11. Serino L., Sanniti di Baja G. A new strategy for skeleton pruning // Pattern Recognition Letters. 2016. 76. P. 41–48.
12. Siyu G., Pingping H., Zhigang L., [et. al] A skeleton pruning method based on saliency sorting // 14th IEEE International Conference on Electronic Measurement and Instruments. 2019. P. 593–599.

УДК 621.391

ВЫБОР АРХИТЕКТУРЫ НЕЙРОННОЙ СЕТИ ДЛЯ ЗАДАЧИ РАСПОЗНАВАНИЯ ФИЗИЧЕСКОЙ АКТИВНОСТИ

П.С. МАРИНИЧ, А.А. БОРИСКЕВИЧ

*Белорусский государственный университет информатики и радиоэлектроники, Республика Беларусь**Поступила в редакцию 8 марта 2021*

Аннотация. Выбрана оптимальная архитектура нейронной сети для распознавания физической активности. Описана архитектура работы сверточной нейронной сети. Обоснован выбор среды разработки.

Ключевые слова: нейронные сети, изображение.

Введение

Одним из основных навыков, которыми человек овладевает в первые годы жизни, является распознавание образов, которые предстают перед ним. Эволюция заложила в человека эту функцию для моментального распознавания уже знакомых образов, будь то убежище, пища или угроза. Чем больше мы сталкиваемся с тем или иным явлением или объектом, тем быстрее мы сможем охарактеризовать их, запомнить определенные паттерны и дать соответствующие ярлыки происходящему вокруг. Однако мир становится все сложнее, объемы информации увеличиваются, а задачи, которые ставит перед собой человечество – (сложнее), и, следовательно, требуются все более точные системы, которые не только помогают естественным чувствам, но и зачастую их заменяют.

Выбор и описание архитектуры

Наилучшей архитектурой искусственных нейронных сетей, направленных на распознавание образов на изображениях, является архитектура сверточных нейронных сетей. Она использует некоторые особенности работы человеческого мозга, в частности, зрительной коры. В 1962 году в ходе экспериментов группы американских нейрофизиологов выяснилось, что отдельные мозговые нервные клетки реагировали (или активировались) только при визуальном восприятии границ определенной ориентации. Например, некоторые нейроны активировались, когда воспринимали вертикальные границы, а некоторые – горизонтальные или диагональные. Нейроны, ответственные за это, сосредоточены в виде стержневой архитектуры и вместе формируют визуальное восприятие. Самый простой пример из повседневной жизни – выявление отличительных признаков и представление на его основе выводов об объекте: мы видим две круглые детали внизу объекта – это велосипед. Эта идея легла в основу работы французского ученого Яна Лекуна, который в 1988 году и представил архитектуру сверточных нейронных сетей.

На вход сверточной нейронной сети подается матрица чисел (если используется изображение в оттенках серого), которые характеризуют каждый пиксель изображения, либо набор матриц (изображение, разделенное на три слоя – R-, G- и B-каналы).

За распознавание определенных характеристик изображения отвечает фильтр – также матрица, но уже с искомыми значениями.

Выбор архитектуры сверточной нейронной сети для распознавания активности должен определяться сложностью обработки слоев, их количества, а также скоростью получения результата, поскольку на вход системы поступает видеопоток, состоящий из отдельных изображений, с частотой порядка 20–30 кадров в секунду. Исходя из этого следует выбрать оптимальный в количестве и скорости выполняемых шагов алгоритм с минимальными потерями в качестве результата.

Сверточные нейронные сети состоят из нескольких слоев, количество которых зависит от ее архитектуры и решаемой задачи [1]. Первым слоем всегда выступает сверточный слой, от которой эта сеть и получила название. Он представляет собой так называемую карту признаков – результирующую матрицу, являющуюся результатом операции свертки матрицы фильтра и части матрицы изображения, сопоставимой размерами с матрицей фильтра (рис. 1). Каждый элемент результата вычисляется как скалярное произведение матрицы фильтра и подматрицы такого же размера (части изображения) с помощью выражения:

$$C_{i,j} = \sum_{u=0}^{m_x-1} \sum_{v=0}^{m_y-1} A_{i+u,j+v} B_{u,v},$$

где A и B – матрицы размера $n_x \times n_y$ и $m_x \times m_y$ соответственно; C – матрица размера $(n_x - m_x + 1) \times (n_y - m_y + 1)$.

Большее числовое значение результирующей матрицы говорит о большем подобии части исходного изображения к фильтру (рис. 1).

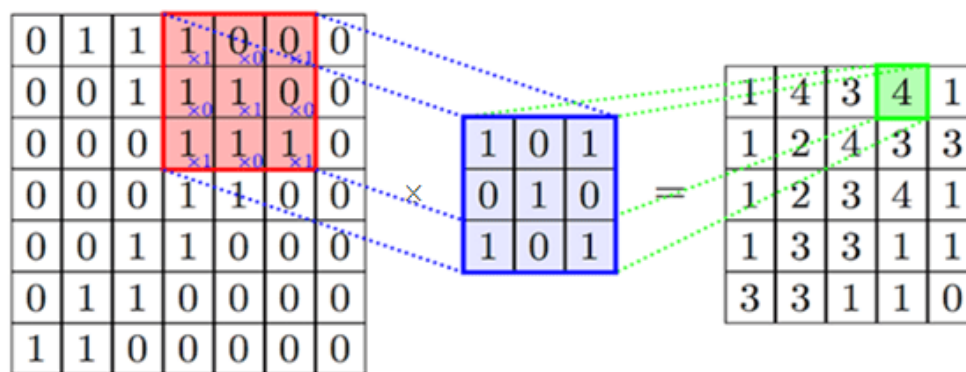


Рис. 1. Пример свертки двух матриц размером 7×7 и 3×3

Важно понимать, что количество каналов фильтра должно соответствовать количеству каналов исходного изображения. То есть при работе с тремя матрицами изображения необходимо иметь три матрицы фильтра для каждого из каналов.

Количество карт определяется требованиями к задаче, если взять большое количество карт, то повысится качество распознавания, но увеличится вычислительная сложность [2]. В большинстве случаев предлагается брать соотношение один к двум: каждая карта предыдущего слоя (например, у первого сверточного слоя, предыдущим является входной слой) связана с двумя картами сверточного слоя (рис. 2). Количество карт – 6, их размер рассчитывается в соответствии с вышеприведенной формулой.

Чем больше сверточных слоев проходит изображение и чем дальше оно движется по сети, тем более сложные характеристики выводятся в картах активации.

Подвыборочный слой также, как и сверточный, имеет карты, но их количество совпадает с предыдущим (сверточным) слоем – таких карт также 6. Цель слоя – уменьшение размерности карт предыдущего слоя. Если на предыдущей операции свертки уже были выявлены некоторые признаки, то для дальнейшей обработки настолько подробное изображение уже не нужно, и оно уплотняется до менее подробного. К тому же фильтрация уже ненужных деталей помогает не переобучаться.

В процессе сканирования ядром подвыборочного слоя карты предыдущего слоя, сканирующее ядро не пересекается в отличие от сверточного слоя. Зачастую каждая карта имеет ядро размером 2×2 , что позволяет уменьшить предыдущие карты сверточного слоя в 2 раза. Вся карта признаков разделяется на ячейки 2×2 элемента, из которых выбираются максимальные по значению. Также чаще всего в подвыборочном слое применяется функция активации ReLU [3]. Благодаря ней картина, формируемая с помощью операции свертки, получает некоторое искажение, позволяющее нейронной сети более ясно оценивать ситуацию.

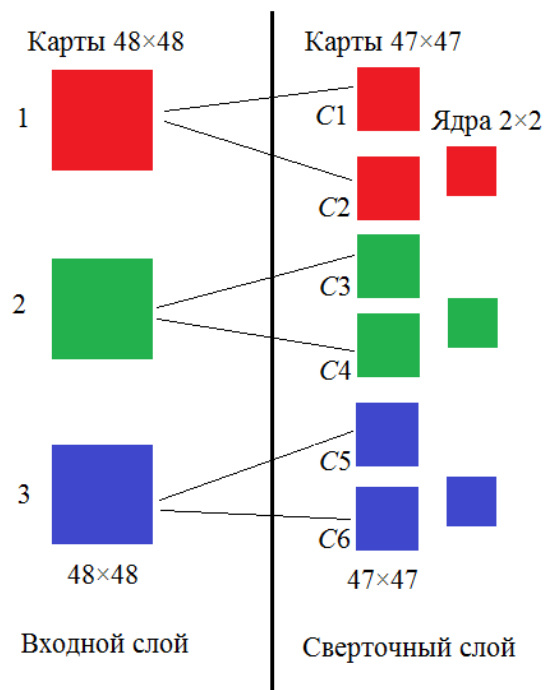


Рис. 2. Связь количества карт между сверточным и входным слоями

Последний из типов слоев – полносвязный слой. Чаще всего его представляет слой многослойного перцептрона. Цель слоя – классификация, которая моделирует сложную нелинейную функцию, оптимизируя которую улучшается качество распознавания. После нескольких проходов свертки изображения и уплотнения с помощью подвыборки система перестраивается от конкретной сетки пикселей с высоким разрешением к более абстрактным картам признаков, как правило на каждом следующем слое увеличивается число каналов и уменьшается размерность изображения в каждом канале. В конце концов остается большой набор каналов, хранящих небольшое число данных (даже один параметр), которые интерпретируются как самые абстрактные понятия, выявленные из исходного изображения.

Эти данные объединяются и передаются на обычную полносвязную нейронную сеть, которая тоже может состоять из нескольких слоев. При этом полносвязные слои уже утрачивают пространственную структуру пикселей и обладают сравнительно небольшой размерностью (по отношению к количеству пикселей исходного изображения).

На рис. 3 представлена обобщенная схема работы нейронной сети с многослойным перцептроном.

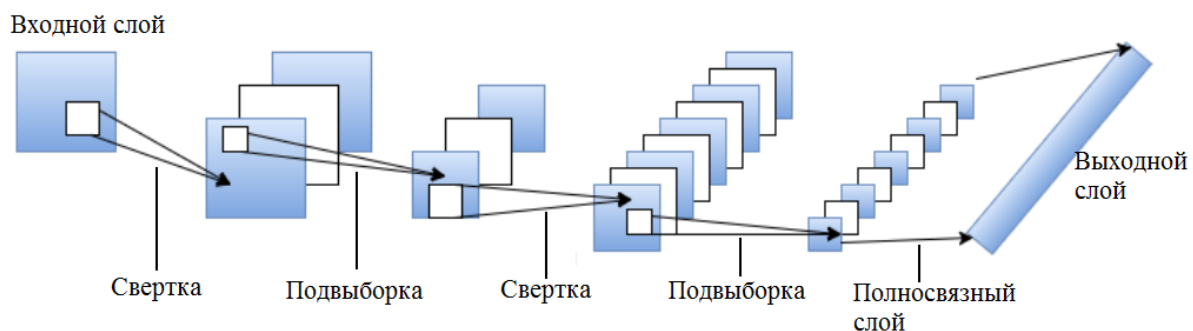


Рис. 3. Обобщенная схема работы сверточной нейронной сети с использованием многослойного перцептрона

Выходной слой отвечает за формирование вероятностей принадлежности входного образа тому или иному классу (некоторому числу). Чтобы добиться этого, выходной слой должен

содержать количество нейронов, соответствующих количеству классов. Взвешенные и просуммированные сигналы далее модифицируются с помощью функции активации. Обучение проходит классическим способом обратного распространения ошибки.

Заключение

Начиная с 2010 года, на основе базы данных проекта ImageNet ежегодно проводится тестирование методов распознавания образов и машинного зрения. Данное испытание считается одним из самых авторитетных в сфере машинного обучения, на результаты которого равняются многие разработчики по всему миру. В большинстве случаев его лидером становится нейронная сеть, разработанная на принципах свертки, показатель точности которых на сегодняшний день составляет не менее 95 %.

Оптимальным выбором для решения задачи распознавания активности выбрана нейронная сеть PNASNet-5, считающаяся на данный момент самой быстрой в принятии решений сверточной нейронной сетью [4]. Нейронная сеть PNASNet-5 реализована на языке программирования Python с использованием библиотеки TensorFlow, вошедшей в дистрибутив Anaconda, зарекомендовавшей себя как универсальная среда, помогающая достигать качества человеческого восприятия.

SELECTING NEURAL NETWORK ARCHITECTURE FOR TASK OF PHYSICAL ACTIVITY RECOGNIZING

P.S. MARYNICH, A.A. BORISKEVICH

Abstract. The optimal neural network architecture for recognizing physical activity is selected. The architecture of the convolutional neural network is described. The choice of the development environment is justified.

Keywords: neural networks, image.

Список литературы

1. Yann LeCun. Gradient-Based Learning Applied to Document Recognition, 1998.
2. Alejandro Escontrela. Convolutional Neural Networks from the ground up, 2018.
3. Alex Krizhevsky, Ilya Sutskever, Geoffrey E. Hinton. ImageNet Classification with Deep Convolutional Neural Networks, University of Toronto, 2010.
4. Chenxi Liu. Progressive Neural Architecture Search, 2018.

УДК 621.391

СРАВНИТЕЛЬНЫЙ АНАЛИЗ АЛГОРИТМОВ АВТОМАТИЧЕСКОГО ОБОБЩЕНИЯ ТЕКСТА

Е.К. ВАШКЕВИЧ, И.А. БОРИСКЕВИЧ

*Белорусский государственный университет информатики и радиоэлектроники, Республика Беларусь**Поступила в редакцию 8 марта 2021*

Аннотация. Проведен сравнительный анализ различных алгоритмов извлечения ключевых предложений для автоматического реферирования текста на наборах текстовых данных новостных статей на английском языке. Рассмотрено тринадцать различных алгоритмов реферирования, а именно TextRank, LexRank, Luhn, LSA, Edmundson, ChunkRank, TGraph, UniRank, NN-ED, NN-SE, FE-SE, SummaRuNNer и MMR-SE, и произведена оценка их эффективности с использованием нескольких показателей производительности, таких как точность, отзыв, F_1 на пяти различных уровнях отсечения суммарной длины для разных n -грамм.

Ключевые слова: обобщение текста, алгоритмы на основе графиков, алгоритмы на основе нейронных сетей, алгоритмы на основе метаэвристики, скрытый семантический анализ, ROUGE.

Введение

Задача суммаризации текстов (автореферирование) – одна из ключевых, широко обсуждаемых задач NLP (Natural Language Processing – Обработка естественного языка). Она состоит в сжатии больших объемов текста до связного краткого содержания, отражающего только основные идеи.

Сейчас доступно огромное количество текстовой информации по любой теме. Чтобы сократить время на ознакомление с интересующей информацией используют алгоритмы суммаризации текстов. Их задача – выделить из потока текстовых данных главные идеи и создать на их основе сокращенный читаемый текст. Так, суммаризация может помочь понять содержание той или иной научной статьи, получить свежие выдержки из новостей или облегчить понимание юридического заключения или финансового отчета. Автореферирование актуально практически во всех областях, так как существенно сокращает время чтения.

Алгоритмы обобщения текста на основе графов, метаэвристики и нейронной сети

Алгоритмы LSA, NN-ED, SummaRuNNer и NN-SE – основаны на семантическом анализе [1–2]. Среди них LSA – единственная модель, не основанная на нейронной сети, в которой семантическое сходство между двумя векторами слов было вычислено путем реализации инструмента уменьшения размерности SVD, который используется для проецирования каждого слова в подпространстве заранее определенного размера. Другие подходы были обработаны с помощью вложений, которые используют статистические свойства текстовой структуры для встраивания слов в векторное пространство. Внедрение слова – это модель, основанная на прогнозировании, тогда как LSA – это модель на основе счетчиков. Подходы, основанные на встраивании, превосходят LSA из-за того, что они лучше разбираются в словах, чем LSA (например, слово "Obama" тесно связано со словом "president" при встраивании слов). Производительность встраивания постепенно меняется в зависимости от размера обучающих данных. Если размер обучающих данных невелик, то он не укладывается во все параметры встраивания и имеет тенденцию к ухудшению производительности. Результат работы алгоритма NN-SE содержит слишком много предложений, что часто приводит к включению незначительных предложений, так же включение несущественных предложений снижает производительность системы. Алгоритм SummaRuNNer учитывает четыре особенности текста

хорошего резюме. Новизна предложений может удалить значимые предложения и вызвать снижение производительности системы. Алгоритм NN-DE просто основан на подобию уровня документа с использованием встроенных распределений. Он генерирует резюме с более похожими предложениями и имеет тенденцию уменьшать заметность и новизну резюме.

TextRank, LexRank, ChunkRank, TGRaph и UniRank – это подходы, основанные на графах [3–5]. TextRank использует функцию сходства косинусов на основе перекрытия слов, чтобы найти сходство между двумя предложениями, в то время как LexRank использует функцию сходства косинусов на основе векторов TF-IDF, чтобы найти сходство между двумя предложениями. В LexRank подобие учитывает фактор значимости, что приводит к увеличению его производительности по сравнению с TextRank. ChunkRank создает чанки, используя описанные правила, которые вызывают чанки с разными темами. Он использует расстояние Левенштейна, чтобы найти сходство между двумя предложениями. Затем текст ранжируется на основе схожести предложений и веса предложения. TGRaph обрабатывает тематическое моделирование с использованием двудольного графа, а предложения ранжируются с использованием алгоритма HITS. Помимо этого, он также опирается на два основных фактора: согласованность и отсутствие избыточности в резюме, что делает его более подходящим. UniRank резюмирует текст, основываясь на локальной и глобальной значимости. Он использует матрицу аффинности косинусного сходства, чтобы найти взаимосвязь между предложениями. Концепции локальной и глобальной значимости превосходят другие алгоритмы.

MMR-SE [6] резюмирует документ на основе содержания темы и новизны, которые достигаются максимальной предельной релевантностью. Алгоритм извлекает предложения с минимальным сходством из предыдущего предложения, что снижает удобочитаемость и последовательность в резюме. Фактор релевантности решает проблему низкой согласованности, но низкая читаемость снижает ее производительность. Алгоритм Луна [7] просто основан на термине "частота". Он рассматривает наиболее часто встречающийся термин как наиболее значимый термин. Алгоритм Эдмундсона [8] основан исключительно на четырех характеристиках текста: ключевой фразе, ключевой фразе, заголовке и позиции. Вес текстовых функций помогает назначить приоритет одной функции над другими. Алгоритм не находит оптимальных весов текстовых функций, что снижает его эффективность. Последний алгоритм, FE-SE, представляет собой нечеткий эволюционный подход, который анализирует особенности текста на уровне слова и предложения. Признак сходства предложений оценивается с использованием метаэвристического подхода. Данный алгоритм также находит оптимальные веса текстовых функций с использованием гибридного метаэвристического подхода и извлекает предложения в соответствии с описанными функциями. Данные функции делают алгоритм FE-SE лучшим среди всех, но у него также есть некоторые недостатки, которые могут снизить производительность. Во-первых, метаэвристический подход не гарантирует получения оптимального решения. Во-вторых, производительность алгоритма FE-SE зависит от набора обучающих данных. Эти два недостатка снижают эффективность алгоритма.

Наборы данных для сравнительной оценки эффективности алгоритмов

Эффективность алгоритмов обобщения текста зачастую зависит от типа данных. В данной работе используются наборы данных, представляющие собой новости на английском языке. Каждый набор данных, используемый для экспериментов, содержит документы из различного типа новостей, таких как политика, спорт, криминал, награды, конференции и т.д.

Для сравнения из каждого набора данных случайным образом выбрано 500 документов примерно одинаковой длины (100 предложений в каждом документе). В табл. 1 содержится подробная информация о наборах данных.

Табл. 1. Набор данных

Набор данных	Размер	Язык	Тип	Количество документов	Ср. кол-во предложений в документе	Ср. длина предложения (слова)
FIRE data (2019)	397,8 MB	English	News article	500	100	22

Оценка эффективности алгоритмов автоматического обобщения текста на основе графов, метаэвристики и нейронной сети

Анализ результатов проводился на основе следующих методов: анализ на основе алгоритма обобщения, анализ по n -граммам, анализ на основе суммарной длины.

Табл. 2. Оценка точности по ROUGE-1

Алгоритмы	Длина резюме, %				
	10	20	30	40	50
TextRank	0,8312	0,7734	0,7137	0,6599	0,5899
LexRank	0,9134	0,7639	0,6905	0,6498	0,5859
LSA	0,7543	0,7250	0,6897	0,6405	0,5876
Luhn	0,7254	0,7146	0,6803	0,6516	0,5901
Edmundson	0,9122	0,76621	0,6993	0,6428	0,5925
ChunkRank	0,7316	0,7190	0,6837	0,6513	0,5911
TGRAPH	0,9523	0,8678	0,7813	0,6956	0,6582
UniRank	0,9562	0,8697	0,7884	0,7003	0,6632
SummaRuNNer	0,9015	0,7851	0,7351	0,6195	0,5464
NN-ED	0,9112	0,7987	0,7615	0,6465	0,5591
FE-SE	0,9769	0,8938	0,8027	0,7419	0,6830
NN-SE	0,9187	0,8021	0,7843	0,6501	0,5637
MMR-SE	0,9219	0,8106	0,7414	0,6771	0,6234

В области лингвистической обработки n -грамма – это последовательность из n элементов любого текста. Оценка производилась до 4 граммов. При использовании 1-граммового метода (табл. 2) каждый член предложения сравнивается со справочной сводкой. В 2-граммовом методе (табл. 3) одновременно рассматриваются два термина, которые должны соответствовать справочной сводке. Если обнаруживается точная последовательность из двух терминов, то говорят, что они совпадают и показатели ROUGE-1 лучше, как показано в табл. 2, 6, 8. С увеличением n -значения ROUGE скорость перекрывающихся членов уменьшается. Таким образом, оценка постепенно снижается с ROUGE-1 до ROUGE-4.

Оценка точности: показатели точности постоянно снижаются с постепенным увеличением уровня отсечения. Из табл. 2–5 видно, что существуют лишь незначительные различия в производительности различных алгоритмов. Некоторые алгоритмы хорошо работают при малой краткой длине, другие – при большой краткой длине. В случае сравнения оценок точности LexRank и LSA при 10 % LexRank работает лучше, чем LSA, а при 50 % LSA работает лучше, чем LexRank. Это также указывает на то, что производительность системы зависит от суммарной длины.

Табл. 3. Оценка точности по ROUGE-2

Алгоритмы	Длина резюме, %				
	10	20	30	40	50
TextRank	0,7361	0,6462	0,5936	0,5621	0,5120
LexRank	0,8911	0,6397	0,5587	0,5339	0,4819
LSA	0,6142	0,5809	0,5582	0,5304	0,4998
Luhn	0,5755	0,5585	0,5425	0,5396	0,5001
Edmundson	0,8919	0,6492	0,5689	0,5234	0,4881
ChunkRank	0,5963	0,5610	0,5467	0,5401	0,5014
TGRAPH	0,9122	0,8321	0,7455	0,6627	0,6192
UniRank	0,9173	0,8388	0,7504	0,6664	0,6220
SummaRuNNer	0,8662	0,7442	0,6665	0,5794	0,5053
NN-ED	0,8710	0,7668	0,6741	0,6035	0,5180
FE-SE	0,9564	0,8791	0,7914	0,7046	0,6690
NN-SE	0,8791	0,7680	0,6775	0,6119	0,5185
MMR-SE	0,9035	0,7514	0,6931	0,6047	0,5737

Табл. 4. Оценка точности по ROUGE-3

Алгоритмы	Длина резюме, %				
	10	20	30	40	50
TextRank	0,7076	0,6228	0,5750	0,5481	0,5021
LexRank	0,8787	0,6217	0,5402	0,5176	0,4678
LSA	0,5881	0,5568	0,5375	0,5120	0,4856
Luhn	0,5372	0,5324	0,5224	0,5226	0,4873
Edmundson	0,8793	0,6287	0,5508	0,5072	0,4738
ChunkRank	0,5468	0,5433	0,5281	0,5230	0,4888
TGRAPH	0,8906	0,7913	0,7220	0,6315	0,5926
UniRank	0,8961	0,7959	0,7272	0,6338	0,6013
SummaRuNNer	0,8460	0,7089	0,6452	0,5417	0,4782
NN-ED	0,8472	0,7355	0,6401	0,5884	0,4812
FE-SE	0,9445	0,8386	0,7625	0,6893	0,6438
NN-SE	0,8486	0,7367	0,6450	0,5890	0,4942
MMR-SE	0,8912	0,7714	0,6763	0,5914	0,5222

Табл. 5. Оценка точности по ROUGE-4

Алгоритмы	Длина резюме, %				
	10	20	30	40	50
TextRank	0,7002	0,6179	0,5717	0,5461	0,5008
LexRank	0,8753	0,6187	0,5385	0,5167	0,4667
LSA	0,5839	0,5521	0,5321	0,5077	0,4821
Luhn	0,5260	0,5269	0,5179	0,5193	0,4854
Edmundson	0,8785	0,6255	0,5485	0,5064	0,4727
ChunkRank	0,5452	0,5374	0,5241	0,5231	0,4862
TGRAPH	0,8891	0,7881	0,7013	0,6253	0,5915
UniRank	0,8943	0,7915	0,7086	0,6321	0,5958
SummaRuNNer	0,8432	0,7043	0,6445	0,5408	0,4771
NN-ED	0,8458	0,7324	0,6319	0,5848	0,4806
FE-SE	0,9381	0,8227	0,7339	0,6738	0,6260
NN-SE	0,8461	0,7357	0,6431	0,5870	0,4931
MMR-SE	0,8831	0,7236	0,6416	0,5723	0,5241

Оценка отзыва: оценка отзыва постепенно увеличивается по мере увеличения порогового уровня. Также скорость извлечения релевантных терминов уменьшается по мере увеличения суммарной длины. Из табл. 6–7 видно, что, как и при оценке точности, некоторые алгоритмы лучше работают при малой суммарной длине, а другие – при большой суммарной длине. Во всех этих случаях показана зависимость от суммарной длины.

