

*Журнал «Научное обозрение.
Технические науки»
зарегистрирован Федеральной службой
по надзору в сфере связи, информационных
технологий и массовых коммуникаций.
Свидетельство ПИ № ФС77-57440
ISSN 2500-0799*

**Двухлетний импакт-фактор РИНЦ – 0,887
Пятилетний импакт-фактор РИНЦ – 0,350**

*Учредитель, издательство и редакция:
ООО НИЦ «Академия Естествознания»,
Почтовый адрес: 105037, г. Москва, а/я 47
Адрес редакции и издателя: 410056, Саратовская
область, г. Саратов, ул. им. Чапаева В.И., д. 56*

**Founder, publisher and edition:
LLC SPC Academy of Natural History,
Post address: 105037, Moscow, p.o. box 47
Editorial and publisher address: 410056,
Saratov region, Saratov, V.I. Chapaev Street, 56**

*Подписано в печать 30.12.2022
Дата выхода номера 31.01.2023
Формат 60×90 1/8*

*Типография
ООО НИЦ «Академия Естествознания»,
410035, Саратовская область,
г. Саратов, ул. Мамонтовой, д. 5*

**Signed in print 30.12.2022
Release date 31.01.2023
Format 60×90 8.1**

**Typography
LLC SPC «Academy Of Natural History»
410035, Russia, Saratov region,
Saratov, 5 Mamontovoi str.**

*Технический редактор Доронкина Е.Н.
Корректор Галенкина Е.С., Дудкина Н.А.*

*Тираж 1000 экз.
Распространение по свободной цене
Заказ НО 2022/6
© ООО НИЦ «Академия Естествознания»*

Журнал «НАУЧНОЕ ОБОЗРЕНИЕ» выходил с 1894 по 1903 год в издательстве П.П. Сойкина. Главным редактором журнала был Михаил Михайлович Филиппов. В журнале публиковались работы Ленина, Плеханова, Циолковского, Менделеева, Бехтерева, Лесгафта и др.

Journal «Scientific Review» published from 1894 to 1903. P.P. Soykin was the publisher. Mikhail Filippov was the Editor in Chief. The journal published works of Lenin, Plekhanov, Tsiolkovsky, Mendeleev, Bekhterev, Lesgaft etc.



М.М. Филиппов (M.M. Philippov)

**С 2014 года издание журнала возобновлено
Академией Естествознания**

**From 2014 edition of the journal resumed
by Academy of Natural History**

**Главный редактор: М.Ю. Ледванов
Editor in Chief: M.Yu. Ledvanov**

Редакционная коллегия (Editorial Board)
А.Н. Курзанов (A.N. Kurzanov)
Н.Ю. Стукова (N.Yu. Stukova)
М.Н. Бизенкова (M.N. Bizenkova)
Н.Е. Старчикова (N.E. Starchikova)
Т.В. Шнуровозова (T.V. Shnurovozova)

НАУЧНОЕ ОБОЗРЕНИЕ • ТЕХНИЧЕСКИЕ НАУКИ

SCIENTIFIC REVIEW • TECHNICAL SCIENCES

www.science-education.ru

2022 г.



***В журнале представлены научные обзоры,
статьи проблемного
и научно-практического характера***

***The issue contains scientific reviews,
problem and practical scientific articles***

СОДЕРЖАНИЕ

Технические науки

СТАТЬИ

ИСПОЛЬЗОВАНИЕ МЕТОДА МОНТЕ-КАРЛО В ОБУЧЕНИИ С ПОДКРЕПЛЕНИЕМ	
<i>Акимов А.А., Малагина Е.С.</i>	5
МЕТОД ВЫБОРА КОЛИЧЕСТВА УРОВНЕЙ СИГНАЛЬНОГО СОЗВЕЗДИЯ СИГНАЛОВ С КВАДРАТУРНОЙ АМПЛИТУДНОЙ МОДУЛЯЦИЕЙ С УЧЕТОМ ВЗАИМНОГО ВЛИЯНИЯ СИГНАЛОВ В МНОГОКАНАЛЬНОЙ ТЕЛЕКОММУНИКАЦИОННОЙ СИСТЕМЕ	
<i>Власов С.В.</i>	12
ИСПОЛНЕНИЕ И АВТОРИЗАЦИЯ EMV-ТРАНЗАКЦИЙ	
<i>Ермаков К.С., Сидякин И.М.</i>	17
ГОСУДАРСТВЕННЫЕ ИНФОРМАЦИОННЫЕ СИСТЕМЫ. ВОПРОСЫ БЕЗОПАСНОСТИ	
<i>Кечкина Н.И., Наумова Е.Г.</i>	22
НАУЧНЫЙ ОБЗОР	
АНАЛИЗ, ПРОБЛЕМЫ И ВОЗМОЖНОСТИ ИСПОЛЬЗОВАНИЯ БОЛЬШИХ ДАННЫХ В ГОРОДСКОМ ХОЗЯЙСТВЕ	
<i>Массеров Д.Д., Массеров Д.А.</i>	27
СТАТЬИ	
ИСПОЛЬЗОВАНИЕ ГРАФА СОАВТОРСТВА ДЛЯ ОПРЕДЕЛЕНИЯ АВТОРСТВА ТЕКСТА С НЕСКОЛЬКИМИ АВТОРАМИ	
<i>Батурин М.М., Белов Ю.С.</i>	32
МОДИФИКАЦИИ АРХИТЕКТУРЫ WAVENET ДЛЯ РЕАЛИЗАЦИИ ВОКОДЕРА В ГЕНЕРАТИВНОЙ МОДЕЛИ ПРЕОБРАЗОВАНИЯ ТЕКСТА В РЕЧЬ	
<i>Белоножко П.Е., Белов Ю.С.</i>	37