Табл. 6. Оценка отзыва по ROUGE-1

Алгоритмы	Длина резюме, %				
	10	20	30	40	50
TextRank	0,1979	0,4100	0,5603	0,6803	0,7608
LexRank	0,1974	0,3578	0,4800	0,5914	0,6845
LSA	0,1980	0,4068	0,5538	0,6797	0,7706
Luhn	0,1569	0,3436	0,4919	0,6463	0,7409
Edmundson	0,1854	0,3512	0,4770	0,5919	0,6815
ChunkRank	0,1824	0,3717	0,5143	0,6567	0,7524
TGRAPH	0,2434	0,4768	0,5917	0,6943	0,7995
UniRank	0,2461	0,4782	0,5952	0,7011	0,8030
SummaRuNNer	0,1962	0,4079	0,5481	0,6147	0,7212
NN-ED	0,2019	0,4124	0,5490	0,6329	0,7381
FE-SE	0,27672	0,5073	0,6262	0,7470	0,8321
NN-SE	0,2118	0,4315	0,5543	0,6580	0,7441
MMR-SE	0,2119	0,4135	0,5412	0,6633	0,7218

Табл. 7. Оценка отзыва по ROUGE-4

Алгоритмы	Длина резюме, %				
	10	20	30	40	50
TextRank	0,1446	0,2946	0,4079	0,5097	0,5844
LexRank	0,1676	0,2552	0,3306	0,4168	0,4838
LSA	0,1405	0,2862	0,3930	0,4917	0,5763
Luhn	0,1048	0,2275	0,3365	0,465982	0,5520
Edmundson	0,1585	0,2523	0,3291	0,4118	0,4809
ChunkRank	0,1227	0,2558	0,3636	0,4741	0,5631
TGRAPH	0,1949	0,3019	0,4258	0,5479	0,6211
UniRank	0,1972	0,3042	0,4288	0,5494	0,6240
SummaRuNNer	0,1638	0,2730	0,3926	0,5059	0,5726
NN-ED	0,1720	0,2821	0,4017	0,5102	0,5801
FE-SE	0,2367	0,3382	0,4562	0,5672	0,6321
NN-SE	0,1721	0,2910	0,4018	0,5153	0,5910
MMR-SE	0,1764	0,2819	0,3761	0,4928	0,5712

Табл. 8. Оценка F_1 по ROUGE-1

Алгоритмы	Длина резюме, %				
	10	20	30	40	50
TextRank	0,3134	0,5299	0,6215	0,6657	0,6613
LexRank	0,3185	0,4817	0,5612	0,6157	0,6286
LSA	0,3076	0,5145	0,6068	0,6545	0,6617
Luhn	0,2532	0,4595	0,5656	0,6433	0,6525
Edmundson	0,3042	0,4756	0,5622	0,6134	0,6317
ChunkRank	0,2919	0,4901	0,5870	0,6539	0,6621
TGRAPH	0,3877	0,6154	0,6735	0,6949	0,7220
UniRank	0,3914	0,6171	0,6783	0,7007	0,7264
SummaRuNNer	0,3223	0,5369	0,6279	0,6170	0,6217
NN-ED	0,3306	0,5439	0,6380	0,6396	0,6362
FE-SE	0,4312	0,6472	0,7035	0,7444	0,7502
NN-SE	0,3442	0,5611	0,6495	0,6540	0,6415
MMR-SE	0,3446	0,5476	0,6257	0,6701	0,6690

Табл. 9. Оценка F_1 по ROUGE-4

Алгоритмы	Длина резюме, %				
	10	20	30	40	50
TextRank	0,2353	0,3948	0,4713	0,5239	0,5365
LexRank	0,2763	0,3569	0,4059	0,4588	0,4728
LSA	0,2224	0,3721	0,4463	0,4955	0,5207
Luhn	0,1719	0,3152	0,4038	0,4866	0,5128
Edmundson	0,2652	0,3546	0,4073	0,4521	0,4750
ChunkRank	0,2003	0,3466	0,4293	0,4974	0,5218
TGRAPH	0,3197	0,4365	0,5298	0,5840	0,6059
UniRank	0,3231	0,4394	0,5343	0,5878	0,6095
SummaRuNNer	0,2743	0,3935	0,4879	0,5227	0,5205
NN-ED	0,2859	0,4073	0,4911	0,5449	0,5256
FE-SE	0,3780	0,4793	0,5626	0,6159	0,6290
NN-SE	0,2860	0,4170	0,4945	0,5488	0,5376
MMR-SE	0,2941	0,4057	0,4742	0,5296	0,5466

Оценка английского языка F_1 : из оценки английского языка F_1 в табл. 8–9 видно, что производительность некоторых алгоритмов, таких как TextRank, SummaRuNNer, NN-ED и NN-SE, снижается на 50 %, тогда как производительность других алгоритмов плавно увеличивается по мере увеличения суммарной длины.

Заключение

Исследованы тринадцать алгоритмов автоматического реферирования с аналогичными настройками для новостных наборов данных на английском языке. Произведена оценка производительности с использованием показателей точности, отзыва и показателей F_1 на пяти различных уровнях отсечения суммарной длины для разных n -грамм. Обнаружено, что показатель точности уменьшается с увеличением длины сводки, а также с увеличением значений слова n -грамм. Оценки запоминания увеличиваются с увеличением длины сводки, но уменьшаются по сравнению с n -значениями в n -граммах слов, максимальные значения наблюдались при длине сводки 40 %. Ограничение данных алгоритмов обусловлено тем, что их эффективность зависит от суммарной длины. Непосредственная близость значений F_1 для алгоритмов SummaRuNNer, NN-SE и NN-ED обусловлена тем, что они являются нейронными сетями, основанными на инструменте word2vec. При этом алгоритмы на основе нейронных сетей не показали лучшей производительности по сравнению с алгоритмами на основе графов. Также было показано, что почти все алгоритмы генерируют избыточные, удобочитаемые, значимые резюме.

COMPARATIVE ANALYSIS OF TEXT SUMMARIZATION ALGORITHMS IN ENGLISH LANGUAGE

E.K. VASHKEVICH, I.A. BARYSKIEVIC

Abstract. A detailed comparative study of various extraction algorithms for automatic text summarization on text data sets of news articles in English was carried out. Thirteen different summarization algorithms were considered, namely TextRank, LexRank, Luhn, LSA, Edmundson, ChunkRank, TGraph, UniRank, NN-ED, NN-SE, FE-SE, SummaRuNNer and MMR-SE, and their effectiveness was assessed using several performance metrics such as Accuracy, Recall, F_1 at five different levels of total length cutoff for different n -grams.

Keywords: text summarization, graph based techniques, neural networks based techniques, meta-heuristic based techniques, latent semantic analysis, ROUGE

Список литературы

1. Hayato Kobayashi, Masaki Noguchi, Taichi Yatsuka. // Summarization Based on Embedding Distributions, 2015. In EMNLP. 1984–1989.
2. Ramesh Nallapati, Feifei Zhai, Bowen Zhou. // SummaRuNNer: A recurrent neural network based sequence model for extractive summarization of documents, 2017
3. Rada Mihalcea, Paul Tarau. // TextRank: Bringing order into texts. Association for Computational Linguistics, 2004.
4. Günes Erkan, Dragomir R Radev. // LexRank: Graph-based lexical centrality as salience in text summarization. Journal of Artificial Intelligence Research, 2004, P. 457–479.
5. Daraksha Parveen, Mohsen Mesgar, Michael Strube. // Generating Coherent Summaries of Scientific Articles Using Coherence Patterns, 2016. In EMNLP. P. 772–783.
6. Jade Goldstein, Jaime Carbonell. // Summarization: using MMR for diversity-based reranking and evaluating summaries. In Proceedings of a workshop on held at Baltimore, Maryland: October 13–15, 1998. Association for Computational Linguistics, P. 181–195.
7. Hans Peter Luhn. // The automatic creation of literature abstracts. IBM Journal of research and development, 1958, P. 159–165.
8. Harold P Edmundson. // New methods in automatic extracting. Journal of the ACM (JACM) 16, 1969, P. 264–285.

УДК 004.056.57:032.27

ПОДДЕРЖКА ИНФОРМАЦИОННОГО УПРАВЛЕНИЯ В ОБРАЗОВАНИИ С ИСПОЛЬЗОВАНИЕМ БЛОКЧЕЙН

Д.А. КАЧАН, В.А. ВИШНЯКОВ

Белорусский государственный университет информатики и радиоэлектроники, Республика Беларусь

Поступила в редакцию 4 марта 2021

Аннотация. Предложена концепция и модель информационной поддержки управления в образовании с использованием блокчейн. Представлена модель смарт-контракта. Рассмотрен подход и модель для представления объектного идентификатора.

Ключевые слова: информационное управление в образовании, блокчейн, смарт-контракт, идентификатор

Введение

Сложности построения процессов управления и принятия управленческих решений в образовании обусловлены рядом особенностей, присущих этой области:

многоаспектность происходящих процессов (экономических, социальных и т.п.) и их взаимосвязанность. Это зачастую не позволяет вычлениить и детально исследовать отдельные явления – все происходящие процессы должны рассматриваться в совокупности:

– отсутствие достаточной количественной информации о динамике процессов, что вынуждает переходить к качественному анализу процессов;

– изменчивость характера процессов во времени и т.д. Для рассматриваемой слабоструктурированной системы число факторов в различных ситуациях может измеряться десятками;

– особенностями системы управления отраслью.

Существуют различные модели и средства представления данных для лиц, принимающих решение в процессе управления [1]. В качестве одной из перспективных технологий для поддержки управления в образовании предлагается использование технологии распределенных реестров [2].

Концепция и модель информационного управления в образовании с использованием блокчейн

В настоящее время актуальным является использование технологии блокчейн для подтверждения авторства и прав на объекты интеллектуальной собственности [2]. В рамках статьи понятие объекта интеллектуальной собственности расширяется на документы, выдаваемые учреждениями образования. Внедрение подобного подхода определяет необходимость разработки информационной поддержки в образовании с использованием блокчейн (ИПвОБ). Под концепцией ИПвОБ будем понимать реализацию, при которой имеется способность автоматического формирования электронного документа и независимая его верификация в местах предъявления. Система должна иметь возможность принять и обработать заявку на формирование электронного нормализованного документа установленного образца и формата, вычислить хэш-значение полученного документа включая уникальный идентификатор учреждения образования, добавить информацию контрольной величины хэш-функции в блок цепей транзакций блокчейн для последующей проверки. В рамках этой концепции сформулируем две задачи [3].

Постановка первой задачи: необходимо разработать технологию для обеспечения подтверждения достоверности документов об образовании на основе технологии распределенных реестров с использованием смарт-контрактов. Для второй задачи в ИПвО надо

разработать модель, чтобы сбалансировать потребности промышленности и выпускников и поддержать это технологией блокчейн.

Тогда обобщенную модель M_{mse} для ИПвОБ представим в виде:

$$M_{mse} = (M_{edd}, M_{sdd}, M_{ebe}),$$

где M_{edd} – модель генерации цифрового документа в образовании, M_{sdd} – модель подтверждения цифрового документа в образовании; M_{ebe} – модель сбалансирования потребностей экономики и выпуска специалистов в образовании.

Подтверждение достоверности является одной из ключевых характеристик технологии распределенных реестров. Технология TRP позволяет осуществлять подтверждение достоверности (существования) записи в виде хэш-суммы интересующего документа. Это позволяет построить сервис для сравнения предоставленного документа с хранимым в сети блокчейн, не нарушая конфиденциальность данных – сами документы в сеть не попадают, сравнение осуществляется только на основе хэш-значений.

Принцип работы публикации документа следующий – производится вычисление значения хэш-суммы документа, проводится транзакция в сети блокчейн, содержащая вместо адреса получателя данной транзакции полученное хэш-значение с используется функций OP_RETURN для блокчейн-сети Bitcoin. Функция OP_RETURN имеет ограничение на длину данных в 40 байт, что является оптимальным значением ввиду используемого алгоритма вычисления хэш-функции SHA – полученное значение хэш-функции имеет длину в 32 байта [4]. Для блокчейн-сетей Ethereum транзакция имеет дополнительный атрибут, что позволяет использовать сеть для хранения данных без ограничений по длине [2].

Модель смарт-контракта

Для публикации документа в сети блокчейн предполагается использовать возможности смарт-контрактов. Формальное представление смарт-контракта может быть отражено в виде математической модели конечного автомата, представляющего в теории алгоритмов математическую абстракцию или модель дискретного устройства, имеющего один вход, один выход и в каждый момент времени находящегося в одном состоянии из множества возможных:

$$M = (Q, \Sigma, \delta, s_0, F),$$

где Q – конечное множество всех возможных состояний смарт-контракта; Σ – набор всех входных событий смарт-контракта; δ – множество переходных функций смарт-контракта; $\delta: Q \cdot \Sigma \rightarrow Q$ – конечное состояние смарт-контракта, F – конечное состояние смарт-контракта $F \in Q$; s_0 – начальное состояние смарт-контракта $s_0 \in Q$.

Обозначив начальное состояние блокчейн-сети γ получаем переход сети в новое состояние при условии совершения успешной транзакции:

$$\gamma \xrightarrow{T_x} \gamma'.$$

Новое состояние сети блокчейн влияет в разной степени на многие учетные записи в сети, а также на другие смарт-контракты, которые в свою очередь оказывают влияние на данные в цепочке блоков.

Процедура проверки осуществляется по следующему алгоритму – проверяющей стороной вычисляется хэш-значение электронной версии документа и сравнивается полученное значение со значением, указанным в первой транзакции, когда данные были отправлены в блокчейн. На основании сравнения принимается решение о достоверности документа.

Транзакции, связанные с механизмами подтверждения авторства или достоверности с помощью цифрового отпечатка, применяются для предъявления доказательства одной стороны другой, когда проверяющая сторона сверяет хэш-значение, временную метку транзакции и, что наиболее важно, подлинность (принадлежность) криптовалютного "кошелька" предъявителя.

Механизм для автоматизированного подтверждения достоверности документа на основе использования ТРР охватывают лишь две стороны (предъявитель и проверяющий), что недостаточно в случае официальных документов, эмитент которых обязательно должен присутствовать в модели в качестве доверенной третьей стороны (ДТС). Модель подтверждения должна устанавливать не только принадлежность документа эмитенту, но и подтверждать полномочия эмитента на осуществление данного вида деятельности и дополнительные сведения (например, для сферы образования перечни специальностей подготовки в определенный период времени в соответствии с лицензией).

Модель объектного идентификатора

Для реализации модели идентификаторов предлагается использование приватной сети блокчейн в части создания и ведения регистра записей, а также публичной сети, для обеспечения доступа третьей стороны при необходимости подтверждения достоверности документа. Приватная сеть блокчейн обеспечивает хранение полных копий распределенного реестра транзакций, обеспечивая их сохранность и достоверность хранимых данных. Реестр содержит записи о документах об образовании и программные модули (смарт-контракты), которые позволяют участникам распределенной сети осуществлять взаимодействие с реестром и друг с другом.

Взаимодействие ответственных за ведение реестра, обучающихся, представителей третьих сторон, заинтересованных в проверке достоверности данных, осуществляется на основании смарт-контрактов на выполнение конкретных действий, которые позволяют участникам вызывать различные события для управления записями в сети.

Использование стандартизированного международного объектного идентификатора позволяет решить проблему подтверждения достоверности и существования эмиссионного центра, издавшего рассматриваемый документ об образовании, в том числе позволит проверить данные о существовавших ранее эмиссионных центрах, а также о внесенных изменениях, связанных с их функционированием (наименование, уровни подготовки, местонахождение).

Предложенное решение данной проблемы основано на использовании реестра Международного регистрационного органа, в качестве которого выступает совместный орган Международного союза электросвязи ITU-T и Международной организации по стандартизации ISO, ответственного за назначение идентификаторов объектов верхнего уровня с первичным целочисленным значением (метка JOINT-ISO-ITU-T) [5]. Международный объектный идентификатор (OID) представляет собой запись, состоящую из комбинации последовательных идентификаторов, сформированных в соответствии с принятыми правилами.

Для рассматриваемого случая идентификатор Республики Беларусь, состоящий из цепочки идентификаторов "первичный международный идентификатор – идентификатор группировки государств – идентификатор государства Республика Беларусь", может быть выражен тремя различными способами:

1. В нотации ASN.1: {joint-iso-itu-t(2) country(16) by(112)}.
2. В нотации dot: .16.112.
3. В нотации OID-IRI: /Country/BY.

При анализе предложенных нотаций очевидно, что нотации 1 и 3 содержат страновой идентификатор "BY". При обозначении стран используются идентификаторы Alfa-2 code (состоит из двух романских символов), Alfa-3code (состоит из трех романских символов), числовой код (numeric code, для Республики Беларусь принят равным 112).

В связи с этим предлагается использование dot-нотации, представляющей собой последовательности числовых идентификаторов, разделенных точками.

Каждое число, отделяемое от последующего (и предыдущего для не первого) точкой, условно обозначим уровнем идентификатора. Таким образом, получаем следующее базовое дерево идентификаторов, в которых уровни 1–3 уже определены:

1. Уровень 1 – совместный идентификатор ITU-T и ISO, значение – 2.
2. Уровень 2 – страновой идентификатор (country), значение – 16.
3. Уровень 3 – Республика Беларусь, значение – 112.
4. Уровень 4 – территориальные единицы Республики Беларусь.

Формирование данного уровня дерева OID-идентификаторов осуществим на основе международных обозначений территориального деления Республики Беларусь, принятых в [6, 7]. Численные обозначения регионов целесообразно принять в соответствии с Общегосударственным классификатором Республики Беларусь "Система обозначений объектов административно-территориального деления и населенных пунктов" (классификатор СОАТО, уровень 1). Для уровня 4 предлагается введение территориального деления Республики Беларусь. Уровень 5 – предлагается включить рассматриваемую отрасль (образование). Присвоим идентификатору отрасли образования значение 1.

Подобная логика построения OID-идентификатора позволит охватить все отрасли народного хозяйства при использовании результатов работы за рамками системы образования. OID Белорусского государственного университета информатики и радиоэлектроники, выраженный в формате dot-нотации, будет представлять собой последовательность 2.16.112.5.1.1, при обозначении направления подготовки с присвоением квалификации "бакалавр" в рамках данного учреждения образования – 2.16.112.5.1.1.1, а степени "магистр" – 2.16.112.5.1.1.2.

Заключение

Предложена концепция информационной поддержки в образовании с использованием технологии блокчейн. Представлена интегрированная модель для концепции, включающая модель для подтверждения достоверности документов об образовании (эмиссия цифрового документа об образовании, его проверка) и модель оптимизации выпуска специалистов под потребности отраслей промышленности. Формальное представление смарт-контракта отражено в виде математической модели конечного автомата. Расширен идентификатор для цифровых документов в образовании до шести уровней.

THE SUPPORT OF INFORMATION CONTROL IN EDUCATION USING BLOCKCHAIN

D.A. KACHAN, U.A. VISHNYAKOU

Abstract. The concept and model of information support of management in education using blockchain is proposed. A smart contract model is presented. An approach and a model for representing an object identifier are considered.

Keywords: information management in education, blockchain, smart contract, identifier.

Список литературы

1. Губанов Д.А., Новиков Д.А., Чхартишвили А.Г. Социальные сети: модели информационного влияния, управления и противоборства. М.: Институт проблем управления им. В.А. Трапезникова, ФИЗМАТЛИТ, 2010.
2. Свон М. Блокчейн: схема новой экономики. М.: Олимп-Бизнес, 2017.
3. Качан Д.А., Вишняков В.А. // Докл. БГУИР. 2020. № 7.С. 14–23.
4. Crespo A.S. Blockchain Timestamping Architecture (BTA) [Электронный ресурс]. URL: <http://arxiv.org/abs/1711.04709v0>.
5. ITU-T. ITU-T X.660 Information technology – Procedures for the operation of object identifier registration authorities: General procedures and top arcs of the international object identifier tree, 2011.
6. ISO3166-1:2013. Codes for the representation of names of countries and their subdivisions. Part 1: Country codes, 2013.
7. ISO3166-2:2013. Codes for the representation of names of countries and their subdivisions. Part 2: Country subdivision code, 2013.

УДК 621.391

СЕТЕВОЕ КОДИРОВАНИЕ КОДОМ РИДА-СОЛОМОНА С УЧЕТОМ ЛИДЕРОВ СТОХАСТИЧЕСКОЙ СЕТИ

С.Б. САЛОМАТИН, А.Э. АЛЕКСЕЕНКО, А.П. ТУРЛАЙ

Белорусский государственный университет информатики и радиоэлектроники, Республика Беларусь

Поступила в редакцию 8 марта 2021

Аннотация. Рассмотрены алгоритмы сетевого кодирования в линейной стохастической сети с использованием лидеров сети. Приведены модели поиска лидеров, алгоритмы кодирования и декодирования кода Рида-Соломона с исправлением ошибок и стираний. Приведены результаты моделирования.

Ключевые слова: лидеры стохастической линейной сети, алгоритмы поиска, код Рида-Соломона, декодирование ошибок и стираний

Введение

Известные методы и алгоритмы сетевого кодирования [1–2] позволяют улучшить надежность работы стохастических сенсорных сетей. В гомогенных сетях с разными состояниями узлом сети этого может оказаться недостаточно.

Одним из путей повышения эффективности работы стохастических сенсорных сетей является применение алгоритма кодирования пакетов сети с указанием лидера сети, имеющего наилучшее управление пространством состояний, и кодов, корректирующих ошибки и стирания.

1. Модель динамической стохастической сенсорной сети

Мы рассматриваем сети, в которых каждый узел обновляет скалярное состояние ψ_i , $\psi = u_i + w_i$, $i = 1, \dots, n$, где u_i – управляющий вход и w_i – белое стохастическое возмущение с нулевым средним и единичной дисперсией. Узел является последовательным, если для формирования управляющего действия используется только относительный обмен информацией с соседями $u_i = -\sum_{j \in N_i} (\psi_i - \psi_j)$.

Узел является *лидером*, если, помимо относительного обмена информацией с соседями, он также имеет доступ к собственному состоянию

$$u_i = -\sum_{j \in N_i} (\psi_i - \psi_j) - k_i \psi_i,$$

где k_i – положительное число и N_i является набором всех узлов, с которыми связывается узел i . Таким образом, представление пространства-состояния лидера согласованной сети задается уравнением

$$\dot{\psi} = -(L + \text{diag}(\mathbf{k}) \text{diag}(\mathbf{x})) \psi + w,$$

где $\mathbf{x}$ – вектор с логическим значением с i -ого элемента $x_i \in \{0, 1\}$, указывающей, что узел i является лидером, если $x_i = 1$ и что узел i является последовательным, если $x_i = 0$.

Ковариационная матрица в установившемся режиме имеет вид $\sigma = \lim_{t \rightarrow \infty} M(\psi(t) \psi^T(t))$ и может быть определена с помощью уравнения Ляпунова:

$$(L + \text{diag}(k) \text{diag}(x))\sigma + \sigma(L + \text{diag}(k) \text{diag}(x)) = I.$$

где $M(\cdot)$, $M(w(t)w^T(\tau)) = I\delta(y = \tau)$ – оператор математического ожидания,
 $\sigma = \frac{1}{2}(L + \text{diag}(k) \text{diag}(x))^{-1}$.

Дисперсия стационарного режима в общем виде определяется как $\text{trace}(\sigma) = \frac{1}{2} \text{trace}((L + \text{diag}(\mathbf{k}) \text{diag}(\mathbf{x}))^{-1})$ и количественно определяет величину отклонение от согласованного режима стохастических сетей. Таким образом, проблема идентификации лидеров сети, сводится к решению задачи минимизации среднеквадратического отклонения по алгоритмам МНК следующего вида

$$\text{minimize } J(x) = \text{trace}((L + \text{diag}(k) \text{diag}(x))^{-1}),$$

$$\text{subject to } x_i \in \{0,1\}, \quad i=1,\dots,n,$$

$$\sum_{i=1}^n x_i = N_l,$$

где граф Лапласиана L и вектор k с положительными элементами являются исходными данными задачи, а логически значимый вектор x с кардинальностью является переменной оптимизации.

Для сенсорных сетей МНК является эквивалентным задаче выбора абсолютных измерений положения среди n датчиков с минимальной дисперсией ошибки оценки [2].

Пример работы алгоритма показан на рис. 1.

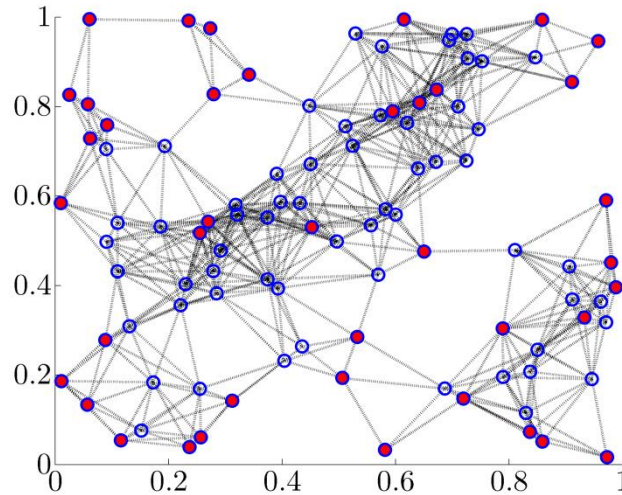


Рис. 1. Выбор лидеров стохастической сети, $J = 0,95$

2. Модель сетевого кодирования

Модель сети состоит из источника, нескольких приемников, и промежуточных узлов. Источник осуществляет кодирование, а приемники декодирование сетевого кода. Вычисления выполняются в конечном поле $GF(q^m)$.

Исходная информация представляется в виде информационных векторов длины k над полем Галуа $GF(q^m)$, где q – степень простого числа. Каждый вектор поступает от источника на вход кодера сетевого канала. В процессе кодирования формируется набор векторов длины n .

Модель сетевого канала предполагает, что сетевой канал может порождать ошибки в виде передачи базисного вектора и стирания – исчезновение базисного вектора. В сетевых узлах производятся неизвестные линейные комбинации над полем $GF(q)$.

Декодер преобразует (декодирует) принятый набор векторов и определяет информационный вектор длины k .

Кодер сопоставляет информационному вектору подпространство, и передает по сети композицию базисных векторов. Предполагается, что линейные комбинации в узлах не выводят векторы из подпространства. Линейные комбинации в узлах случайные, что может приводить к получению порождающей системы вложенного в исходное подпространства (рис. 2).

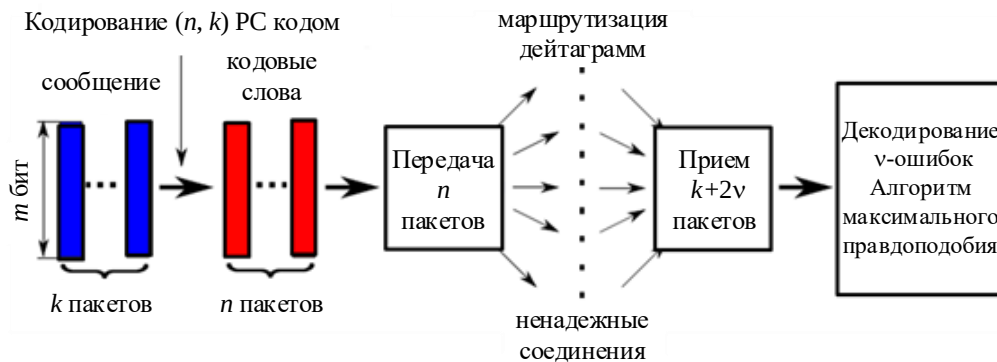


Рис. 2. Схема сетевого кодирования с использованием кода Рида-Соломона

Рассмотрим конечное поле $GF(q)$ и его расширение $GF(q^m)$, которое иногда будет удобно представлять как векторное пространство $GF(q^m)^n$. Выберем некоторое множество $A = \{\alpha_1, \alpha_2, \dots, \alpha_l\}$ линейно независимых над $GF(q)$ элементов $GF(q^m)^n$. Кодирование будет выглядеть следующим образом.

Пусть $\mathbf{a} = (a_0, a_1, \dots, a_{k-1}) \in GF(q^m)^k$ – информационное сообщение.

Вектору $\mathbf{a}$ ставится в соответствие линейаризованный полином

$$a(x) = \sum_{i=0}^{k-1} a_i x^{q^i} \in GF(q^m)[x].$$

Для полинома $a(x)$ во всех точках $\{\alpha_i\}$ множества A вычисляется его значение β_i . Пары (α_i, β_i) можно рассматривать как элементы $GF(q)^{2m}$ пространства $W = A \oplus F_q^m$ размерности $(l + m)$. Множество пар $\{(\alpha_1, \beta_1), (\alpha_2, \beta_2), \dots, (\alpha_l, \beta_l)\}$ являются линейно независимыми над полем $GF(q)$ и образуют l -мерное пространство.

Режим исправления ошибок и стираний. В режиме стирания фиксируются не сами оценки принятых символов, а их местоположение и им присваивается статус стертого символа. Суть исправления и стирания состоит в том, что после декодирования можно провести восстановление стертых символов, используя алгоритмы интерполяции. Если при декодировании использовать l стертых символов, тогда два кода будут отличаться друг от друга по меньшей мере на $(d - l)$ позиций, где d – кодовое расстояние. Тогда в дополнение к стиранию можно будет исправлять $t_m = \lfloor (d - l - 1) / 2 \rfloor$ ошибок, где $\lfloor x \rfloor$ – это целая часть числа x . Код может исправлять все комбинации из v ошибок и l стираний в канале, для которого $2v + l < d$.

Для восстановления одного стертого символа необходим только один проверочный символ – для восстановления значения, т.к. позиция стирания принимающей стороне известна.

Например, для поля Галуа $GF(256)$ РС-код (94, 88) из 8-битовых символов имеет $n = 94$, $k = 88$ и может исправить до 3 ошибок ($t = 3$) и восстановить до 6 стертых символов.

Алгоритм исправления стираний. Предположим, что выполнено f стираний при приеме кодового слова, в котором имеется v ошибок.

Обозначим локаторы ошибок как $X_1 = \alpha^{i_1}, X_2 = \alpha^{i_2}, \dots, X_v = \alpha^{i_v}$, а стираний – как $Y_{c,1} = \alpha^{j_1}, Y_{c,2} = \alpha^{j_2}, \dots, Y_{c,f} = \alpha^{j_f}$. Декодирование ведется в следующем порядке:

1. Вычисляется полином локаторов стираний:

$$\Gamma(x) = \prod_{l=1}^f (1 - Y_{c,l} \cdot x).$$

2. В декодируемом векторе заменяют символы с координатами стираний на нулевые символы. Для нового вектора находится полином синдрома стираний $s(x)$.

3. Определяется модифицированный полином синдрома

$$SE(x) = (\Gamma(x)[1 + s(x)] - 1) \bmod x^{2t+1}.$$

4. Вычисляется полином локаторов ошибок $\sigma(x)$, используя для этого алгоритм Берлекемпа-Мессии и значения модифицированного полинома $SE_i, i = f + 1, \dots, 2t$.

5. Определяются корни уравнения $\sigma(x) = 0$ и координаты ошибок.

6. Составляется ключевое уравнение $\omega(x) = \sigma(x)[1 + SE(x)] \bmod x^{2t+1}$ и определяется полином локаторов ошибок-стираний $\psi(x) = \sigma(x)\Gamma(x)$.

7. Оцениваются значения ошибок и стираний. Значения ошибок вычисляют по формуле

$$Q_{i_k} = \frac{-X_k \omega(X_k^{-1})}{\psi'(X_k^{-1})},$$

где ψ' – формальная производная.

$$\text{Значения стираний вычисляют по формуле } F_{i_k} = \frac{-Y_k \omega(Y_k^{-1})}{\psi'(Y_k^{-1})}.$$

3. Моделирование

Алгоритмы кодирования и поиска лидеров стохастической сети были промоделированы в среде Matlab. Результаты моделирования показали, что применение предлагаемых алгоритмов позволяет уменьшить вероятность ошибки при приеме пакетов сети. Такт для сети с параметрами $m = 4000$ – длина пакета, $k = 80$ пакетов на сообщение, $p = 0,03$ вероятности приема поврежденного пакета, вероятность ошибки $P_{err} = 10^{-4}$, тогда как для сети без кодирования $P_{err} = 0,9$.

Заключение

Применение сетевого кодирования с использованием лидеров стохастической сети и кода Риды-Соломона позволяет повысить надежность работы стохастической сети.

NETWORK CODING BY REED-SOLOMON CODE WITH STOCHASTIC NETWORK LEADING

S.B. SALOMATIN, A.E. ALEKSEENKO, A.P. TURLAY

Annotation. The algorithms of network coding in a linear stochastic network with the use of network leaders were considered. Lead search models, Reed-Solomon code coding and decoding algorithms with error correction and erasure were given. Simulation results were presented.

Keywords: stochastic line network leaders, search algorithms, Reed-Solomon code, error decoding and erasure

Список литературы

1. Martinez-Penas U., Kschischang F.R. Reliable and Secure Multishot Network Coding using Linearized Reed-Solomon Codes [Электронный ресурс]. URL: [https://arXiv:1805.03789v3\[csIT\]](https://arXiv:1805.03789v3[csIT]).
2. Lin F., Fardad M., Jovanović R. // IEEE Transactions on automatic control. 2014. Vol. 59, № 7.

УДК 621.391

ОЦЕНКА ИСКАЖЕНИЙ ИЗОБРАЖЕНИЙ ПРИ СУБПИКСЕЛЬНОМ СДВИГЕ КАМЕРЫ

А.В. ЗАХАРЕНКО, О.Г. ШЕВЧУК, В.Ю. ЦВЕТКОВ

*Белорусский государственный университет информатики и радиоэлектроники, Республика Беларусь**Поступила в редакцию 9 марта 2021*

Аннотация. Разработан алгоритм формирования изображений низкого пространственного разрешения. Алгоритм основан на использовании изображения высокого пространственного разрешения для генерирования изображений низкого пространственного разрешения, смещенных друг относительно друга на доли пикселя в низком разрешении. Произведена оценка искажений изображений при субпиксельном сдвиге с использованием среднеквадратичной ошибки.

Ключевые слова: субпиксельный сдвиг, изображения, среднеквадратичная ошибка.

Введение

В ряде случаев необходимо накопление изображений для повышения отношения сигнал-шум при передаче данных. Для этого требуется высокая точность пространственного совпадения изображений, полученных в различные моменты времени. В этой связи актуальной задачей является оценка искажений изображений при субпиксельном смещении видеокамеры. Данная задача может быть решена с использованием методов физического или программного моделирования. Один из вариантов программного моделирования основан на использовании изображения высокого пространственного разрешения с достаточно высоким числом деталей для получения изображений низкого пространственного разрешения с помощью смещения окна аппроксимации, определяющего значения пикселей в изображении низкого разрешения. При этом могут использоваться следующие алгоритмы: ближайшего соседа [1]; билинейная интерполяция [2]; бикубическая интерполяция [3]; свертка [4]; суперсемплинг [5].

В алгоритме ближайшего соседа пиксель итогового изображения низкого разрешения формируется путем копирования одного наиболее близкого по положению пикселя исходного изображения высокого разрешения. Недостатком алгоритма является потеря мелких деталей и зернистость итогового уменьшенного изображения. В алгоритме билинейной интерполяции для формирования каждого пикселя уменьшенного изображения используются 2×2 пикселей исходного изображения. Алгоритм нельзя применять для уменьшения более чем в 2 раза. При уменьшении до двух раз имеется заметный алиасинг. Бикубическая интерполяция работает по тому же принципу, что и билинейная, но для каждого пикселя уменьшенного изображения используются 4×4 пикселей исходного изображения. Алгоритм, использующий свертку, похож на алгоритмы билинейной и бикубической интерполяции, но вместо фиксированного количества пикселей используется количество, пропорциональное масштабу. Вклад каждого исходного пикселя в конечный определяется фильтром. При суперсемплинге исходное изображение разбивается на участки, равные количеству пикселей итогового изображения. Для формирования одного пикселя изображения низкого разрешения рассчитывается сумма всех пикселей, вошедших в площадь участка. Разбиение на участки может быть с округлением до ближайшего целого числа пикселей и без округления.

Рассмотренные алгоритмы позволяют изменить пространственное разрешение изображений, сохранив их структуру с необходимой точностью. Качество сформированного изображения будет тем выше, чем больше информации в него попадет из исходного изображения.

Целью работы является разработка алгоритма и программной модели формирования изображений с субпиксельным сдвигом и оценка его влияния на структуру изображений.

Алгоритм формирования изображений низкого пространственного разрешения

Предложенный алгоритм формирования изображений низкого пространственного разрешения, смещенных друг относительно друга на доли пикселя в низком разрешении, может использоваться для уменьшения разрешения изображений в $S = 2^N$ раз, где S – коэффициент масштабирования, а N – целое положительное число. При этом количество получаемых изображений совпадает с коэффициентом масштабирования. Для формирования серии перекрывающихся изображений низкого пространственного разрешения из изображения высокого пространственного разрешения коэффициент масштабирования S должен округляться до ближайшего целого с избытком. Блок-схема алгоритма представлена на рис. 1.

Исходными данными для алгоритма формирования изображений низкого пространственного разрешения является изображение высокого разрешения $I = \{i(y, x)\}_{y=0, \overline{Y-1}, x=0, \overline{X-1}, Y=X}$.

Предложенный алгоритм состоит из следующих основных шагов:

1. Обработка исходного изображения $I = \{i(y, x)\}_{y=0, \overline{Y-1}, x=0, \overline{X-1}, Y=X}$.

В первом шаге происходит загрузка исходного изображения $I = \{i(y, x)\}_{y=0, \overline{Y-1}, x=0, \overline{X-1}, Y=X}$ в формате bmp и определение размерности матрицы, содержащей исходное изображение.

2. Выбор пикселя $i(y, x)_{y=0, \overline{Y-1}, x=0, \overline{X-1}}$.

Пиксель $i(y, x)_{y=0, \overline{Y-1}, x=0, \overline{X-1}}$ выбирается в левом верхнем углу матрицы $I = \{i(y, x)\}_{y=0, \overline{Y-1}, x=0, \overline{X-1}, Y=X}$, содержащей исходное изображение.

3. Формирование матрицы $I_1 = \{i_1(y, x)\}_{y=0, \overline{\frac{Y}{S}-1}, x=0, \overline{\frac{X}{S}-1}}$, содержащей изображение низкого пространственного разрешения.

От выбранного пикселя строится матрица скользящего окна $I' = \{i'(y, x)\}_{y=0, \overline{S-1}, x=0, \overline{S-1}}$. Вычисление значения суперпикселя, которое является средним арифметическим значением яркости пикселей скользящего окна, выполняется по формуле $P = \frac{\sum_{n=0}^{S-1} \sum_{n=0}^{S-1} i'(y, x)}{S^2}$.

Полученное значение суперпикселя P заносится в матрицу $I_1 = \{i_1(y, x)\}_{y=0, \overline{\frac{Y}{S}-1}, x=0, \overline{\frac{X}{S}-1}}$. Для заполнения всей матрицы I_1 выполняется сдвиг матрицы I' на S пикселей слева направо сверху вниз, пока матрица I_1 не будет заполнена до конца.

4. Сохранение матрицы $I_1 = \{i_1(y, x)\}_{y=0, \overline{\frac{Y}{S}-1}, x=0, \overline{\frac{X}{S}-1}}$.

Матрица $I_1 = \{i_1(y, x)\}_{y=0, \overline{\frac{Y}{S}-1}, x=0, \overline{\frac{X}{S}-1}}$, содержащая изображение низкого пространственного разрешения сохраняется в виде изображения в формате bmp.

5. Формирование серии перекрывающихся изображений низкого пространственного разрешения $I_s = \{i_s(y, x)\}_{y=0, \overline{\frac{Y}{S}-1}, x=0, \overline{\frac{X}{S}-1}}$, при $s = \overline{2, S}$.

Для формирования серии перекрывающихся изображений низкого пространственного разрешения $I_s = \{i_s(y, x)\}_{y=0, \overline{\frac{Y}{S}-1}, x=0, \overline{\frac{X}{S}-1}}$ при $s = \overline{2, S}$ количеством S штук от выбранного пикселя

$i(y, x)_{y=0, \overline{Y-1}, x=0, \overline{X-1}}$ производится сдвиг на 1 пиксель по диагонали. Далее снова повторяется шаг 3,

с помощью которого формируется новое изображение низкого разрешения, сдвинутое от предыдущего на 1 пиксель. Алгоритм заканчивается, когда формируется S матриц низкого разрешения размерностью $\frac{Y}{S} \times \frac{X}{S}$ пикселей.

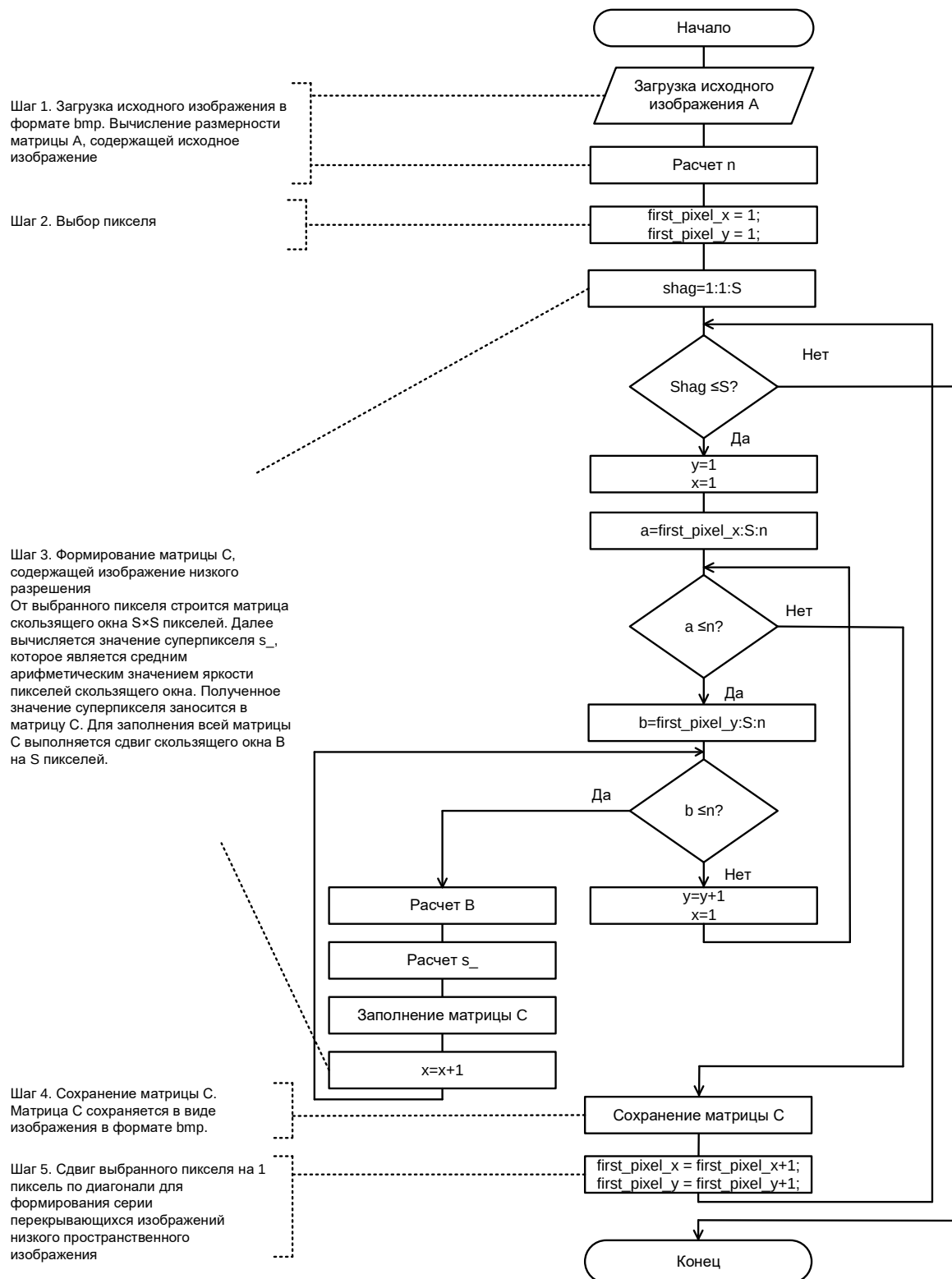


Рис. 1. Блок-схема алгоритма формирования изображений низкого пространственного разрешения

5. Формирование серии перекрывающихся изображений низкого пространственного разрешения $I_s = \{i_s(y, x)\}_{y=0, \frac{Y}{S}-1, x=0, \frac{X}{S}-1}$, при $s = \overline{2, S}$.

Для формирования серии перекрывающихся изображений низкого пространственного разрешения $I_s = \{i_s(y, x)\}_{y=0, \frac{Y}{S}-1, x=0, \frac{X}{S}-1}$ при $s = \overline{2, S}$ количеством S штук от выбранного пикселя

$i(y, x)_{y=0, Y-1, x=0, X-1}$ производится сдвиг на 1 пиксель по диагонали. Далее снова повторяется шаг 3, с помощью которого формируется новое изображение низкого разрешения, сдвинутое от предыдущего на 1 пиксель. Алгоритм заканчивается, когда формируется S матриц низкого разрешения размерностью $\frac{Y}{S} \times \frac{X}{S}$ пикселей.

Оценка эффективности алгоритма преобразования изображения высокого пространственного разрешения в серию перекрывающихся изображений низкого пространственного разрешения

Разработанный алгоритм реализован в среде Matlab. Для сравнительной оценки использованы алгоритмы билинейной и бикубической интерполяции. Эксперимент проведен на ЭВМ со следующими техническими характеристиками: процессор Intel(R) Core(TM) i5-3340M @ 3,4 ГГц; ОЗУ – 8 ГБ; тип системы – 64-разрядная операционная система Windows 10.

Для тестирования алгоритмов использовались изображения, представленные на рис. 2.

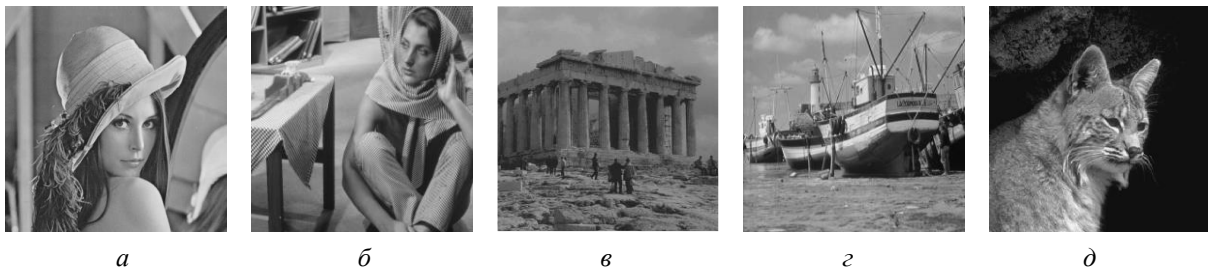


Рис. 2. Тестовые изображения высокого разрешения: a – изображение 1; $б$ – изображение 2; $в$ – изображение 3; $г$ – изображение 4; $д$ – изображение 5

В качестве критерия эффективности используется значение среднеквадратической ошибки (MSE). Значение среднеквадратической ошибки рассчитывается по следующим формулам:

1. Для оценки подобию изображений с разным сдвигом при использовании разработанного алгоритма:

$$MSE = \frac{1}{N^2} \sum_{i=0}^{n-1} (X_i - X'_i)^2,$$

где X_i – значение i -го пикселя изображения, сформированного предложенным алгоритмом, X'_i – значение i -го пикселя изображения, сдвинутого относительно изображения X_i на пиксель, N^2 – общее число пикселей изображения.

2. Для сравнительной работы тестируемых алгоритмов:

$$MSE = \frac{1}{N^2} \sum_{i=0}^{n-1} (X_i - X'_i)^2,$$

где X_i – значение i -го пикселя изображения, сформированного предложенным алгоритмом, X'_i – значение i -го пикселя изображения, сформированного алгоритмом билинейной или бикубической интерполяции, N^2 – общее число пикселей изображения.

На рис. 3 представлено исходное изображение высокого разрешения и изображение низкого пространственного разрешения, сформированное с помощью алгоритма формирования изображений низкого пространственного разрешения.



Рис. 3. Результат работы алгоритма формирования изображений низкого пространственного разрешения: a – исходное изображение высокого разрешения; b – изображение низкого пространственного разрешения

Так как предложенный алгоритм формирует серию перекрывающихся изображений низкого пространственного разрешения, то наличие перекрытия позволяет оценить подобие изображений. Для этого может использоваться значение среднеквадратичной ошибки.

Значение среднеквадратичной ошибки вычисляется от изображений с разным сдвигом. На рис. 4 представлены изображения низкого пространственного разрешения, сдвинутые на 1, 16 и 32 пикселя.

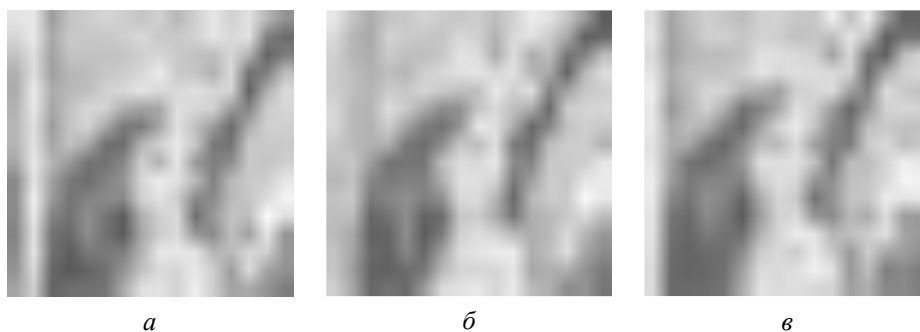


Рис. 4. Изображения низкого пространственного разрешения, сгенерированные предложенным алгоритмом, с разной величиной сдвига: a – 1 пикселей; b – 16 пикселей; v – 32 пикселя

Зависимость среднеквадратичной ошибки MSE серии перекрывающихся изображений от величины диагонального сдвига между ними (n) приведена на рис. 5.

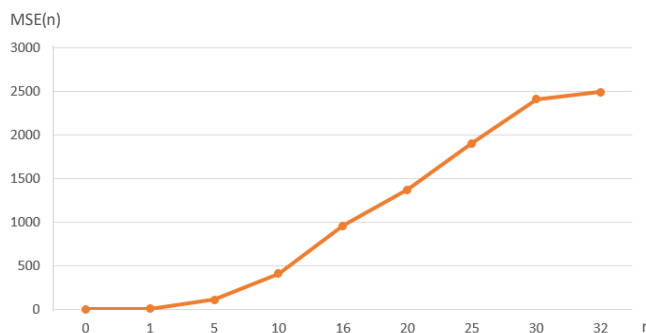


Рис. 5. Зависимость среднеквадратичной ошибки изображений от величины диагонального сдвига

При сдвиге полученных изображений низкого разрешения на 1 пиксель значение среднеквадратичной ошибки равняется 4,4727, что свидетельствует о высокой степени идентичности изображений. Увеличение сдвига приводит к увеличению значения среднеквадратичной ошибки и, следовательно, к снижению схожести изображений.

На рис. 6 представлены результаты обработки изображений с помощью алгоритма преобразования изображения высокого пространственного разрешения в серию перекрывающихся изображений низкого пространственного разрешения и алгоритмов билинейной и бикубической интерполяции.

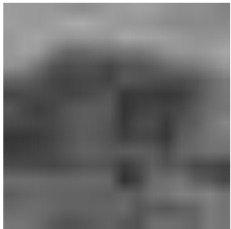
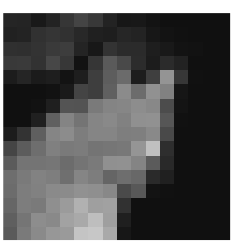
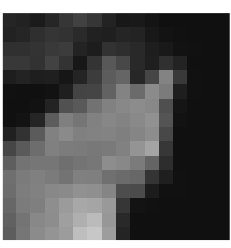
Исходное изображение ВР	Предложенный алгоритм	Билинейная интерполяция	Бикубическая интерполяция
			
			
			
			
			

Рис. 6. Результаты обработки изображений

В таблице представлены полученные значения MSE при сравнении изображений, сгенерированных с помощью алгоритма формирования изображений низкого пространственного разрешения, относительно изображений, полученных при использовании алгоритмов билинейной и бикубической интерполяции.

Рассчитанное значение MSE для тестовых изображений

Тестовое изображение	Алгоритм интерполяции	
	Билинейная интерполяция	Бикубическая интерполяция
Изображение 1	48,4570	17,8945
Изображение 2	44,7461	19,5312
Изображение 3	10,4102	8,8906
Изображение 4	29,1563	15,6523
Изображение 5	16,5469	8,1133

Среднее значение MSE при сравнении предложенного алгоритма с алгоритмом билинейной интерполяции составило 29,8633. При сравнении предложенного алгоритма с алгоритмом бикубической интерполяции среднее значение MSE равняется 14,0164. Сформированное предложенным алгоритмом изображение близко по качеству к изображению, полученному алгоритмом бикубической интерполяции, как по значению MSE, так и по визуальной оценке.

Заключение

Предложен алгоритм формирования изображений низкого пространственного разрешения. Показано, что при формировании серии изображений низкого пространственного разрешения, при увеличении сдвига между изображениями, их идентичность в среднем снижается в 2 раза. При сравнении с алгоритмами билинейной и бикубической интерполяции, предложенный алгоритм формирует изображения наиболее близкие по значению MSE к методу бикубической интерполяции.

ESTIMATION OF IMAGE DISTORTIONS
WITH SUB-PIXEL SHIFT OF THE CAMERA

A.V. ZAKHARENKO, O.G. SHEVCHUK, V.Yu. TSVIATKOU

Abstract. An algorithm for the formation of images of low spatial resolution has been developed. The algorithm is based on the use of high spatial resolution images to generate low spatial resolution images offset from each other by fractions of a pixel in low resolution. The estimation of distortions of images at subpixel shift using the root mean square error has been made.

Keywords: subpixel shift, images, MSE.

Список литературы

1. Blu T. // IEEE Transactions on Image Processing. 2004. Vol. 13. P. 710–719.
2. Thevenaz P., Blu T. // IEEE Transactions on Medical Imaging. 2000. Vol. 19. P. 739–758.
3. Wilhelm B., Burge M. J. // Springer-Verlag. 2009. Vol. 7. P. 327
4. Turkowski K. // Graphics gems. 1990. P. 147–165.
5. Yang W., Zhang X., // IEEE Transactions on Image Processing. 2018. Vol. 20. P. 512– 520.

УДК 621.391

АЛГОРИТМЫ ОБРАБОТКИ ИЗОБРАЖЕНИЙ В КОМБИНИРОВАННЫХ ОПТИКО-ЭЛЕКТРОННЫХ СИСТЕМАХ ЗАЩИТЫ ОБЪЕКТОВ ОТ БЕСПИЛОТНЫХ ЛЕТАТЕЛЬНЫХ АППАРАТОВ

О.А. АЗАРЧИК, Е.А. МАШКОВ, Н.С. ДАВЫДОВА

*Белорусский государственный университет информатики и радиоэлектроники, Республика Беларусь**Поступила в редакцию 6 марта 2021*

Аннотация. Рассматриваются подходы к разработке алгоритмического обеспечения оптико-электронных систем в части технического зрения. Целью данной работы является обзор алгоритмов обработки изображений комбинированных оптико-электронных систем защиты объектов от беспилотных летательных аппаратов для оценки и анализа эффективности детектирования объектов и расширения области применения подобных систем.

Ключевые слова: обработка изображений, беспилотный летательный аппарат, защита объектов.

Введение

Стремительное развитие технологий беспилотной авиации значительно повысило автономность, дальность действия и спектр решаемых беспилотными летательными аппаратами (БЛА) задач. В связи с этим особенно остро встает вопрос защиты важных объектов от несанкционированного проникновения БЛА [1]. При этом малый размер, высокая маневренность и автономность БЛА создают необходимость в разработке сложных многоступенчатых автоматических комплексов обнаружения, распознавания и подавления бортовых систем БЛА.

Комбинированная оптико-электронная система (КОЭС) предназначена для автоматического обнаружения, селекции, распознавания и автосопровождения объектов. Система может выполнять обнаружение и селекцию БЛА как автономно, так и по пеленгу от системы радиопеленгации, дальность обнаружения зависит от размеров БЛА. КОЭС состоит из широкоугольной панорамной (ОЭС ШУ) и узкоугольной подсистем (ОЭС УУ) оптико-электронного мониторинга. ОЭС ШУ при обнаружении несанкционированных БЛА в охраняемой зоне передает угловые координаты всех обнаруженных объектов ОЭС УУ. ОЭС УУ предназначена для автоматического обнаружения, распознавания и сопровождения объектов. ОЭС УУ принимает угловые координаты от ОЭС ШУ, производит поиск объектов по предоставленным пеленгам, их автоматическое распознавание и, в случае классификации объекта как БЛА, переходит в режим автоматического сопровождения (рис. 1).

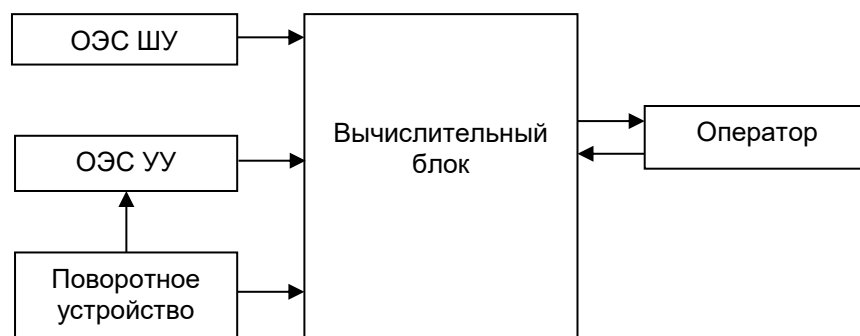


Рис. 1. Структурная схема КОЭС

1 Алгоритм формирования панорамного изображения

Данный алгоритм осуществляет совмещение изображений, получаемых от камер кругового обзора в единое бесшовное панорамное изображение, и состоит из нескольких этапов:

- устранение дисторсий на получаемых с камер изображениях;
- автоматическая коррекция яркости и контраста изображений;
- объединение исправленных изображений в единое панорамное изображение.

Для устранения дисторсий используется стандартная модель аппроксимации радиальной дисторсии, описываемая выражением

$$\begin{pmatrix} x - x_c \\ y - y_c \end{pmatrix} = L(r) \begin{pmatrix} x - x_c \\ y - y_c \end{pmatrix}, \quad (1)$$

где (x, y) – координаты искаженной точки изображения, (x_c, y_c) – координаты исправленной точки изображения; (x_c, y_c) – центр модели дисторсии фотоприемника, обычно расположенный в центре изображения. Дистанция от каждой точки изображения до центра модели радиальной дисторсии фотоприемника определяется по формуле

$$r = \sqrt{(x - x_c)^2 + (y - y_c)^2}. \quad (2)$$

$L(r)$ – функция, определяющая форму дисторсий в применяемой модели, обычно аппроксимируется рядом Тейлора как

$$L(r) = k_0 + k_1 r + k_2 r^2 + k_3 r^3 + \dots, \quad (3)$$

где набор $k = (k_0, k_1, \dots, k_N)^T$ определяет параметры дисторсии. Точность модели определяется числом используемых членов аппроксимирующего ряда.

Для определения параметров модели k (калибровки) используется общепринятый подход [2], в основе которого лежит наложение требований на проекции трехмерных линий, которые должны быть прямыми линиями на изображениях.

Калибровка выполняется один раз, после чего одни и те же калибровочные параметры используются для устранения дисторсий на каждом изображении, поступающем с откалиброванного фотоприемника. После определения параметров модели дисторсии выражение (1) применяется для формирования исправленного изображения.

Для всех исправленных изображений выполняется эквализация их гистограмм [3], а также применяются другие алгоритмы улучшения изображений в случае необходимости, после из них формируется результирующее панорамное изображение.

2 Алгоритм автоматического обнаружения объектов на панорамном изображении

Для автоматического обнаружения малоразмерных объектов на панораме, формируемой со стационарно установленной оптико-электронной системы, применяется комплексный двухэтапный алгоритм. На первом этапе алгоритма производится вычитание усредненного панорамного изображения фоновой обстановки из текущего панорамного изображения. По результатам вычитания формируется бинарное разностное изображение, на котором нули означают фон, а единицы – подозрительные на объект зоны. Усредненное панорамное изображение формируется по следующей формуле:

$$bg_i = bg_{i-1} \cdot \lambda + Im_i \cdot (1 - \lambda), \quad (4)$$

где bg_i – усредненное панорамное изображение на текущем кадре; bg_{i-1} – усредненное панорамное изображение на предыдущем кадре; Im_i – текущее панорамное изображение, а $\lambda \in [0; 1]$ устанавливает время накопления фона, что в свою очередь влияет на соотношение (точность обнаружения) / (число ложных срабатываний). Абсолютная разность между текущим панорамным изображением и усредненным панорамным изображением определяется как

$$d = |bg_i - Im_i|. \quad (5)$$

Абсолютная разность вычисляется для каждого пикселя панорамного изображения и если полученная разница выше порогового значения, $d \geq t$, то в результирующее бинарное изображение записывается единица, иначе – ноль.

Для определения порогового значения вычисляется средняя абсолютная разность между текущим панорамным изображением и усредненным панорамным изображением:

$$d_{avg} = \sum_{i=0, j=0}^{i=M, j=N} d_{ij}, \quad (6)$$

где M – ширина изображения; N – высота изображения; d_{ij} – модуль абсолютной разности между текущим панорамным изображением и усредненным панорамным изображением в точке (i, j) .

Таким образом, порог минимальной абсолютной разности можно определить как

$$t = d_{avg} \cdot w, \quad (7)$$

где w – корректирующий коэффициент, определяемый на этапе настройки системы.

Полученное разностное бинарное изображение, после обработки с помощью морфологических операций [4], является входными данными для второго этапа алгоритма.

На втором этапе применяется алгоритм автоматического обнаружения целей, подробно описанный в [5]. Данный алгоритм обнаружения целей основан на построении модели ключевых объектов сцены с использованием данных о количестве, положении и характеристиках объектов, получаемых при последовательной обработке каждого кадра видеопоследовательности. Этот способ предполагает два этапа анализа – пространственный и временной. На этапе пространственного анализа происходит обработка текущего кадра видеопоследовательности, а его результатом является некоторый список подозрительных объектов. На этапе временного анализа результаты пространственного анализа сравниваются с текущей моделью ключевых объектов сцены, после чего модель уточняется и обновляется. Результатом работы двухэтапного алгоритма автоматического обнаружения объектов на панорамном изображении является список объектов, угловые координаты которых передаются на дополнительный анализ в систему ОЭС УУ.

3 Алгоритмическое обеспечение ОЭС УУ

Получив угловые координаты от ОЭС ШУ, ОЭС УУ разворачивается по полученным угловым координатам и осуществляет автоматический поиск, распознавание и захват на сопровождение объекта в указанных координатах.

Алгоритмическое обеспечение ОЭС УУ состоит из нескольких ключевых частей:

- алгоритм автоматического обнаружения объектов в телевизионном (ТВ) и инфракрасном (ИК) каналах;
- алгоритм автоматического распознавания объектов;
- алгоритм мультиспектрального комплексирования ТВ- и ИК-каналов;
- алгоритм автоматического сопровождения объектов в ТВ- и ИК-каналах.

Алгоритм автоматического обнаружения объектов в ТВ- и ИК-каналах. Данный алгоритм схож со вторым этапом алгоритма автоматического обнаружения целей ОЭС ШУ [5], описанным ранее, с той лишь разницей, что бинарное изображение для анализа строится с применением адаптивного алгоритма бинаризации на основе анализа локальных контрастов [6–8]. Подобный алгоритм бинаризации позволяет применять его как для обработки изображений, полученных в видимом диапазоне, так и для инфракрасных изображений. Его отличительным свойством является инвариантность к четкости контуров объекта и его освещенности, необходимо лишь, чтобы объект обладал достаточным локальным контрастом на изображении.

Алгоритм автоматического распознавания объектов. Каждый обнаруженный на предыдущем этапе обработки объект подвергается процедуре автоматического распознавания.

Для автоматического распознавания используется алгоритм с применением сверточных нейронных сетей [9–10].

Сети сверточной архитектуры обычно состоят из нескольких идущих друг за другом слоев. Существует множество типов слоев в нейронной сети, наиболее часто используемыми являются сверточные слои, слои активации, слои субдискретизации, слои батч-нормализации, полносвязные слои и т.д.

Основой сверточных нейронных сетей являются слои свертки. Отличие данного слоя от полносвязного слоя заключается в том, что каждый нейрон связан только с ограниченным числом соседних нейронов, что позволяет существенно уменьшить число параметров сети. Веса данного слоя представляют собой набор многомерных изображений-фильтров с числом каналов, равным числу каналов входного изображения. В процессе прямого прохода данный фильтр двигают по входному изображению и вычисляют сумму попарных произведений элементов фильтра и значений входного изображения. В итоге на выходе блока получается изображение с количеством каналов, равным количеству фильтров блока.

Операцию свертки можно свести к операции матричного умножения, если предварительно преобразовать входное изображение в матрицу, в которой соседние элементы изображения оказываются в одном столбце. Данное преобразование называется "image to column".

Алгоритм мультиспектрального комплексирования ТВ- и ИК-каналов. Данный алгоритм осуществляет геометрическое и информационное совмещение изображений, получаемых в двух различных диапазонах видимости, с целью улучшения восприятия оператором (в особенности при плохой освещенности или сложных погодных условиях). Метод мультиспектрального комплексирования подробно описан в [11] и основан на использовании вейвлет преобразования [12]. В отличие от Фурье-преобразований, вейвлет-базисные функции являются хорошо локализованными, что дает возможность проводить локальный спектральный анализ [13]. Спектральные вейвлет-коэффициенты соответствуют не только амплитудам различных частот, но и различным пространственным участкам на изображении.

Для осуществления быстрого алгоритма вычисления вейвлет-преобразования в качестве базиса применяется система вейвлетов Хаара [14]. Данная система вейвлетов требует минимум вычислений, что немаловажно в условиях требований выполнения операций совмещения изображений в реальном масштабе времени.

Алгоритм автоматического сопровождения объектов в ТВ- и ИК-каналах. Для реализации устойчивого автоматического сопровождения малоразмерных подвижных объектов используется комбинированный корреляционно-признаковый алгоритм. Алгоритм основан на параллельном применении двух различных подходов к сопровождению.

В качестве результата автоматического обнаружения и распознавания формируется эталон цели, получаемый из изображения в видимом или инфракрасном спектре либо на основе комплексированного мультиспектрального изображения (в зависимости от времени суток и установок оператора). Полученный эталон используется в качестве шаблона для корреляционно-экстремального алгоритма поиска с применением нормализованной кросс-корреляции [15, 16] в качестве меры сходства шаблона (эталона) и участка изображения в проверяемых координатах. Поскольку объект претерпевает эволюции в процессе сопровождения (изменения масштаба и ракурса, развороты), необходимо периодически обновлять используемый эталон с заданным периодом либо по значению коэффициента корреляции. Для перезаписи используемого эталона и уточнения границ и размеров сопровождаемого объекта применяется признаковый алгоритм автоматического обнаружения и сопровождения, описанный ранее [5], который периодически определяет и уточняет строб сопровождения цели (описывающий прямоугольник). По уточненному стробу сопровождения производится перезапись эталона для корреляционного сопровождения.

Заключение

В статье был проведен обзор алгоритмов обработки изображений комбинированных оптико-электронных систем защиты объектов от БЛА для оценки и анализа эффективности детектирования объектов и расширения области применения подобных систем. Рассмотренные алгоритмы частично либо полностью реализованы в настоящее время в следующих комплексах

защиты объектов от БЛА: комплекс для защиты стратегических объектов от мультикоптеров "Гроза-3" производства КБ "Радар" (Республика Беларусь), комплекс "Силок" производства ООО "Специальный технологический центр" (Российская Федерация) и др.

Таким образом, комбинированные оптико-электронные системы позволяют эффективно обнаруживать и сопровождать воздушные объекты, что особенно актуально в случаях, когда БЛА совершают полет в режиме "радиомолчания" без обмена сигналами с пультом оператора.

IMAGE PROCESSING ALGORITHMS IN COMBINED OPTOELECTRONIC SYSTEMS FOR PROTECTING OBJECTS FROM UNMANNED AERIAL VEHICLES

O.A. AZARCHIK, E.A. MASHKOV, N.S. DAVYDOVA

Abstract. Approaches to the development of algorithmic support for optoelectronic systems in terms of technical vision are considered. The purpose of this work is to review the image processing algorithms of combined optoelectronic systems for protecting objects from unmanned aerial vehicles to evaluate and analyze the effectiveness of object detection and expand the scope of such systems.

Keywords: image processing, unmanned aerial vehicle, object protection.

Список литературы

1. Мясников Е.В. Угроза терроризма с использованием беспилотных летательных аппаратов: технические аспекты проблемы. 2004.
2. Luis Alvarez, Luis Gomez, and J. Rafael Sendra. Algebraic Lens Distortion Model Estimation. 2010.
3. Gonzalez R.C., Woods R.E. Digital Image Processing. 2002. P. 91–94.
4. Фурман Я.А., Юрьев А.Н., Яншин В.В. Цифровые методы обработки и распознавания бинарных изображений. 1992.
5. Тупиков В.А., Павлова В.А., Бондаренко В.А., [и др.] Способ автоматического обнаружения объектов на морской поверхности в видимом диапазоне. 2016. № 11. С. 105–121.
6. Sauvola J., Pietikainen M. Adaptive document image binarization. 2000.
7. Bradley D., Roth G. Adaptive Thresholding Using Integral Image. 2007. Vol. 12. P. 13–21.
8. Shafait F., Keysers D., Breuel T.M. Efficient implementation of local adaptive thresholding techniques using integral images. 2008.
9. Goodfellow I., Bengio Y., Courville A. Deep Learning. 2016.
10. Deng L. and Yu. D. Deep Learning: Methods and Applications. 2013. Vol. 7. P. 197–387.
11. Фролов В.Н., Тупиков В.А., Павлова В.А., [и др.] Методы информационного совмещения изображений в многоканальных оптико-электронных системах. 2016.
12. Тетерин В.В. [и др.] Метод комплексирования информации от многоканальной системы с использованием вейвлет-спектров. 2006.
13. Петухов А.П. Введение в теорию базисов всплесков. 1999.
14. Воробьев В.И., Грибунин В.Г. Теория и практика вейвлет-преобразования. 1999.
15. Баклицкий В.К. Корреляционно-экстремальные методы навигации и наведения. 2009.
16. Алпатов Б.А., Бабаян П.В., Балашов О.Е., [и др.] Методы автоматического обнаружения и сопровождения объектов. 2008.

ЗАЩИТА ИНФОРМАЦИИ НА ОСНОВЕ СПЕКТРАЛЬНО-ПРОСТРАНСТВЕННОГО КОДИРОВАНИЯ

А.И. МИТЮХИН

Белорусский государственный университет информатики и радиоэлектроники, Республика Беларусь

Поступила в редакцию 9 марта 2021

Аннотация. Рассматривается подход решения задачи защиты информации путем применения методов частотного преобразования и помехоустойчивого кодирования. Показано, что предлагаемый вычислительный алгоритм на основе линейного ортогонального дискретного преобразования и низкоскоростного кодирования над полем целых чисел обеспечивает высокую степень скрытности сигнала. Метод защиты информации основывается на практическом подходе из теории информации.

Ключевые слова: обнаружение, кодирование, спектр, дискретное преобразование, скрытность, корреляция, источник, шум, полином.

Введение

Одним из основных требований при проектировании специальных инфокоммуникационных систем является эффективная и надежная защита пространственно-временного сигнала от перехвата. Современные системы перехвата включают в себя аппаратно-программные средства для спектрального анализа передаваемого сигнала и высокопроизводительные средства радиоразведки параметров сигнала. К таким основным параметрам относятся период излучаемого кодированного сигнала, вид модуляции, класс кодов, конкретная структура порождающей / проверочной матрицы или полинома и др. Имея достаточно мощные специальные аппаратно-программные средства, можно с высокой достоверностью решать задачи, связанные с перехватом информации. При этом необходимо учитывать вид (класс) наблюдаемой системы, когда успешному перехвату, может предшествовать решение задачи обнаружения, различения, распознавания, декодирования сигнала. Задача перехвата существенно усложняется, когда на получение необходимых энергетических, пространственных, временных, спектральных, поляризационных признаков сигналов требуются временные затраты, превосходящие сеанс перехватываемой связи.

Известно [1], требование надежной защиты может выполняться при условии применения низкоскоростных помехоустойчивых кодов. В этом случае можно уменьшить спектральную плотность мощности сигнала до величины, обеспечивающей его энергетическую скрытность. Кроме того, вероятность обнаружения сигнала усложняется, если применять сигналы, модулированные кодом с равномерным спектром в полосе частот канала.

В статье описывается вычислительный алгоритм защиты информации, где используется спектральное описание сигнала, а затем кодирование квантованных коэффициентов преобразования низкоскоростным кодом.

Теоретические принципы

Пусть сигнал с шириной спектра W передается в течении временного интервала T . Канальный параметр $q = \frac{S}{N}$ характеризует отношение средней мощности S сигнала к средней мощности $N = N_0 W$ шума на входе канала перехватчика. Шум описывается равномерным распределением плотности мощности N_0 в полосе W . Физические параметры T , W , S и N описывают основной канал. Обобщенно свойство канала можно представить в виде его геометрической интерпретации – объема

$$V = WT \frac{S}{N} = WTq. \quad (1)$$

Канал перехвата также характеризуется физическими параметрами. Перехват усложняется, если временные затраты T_c на обнаружение сигнала, выявление способов модуляции, декодирования и пр. превышают T , т.е. $T_c > T$. Поиск сигнала по частоте потребует производить частотное сканирование по несущей и тактовой частоте в канале с шумами при отношении

$$\frac{S}{N_0} \ll 1. \quad (2)$$

Очевидно, при отсутствии априорных знаний о частотных параметрах: несущей частоты и ширины спектра W (тактовой частоты) возникают дополнительные временные затраты на прием сигнала.

По аналогии с (1) свойство канала перехвата представим в виде величины его объема

$$V_c = W_c T_c q_c.$$

Для надежной передачи информации по основному каналу необходимо, чтобы выполнялось условие $V_c < V$. Возможное выполнение этого неравенства можно рассмотреть, используя два подхода.

1. Практическая стратегия надежной защиты инфокоммуникационной системы должна строиться на ограничении времени t передачи открытой информации, когда $t \ll T_c$. Решение этой задачи возможно, применяя алгоритмы эффективного кодирования. Например, такие как энтропийный, универсальный, арифметический, спектральный. Недостатком первых трех является сравнительно низкая эффективность.

Пусть качестве исходного информационного источника (1D или 2D) рассматривается дискретный источник информации без памяти, где множество $\{x_1, x_2, \dots, x_m\} \in X$ – это алфавит источника. Источник описывается распределением вероятностей $P = \{p_1, p_2, \dots, p_m\}$, где p_i – вероятность появления символа x_i . Примем, что i -у символу источника X соответствует X_i -е информационное слово источника одиночных символов на множестве $\{1, -1\}$. Длина слова этого равномерного кода равна k . Если учитывать статистические характеристики источника (априорное знание распределения P) появляется возможность более эффективно использовать сигнал объемом V .

2. Уменьшая временной параметр объема V , появляется возможность применения метода широкополосного кодирования, т.е. увеличения значения составляющей W объема V . Свойство

$$\frac{k}{n} \ll 1 \quad (3)$$

широкополосного кода на втором этапе кодирования источника может обеспечить высокую степень информационной безопасности за счет необходимости при перехвате решать задачу обнаружения сигнала. Очевидно, увеличение объема V_c приведет к усложнению выполнения условия $V_c < V$ и решению задачи перехвата с учетом условий (2) и (3).

Спектральное кодирование. С целью обеспечения экономии вычислительных и временных ресурсов кодер источника строится с использованием ортогонального действительного дискретного преобразования Хартли (ДПХ). Ортогональный базис дискретных функций Хартли на интервале из N точек выражается числами вида [2]

$$h_v = \cos\left(\frac{2\pi}{N}nv\right), \quad n, v \in \{0, 1, \dots, N-1\},$$

где N – период дискретного сигнала $g_n = (g_0, g_1, \dots, g_{N-1})$ источника. Аргументы n и v определяют соответственно временной (пространственный) и частотный (частотно-

пространственный) параметры сигнала. 1-D спектральное преобразование по всей пространственной области выполняется как [2]

$$\hat{g}_v = \frac{1}{N} \sum_{n=0}^{N-1} g_n \cos\left(\frac{2\pi}{N} nv\right), v \in \{0, 1, \dots, N-1\}. \quad (4)$$

Вычисление (4) выполняется с помощью быстрого алгоритма типа БПФ. Сокращение времени передачи информации сводится к отбору коэффициентов преобразования (4) с наибольшей дисперсией. Для того, чтобы не передавать служебную информацию о номере v отобранных коэффициентов $\hat{g}_v$, предлагается использовать зональный метод отбора (фильтрации) [3]. Учитывая знание распределения вероятностей P источника X , вид ковариационной матрицы реального сигнала (например, речи, контура изображения объекта и пр.), зона фильтрации определяется через вычисление 1D- или 2D-функции распределения дисперсий компонент спектра [4]. В спектре случайного сигнала источника, как правило, имеется составляющая, соответствующая значению его математического ожидания. Компонента $\hat{g}_v$ на частоте $v=0$ проявляется в большей степени, чем остальные составляющие спектрального образа сигнала. Поэтому преобразование (4) выполняется для процесса с нулевым математическим ожиданием.

Пространственное кодирование. Второй этап кодирования связан с выбором $[n, k, d]$ -кода, обеспечивающего получение псевдослучайного потока двоичных символов и тем самым структурную скрытность. В спектральной области выбор класса кода связан с необходимостью иметь сравнительно равномерный спектр для всех форм кодовых слов кода. Возможность обнаружения кодированного сигнала сводится к минимуму, если в наблюдаемой полосе частот отсутствуют периодически повторяющиеся спектральные точки. Известно [5], чем меньше боковые остатки функция автокорреляции сигнала (АКФ), тем более равномерным амплитудным спектром описывается сигнал. В наибольшей степени этому условию удовлетворяет сигнал, модулированный симплексным кодом Рида-Маллера первого порядка $RM(1, k)$ [6]. Для двужначного алфавита $(1, -1)$ АКФ $r(\tau)$ имеет только два значения: $r(\tau) = 1, \tau = 0$ и $r(\tau) = -\frac{1}{n}, 0 \leq \tau \leq n-1$. Возможность эффективного обнаружения и декодирования сигнала с малой вероятностью ошибки при энергетическом условии (2) в полосе приема практически сложно осуществить. Выбор этого кода обусловлен и тем, что при условии (3) практически восстановить порождающий полином кода не представляется возможным. Структуру полинома можно раскрыть путем решения системы линейных уравнений. Однако, это возможно при условии последовательного правильного приема всех символов сегмента кода длиной $2k$. Например, в основном канале выбран для применения $RM(1, 7)$ -код, вероятность ошибки в канале перехвата равна $p = 0,1$. Для этого условия вектор ошибок $\mathbf{E}$ имеет вес $wt(\mathbf{E}) \equiv 13$. Вероятность последовательного правильного приема 14-и символов близка к нулевому значению. Заметим, требование надежной передачи информации в реальных энергетически скрытных каналах предполагает работу с уровнем шума, при котором на входе приемного устройства перехвата вероятность ошибки $p \geq 0,1$.

Экспериментальные исследования

В качестве исходных данных, требующих эффективного описания и защиты от перехвата, использовалось бинарное изображение границы некоторого объекта интереса. После аналого-цифрового преобразования и бинаризации изображения получались пространственные данные в виде целочисленных пар (точек) декартового произведения. Число точек, описывающих границу, составляло $N = 16$. В пространственной области для описания границы требовалось 32 числа. Понижение размерности входа кодера $RM(1, k)$ -кода за счет выполнения процедуры эффективного кодирования привело к значению $k = 7$. Проверка восстановления исходных данных путем вычисления обратного спектрального преобразования по 7-и коэффициентам

Хартли подтвердила получение всех 32 чисел с нулевым значением СКО. Далее осуществлялось кодирование 7-и информационных символов произвольно выбранной псевдослучайной последовательностью [127, 7, 64]-кода. Проведенные в работе экспериментальные исследования спектральных особенностей кода показали, что, начиная от значения $n \geq 63$ во всей полосе частот интенсивность коэффициентов Фурье равномерно уменьшается без явно проявляемых выбросов. В канале перехвата формировалась случайная аддитивная смесь сигнала и помехи в виде гауссовского процесса. При этом величина мощности помехи в полосе сигнала превышала на $(0 \div 10)$ дБ мощность сигнала. Экспериментальная оценка вероятности ошибки в канале перехвата давала значения $(2 \cdot 10^{-1} \div 4 \cdot 10^{-1})$.

Заключение

Предложенный вычислительный алгоритм защиты данных на основе кодирования в спектральной и пространственной областях позволяет получить определенную степень информационной безопасности. Не располагая сведениями о структурах порождающих полиномов кодов, перехватчик не может осуществить надежный прием информации, используя неоптимальные алгоритмы декодирования кодов. Исследования показали, что увеличение величины длины кода не приводит к резкому увеличению вероятности ошибки в канале перехвата, но приводит к уменьшению вероятности обнаружения сигнала в смеси сигнал/шум. Анализ формы спектра входного случайного процесса в канале перехвата показывает, что вероятность перехвата уменьшается при большем отношении сигнал/шум, но с условием увеличения длины кода. Дальнейшие исследования свойств алгоритма защиты могут быть продолжены, если для пространственного кодирования использовать сингулярные преобразования $RM(1, k)$ -кода.

PROTECTION OF INFORMATION BASED ON SPECTRAL-SPATIAL CODING

A.I. MITSUKHIN

Abstract. We are considering an approach to solving the problem of information protection through the use of frequency conversion methods and interference-resistant coding. It is shown that the proposed computational algorithm based on linear orthogonal discrete conversion and low-grown code over the field of whole numbers provides a high degree of stealth signal. The method of information protection is based on a practical approach from the theory of information.

Keywords: detection, coding, spectrum, discrete conversion, stealth, correlation, source, noise, polynomial.

Список литературы

1. Ipatov V. // Spread Spektrum and CDMA. John Wiley & Sons, Ltd, 2005.
2. Митюхин А. // Докл. БГУИР, 2018, № 7 (117). С. 74–79.
3. Гонсалес Р, Вудс Р. Цифровая обработка изображений. М., 2005.
4. Mitsukhin A., Pachynin V, Petrovskaya E. // Proc. 52. IWK. 2007. Vol. 2. P. 321–325.
5. Пестряков В. Б., Афанасьев В.П., Гурвиц В. Л. [и др.] Шумоподобные сигналы в системах передачи информации. М., 1973.
6. Мак-Вильямс Ф.Дж., Слоэн Н.Дж. Теория кодов, исправляющих ошибки. М., 1979.

METHOD FOR GENERATING TWO-DIMENSIONAL DEPENDENT ERRORS

REN XUNHUAN, V.K. KONOPELKO, V.Yu. TSVIATKOU

*Belarusian State University of Informatics and Radioelectronics, Republic of Belarus**Submitted 12 March 2021*

Abstract. The method of forming patterns of dependent two-dimensional errors is considered. It is shown that the algorithm works much faster than the known algorithms for generating dependent two-dimensional errors.

Keywords: dependent error, error pattern set, error pattern generation.

Introduction

Currently, error correcting code (ECC) is widely used in many fields, such as in information processing systems, in computer memory devices, in telecommunications systems, in missile control systems, in information compression systems, etc. Decoding by the syndrome is used to correct random errors of small multiplicity, the length of the verification characters will be increased with the multiplying errors. Obviously, the dependent errors can't fix decoding by the syndrome. In the theory of two-dimensional coding, the high efficiency of using the error correction pattern library is noted.

In [1, 3] proposed a method generate dependent errors patterns which is based on the process of calculating the pattern of the $t \times t$ type library and applying the identification parameters to recognise the errors. The method classified all of the forming error patterns in typical and non-typical, but with the increasing of t , the computational complexity increases rapidly. Based on this defect, a new method of generate the two-dimensional dependent errors patterns is proposed.

Method for generating two-dimensional dependent errors

The error, caused by the unknown interference in the channel, appeared in the receive device can be divided into random and dependent. Since the random errors are independently happened, the multiplicity of errors t over the length of the message are usually used to evaluate this type of error. Dependent errors can be further divided into batch errors and modular errors according to the degree of the generality. In fact, the modular errors are the special case of the batch errors. Batch errors may cause by the events such as the aging in communication cable, interference from periodic noise, dust particles on magnetic tape, etc. The modular errors are phased burst errors, which can be observed in memory systems built on multi-bit LSI memory. A good example is in a 16-bit memory system implemented on microcircuits with four-bit outputs, modular errors of length due to failures of individual memory LSIS [4]. The length of the modular errors is denoted by b in this paper.

In addition, batch errors and modular errors can be either single or multiple. Supposing there are some codes whose n is 16 and b is 4. Then supposing there are two error pattern which we detected which are $E = (0000\ 0000\ 1101\ 0000)$ and $E = (0000\ 1111\ 0000\ 1100)$. It is easy to find that the former error pattern has one error segment, in which there are four bits. Whereas, in the latter one, there are two error segments [4].

If two error patterns can become the same pattern by conducting a special rules set which given later to filter the similar error patterns, we will categorize them into one type. Based on our method, the following 8 error patterns – $E = \{(11\ 00), (00\ 11), (10\ 10), (01\ 01), (10\ 01), (01\ 10), (11\ 10), (11\ 11)\}$, which are the all-possible potential error patterns $2 \leq t \leq 4$, will reduced to five patterns. The procedure of the reducing presented in fig. 1.

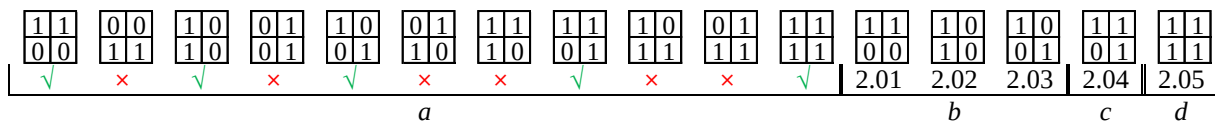


Fig. 1. procedure of the reducing size 2×2 : *a* – original error patterns; *b* – dependent error patterns when $t = 2$; *c* – two dependent error patterns when $t = 3$; *d* – two dependent error patterns when $t = 4$

Since special rules set eliminate the redundant error patterns, the leftover patterns of error can be regarded as the basis set, which will use to extend the higher order error patterns. This procedure will first extend the row and column of each element in basis set and add "1" (error) to each site where the value is "0". After that, we again applied our rule set to those extended error pattern to acquire the higher order basis set. The whole procedure can be viewed in the fig. 2.



Fig. 2. Procedure of extending basic set of size 2×2 to pattern errors of all size 3×3 and the procedure of reducing error patterns of order 3 to basis set of order $t = 3:5$

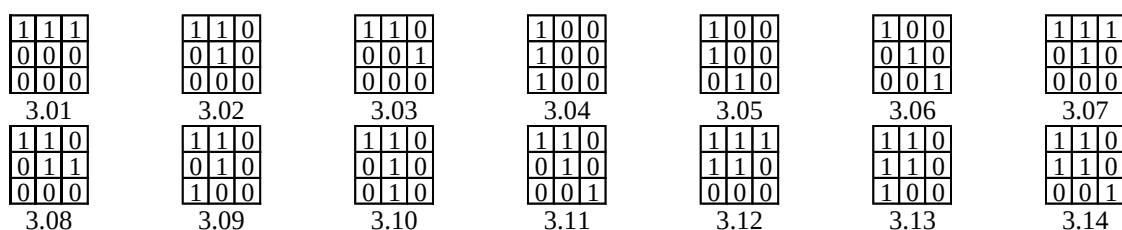


Fig. 3. From the basic set of All the basis set of order $t \times t = 2 \times 2$ chose the library of $t \times t = 3 \times 3$

From the discussion above, it is intuitive to noticed that the basis set of order 2×2 can derived 14 distinct error patterns of order 3×3 , which have 5 maximum errors.

Special rules set for two-dimensional dependent errors

The rules set are worked according to the parameters of the error pattern. If two error patterns have the same parameters, one of them will be removed.

These parameters can be calculated as follows:

1. Calculate the total number of units for each row (*R*) and column (*C*).
2. Counting the number of point of intersections (*PI*) by rows and columns.

3. Calculation of the coordinate of the point of intersection (CPI).

4. Calculation of sum (S) and the difference (D) of the point of intersection (CPI).

For example, in table 1 shows the values of the parameters correspond to the above rules for fig. 3.

Table 1 shows that all of the error patterns for the size of $t \times t = 3 \times 3$ have the distinct identification parameters which can represent the other error patterns in the error space of $t \times t$. Therefore, based on these special error patterns, the higher error order of t can be acquired by extending the row and column.

Table 1. Identification parameters for two-dimensional dependent error patterns when $t \times t = 3 \times 3$

№	R	C	PI	CPI	(S, D)
3.01	(300)	(111)	0	–	–
3.02	(210)	(210)	1	(2,1)	(3,1)
3.03	(210)	(111)	0	–	–
3.04	(111)	(300)	0	–	–
3.05	(111)	(210)	0	–	–
3.06	(111)	(111)	0	–	–
3.07	(310)	(121)	1	(3,2)	(5, 1)
3.08	(220)	(121)	2	(2,2), (2,2)	(4,0), (4,0)
3.09	(211)	(220)	2	(2,2), (2,2)	(4,0), (4,0)
3.10	(211)	(130)	1	(2,3)	(5, –1)
3.11	(211)	(121)	1	(2,2)	(4,0)
3.12	(320)	(221)	4	(3,2), (3,2), (2,2), (2,2)	(5,1), (5,1), (4,0), (4,0)
3.13	(221)	(320)	4	(2,3), (2,3), (2,2), (2,2)	(5, –1), (5, –1), (4,0), (4,0)
3.14	(221)	(221)	4	(2,2), (2,2), (2,2), (2,2)	(4,0), (4,0), (4,0), (4,0)

Evaluation of the efficiency of the algorithm for generating two-dimensional dependent error patterns

Table 2. shows the values of the average time for calculating two-dimensional dependent errors when conducting the experiment, a quad-core platform was used and the program was developed in the environment of the windows 10 operating system, and Matlab was used for the program.

Table 2. The average generating time of two-dimensional dependent error patterns for $t = 2:6$

Methods	Average time for generating the library of dependent error patterns				
	2	3	4	5	6
Method [3]	0,008 с	0,03 с	0,1 с	4 с	800 с
Proposed	–	0,02 с	0,02 с	0,5 с	2 с

The analysis of the table data shows that the time of generating libraries of dependent errors under the proposed method is many times faster than the convoluted method.

Conclusion

We have proposed a novel fast error pattern generating algorithm based on the idea of gradually extending the lower order error pattern to higher order error pattern. The experiments have proved that the efficiency of the proposed algorithm is much higher than that of the algorithm in [3]. The reason for this result is that our method adopts the method of extending the $t \times t$ to $(t + 1) \times (t + 1)$ rather than the brute force search used in [3].

References

1. Конопелько В.К., Смолякова О.Г. // Докл. БГУИР. 2008. С. 19–28.
2. Конопелько В.К., Липницкий В.А., Смолякова О.Г. [и др.] // Докл. БГУИР. 2010. № 5. С. 40–47.
3. Конопелько В.К., Макейчик Е.Г., Смолякова О.Г. // Докл. БГУИР. 2009. № 5. С. 57–64.
4. Конопелько В.К., Липницкий В.А., Дворников В.Д [и др.] Теория прикладного кодирования. БГУИР. 2004.

УДК 621.391.13

**ПОМЕХОУСТОЙЧИВОСТЬ СИСТЕМ СВЯЗИ С КАСКАДНЫМ
КОДИРОВАНИЕМ И МНОГОПОЗИЦИОННОЙ МОДУЛЯЦИЕЙ**

Э.Б. ЛИПКОВИЧ, Е.А. БЕЛОКОНЬ

*Белорусский государственный университет информатики и радиоэлектроники, Республика Беларусь**Поступила в редакцию 9 марта 2021*

Аннотация. Предложены аналитические соотношения расчета помехоустойчивости и энергетической эффективности систем связи с многопозиционными видами модуляции и каскадным кодированием без необходимости использования процедур компьютерного моделирования. Выполнены исследования для двух- и трехкаскадных конструкций, использующих сверточные коды с мягким декодированием.

Ключевые слова: каскадное кодирование, энергетическая эффективность, помехоустойчивость, достоверность приема.

Введение

Одним из способов повышения помехоустойчивости и эффективности систем связи с многопозиционными видами модуляции является каскадное кодирование и декодирование сигналов. Оно основано на последовательном включении на стороне передачи нескольких ступеней кодирования и соответствующего числа ступеней декодирования на стороне приема [1]. Наибольшее распространение среди известных решений получили составные конструкции с двухкаскадным кодированием и декодированием, в состав которых могут входить блочные, сверточные, двоичные и недвоичные коды. Например, в системах наземного телевизионного вещания стандарта DVB-T и спутникового мультимедийного вещания стандарта DVB-S в качестве внутреннего кода используется несистематический сверточный код (СК) с мягким декодированием, в качестве внешнего – укороченный недвоичный код Рида-Соломона (РС) с жестким решением [2, 3]. Теоретические результаты исследования этих структур обычно приводятся в виде табличных или графических построений, полученных с помощью средств компьютерного моделирования.

Недостатком сложившегося подхода к оценке помехоустойчивости и эффективности систем с кодированием является необходимость представления для исследований значительного многообразия графических зависимостей, учитывающих формат модуляции, способы и параметры кодирования, число и последовательность соединения кодов в каскадной конструкции. В результате исключается замкнутость аналитических исследований базовых характеристик систем (пороговая чувствительность, коэффициент системы, энергетический потенциал радиолиний, информационная эффективность и др.) и теряется понимание имеющихся взаимосвязей между системными показателями и структурой используемых кодов.

Цель статьи состоит в получении достаточно простых и общих аналитических выражений для прямого расчета помехоустойчивости и эффективности систем связи с каскадным кодированием и многопозиционной модуляцией при обеспечении требуемой достоверности приема. Предполагается, что для борьбы с образованием пакетных ошибок в канале связи предусмотрены устройства перемежения и деперемежения данных. Приведенные результаты расчетов даны для двух- и трехкаскадного сверточного кодирования с мягким декодированием.

Математическая модель расчета помехоустойчивости и эффективности систем связи с каскадным кодированием

Основываясь на работе [4], общее уравнение взаимосвязи между вероятностью ошибки P_{bN} на выходе приемного устройства с N -каскадным декодированием и отношением сигнал/шум (ОСШ) h'_k на его входе записывается в следующем виде:

$$P_{bN} = \frac{C_i \sqrt{\mu_{ipN}}}{q_i \sqrt{\pi \cdot h'_k}} \cdot 10^{-\mu_{ipN} \cdot h'_k / 2,3}, \quad (1)$$

$$\mu_{ipN} = \mu_{i1N} \cdot \mu_{i2N} \cdot \dots \cdot \mu_{iNN} = q_i \cdot \prod_{j=1}^N d_{cj} \cdot \beta_{ijN} \cdot R_{jN}; R_{jN} = \prod_{j=1}^N R_j, \quad (2)$$

где C_i – коэффициент, зависящий от формата модуляции; μ_{ipN} – результирующий показатель эффективности N -каскадного декодирования; μ_{ijN} – показатель эффективности декодирования j -ой ступени; q_i – квадрат коэффициента помехоустойчивости для используемого формата модуляции; $h'_k = E_0 / N_0$ – отношение энергии E_0 , затрачиваемой на передачу бита информации, к спектральной плотности мощности шума N_0 ; d_{cj}, β_{ijN} – свободное расстояние кода и функция взаимосвязи между параметрами j -ой ступени декодирования соответственно; R_{jN} – результирующая кодовая скорость; $R_j = k_j / n_j$ – кодовая скорость j -ой ступени кодирования; k_j, n_j – число информационных символов на входе и выходе j -ой ступени кодирования; N – число ступеней кодирования/декодирования; i – индекс, указывающий на используемый формат модуляции.

Расчетные выражения для определения значений C_i и q_i при использовании в системе многопозиционных сигналов с квадратурной амплитудной (КАМ-М), фазовой (ФМ-М), частотной (ЧМ-М), амплитудной (АМ-М) и относительной фазовой (ОФМ-М) модуляцией приведены в табл. 1, где $m = \log_2 M$.

Табл. 1. Расчетные формулы для определения коэффициентов C_i и q_i

Вид модуляции	C_i	q_i
КАМ-М, $m = 2, 4, 6, \dots$	$C_1 = 2(1 - 1/\sqrt{M}) / m$	$q_1 = 3m / 2(M - 1)$
КАМ-М, $m = 3, 5, 7, \dots$	$C_2 = 2 / m$	$q_2 = 3m / 2(M - 0,5)$
ФМ-2, $m = 1$	$C_3 = 0,5$	$q_3 = 1$
ФМ-М, $m \geq 2$	$C_4 = 1 / m$	$q_4 = m \sin^2(\pi / M)$
ЧМ-М, $m \geq 1$	$C_5 = M / 4$	$q_5 = m / 2$
АМ-М, $m \geq 1$	$C_6 = (M - 1) / mM$	$q_6 = 3m / (M^2 - 1)$
ОФМ-2, $m = 1$	$C_7 = 0,5$	$q_7 = 0,897$
ОФМ-М, $m \geq 2$	$C_8 = 1 / m$	$q_8 = m \sin^2(\pi / M \sqrt{2})$

Если в каскадных конструкциях используются сверточные коды с $R_j = 1/n_j$ или с $R_j = (n_j - 1)/n_j$, то функция взаимосвязи между параметрами отдельных ступеней декодирования определяется на основании следующего соотношения:

$$\beta_{ijN} = \left[1 - \frac{q_i w_j \cdot \lg(R_{jN} d_{cj})}{(n_j - r_j) \cdot (-\lg P_{bj})} \right] / \left[1 + w_j (n_j - 1) \sqrt{R_{jN} \cdot P_{bj}} \right], \quad (3)$$

где $w_j = d_{cj}^*/K_j$ – коэффициент, зависящий от параметров кодера; d_{cj}^* – свободное расстояние кода для $R_j = 1/2$; K_j – длина кодового ограничения j -ой ступени кодирования.

В табл. 2 приведены значения d_{cj} для часто используемых оптимальных кодов в зависимости от длины кодового ограничения K_j .

Табл. 2. Значение свободного расстояния кода

K_j	Расстояние кода d_{cj} при кодовой скорости R_j						
	1/4	1/3	1/2	2/3	3/4	5/6	7/8
5	16	12	8	5	4	3	2
7	20	15	10	7	5	4	3
9	24	18	12	8	6	5	4

Отметим, что большему значению d_{cj} соответствует большее значение β_{ijN} (3) и большая эффективность декодирования μ_{ipN} (2). Однако рост d_{cj} приводит к увеличению избыточности кода (меньше R_j) и усложнению его структуры.

Согласно (3) функция взаимосвязи β_{ijN} зависит как от параметров отдельных ступеней декодирования, так и вероятности ошибки P_{bj} на их выходах. Диапазон изменения β_{ijN} составляет от величин, близких к нулю при большом уровне ошибок ($P_{bj} \geq 5 \cdot 10^{-2}$), до единицы при $P_{bj} \rightarrow 0$. Имеющаяся зависимость β_{ijN} и, следовательно, μ_{ipN} от P_{bj} указывает на необходимость правильного выбора типа, параметров и места размещения кода в каскадной конструкции.

Для лучшего понимания проводимых исследований на рис. 1 приведена упрощенная структурная схема блока демодуляции и каскадного декодирования приемного устройства с указанием на ней скоростей передачи данных, вероятностей ошибок и значений, характеризующих эффективность декодирования, где B_0 – информационная скорость.

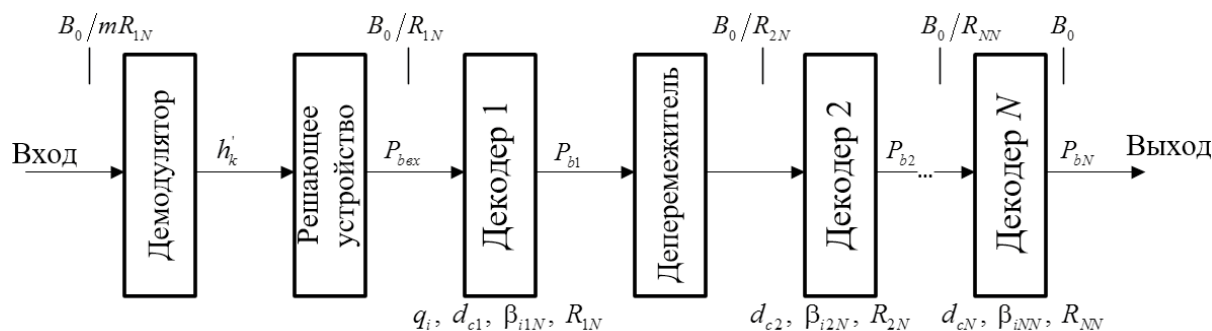


Рис. 1. Структурная схема блока демодуляции и каскадного декодирования

Основываясь на уравнении (1), представляется возможным выразить значение h'_k в зависимости от P_{bN} и получить соотношения для прямого определения помехоустойчивости систем связи с каскадным декодированием и многопозиционными видами сигналов

$$h'_k = \frac{2,3}{\mu_{ipN}} \left[D_{iN} - 0,5 \lg \left[\frac{2,3(D_{iN} - V_{iN})}{\mu_{ipN}} \right] \right]; \quad (4)$$

$$D_{iN} = -\lg P_{bN} - \alpha_i + 0,5 \lg(\mu_{ipN}); \alpha_i = \lg \left(\frac{q_i \sqrt{\pi}}{C_i} \right); V_{iN} = 0,5 \lg \left(\frac{2,3 D_{iN}}{\mu_{ipN}} \right). \quad (5)$$

Из рассмотрения соотношений (4), (5) можно сделать следующие выводы:

– полученные выражения являются достаточно общими для определения помехоустойчивости и энергетической эффективности систем с однокаскадным или N -каскадным декодированием;

– величина ОСШ системы определяется показателем эффективности декодирования первой ступени $\mu_{i1N} = q_i \cdot d_{c1} \cdot \beta_{i1N} \cdot R_{1N}$ и значением ошибки P_{b1} на ее выходе;

– при увеличении числа ступеней декодирования величина ОСШ не изменяется, а повышается исправляющая способность конструкции;

– улучшение помехоустойчивости и достоверности приема связано с увеличением числа ступеней и показателей их эффективности $\mu_{ijN} = d_{cj} \cdot \beta_{ijN} \cdot R_{jN}$;

– увеличение энергетической эффективности (снижение h'_k) обеспечивается увеличением μ_{i1N} первой ступени и выбором режима ее работы в области относительно большого уровня ошибок ($P_{b1} = 10^{-2} \dots 10^{-3}$) в предположении их исправления последующими ступенями.

В соответствии с изложенными выводами устанавливаются требования к параметрам используемых кодов, их структуре и месту размещения в кодовой конструкции.

На основании полученных аналитических соотношений (4), (5) представляется возможным проводить сравнительный анализ энергетического выигрыша от кодирования (ЭВК) при обеспечении равных ошибок на выходах сравниваемых устройств с одинаковыми или разными форматами модуляции. Расчетная формула для определения ЭВК следующая:

$$\Delta G_k = 10 \cdot \lg h'_{kI} / h'_{kII} = h_{kI} - h_{kII}, \text{ дБ}, \quad (6)$$

где h_{kI} , h_{kII} – значения ОСШ сравниваемых устройств, дБ.

При расчете ЭВК в случае сравнения режимов с кодированием и без него значение $h_{kI} = h_0$ определяется по формулам, приведенным в [4, 5].

В табл. 3 и 4 представлены рассчитанные значения h_k и ЭВК (по сравнению с режимом без кодирования) при использовании двух- и трехступенчатого кодирования и модуляции формата КАМ-М. Принято, что длина кодового ограничения $K=9$, вероятность ошибок на выходах устройств составляет 10^{-6} и 10^{-10} и используется несколько схем кодовых конструкций. При двухкаскадном кодировании во внешней ступени $R_2 = 7/8$, во внутренней R_1 принимает два значения, равные $1/2$ и $3/4$. Для трехкаскадной кодовой конструкции $R_2 = R_3 = 7/8$ и R_1 соответственно $1/2$ и $3/4$.

Табл. 3. Эффективность декодирования для двухкаскадной кодовой структуры

Схема конструкции кода	$P_{b2} = 10^{-6}$		$P_{b2} = 10^{-10}$	
	h_k , дБ	ΔG_k , дБ	h_k , дБ	ΔG_k , дБ
$1/2 + 7/8$	1,84	8,70	2,76	10,30
$3/4 + 7/8$	3,15	7,39	3,84	9,22

Табл. 4. Эффективность декодирования для трехкаскадной кодовой структуры

Схема конструкции кода	$P_{b3} = 10^{-6}$		$P_{b3} = 10^{-10}$	
	h_k , дБ	ΔG_k , дБ	h_k , дБ	ΔG_k , дБ
$1/2 + 7/8 + 7/8$	1,60	8,94	1,70	11,36
$3/4 + 7/8 + 7/8$	2,75	7,79	2,87	10,19

Из анализа данных таблиц следует, что ЭВК для приведенных схем при обеспечении $P_b = 10^{-6}$ находится в пределах 7,4...8,7 дБ и 7,8...8,95 дБ для двух и трех ступеней декодирования соответственно. При реализации в системе квазибезошибочного приема с

$P_b = 10^{-10}$ выигрыш в ЭВК от использования трехкаскадной конструкции по сравнению с двухкаскадной составляет 1 дБ и растет по мере снижения P_b .

Заключение

Предложены расчетные соотношения для оценки помехоустойчивости и эффективности систем связи с многопозиционными видами модуляции и каскадным кодированием. Показано, что энергетическая эффективность такой системы определяется первой ступенью декодирования, а требуемая помехоустойчивость достигается за счет последующих ступеней в конструкции, обеспечивающих исправление ошибок. Результаты расчетов энергетического выигрыша от кодирования для двух- и трехкаскадных конструкций, использующих сверточные коды и мягкое декодирование, указывают на целесообразность перехода к трехкаскадным устройствам в случаях необходимости обеспечения высокой помехоустойчивости при сверхнизкой вероятности ошибок ($P_b \leq 10^{-10}$) на выходах систем связи.

IMMUNITY OF COMMUNICATION SYSTEMS WITH CASCADE CODING AND MULTI-POSITION MODULATION

E.B. LIPKOVICH, E.A. BELAKON

Abstract. Analytical relationships are proposed for calculating the noise immunity and energy efficiency of communication systems with multi-position modes of modulation and cascade coding without the need to use computer modeling procedures. Research has been carried out for two- and three-stage constructions using soft decoding convolutional codes.

Keywords: concatenated coding, energy efficiency, noise immunity, reception reliability.

Список литературы

1. Блох Э.Л., Зябков В.В. Линейные каскадные коды. М., 1982.
2. ETSI TR 101 290. Digital Video broadcasting (DVB); Measurement guidelines for DVB systems [Электронный ресурс]. URL: <http://www.etsi.org>.
3. ETSI 300 744. Digital Video broadcasting (DVB); Framing structure, channel coding and modulation for digital terrestrial television [Электронный ресурс]. URL: <http://www.etsi.org>.
4. Липкович Э.Б., Серченя А.А. // Электросвязь. 2020. № 10. С. 62–66.
5. Липкович Э.Б. Системы и устройства спутникового мультимедийного вещания и интерактивной связи. Минск: БГУИР, 2020.

ОЦЕНКА СЛОЖНОСТИ АЛГОРИТМА ОБУЧЕНИЯ С ОШИБКАМИ
АЛГЕБРАИЧЕСКИХ РЕШЕТЧАТЫХ КОДОВ

М.А. АЛИСЕЕНКО, С.Б. САЛОМАТИН

*Белорусский государственный университет информатики и радиоэлектроники, Республика Беларусь**Поступила в редакцию 10 марта 2021*

Аннотация. Рассмотрен алгоритм обучения с ошибками LWE для алгебраических решетчатых кодов. Приведена реализация алгоритма LWE на языке программирования Python. Показана временная сложность вычисления алгоритма LWE для различных длин открытого сообщения и наличия ошибок.

Ключевые слова: алгоритм обучения с ошибками, LWE, алгебраические решетки.

Введение

Одной из задач разработки алгоритмов защиты данных является их потенциальная способность противостоять различного вида атакам, в том числе на основе пост-квантовых и параллельных вычислений. Применение алгоритмов теории многомерных алгебраических решеток предоставляют возможность формирования пространственно-временного многообразия кодовых криптографических структур [1–4]. В данной работе проведена оценка временной сложности алгоритма обучения с ошибками LWE на основе теории решеток.

Моделирование алгоритма обучения с ошибками алгебраических решетчатых кодов

Алгоритм обучения с ошибками LWE [5]:

1. Алгоритм генерация ключа LWE. Вход: $LWE = n, m, l, q$ – целые числа.

1.1. Выбрать секретный ключ $S \in \mathbb{Z}_q^{n \times l}$ случайным образом.

1.2. Выбрать открытый ключ $A \in \mathbb{Z}_q^{m \times n}$ случайным образом.

1.3. Выбрать $E \in \mathbb{Z}_q^{m \times l}$ согласно χ .

1.4. Вычислить открытый ключ $P = AS + E(\bmod q)$, где $P \in \mathbb{Z}_q^{m \times l}$.

Выход: закрытый ключ S и открытый ключ $(A; P)$.

2. Алгоритм шифрования. Вход: целые числа n, m, l, t, r, q , открытый ключ $(A; P)$, открытый текст $M \in \mathbb{Z}_t^{l \times 1}$.

2.1. Выбрать $a \in [-r, r]^{m \times 1}$ случайным образом.

2.2. Вычислить шифротекст $u = A^T a(\bmod q) \in \mathbb{Z}_q^{n \times 1}$.

2.3. Вычислить шифротекст $c = P^T a + [Mq / t](\bmod q) \in \mathbb{Z}_q^{l \times 1}$.

Выход: шифротекст (u, c) .

3. Алгоритм расшифрования. Вход: целые n, m, l, t, r, q , секретный ключ S , шифротекст (u, c) .

3.1. Вычислить $v = c - S^T u$ и $M = [tv / q]$.

Выход: открытый текст M .

Алгоритм реализован на Python 3.8.7 с использованием модуля numpy. Среднее время вычислений рассчитано из 20 измерений на каждую длину открытого текста с использованием встроенного модуля time. При вычислениях использовались следующие параметры:

$n=3$, $m=3$, $t=10$, $r=9$, $q=23$, длина сообщения l , состоящего из случайных целых чисел от 0 до r , варьировалась от 2^1 до 2^{20} . Код функции алгоритма представлен ниже.

```
import numpy as np
import random
import time
def algorithm_lwe(n, m, l, t, r, q, e):
    t0 = time.time()
    lwe_message = np.array([random.randint(0, r) for index in range(l)])
    secret_key = np.random.randint(q, size=(n, l))
    public_key_a = np.random.randint(q, size=(m, n))
    errors = np.zeros((m, l))
    errors = add_errors(errors, e)
    public_key_p = np.mod(((np.dot(public_key_a, secret_key)) + errors), q)
    a_column = np.array([random.randint(-r, r) for index in range(m)])
    ciphertext_u = np.mod(np.dot(public_key_a.transpose(), a_column), q)
    c = np.mod(np.dot(public_key_p.transpose(), a_column) + np.dot(lwe_message, q) / t, q)
    v = np.mod((c - np.dot(secret_key.transpose(), ciphertext_u)), q)
    decoded_message = np.dot(t, v) / q
    message_verification = lwe_message - np.around(decoded_message)
    t1 = time.time()
    print(t1-t0)
```

Результаты вычислений представлены в таблице и на рис. 1 (график построен использованием модуля matplotlib). Сложность вычислений растет экспоненциально.

Длина открытого текста и среднее время вычислений

Длина открытого текста	Среднее время вычислений, с					
	0 ошибок	1 ошибка	2 ошибки	3 ошибки	4 ошибки	5 ошибок
2^1	менее 0,00077692	менее 0,00078103	менее 0,0007809	менее 0,00078601	менее 0,00078105	0,00078129
2^2	менее 0,00077692	менее 0,00078103	0,00078129	менее 0,00078601	менее 0,00078105	0,0007811
2^3	менее 0,00077692	менее 0,00078103	0,0007809	0,00078601	0,00078105	0,00078131
2^4	менее 0,00077692	менее 0,00078103	0,00078121	0,00078114	0,00078108	0,00156629
2^5	0,00077692	0,00078103	0,00078155	0,0	0,00078124	0,0
2^6	0,00078605	0,00078129	0,00156257	0,00078124	0,00078578	0,00077728
2^7	0,00156108	0,00078106	0,00078126	0,00077665	0,00077662	0,00078146
2^8	0,00156344	0,00156243	0,00312494	0,0015626	0,00078135	0,0023437
2^9	0,00233914	0,00156246	0,00390598	0,00234364	0,0023435	0,00234365
2^{10}	0,00547355	0,00625469	0,00859376	0,00547352	0,0046881	0,00781653
2^{11}	0,0124952	0,01093295	0,01406243	0,00858892	0,00859332	0,01249608
2^{12}	0,01718754	0,01796856	0,0367188	0,02500415	0,01797349	0,02812887
2^{13}	0,04140624	0,04531261	0,05703115	0,0546874	0,04062026	0,053066
2^{14}	0,08984395	0,0851567	0,10781269	0,11249609	0,08203155	0,1195312
2^{15}	0,17187498	0,17500001	0,22109404	0,21796908	0,18984798	0,1859421
2^{16}	0,34599365	0,3664056	0,40843289	0,45463223	0,3748706	0,3507766
2^{17}	0,66953600	0,7117192	0,87031239	0,8585941	0,75546863	0,7506607
2^{18}	1,37562145	1,5239129	1,66620113	1,64004219	1,5120651	1,5125001
2^{19}	2,79823248	3,0875399	3,4039516	2,93371499	3,05841386	3,1083119
2^{20}	5,66675596	6,3474723	6,96240778	6,06066538	6,66182725	6,2066043

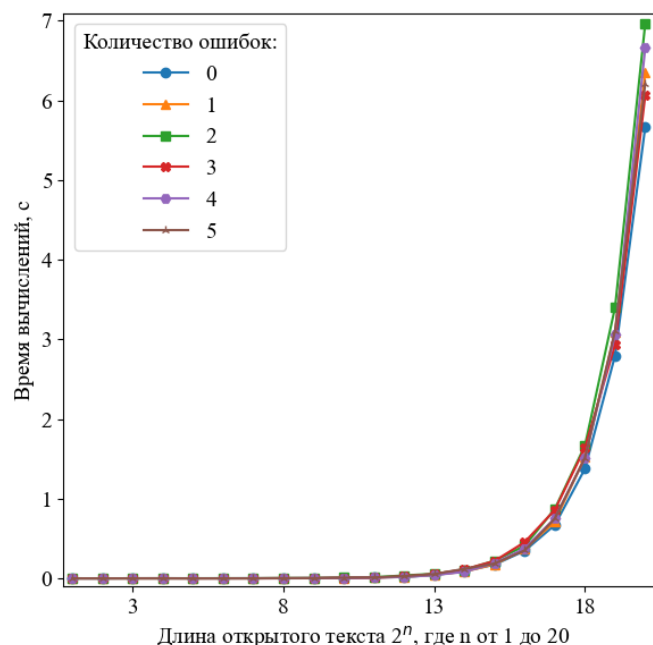


Рис. 1. График зависимости среднего времени вычислений в секундах от количества символов открытого текста и наличия ошибок

Заключение

В настоящее время важность задачи разработки алгоритмов защиты данных определена их потенциальной способностью противостояния к разного рода атакам, таким как постквантовые и параллельные вычисления. Актуальной проблемой становится поиск быстрых алгоритмов шифрования с минимальной вычислительной сложностью для систем связи в сенсорных сетях. Временная сложность алгоритма обучения с ошибками LWE алгебраических решетчатых кодов растет экспоненциально по мере увеличения длины открытого текста, что необходимо учитывать в аппаратном обеспечении устройств интернета вещей.

COMPLEXITY ESTIMATION OF THE LEARNING WITH ERRORS ALGORITHM FOR ALGEBRAIC LATTICE CODES

M.A. ALISEYENKA, S.B. SALOMATIN

Abstract. The LWE (Learning with errors) algorithm for algebraic lattice codes is considered. The implementation of the LWE algorithm in Python is given. The time complexity of calculating the LWE algorithm for different open message lengths and the presence of errors is shown.

Keywords: learning with errors, LWE, algebraic lattice.

Список литературы

1. Ferdinand N.S. // Low complexity lattice codes for communication networks. University of Oulu Graduate School, 2016. P. 178.
2. Olds C.D. // The Geometry of Numbers. Mathematical Association of USA, 2012. P. 192.
3. Johnson N.W., Weiss A.I. // Canadian Journal of Mathematics. 1999.
4. Stallings W. Cryptography and Network Security: Principles and Practic. Prentice-Hall, Upper Saddle River, New-Jersey, fifth edition, 2006.
5. Алисеенко М.А., Саломатин С.Б. // Телекоммуникации: сети и технологии, алгебраическое кодирование и безопасность данных: материалы междунар. науч.-техн. семинара (Республика Беларусь, Минск, ноябрь – декабрь 2020 г.). Минск: БГУИР, 2020. С 23–27.

УДК 654.16

ПРОГНОЗИРОВАНИЕ ЭКСТРЕМАЛЬНОГО ВСПЛЕСКА В СИГНАЛЕ С OFDM

В.А. АКСЕНОВ, С.В. СМОЛЯК

*Белорусский государственный университет информатики и радиоэлектроники, Республика Беларусь**Поступила в редакцию 15 марта 2021*

Аннотация. На основе детерминизма алгоритма формирования сигнала с OFDM рассмотрен метод вычисления временной точки на длительности символа, где может сформироваться экстремальный по величине, или близкий к нему, всплеск в структуре сигнала. Этот метод рассматривается, как составная часть группы более общих методов снижения пик-фактора сигнала OFDM без внесения искажений. Наиболее простым кажется применение предлагаемого подхода к формированию радиоблока узкополосного интернета вещей NB-IoT, где присутствует всего лишь 12 поднесущих.

Ключевые слова: OFDM, пик-фактор, PARP, ортогональные поднесущие, текущая фаза, экстремальный всплеск, радиоблок NB-IoT.

Введение

Ключевым недостатком сигналов с OFDM является их большой пик-фактор. Это означает, что сумма гармонических поднесущих в структуре такого сигнала склонна давать очень большие амплитудные всплески на фоне относительно небольшого среднего значения сигнала на длительности информационного символа. Большой пик-фактор создает проблемы в разработке и применении выходных радиоусилителей и антенно-фидерных устройств в целом. В мире отмечается рост критики применения таких радиосигналов с точки зрения их влияния на здоровье пользователей.

Показательно, что в рамках программы разработки радиосистем шестого поколения 6G, именуемой FG-NET2030, чуть-ли не единственной задачей при разработке нового радиоинтерфейса является поиск OFDM-подобных сигналов, но с существенно сниженными пиковыми выбросами [1].

В научной литературе и реальной инженерной практике предложены и используются несколько методов снижения пик-фактора сигналов с OFDM [2–4]. К сожалению, ни один из этих методов не может быть назван наилучшим по всей совокупности предъявляемых к нему критериев.

Формирование экстремального всплеска в сумме ортогональных поднесущих

Все упомянутые выше методы снижения пик-фактора сигналов с OFDM можно разделить на три группы. К первой относятся методы компрессии (ограничения) по уровню уже сформированного сигнала без относительно того, есть в нем "опасные" всплески или нет.

Ко второй группе можно отнести потоковые регулярные методы перекодирования, предкодирования, перемежения, перестановки как входного потока битов, так и отсчетов "спектра", формируемого перед операцией обратного преобразования Фурье. Такие методы рассматривают входной поток битов и сформированный сигнал, как чисто случайные процессы и гарантированно в среднем снижают вероятность появления всплесков.

Третья группа методов исходит из того, что действительно случайному текущему сочетанию битов на входе модулятора будет соответствовать совершенно детерминированный сигнал на его выходе. Его пик-фактор можно оценить. Если он нас не устраивает, текущее сочетание входных битов, например, исключается из передачи. Как рафинированный пример такого подхода в [5] предлагается заранее просчитать все возможные формы сигнала под все возможные сочетания входного потока битов. Выявить "опасные" сочетания, запомнить их и затем по той или иной схеме избегать их прямой передачи. Вычислительная сложность такого

одноразового подхода, особенно для систем с небольшим количеством поднесущих типа Wi-Fi, NB-IoT, LTE-MTC и др., не является сдерживающим обстоятельством. Затратными оказываются способы передачи – а их надо передать – выявленных "опасных" сочетаний.

С ориентацией на методы третьей группы попытаемся оценить вычислительную сложность предсказания наличия в символе сигнала с OFDM экстремального всплеска, когда достигается максимально возможное значение амплитуды такого сигнала.

В простейшем случае на длительности одного символа T сигнал OFDM может быть представлен суммой гармонических поднесущих

$$S(t) = \sum_{n=0}^{N-1} \sin((\omega_0 + n\Delta\omega)t + \varphi_n), \quad (1)$$

где ω_0 – начальная круговая частота ряда поднесущих; $\Delta\omega$ – разнос по частоте (круговой) между поднесущими; φ_n – начальная фаза поднесущей; $t \in [0, T]$ – текущее время; N – количество поднесущих в сигнале.

Функция текущей фазы для каждой из поднесущих с учетом их ортогональности может быть записана в виде

$$\Phi_n(t) = (\omega_0 + n\Delta\omega)t + \varphi_n = \left(2\pi \frac{m}{T} + 2\pi \frac{n}{T}\right)t + \varphi_n = 2\pi \frac{(m+n)}{T}t + \varphi_n, \quad (2)$$

где m – количество периодов начальной поднесущей на длине символа T (обычно – целое число); $n \in [0, (N-1)]$ – номер поднесущей в ряду; $t \in [0, T]$ – текущее время.

В соответствии с выражением для текущей фазы поднесущей (2) на рис. 1 показано поведение этих функций на длительности T символа OFDM для трех соседних поднесущих. Для упрощения все три поднесущие имели нулевую начальную фазу, поэтому их графики стартуют в точке начала координат. По вертикальной оси графиков отложены значения в долях от π радиан. Другими словами, значение 0,5 соответствует фазе $\pi/2$ и так далее.

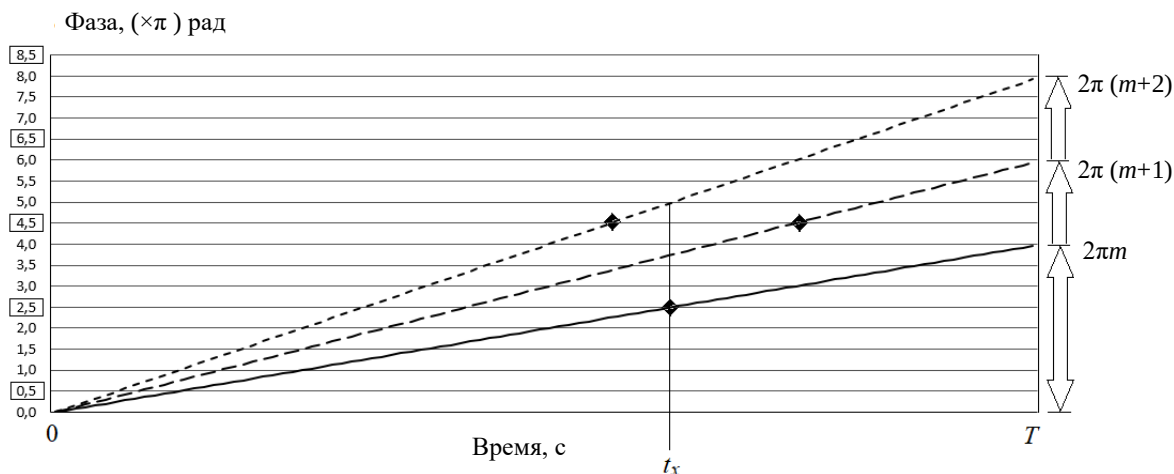


Рис. 1. Поведение текущей фазы для трех гармонических поднесущих на длине символа OFDM

На конце длительности символа T линейная функция для самой низкочастотной поднесущей достигает значения $2\pi m$. Свойство ортогональности поднесущих предполагает, что во временной области всякая следующая поднесущая имеет ровно на один период больше своих колебаний, чем предыдущая [6]. Это означает, что график следующей по порядку поднесущей должен достигать значения на 2π больше, чем график фазы предыдущей поднесущей, как это и показано на рис. 1.

При модуляции поднесущая будет иметь некоторую начальную фазу в интервале от 0 до 2π . В этом случае ее линейный график смещается вверх на указанное значение начальной фазы параллельно самому себе относительно положения на рис. 1.

В предлагаемом графическом представлении все поднесущие будут достигать, например, локального положительного амплитудного значения в моменты времени, когда их текущие фазы пересекают значения вертикальной сетки со значениями 0,5, 2,5, 4,5 и т.д., как показано на рис. 1. Понятно, что сигнал OFDM достигнет экстремального всплеска только в случае, если в некоторый момент времени t_x все поднесущие просуммируются точно в фазе своих локальных амплитудных значений. На графиках рис. 1 это будет соответствовать ситуации, когда в момент t_x все линейные функции текущей фазы будут иметь значения 0,5, либо 2,5, либо 4,5 и т.д. Другими словами, вертикальная линия при аргументе t_x будет пересекать все линейные функции только в указанных значениях фазы.

На рис. 1 представлен прямо противоположный случай. В момент t_x значения 2,5, обозначенного маркером, достигает первая поднесущая в структуре сигнала. Однако две другие поднесущие имеют требуемые для формирования экстремального всплеска фазы (4,5 на графиках) совершенно в другие моменты времени.

Понятно, что формирование экстремального всплеска с отрицательным знаком будет происходить при аналогичных условиях, но для значений текущих фаз из ряда 1,5, 3,5, 5,5 и т.д.

Алгоритм выявления экстремального всплеска

На основе сказанного выше может быть записана система из N уравнений следующего вида, определяющая наличие экстремального положительного всплеска

$$2\pi \frac{(m+n)}{T} t_k^n + \varphi_n = \frac{\pi}{2} + 2\pi k, \quad (3)$$

где $k \in [0, (m+n)]$ – порядковый номер периода поднесущей, где отыскивается всплеск; $n \in [0, (N-1)]$ – номер поднесущей в ряду; t_k^n – момент времени достижения положительного амплитудного значения текущей поднесущей n в ее периоде k .

Разрешение такой системы в общем виде не требуется. Может быть использован упрощенный алгоритм. Для поднесущей с $n=0$, имеющей наименьшее количество периодов на длине символа, определяется момент времени достижения ею первого амплитудного значения t_1^0 . Это значение подставляется в уравнение (3) для $n=1$ и проверяется его непротиворечивость с перебором всех k . Если при каком-либо k присутствует непротиворечивость, то повторяем проверку для следующего $n=2$ и т.д. Если уравнение противоречиво (не может быть разрешено), то отбрасываем момент времени t_1^0 и переходим к моменту t_2^0 . Для этого значения повторяем проверки. Такой алгоритм вытекает из условия, что все поднесущие должны достигать локального амплитудного значения в один и тот-же момент времени для формирования общего экстремального всплеска в сигнале.

Из уравнения (3) видно, что при значении времени $t_k^n = 0$ пропадает зависимость от номера поднесущей и для всех поднесущих при $k=0$ получается взаимно не противоречивая система из уравнений вида

$$\varphi_n = \frac{\pi}{2}, \quad (4)$$

где $n \in [0, (N-1)]$.

Другими словами, экстремальный всплеск будет сформирован в первой же точке сигнала, т.к. все поднесущие начинаются с амплитудного значения (начальная фаза $\pi/2$).

Аналогично, при $t_k^n = T$ уравнения в системе приобретут вид

$$2\pi(m+n) + \varphi_n = 2\pi k + \frac{\pi}{2}, \quad (5)$$

что при той же начальной фазе всех поднесущих $\varphi_n = \frac{\pi}{2}$ и $k = (m+n)$ обеспечит экстремальный всплеск в последней точке сигнала.

Описанной частной ситуации с одинаковыми для всех поднесущих начальными фазами $\pi/2$ будет соответствовать семейство графиков, показанное на рис. 2. На графиках можно видеть, что колебания достигают требуемых значений фаз, кратных $\pi/2$ и обозначенных маркерами, в одинаковые моменты времени.

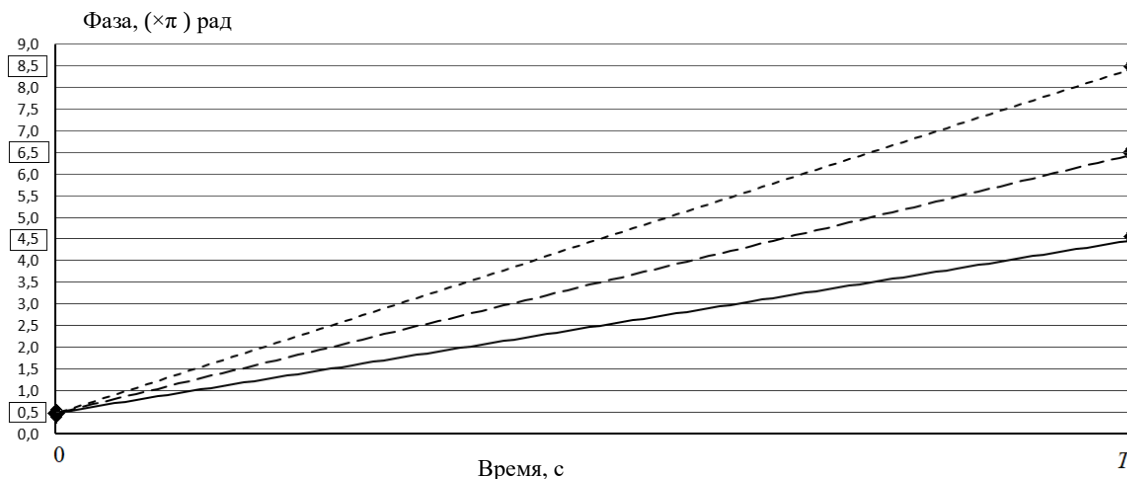


Рис. 2. Текущая фаза для трех гармонических поднесущих при их начальных фазах $\pi/2$

Данный метод наиболее подходит для сигналов OFDM с небольшим количеством поднесущих, например, как упомянутый выше узкополосный интернет вещей NB-IoT, где при передаче в DL имеется только 12 поднесущих в полосе 180 кГц. Даже при небольшом количестве поднесущих, пик-фактор такого сигнала, понимаемый как PAPR (англ. Peak-to-Average Power Ratio), достигает при наличии экстремального всплеска относительно большого значения [7]

$$PAPR = 10\lg(N) = 10\lg(12) = 10,8 \text{ (дБ)}. \quad (6)$$

При использовании протокола передачи с возможными повторениями до 2048 раз, неудачная по пиковой структуре реализация сигнала может оказаться опасной для выходных каскадов передатчика и антенно-фидерного устройства.

Обычно сигналы с OFDM формируют на так называемой "нулевой частоте", когда середина спектра ОБПФ соответствует нулю герц. В случае NB-IoT это полоса ± 90 Гц. Далее сигнал переносится в смесителе на среднюю частоту выбранного радиоканала. При этом временная структура огибающей сигнала сохраняется. Понятно, что описанный метод целесообразно применять именно на "нулевой частоте", что существенно минимизирует количество переборов по коэффициенту k .

Заключение

Представленный метод прогнозирования экстремальных всплесков, ориентированный на использование в NB-IoT, может оказаться простым инструментом для повышения эффективности передачи подобных систем и снижения требований, соответственно, и стоимости выходных каскадов радиооборудования.

EXTREMAL PULSE FORECASTING IN THE SIGNAL WITH OFDM

V.A. AKSYONOV, S.V. SMOLYAK

Abstract. On the basis of the determinism of the signal generation algorithm with OFDM, a method is considered for calculating the time point on the symbol duration, where an extreme in magnitude, or close to it, burst in the signal structure can be formed. This method is considered as part of a group of more general methods for reducing the crest factor of an OFDM signal without introducing distortion. The simplest seems to be the application of the proposed approach to the formation of a radio block of the narrowband Internet of things NB-IoT, where there are only 12 subcarriers.

Keywords: OFDM, crest factor, PARP, orthogonal subcarriers, current phase, extremal pulse, radio block NB-IoT.

Список литературы

1. ITU-T Technical Report FG-NET2030-Sub-G1(2020). Representative use cases and key network requirements for Network 2030. [Электронный ресурс]. URL: https://www.itu.int/dms_pub/itu-t/opb/fg/T-FG-NET2030-2020-SUB.G1-PDF-E.pdf
2. Sandeep Bhada, Pankaj Gulhane, A.S. Hiwalec. PAPR reduction scheme for OFDM. [Электронный ресурс]. URL: <https://www.sciencedirect.com/science/article/pii/S2212017312002940/pdf?md5=218c5c472db45e3b38ec5119000507c7&pid=1-s2.0-S2212017312002940-main.pdf>
3. Пик-фактор OFDM модулированного сигнала. [Электронный ресурс]. URL: https://bester-ltd.ru/articl/teoriya_praktika/pik_faktor_ofdm_modulyacii/pik_faktor_ofdm_modulyacii.html
4. Zaidi A., Athley F., Medbo J. [et. al] Physical Layer. Principles, Models and Technology Components. 2018, P. 159–198.
5. Рухлин С.Н. // III Всероссийские Армановские чтения. 2013. С. 201–207.
6. Передача сигналов ДЗЗ по технологии OFDM. Тема 5. Лекция № 15. [Электронный ресурс]. URL: <http://www.its.kpi.ua/subjects/22/Documents/%D0%9B%D0%B5%D0%BA%D1%86%D0%B8%D1%8F%2015.pdf>
7. Sadek Ali, Yu Li, Khalid Hossain Jewel [et. al]. // Channel Estimation and Peak-to-Average Power Ratio Analysis of Narrowband Internet of Things Uplink Systems. [Электронный ресурс]. URL: <https://doi.org/10.1155/2018/2570165>.

УДК 621.391

КОДИРОВАНИЕ ЦИФРОВОГО ВИДЕОПОТОКА В СИСТЕМАХ ВИДЕОНАБЛЮДЕНИЯ

Е.А. МАШКОВ, О.А. АЗАРЧИК, А.В. ТЮТРЮМОВА

Белорусский государственный университет информатики и радиоэлектроники, Республика Беларусь

Поступила в редакцию 6 марта 2021

Аннотация. Рассматриваются основные виды и стандарты сжатия цифрового видеопотока, их основные преимущества и недостатки. Приводится обоснование выбора стандарта сжатия видеопотока в системах видеонаблюдения.

Ключевые слова: стандарт сжатия, H.264, H.265, H.266, видеонаблюдение.

Введение

Под кодированием цифрового видеопотока в системах видеонаблюдения принято понимать сжатие цифрового видеопотока, данный вид обработки видеоизображения имеет большое значение в системах видеонаблюдения.

Цифровой способ съемки видео впервые был применен в 1980-ых годах. Тогда это было несжатое видео, требовавшее большого количества памяти для его сохранения и дальнейшего использования. В связи с тем, что устройства хранения информации имели небольшой объем, были разработаны алгоритмы сжатия цифрового видеопотока, называемые кодеками.

Процесс сжатия цифрового видеопотока заключается в потоковом анализе получаемой видеоинформации и автоматическом удалении фрагментов данных, на которых не зафиксировано каких-либо изменений с предыдущим фрагментом, что помогает значительно уменьшить объем видеофайла, либо удалении области из фрагмента по тому же признаку.

Виды и стандарты сжатия цифрового видеопотока

На сегодняшний день существует два вида сжатия видеоинформации: внутрикадровое (покадровое) и межкадровое кодирование.

Внутрикадровое сжатие обрабатывает каждый кадр видеозаписи как отдельное неподвижное изображение, наподобие фотографии. Данная технология позволяет получить видео хорошего качества, но при этом размер файла уменьшается незначительно по той причине, что при такой кодировке сохраняются все кадры, даже если в кадре нет никаких изменений. То есть, например, из десятка или сотни одинаковых кадров сохраняются все, хотя достаточно всего одного.

Межкадровое сжатие работает по противоположному принципу: при обработке сигнала, анализируется все видеоизображение, но сохраняются только ключевые изменения, например, движение объекта, при этом фон заднего плана и окружающая объект обстановка остаются теми же. Это позволяет значительно уменьшить размер видеофайла в сравнении с внутрикадровым сжатием.

При сжатии учитываются: разрешение видео, размер файла, способ передачи и загрузки видеофайла, преобладание статичных или динамичных сцен, цветность, контрастность и т.д. Качество и размер видеофайла, полученного в результате видеозаписи, зависит от применяемого кодека [1].

На сегодняшний день, основными стандартами сжатия в системах видеонаблюдения являются H.264, H.265.

Стандарт сжатия M-JPEG (Motion JPEG)

Это нелицензированный стандарт кодирования, созданный и широко используемый в 90-ых годах. Он использует внутрикадровую технологию сжатия. Цифровой видеопоток, полученный с помощью данного кодека, представляет собой массив полновесных JPEG-изображений. Несмотря на то, что данный кодек позволяет использовать ряд понижающих объем файла инструментов, сегодня его используют крайне редко из-за невысокого качества получаемого изображения, а также по причине минимальной степени сжатия видео.

Стандарт сжатия MPEG-4

Лицензированный стандарт кодирования, использующий объектно-ориентированное (межкадровое) сжатие, то есть когда движение каждого объекта в кадре отслеживается отдельно и на основании этих движений фиксируется видеосигнал. Основным плюсом данного кодека является широта настроек степени сжатия, которые можно подобрать под любую скорость передачи данных. Данный формат универсален и используется для просмотра потокового видео в реальном времени. Следует отметить, что этот стандарт уже устарел [2–4].

Стандарты сжатия H.264 и H.264+

Современные лицензированные стандарты кодирования, которые существенно уменьшают объем цифрового видеофайла. Данные кодеки вносят минимальные изменения в качество видеоизображения и предназначены для записи видеосигнала в течение продолжительного времени, так как требуют небольших пропускных способностей сети и места на жестком диске.

Кодек H.264 – является одним из лучших инструментов для работы в системах видеонаблюдения, особенно при съемке с большой частотой кадров и высоким разрешением. Единственный минус – он требует большой вычислительной мощности оборудования для распаковки и просмотра видеоинформации.

H.264+ получил некоторые изменения, которые еще больше повысили эффективность работы оборудования:

- фоновый инструмент для подавления любых шумов во время работы;
- технология интеллектуального кодирования по модели предсказания фона;
- система автоматического изменения битрейта и регулировки [1–4].

Стандарты сжатия H.265 и H.265+

Лицензированный стандарт H.265 предлагает значительно улучшенное сжатие по сравнению с H.264. Данный кодек сжимает видео почти вдвое лучше, чем его предшественник. То есть с H.265 видео, с одинаковым визуальным качеством, займет лишь половину памяти по сравнению с 264 версией кодека, а видео с одинаковым размером файла и скоростью передачи битов, имеет значительно лучшее качество.

В H.265 используется тот же принцип сжатия, что и в H.264. В случае использования фиксированной камеры видеонаблюдения фоновое изображение меняется не часто, поэтому достаточно кодировать и передавать только изменения – движущиеся объекты. Это позволяет значительно уменьшить требования к пропускной способности канала и емкости хранения.

IP-камеры видеонаблюдения сначала снимают необработанное видео в соответствии с заданным режимом записи, а после обработки изображения кодируют его. Основное преимущество в степени сжатия достигается за счет улучшения прогнозирования с компенсацией движения. В то время как у H.264 максимальный размер блока составляет 16×16 пикселей, H.265 использует при обработке информации макроблоки дерева кодирования Coding Tree Unit (CTU) размером до 64×64 пикселей. Такие блоки более эффективны для кодирования кадров больших размеров и при этом позволяют передавать видео 4K+ [1].

Помимо изменения размера блока, H.265 отличается наличием улучшенного сглаживающего фильтра для устранения нестыковок на границах блоков (deblocking filter). Кроме того, используется новый алгоритм прогнозирования вектора движения Motion Vector

Prediction (MVP) для улучшения прогнозирования внутри кадра. Более высокая точность предсказаний достигается, помимо прочего, благодаря тому, что в пределах кадра вместо 8 возможных направлений, как обеспечивается в H.264, рассматривается 36 возможных направлений.

Для ускорения вычислений в кодеке предусмотрена возможность параллельной обработки за счет поддержки расширенного набора инструкций AVX/AVX2 для процессоров Intel/AMD. Квадратные области, на которые разбивается изображение, независимы одна от другой, так что их обработка может выполняться параллельно. Кроме того, H.265 поддерживает волновую параллельную обработку Wavefront Parallel Processing (WPP): своеобразное дерево принятия решений, способствующее повышению производительности сжатия. Тем не менее для его реализации необходим на порядок более мощный процессор, что является одним из его существенных недостатков.

Возможности стандарта кодирования H.265+:

1. Интеллектуальное кодирование. Во время работы кодек разделяет фон и посетителей. Создается модель из одного или нескольких ранее созданных кадров, что позволяет проводить своего рода "прогнозирование", где блоки обработки предсказываются информацией из ранее переданных блоков и того же кадра. Таким образом, сжатие потока осуществляется благодаря проведению трансляции исключительно динамической части кадра.

2. Подавляется цифровой шум. Интегрированная интеллектуальная система анализа способствует тому, что кодек H.265+ имеет возможность различать движущиеся объекты и фоновые изображения таким образом, что каждый из них может быть закодирован под разной стратегией кодирования. Если фон сжимается под высоким уровнем сжатия для подавления шума, то модуль кодирует также визуальный шум в сцене. Это привело к тому, что удается достигать высокого уровня качества, несмотря на небольшой размер видеопотока.

3. Битрейт под долгосрочным контролем. Современные IP-камеры видеонаблюдения умеют различать моменты, когда на выделенном участке наблюдения ничего не происходит, и в это время снижают качество, чтобы уменьшить нагрузку на сеть и место на жестком диске, сохраняя значения битрейта около установленного максимума. Это значительно экономит ресурсы и повышает эффективность работы системы [1, 2].

Стандарт сжатия H.266

Стандарт сжатия H.266, который также называют Versatile Video Codec (VVC), находится в стадии активной разработки. Рабочий Проект кодекса спецификации VVC был выпущен в октябре 2018 года, поэтому сейчас ведется работа по его стандартизации. Тем не менее, модель лицензирования для VVC остается неясной до момента, пока стандарт не будет завершен и не будут финализированы основные функции кодекса.

Стандарт сжатия H.266 определяет технологию кодирования видеопотока. Он был разработан с двумя основными целями. Первая из них заключается в том, чтобы определить технологию кодирования видео с возможностью сжатия, которая существенно превосходит возможности предыдущих поколений таких стандартов, а вторая – в том, чтобы эта технология была весьма универсальной для эффективного использования в более широком диапазоне применений, чем те, которые рассматривались в предыдущих стандартах. Некоторые ключевые области применения этого стандарта, в частности, включают видео сверхвысокой четкости (например, с разрешением изображения 3840×2160 или 7620×4320 и глубиной битов 10 или 12 бит, как указано в Rec. ITU-R BT.2100), видео с высоким динамическим диапазоном и широкой цветовой гаммой (например, с перцептивным квантованием или гибридными характеристиками передачи log-гамма, указанными в Rec. ITU-R BT.2100), и видео для иммерсивных медиа-приложений, таких как 360° всенаправленное видео, проецируемое с использованием общего формата проекции, такого как формат проекции equirectangular или cubemap, в дополнение к приложениям, которые обычно рассматривались предыдущими стандартами кодирования видео [1].

Выбор стандарта сжатия видеопотока в системах видеонаблюдения

Для того что бы выбрать оптимальный стандарт сжатия видеопотока в системах видеонаблюдения необходимо сравнить результаты сжатия одного и того же видеопотока разными кодеками. В связи с тем, что стандарт сжатия цифрового видеопотока H.266 еще окончательно не принят, поэтому рассматривать его нецелесообразно. Также по причине редкого использования в профессиональных системах видеонаблюдения кодеков M-JPEG и MPEG-4 в сравнительную таблицу они включаться не будут.

Результаты сравнение стандартов сжатия цифрового потока будут представлены в таблице, для расчета требуемого объема дискового пространства и суммарного битрейта использованы следующие данные:

- глубина архива – 30 дней;
- количество камер – 10 штук;
- тип записи для первого сравнения – постоянная;
- тип записи для второго сравнения – по движению, процент движения в сутки – 60,0 %;
- разрешение камер видеонаблюдения – 4Мр (2592×1520);
- скорость записи – 25 кадров в секунду.

Результаты сравнения стандартов сжатия цифрового видеопотока

№	Название критерия	H.264	H.264+	H.265	H.265+
1	Требуемый объем дискового пространства, постоянная запись, Тб	23,48	16,38	12,98	11,12
	Суммарный битрейт, Мбит/с	83,6	58,3	46,2	39,6
2	Требуемый объем дискового пространства, запись по движению, Тб	14,09	9,83	7,79	6,67
	Суммарный битрейт, Мбит/с	76	53	42	36

Данные приведенные в таблице получены расчетным путем при помощи онлайн калькулятора требуемого объема дискового пространства и битрейта [5].

Таким образом оптимальным стандартом сжатия видеопотока является H.265+, так как при его использовании требуемый объем дискового пространства составляет на 52,7 % меньше чем у стандарта H.264, на 32,1 % меньше чем у стандарта H.264+ и на 14,3 % меньше чем у стандарта H.265, а также значительно меньший канал связи для передачи данного цифрового видеопотока.

Заключение

В данной статье рассмотрены основные используемые стандарты сжатия цифрового видеопотока, а также перспективный стандарт сжатия H.266. Проведено сравнение стандартов сжатия цифрового видеопотока, при котором оптимальным стандартом для систем видеонаблюдения принят H.265+. Но не стоит забывать про то, что подбор необходимого стандарта сжатия зависит, прежде всего, от актуальности модели видеорегистратора, так как не все модели смогут работать с данным стандартом сжатия в связи с тем, что для работы со стандартом сжатия цифрового видеопотока H.265+ необходимо иметь больше вычислительных мощностей по сравнению с другими стандартами.

В непрофессиональных системах видеонаблюдения, при простых аппаратных характеристиках оборудования, а также для формирования нескольких потоков видеосигнала (например, для передачи видео по сети или удаленном просмотре с помощью мобильного телефона), достаточно использовать кодек MPEG-4, так как он менее требователен к ресурсам системы [6].

В профессиональных системах видеонаблюдения, где используется множество видеокамер и имеются видеорегистраторы или серверы с большим объемом памяти наиболее оптимальным будет выбор стандарта сжатия цифрового видеопотока H.265+.

DIGITAL VIDEO STREAM ENCODING IN VIDEO SURVEILLANCE SYSTEMS

E.A. MASHKOV, O.A. AZARCHIK, A.V. TYUTRYUMOVA

Abstract. The main types and standards of digital video stream compression, their main advantages and disadvantages are considered.

Keywords: compression standard, H.264, H.265, H.266, video surveillance.

Список литературы

1. Дамьяновски В. Библия видеонаблюдения. 2021. С. 219–232.
2. Подразделение научных разработок МВД Великобритании. Руководство по составлению эксплуатационных требований к системам видеонаблюдения. Версия 5.0. 2014. С. 40–70.
3. Торстен Анштедт. Видеоаналитика: Мифы и реальность. 2019. С. 120–124.
4. Попов А. "Моя азбука видеонаблюдения". 2013. С. 30–34.
5. Онлайн калькулятор. [Электронный ресурс]. URL: https://rvigroup.ru/techsupport/calc_bitrate.
6. Лыткин А. IP-видеонаблюдение: наглядное пособие. 2011.

UDC 004.031.43:004.75

MODELING A NETWORK OF THINGS ON THE BASIS OF A CLOUD PLATFORM

U.A. VISHNIAKOU, A.H. AL-MASRI., S.KH. AL-HAJJ

*Belarusian State University of Informatics and Radioelectronics, Republic of Belarus**Submitted 1 March 2021*

Annotation. The analysis of the capabilities of cloud platforms for building IoT networks is presented. The functionality of the IoT platform from the Cisco company is considered. The structure and description of IoT platform is given, an enlarged algorithm for creating an IoT network based on Bluemix platform is given.

Keywords: cloudy IoT platforms, Cisco, Bluemix, building an IoT network.

Introduction

A variety of Internet of Things (IoT) networks are becoming more widespread [1]. IoT is a collection of embedded systems, wireless sensor networks, control systems and automation tools for processing information received from sensors. IoT networks allow at a new level to implement the automation of production processes, management of objects of transport, energy, agriculture, medicine, to create smart stores, "smart" houses, districts and cities as Automatization 4.0.

To automate the creation of IoT systems, the world's leading companies have developed design tools in the form of IoT platforms [2]. IoT platforms are becoming the central pillar of IoT deployments by 023, the IoT platform market will reach 3 billion dol. [3]. Due to the large volume of information received from IoT sensors, analytical tools for processing Big Data are provided within the framework of such platforms. Let's consider these issues in more detail.

In May 020, the analytical company Counterpoint Research named the leading platforms for creating IoT networks and applications in terms of versatility (meeting user needs) and other parameters. The first position was taken by the Microsoft Azure platform, the second position was taken by Amazon Web Services (AWS), the third position was taken by Huawei OceanConnect, the fourth position was taken by PTC ThingWorx, the fifth position was taken by IBM Watson [2, 3]. The top ten also includes platforms from Cisco, Oracle, Intel, and others. This study took into account eight components: distribution, growth rate, ability to integrate and scale, application support, cloud components, edge interaction, data processing from devices and peripheral components [4].

Cisco IoT Solutions

Among the innovations presented by Cisco in the field of IoT [4].

1. New network IoT platforms. Cisco unveiled new Industrial Catalyst switches and Industrial Integrated Services Routers specifically designed for the IoT environment. All equipment operates on the basis of the state-of-the-art IOS XE operating system, which supports the operation of intention-oriented networks at the campus, branch and wide area network (WAN) levels. The new platforms are managed by Cisco DNA Center, which provides customers with a complete, centralized view of events across campus, branch, and IoT environments.

2. Development Tools (IoT Developer Tools). The Cisco DevNet Developer Assistance Program has been updated with new tools to make it easier for customers and partners to innovate at the IoT edge.

3. Extensibility. Customers can implement new technologies such as 5G without replacing their network infrastructure. Cisco Industrial Routers are the industry's first and only 5G and IPv6 capable routers.

4. Information security. Appropriate tools are built into every layer of the IoT stack, from networking hardware to operating software and edge computing applications.

IoT component stack structure and purpose

As part of the Internet of Things network, five components can be distinguished [5]: sensors (devices) and their hardware, software for sensor management, communication tools, cloud platform and mobile applications. Let's consider the purpose of each of these components.

Devices act as an interface between the physical and digital worlds. They are the first layer of the IoT technology stack. Only one sensor may be needed for simple data collection. For more complex data collection, it may need a computer that contains many sensors, a processor, local storage, a gateway, etc. At this level of the IoT technology stack, it is important to understand such parameters of equipment as cost, size, ease of deployment, reliability, useful life, service, etc.

Device software is the second layer of the IoT technology stack. It is a component that turns a device's hardware into a "smart device". Device software includes the concept of "software-defined hardware", which means that a particular hardware device can serve multiple applications depending on the firmware it runs on.

The device software allows communication with the cloud or other local devices. You can perform analytics in real time, collect data from device sensors and monitor parameters. The device software layer consists of two components: the operating system and the device applications. If the device is simple, the OS may not be used. The application can analyze the data from the sensors, comparing them with the boundary values. If the data exceeds the permissible, it is transmitted to the cloud for monitoring by the operator.

Communications – the third level of the IoT stack, includes both physical networks and the protocols that will be used. The implementation of the communication layer can be found in the hardware and software of the device. But in terms of the conceptual model, you can keep communication as a separate layer to facilitate discussion during development.

Choosing the right communication mechanisms is an important part of an IoT product strategy. It will determine not only how data is transferred to and from the cloud (using different subnets and protocols: Wi-Fi, 4G, LoRA, etc.), but also how communication with other devices is organized.

The cloud platform (the fourth layer of the IoT stack) is the basis of the IoT network (project), which provides the infrastructure that supports all components: device management, data collection and management, data analytics, cloud interfaces, information security. Smart devices will transmit information to the cloud. You need to be aware of the type and amount of data that will be collected daily, monthly and annually.

Analytics refers to the ability to process data, find patterns, make predictions, integrate machine learning, and so on. The ability to find information from data makes a solution useful. Analytics can be as simple as aggregating and displaying data, or as complex as using machine learning or artificial intelligence. Cloud interfaces allow clients and managers to either interact with devices or exchange data. You may need separate applications for desktops, mobile devices, and for different categories of users.

Algorithm for connecting sensors with primary processing

For simulated IoT monitoring of parameters, we use a scheme for reading readings of measured physical quantities (indicators of product quality or environmental parameters). Then they are preprocessed and the application is sent to the client via the IoT platform. The final stage is displaying data on the client side. The generalized network design algorithm is represented by three points:

- Select a sensor or a set of sensors from which we will receive the measured data, and a method for processing the received data.
- Determining how we will communicate with sensors, determine the amount of data and understand how we will build interaction.
- Find a suitable client for our network and describe the work with him.

For the sensor and pre-processor, we will choose the SensorTag circuit from Texas Instruments [6]. Another option is Arduino with BLE-shield (Bluetooth low energy, low power) or BLEduino.

Inside the CC2650 chip, the SensorTag core, is a real-time operating system (TI-RTOS), which together with the BLE stack provides reliable control of three different microcontrollers:

- The core of the first microcontroller is Cortex-M3 (it usually runs the custom application we have written).
 - The core of the second Cortex-M0 (responsible for the physical layer, radio communication).
 - A separate controller for sensors (helps to quickly receive data from them).
- Android phones with BLE stack support are widespread, we will use it as a hub and gateway on the way to the cloud.

Generalized algorithm for creating a network by means of the IoT platform

Consider an IoT network modeling algorithm, which is divided into two parts. First, we organize the transfer of information from the sensors using a smartphone (three levels of the IoT stack), then we connect the cloud platform and implement the remaining two levels of the IoT stack.

To send the data received to the gateway (smartphone) via the BLE protocol to the Internet, we use the MQTT protocol and the JSON data transfer format.

MQTT (Message Queue Telemetry Transport) is a simplified network layer protocol for exchanging messages between devices, it runs on top of the TCP/IP stack and is designed to connect sensors, microcomputers, smartphones, tablets. MQTT is a publisher/ subscriber. A publisher (devices of type publishers) sends a message, which is published in a centralized service (a message broker), and a subscriber (devices of type subscriber) receives a message from the broker.

JSON (JavaScript Object Notation) is a textual data interchange format based on JavaScript and used with this particular language. The format is considered language independent and can be used with any programming language. For this, there is a ready-made code for creating and processing data in JSON format.

JSON text is (encoded) one of two structures.

- A set of key-value pairs. In various languages, this is implemented as an object, record, structure, dictionary, hash table, keyed list, or associative array. The key can only be a string (case-sensitive: names with letters in different cases are considered different), the value can be any form.
- An ordered set of values. In many languages, this is implemented as an array, vector, list, or sequence.

These are universal data structures: as a rule, any modern programming language supports them in one form or another. They formed the basis for JSON, as it is used to exchange data between different programming languages.

A common implementation of the MQTT protocol is the Paho MQTT library, which is implemented for common programming languages: C/C++, Java, JavaScript, Python, etc. Let's consider the algorithm of communication between the client and the cloud.

1. We import the Paho MQTT library and the classes we need to work with the MQTT protocol.
2. We indicate the address of the cloud.
3. We set the number of the standard port of the broker of the cloud platform of the MQTT protocol.
4. We send data to the cloud.
5. The hub/gateway device (android phone) generates MQTT packets and transmits them to the cloud for storage and processing.
6. From the cloud, the hub can receive commands for device control or for the gateway.

Then we will create our web service to receive and display the readings of our sensors. To do this, we will use the cloud platform from IBM – Bluemix [7], which is needed to implement storage services, analytical processing and visualization of data received from SensorTag and pumped through an android phone. Bluemix is PaaS (Platform as a Service) open-source cloud offering based on the Cloud Foundry open-source project. The platform is designed for application development and hosting, and it also simplifies infrastructure management.

Building a cloud platform app

In Bluemix terminology, an application is any generated code (source code or executable binaries) that must be run or referenced at runtime. Mobile apps run outside of the Bluemix environment and use the services provided by the apps. In the case of web applications, an application is code uploaded to

the Bluemix platform for hosting purposes. In addition, the platform is capable of hosting the application code that we want to run on an internal server in a container-based environment.

A service is code that runs on the Bluemix platform and offers specific functionality that applications can use. This can be a ready-made service used directly, such as push notifications for mobile apps or elastic caching for a web app. You can create your own services ranging from simple utility functions to complex business logic.

There are three steps to using services in Bluemix:

1. Tell the Bluemix platform that we need a new instance of the service and specify which particular application will use this new instance.
2. Bluemix automatically initializes a new instance of this service and associates it with the application.
3. The application interacts with the service.

Service bundles are collections of APIs used in specific areas. For example, the Mobile Services package includes MobileData, Cloud Code, Push, and Mobile Application Management services. The available services and runtimes are listed in the Bluemix catalog. In addition, you can register your own services.

In the platform, we will choose Node-RED because of its convenience and ease of configuration. It has a convenient graphical programming interface consisting of JS blocks, which can be described by loading a JSON file. You can also use some Node.js.

The data from the MQTT package is sent to the broker as part of the IoT Foundation service. A data subscriber is a Node-RED application that allows you to manipulate data using simple visual aids. Many primitive processing units (nodes) are JavaScript applications connected to each other by streams of data.

Conclusion

To automate the creation of IoT systems, design tools are used in the form of IoT platforms. The main functions of such a platform from Cisco are presented. The structure of the stack in the IoT network is considered. The connection of sensors with means of primary processing, including protocols and data structure, is described. A generalized algorithm for creating a network using the IoT platform Bluemix from IBM is presented.

References

1. Roslyakov A.V. Internet of Things: textbook. manual. Samara. 2019.
2. IoT platform [Electronic resource]. URL: <https://iot.ru/wiki/iot-platforma>.
3. IoT Platforms [Electronic resource]. URL: <http://www.tadviser.ru/index.php/> Article: IoT platforms.
4. Cisco Kinetic [Electronic resource]. URL: [https://www.tadviser.ru/index.php/](https://www.tadviser.ru/index.php/Product:Cisco_Kinetic) Product: Cisco_Kinetic.
5. Elizable D. The 5 Layers of the IoT Technology Stack [Electronic resource]. URL: <https://danielelizalde.com/iot-primer>.
6. How to create an IoT for yourself. Learning to do the Internet of things on Android and hardcore hardware [Electronic resource]. URL: <https://xakep.ru/2016/04/28/iot-android-sensortag/>.
7. Platform description Bluemix [Electronic resource]. URL: <https://searchcloudcomputing.techtarget.com/definition/IBM-Bluemix>.

УДК 681.382.473

ИССЛЕДОВАНИЕ ОПТИЧЕСКИХ ИЗЛУЧАТЕЛЕЙ НА ОСНОВЕ НАНОСТРУКТУРИРОВАННОГО КРЕМНИЯ

А.Ю. КЛЮЦКИЙ, С.К. ЛАЗАРУК

*Белорусский государственный университет информатики и радиоэлектроники, Республика Беларусь**Поступила в редакцию 8 марта 2021*

Аннотация. Рассмотрена перспектива использования светоизлучающих структур на основе кремния в системах кодирования. Исследованы характеристики лавинных светодиодов на основе наноструктурированного кремния. Показана возможность создания оптических систем межсоединений на основе кремния.

Ключевые слова: кремниевые оптические межсоединения, светодиоды на основе наноструктурированного кремния.

Введение

Системы кодирования и цифровой обработки сигнала находят широкое применение во всех областях науки и техники. Принципы кодирования и шифрования информации базируются на использовании в своем алгоритме последовательностей случайных чисел, которые формируются с помощью таких устройств, как генераторы. Существующие программные решения являются генераторами псевдослучайных чисел, что делает их непригодными для задач криптографии. Аппаратные генераторы, в свою очередь, обладают низкой скоростью работы по сравнению с программными, а также подвержены деградации. Поэтому важной задачей является разработка устройств, решающих данные проблемы. Одним из современных направлений аппаратных генераторов случайных чисел является использование однофотонных источников излучения. Конструктивно такие генераторы состоят из квантово-оптической системы, включающей в себя однофотонный источник излучения, лавинный фотодетектор, оптический канал связи, оптическую линию задержки, и устройства преобразования и обработки информации [1]. Однофотонные излучатели, как правило, представляют собой квантовые точки арсенида индия. Однако существует перспектива создания таких излучателей на основе наноструктурированного кремния. Использование наноструктурированного кремния в качестве основного материала светоизлучающего элемента позволяет добиться интеграции всех составляющих системы оптических межсоединений на одном кристалле. Интеграция внутри кремниевого чипа значительно увеличит эффективность таких приборов. Кроме того, светоизлучающие элементы на основе наноструктурированного кремния будут отличаться низкой стоимостью, в связи с тем, что при их изготовлении будут применяться традиционные методы кремниевой технологии, которая в настоящее время достаточно хорошо освоена [2].

Лавинные светодиоды на основе наноструктурированного кремния

Разработка и создание эффективных светоизлучающих устройств может быть реализована на основе наноструктурированного кремния [3]. Их отличительной особенностью является излучение света при обратном смещении p - n переходов или контактов Шоттки в режиме лавинного пробоя. Использование лавинного механизма пробоя обеспечивает высокое быстродействие таких устройств.

В работе [4] представлена методика формирования матрицы лавинных фотодиодов на основе наноструктурированного кремния. Светоизлучающая часть матрицы представляет из себя ячейку встречноключенных лавинных светодиодов, состоящую из двух контактов Шоттки, а также из слоя анодного оксида алюминия, разделяющего алюминиевые электроды. Наноструктурированный кремний встроен в матрицу оксида алюминия в качестве активного

материала. По данной методике была изготовлена матрица светодиодов (рис. 1), состоящая из 35 ячеек.

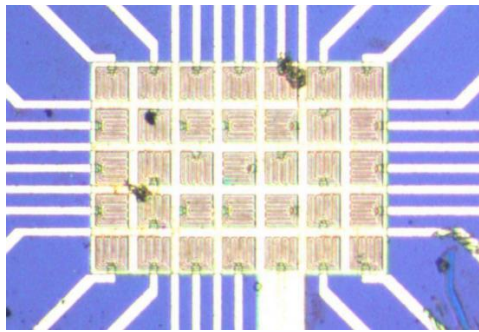


Рис. 1. Матрица лавинных светодиодов

Светодиодные ячейки объединены одним контактом с подложкой таким образом, что образуют структуру, представленную на рис. 2.

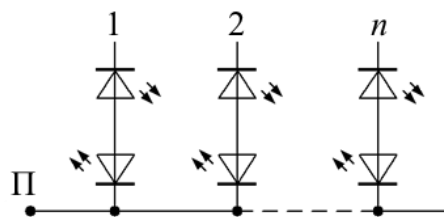


Рис. 2. Электрическая схема матрицы светодиодов: 1, 2, n – выводы подключения светодиодов; П – вывод подключения подложки

Для исследования сформированной матрицы светодиодов был изготовлен тестовый образец оптической системы. Тестовый образец представляет собой металлический корпус для защиты от помех и внешних источников света, в котором размещены светоизлучающая матрица и оптически связанный с ней фотодетектор.

В качестве фотодетектора для тестового образца выбран $p-i-n$ фотодиод типа BPW34. Он обладает высоким быстродействием (до 100 нс), высокой чувствительностью, а также широким спектром принимаемых длин волн (от 430 нм до 1100 нм).

Результаты и их обсуждение

На рис. 3 представлены микрофотографии светодиодной ячейки сформированной матрицы светодиодов. При увеличении питающего напряжения между ячейкой светодиодов и подложкой до напряжения лавинного пробоя обратносмещенного светодиода ячейки, наблюдается излучение света вдоль периметра структуры (рис. 3, б). Соответственно напряжение питания для данной матрицы можно подавать как прямой, так и обратной полярности.

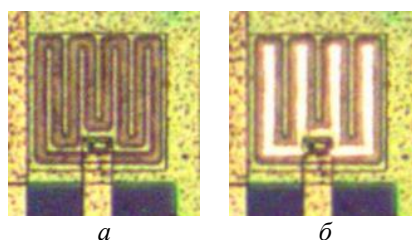


Рис. 3. Микрофотографии сформированной ячейки матрицы светодиодов: а – в отсутствии питающего напряжения; б – при подаче напряжения питания

Для исследования характеристик тестового образца использовался высокоточный двухканальный источник-измеритель Keithley2604B. Источник-измеритель позволяет формировать ступенчатый сигнал с заданной длительностью и амплитудой ступеньки, а наличие

второго канала дает возможность подключать и снимать характеристики светодиода и фотодетектора одновременно. Для поддержания постоянной температуры при измерении характеристик используется термокамера, которая позволяет поддерживать температуру внутри нее в диапазоне от -20 до $+70$ °C с точностью не хуже 1 % от установленного значения.

Вольт-амперные характеристики ячейки светодиодов измерялись при прямом и обратном подключении источника питания и представлены на рис. 4, а, б соответственно.

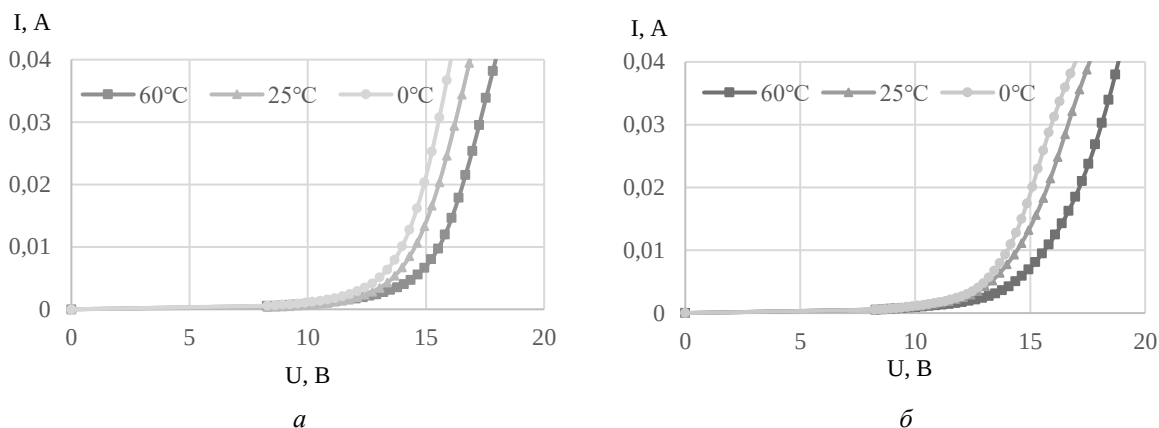


Рис. 4. ВАХ ячейки светодиодов: а – при прямом смещении; б – при обратном смещении

На графиках наблюдается смещение кривой характеристики вправо при увеличении температуры (как при прямом подключении, так и при обратном), что соответствует положительному температурному коэффициенту напряжения, характерному для лавинных пробоев полупроводника. Кривые на рис. 4, б имеют меньшую крутизну, по сравнению с прямым включением, однако характеристики практически идентичны, что можно объяснить одинаковыми параметрами двух светодиодов в данной ячейке.

Измерения сигнала ответа на фотодетекторе проводились в фотодиодном режиме работы (при подаче постоянного обратного смещения на фотодиод); в режиме короткого замыкания фотодиода на измеритель и в фотогальваническом режиме работы фотодиода при нагрузке на измеритель. Полученные данные ответа фотодиода имеют достаточно малые абсолютные значения и нагляднее всего отображаются при фотогальваническом режиме работы фотодиода, поэтому передаточные характеристики тестового образца представлены графиками зависимостей ЭДС фотодиода от тока светодиодной матрицы для прямого и обратного подключения (рис. 5).

На графиках можно наблюдать характерное для фотодиодов нелинейное увеличение фото-ЭДС. Также с увеличением температуры системы фото-ЭДС падает, что может быть связано с увеличением собственных шумов фотодиода. Наиболее эффективным режимом работы для тестового образца является импульсный, так как он позволяет уменьшить саморазогрев системы.

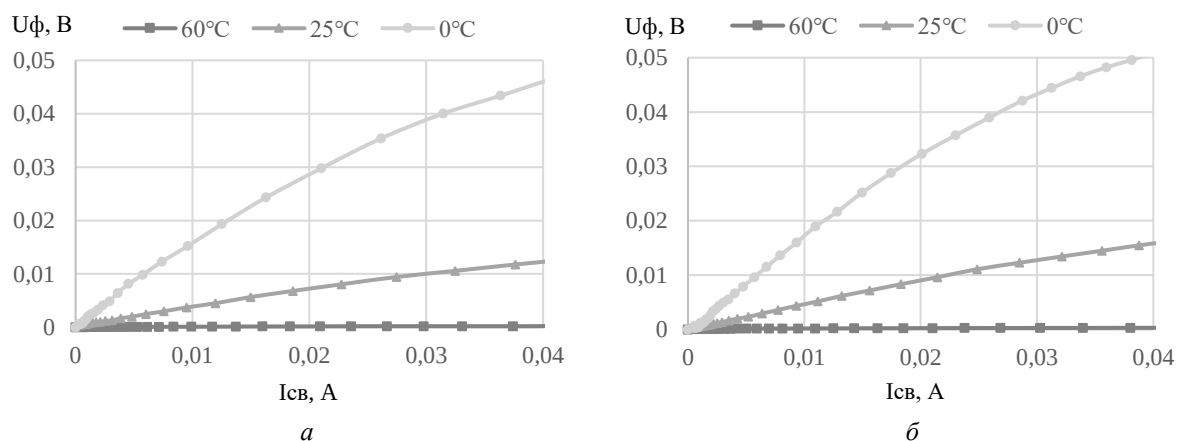


Рис. 5. Передаточные характеристики тестового образца: а – при прямом смещении; б – при обратном смещении

Заключение

Проведенные исследования показали возможность создания устройств, включающих в свой состав светоизлучающие структуры на основе кремния. Полученные результаты можно использовать при разработке устройств оптоэлектроники на основе наноструктурированного кремния, что позволяет добиться интеграции всех составляющих оптических систем на одном кристалле. В частности, формирование квантовых точек на основе кремния позволит в перспективе создать однофотонный генератор случайных чисел, который может быть интегрирован с другими устройствами, чем значительно увеличит эффективность и быстродействие таких устройств.

STUDY OF OPTICAL LIGHT EMITTERS BASED ON POROUS SILICON

A.Yu. KLIUTSKI, S.K. LAZAROUK

Abstract. The prospect of using silicon-based light-emitting structures in coding systems are considered. The characteristics of avalanche LEDs based on nanostructured silicon have been investigated. The possibility of creating optical interconnection systems based on silicon are shown.

Keywords: silicon optical interconnects, LEDs based on por-Si.

Список литературы

1. Балыгин К.А. // Письма в ЖЭТФ. 2017. № 106. С 451–458.
2. Трегулов В.В. Пористый кремний: технология, свойства, применение. Рязань.: Ряз. гос. ун-т им. С.А. Есенина, 2011.
3. Лазарук С.К. [и др.] // Докл. БГУИР. 2004. № 3 (7). С. 27–37.
4. Ле Динь Ви [и др.] // Докл. БГУИР. 2019. № 7–8 (126). С. 165–172.

УДК 621.391.63

МЕЖКАНАЛЬНЫЕ ПЕРЕКРЕСТНЫЕ ПОМЕХИ В ВОСХОДЯЩЕМ НАПРАВЛЕНИИ В СЕТЯХ СЛЕДУЮЩЕГО ПОКОЛЕНИЯ NG-PON2

Н.Н. СЕРГЕЕВ

Белорусский государственный университет информатики и радиоэлектроники, Республика Беларусь

Поступила в редакцию 08 марта 2021

Аннотация. Проведен анализ перекрестных помех в сетях NG-PON2 в восходящем направлении. Показано, такой новый тип сети следующего поколения как NG-PON2 может спокойно функционировать с учетом возможных перекрестных помех сигналов.

Ключевые слова: перекрестные помехи, пассивная оптическая сеть, многоволновая архитектура

Введение

В системах PON характеристика отношения перекрестных помех к сигналу применялись только относительно к ONT. Сеть NG-PON2 представляет собой многоволновую архитектуру [1], и поэтому сигналы с длинами волн NG-PON2, не предназначенные для конкретного приемника ONT либо OLT. Эти сигналы можно рассматривать как мешающие. В статье рассмотрены перекрестные помехи в восходящем направлении передачи.

Отношение перекрестных помех к сигналу в NG-PON2

Рассмотрим упрощенную структурную схему TWDM PON, без линии, для расчета межканальных перекрестных помех в восходящем направлении.

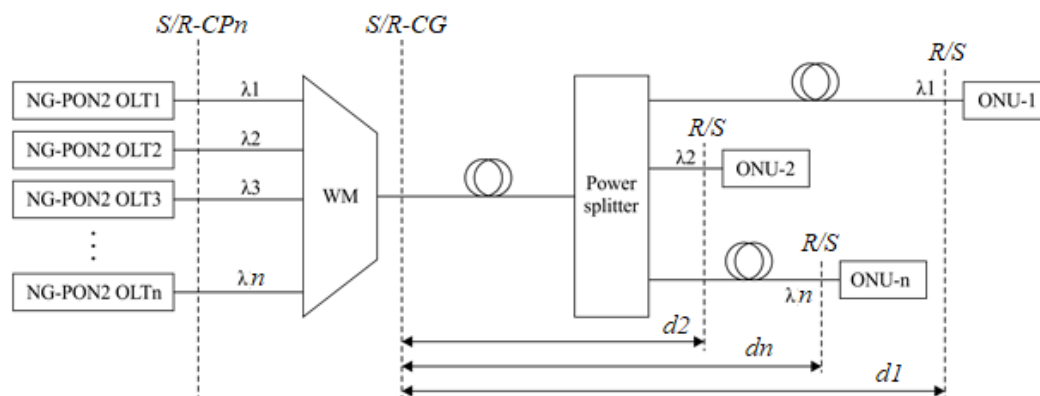


Рис. 1. Упрощенная структурная схема TWDM PON

ONU расположены от OLT на различном расстоянии, чем больше расстояние, тем больше потери в оптическом тракте в восходящем направлении, тогда как порты OLT NG-PON2 расположены рядом друг с другом, что не дает дифференциальных потерь в оптическом тракте в нисходящем направлении. Максимальные дифференциальные потери в оптическом тракте определены как 15 дБ [2].

Спектр мощности сигнала в точке $S/R-CG$ перед WM на центральной АТС и спектр мощности сигнала в точке R/S перед каждым ONU различаются по мощности. Согласно данным рекомендации G.984.3, разница между средней максимальной мощностью и средней минимальной мощностью передатчика OLT для нисходящего направления составляет 4 дБ. Учитывая 1,5 дБ для однородности устройства WM, максимальная разница мощности между

нисходящими сигналами в точке R/S перед ONU может достигать 5,5 дБ. В восходящем направлении потери в оптическом тракте 15 дБ требуется добавить к разнице между максимальной средней мощностью и средней минимальной мощностью передатчика 5 дБ [3]. Максимальная разница в мощности между восходящими сигналами в точке $S/R - CG$ перед устройством WM может достигать 20 дБ.

Разность мощностей восходящего сигнала в точке $S/R - CG$ перед WM выше на 14,5 дБ, чем разность мощности нисходящего направления [4]. Межканальные перекрестные помехи в восходящем направлении могут быть рассчитаны с использованием следующего уравнения:

$$C_c = \Delta P_{ONU} + d_{MAX} + 10 \log_{10} \left(2 \times 10^{-\frac{I_A}{10}} + (N-3) \times 10^{-\frac{I_{NA}}{10}} \right), \text{ дБ} \quad (1)$$

где I_A и I_{NA} – значения для смежного и несмежного канала устройства WM, соответственно, N – общее количество каналов. А при гауссовой аппроксимации, перекрестные помехи P_c можно рассчитать с помощью следующего уравнения [5]:

$$P_c = -5 \log \left(1 - \frac{10^{\frac{2C_c}{10}}}{N-1} Q^2 \left(\frac{ER+1}{ER-1} \right)^2 \right), \quad (2)$$

где $Q = \sqrt{2\text{erfc}^{(-1)}} (2 \times BER)$.

Для восходящего направления $BER = 10^{-4}$, а $Q = 3,72$. Оптические потери при линейном коэффициенте ослабления 0,16 (8 дБ) показаны в зависимости от межканальных перекрестных помех для случаев общего количества каналов $N = 4$ и $N = 8$ на рис. 2. Сплошной линией на рис. 2 обозначена зависимость для 4 каналов, при $Q = 3,72$, а прерывистая обозначает зависимость для 8 каналов при $Q = 3,72$.

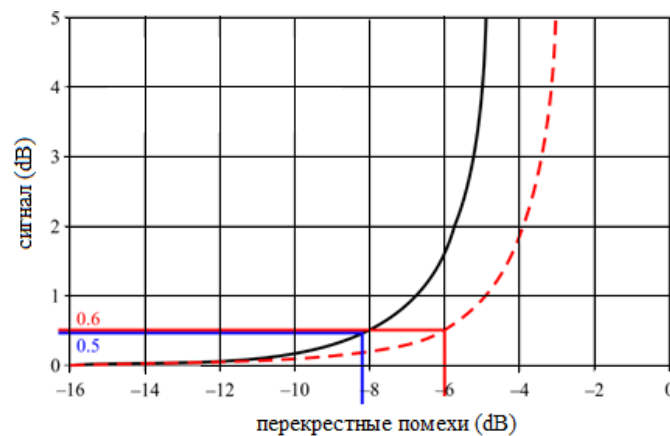


Рис. 2. Оптические потери от межканальных перекрестных помех для $N = 4$ и $N = 8$

Вычисленные значения межканальных перекрестных помех сведены в таблицу для четырехканальных систем и восьмиканальных систем с примерами, существующих WM и фильтров.

Рассчитанные данные межканальных перекрестных помех для различных фильтров (WM)

Тип системы	Межканальные перекрестные помехи	AWG	Фильтр с плотной полосой пропускания	Каскадный фильтр
		$I_A = 23 \text{ дБ},$ $I_{NA} = 30 \text{ дБ}$	$I_A = 26 \text{ дБ},$ $I_{NA} = 33 \text{ дБ}$	$I_A = 32 \text{ дБ},$ $I_{NA} = 36 \text{ дБ}$
4 канальная система	C_c (дБ)	0,4	-2,6	-8,2
	P_c (дБ)	∞	∞	0,5
8 канальная система	C_c (дБ)	1,8	-1,2	-6
	P_c (дБ)	∞	∞	0,6

Заключение

При наихудшем варианте проектирования требуется WM с очень высокими характеристиками изоляции каналов ($I_A = 32 \text{ дБ}$ и $I_{NA} = 36 \text{ дБ}$) для ограничения TWDM PON с оптическими потерями 0,6 дБ, добавленными из-за межканальных перекрестных помех. Используя приведенные значения NG-PON2, сеть может спокойно функционировать с учетом возможных перекрестных помех сигналов.

UPLINK CROSS-CHANNEL CROSSTALK IN NEXT-GENERATION NETWORKS NG-PON2

N.N. SERGEEV

Abstract. The analysis of crosstalk in NG-PON2 networks in the ascending direction is carried out. It is shown that such a new type of next-generation network as NGPON 2 can safely function taking into account possible crosstalk of signals.

Keywords: crosstalk, passive optical network, multi-wave architecture.

Список литературы

1. ITU-T Recommendation G.983.1. Broadband optical access networks based on passive optical networks. 2005. P. 11–20.
2. ITU-T Recommendation G.989.2. NG-PON2: Physical media dependent layer specification. 2019. P. 9 –11.
3. ITU-T Recommendation G.989.3. 40-Gigabit-capable passive optical networks (NG-PON2): Transmission convergence layer specification. 2020. P. 21–28.
4. Nettet D. // Journal of Lightwave Technology. London. 2015. Vol. 33. P. 1136–1143.
5. Giorgi L., Sfameni G., Cavaliere F., [et. al] // Opt. Fiber Commun. Conf. Expo. 2013. P. 91–93.

СВЕДЕНИЯ ОБ АВТОРАХ

- | | | |
|----|-----------------------------------|---|
| 1 | Азарчик Олег Александрович | – магистрант кафедры
инфокоммуникационных технологий
БГУИР |
| 2 | Аксенов Вячеслав Анатольевич | – старший преподаватель кафедры
инфокоммуникационных
технологий БГУИР |
| 3 | Алексееenko Александр Эдуардович | – магистрант кафедры
инфокоммуникационных
технологий БГУИР |
| 4 | Алисеенко Маргарита Александровна | – аспирант кафедры
инфокоммуникационных технологий
БГУИР |
| 5 | Аль-Масри Абдель Хуссейн | – аспирант кафедры
инфокоммуникационных технологий
БГУИР |
| 6 | Аль-Хаджи Слейман Кхалед | – аспирант кафедры
инфокоммуникационных технологий
БГУИР |
| 7 | Белоконь Евгений Анатольевич | – магистрант кафедры
инфокоммуникационных
технологий БГУИР |
| 8 | Борискевич Анатолий Антонович | – д.т.н., профессор кафедры
инфокоммуникационных технологий
БГУИР |
| 9 | Борискевич Илья Анатольевич | – к.т.н., доцент кафедры
инфокоммуникационных
технологий БГУИР |
| 10 | Вашкевич Евгений Кириллович | – магистрант кафедры
инфокоммуникационных
технологий БГУИР |
| 11 | Вишняков Владимир Анатольевич | – д.т.н., профессор кафедры
инфокоммуникационных технологий
БГУИР |
| 12 | Давыдова Надежда Сергеевна | – к.т.н., доцент кафедры
инфокоммуникационных
технологий БГУИР |
| 13 | Жэнь Сюньхуань | – аспирант кафедры
инфокоммуникационных
технологий БГУИР |
| 14 | Захаренко Ангелина Валерьевна | – магистрант кафедры
инфокоммуникационных
технологий БГУИР |
| 15 | Качан Дмитрий Александрович | – соискатель кафедры
инфокоммуникационных
технологий БГУИР |

- | | | |
|----|-----------------------------------|---|
| 16 | Клюцкий Алексей Юрьевич | – аспирант кафедры информационных радиотехнологий БГУИР |
| 17 | Конопелько Валерий Константинович | – д.т.н., профессор кафедры инфокоммуникационных технологий БГУИР |
| 18 | Лазарук Сергей Константинович | – д.ф.-м.н., заведующий НИЛ 4.12 БГУИР |
| 19 | Липкович Эдуард Борисович | – доцент кафедры инфокоммуникационных технологий БГУИР |
| 19 | Ма Цзюнь | – аспирант кафедры инфокоммуникационных технологий БГУИР |
| 20 | Маринич Павел Сергеевич | – магистрант кафедры инфокоммуникационных технологий БГУИР |
| 21 | Машков Евгений Александрович | – магистрант кафедры инфокоммуникационных технологий БГУИР |
| 22 | Митюхин Анатолий Иванович | – доцент кафедры инфокоммуникационных технологий БГУИР |
| 23 | Саломатин Сергей Борисович | – к.т.н., доцент кафедры инфокоммуникационных технологий БГУИР |
| 24 | Сергеев Николай Николаевич | – соискатель кафедры инфокоммуникационных технологий БГУИР |
| 25 | Смоляк Сергей Владимирович | – магистрант кафедры инфокоммуникационных технологий БГУИР |
| 26 | Тютрюмова Анастасия Валерьевна | – аспирант кафедры инфокоммуникационных технологий БГУИР |
| 27 | Цветков Виктор Юрьевич | – д.т.н., заведующий кафедрой инфокоммуникационных технологий БГУИР |
| 28 | Шевчук Оксана Геннадьевна | – к.т.н., доцент кафедры инфокоммуникационных технологий |

This image shows a single sheet of white paper with horizontal blue ruling lines. The lines are evenly spaced and run across the width of the page. There are no margins, text, or other markings on the paper.

[illegible]

This image shows a single sheet of white paper with horizontal ruling lines. The lines are evenly spaced and run across the width of the page. There are no margins, text, or other markings on the paper.