CONTENTS

Technical sciences

ARTICLES

APPLYING THE MONTE CARLO METHOD IN REINFORCEMENT LEARNING

Akimov A.A., Malagina E.S. 5

METHOD FOR SELECTING THE NUMBER OF LEVELS OF A SIGNAL CONSTELLATION OF SIGNALS WITH QUADRATIVE AMPLITUDE MODULATION TAKING INTO ACCOUNT THE MUTUAL INFLUENCE OF SIGNALS IN A MULTI-CHANNEL TELECOMMUNICATION SYSTEM

Vlasov S.V. 12

EMV-TRANSACTION PROCESSING AND AUTHORIZATION

Ermakov K.S., Sidyakin I.M. 17

STATE INFORMATION SYSTEMS. SECURITY QUESTIONS

Kechkina N.I., Naumova E.G. 22

REVIEW

ANALYSIS, CHALLENGES AND OPPORTUNITIES FOR THE USE OF BIG DATA IN URBAN MANAGEMENT

Masserov D.D., Masserov D.A. 27

ARTICLES

USING THE CO-AUTHORSHIP GRAPH TO DETERMINE THE AUTHORSHIP OF A TEXT WITH MULTIPLE AUTHORS

Baturin M.M., Belov Yu.S. 32

MODIFICATIONS OF THE WAVENET ARCHITECTURE FOR THE IMPLEMENTATION OF A VOCODER IN A GENERATIVE MODEL OF TEXT-TO-SPEECH CONVERSION

Belonozhko P.E., Belov Yu.S. 37

СТАТЬИ

УДК 004.023

**ИСПОЛЬЗОВАНИЕ МЕТОДА МОНТЕ-КАРЛО
В ОБУЧЕНИИ С ПОДКРЕПЛЕНИЕМ****Акимов А.А., Малагина Е.С.***ФГБОУ ВО «Башкирский государственный университет», Уфа, e-mail: andakm@rambler.ru*

Обучение с подкреплением – наиболее быстро развивающаяся область, которая связана с созданием искусственных интеллектуальных систем. На данный момент эта отрасль довольно обширна. Многие исследователи по всему миру активно работают с обучением с подкреплением в разнообразных областях: нейробиология, теория управления, психология и многое другое. Целью данной работы является изучение возможностей использования метода Монте-Карло в обучении с подкреплением. Главное при обучении с подкреплением – это зафиксировать основные аспекты реальной проблемы при взаимодействии обучающегося с окружающим миром для достижения своей цели. То есть агент обучения должен иметь цель, связанную с состоянием окружающей среды. Также учащийся должен иметь возможность ощущать среду и совершать действия, влияющие на нее. Формулировка задачи обучения с подкреплением должна учитывать все три аспекта – ощущение, действие и цель – в их наиболее простых формах. Методы Монте-Карло способны решить проблемы обучения с подкреплением, основываясь на усреднении результатов выборки. Чтобы обеспечить доступность четко определенных результатов, определим методы Монте-Карло только для эпизодических задач. Таким образом, методы Монте-Карло могут быть инкрементными на уровне эпизодов.

Ключевые слова: метод Монте-Карло, обучение с подкреплением, искусственный интеллект, изучение-применение, численные методы

**APPLYING THE MONTE CARLO METHOD
IN REINFORCEMENT LEARNING****Akimov A.A., Malagina E.S.***Bashkir State University, Ufa, e-mail: andakm@rambler.ru*

Reinforcement learning is the most rapidly developing area that is associated with the creation of artificial intelligence systems. At the moment, this industry is quite extensive. Many researchers around the world are actively working with reinforcement learning in a variety of fields: neuroscience, control theory, psychology, and more. The purpose of this work is to study the possibilities of using the Monte Carlo method in reinforcement learning. The main thing in reinforcement learning is to fix the main aspects of the real problem in the interaction of the learner with the outside world in order to achieve his goal. That is, the learning agent must have a goal related to the state of the environment. Also, the student should be able to feel the environment and perform actions that affect it. Reinforcement learning problem formulation should take into account all three aspects – feeling, action and goal – in their simplest forms. Monte Carlo methods are able to solve reinforcement learning problems based on averaging the results of a sample. To ensure that well-defined results are available, we define Monte Carlo methods for episodic problems only. Thus Monte Carlo methods can be incremental at the episode level.

Keywords: Monte Carlo method, reinforcement learning, artificial intelligence, study-application, numerical methods

Если задуматься о том, как человек обучается, то, скорее всего, первой мыслью, приходящей в голову, будет идея, что этот процесс происходит при взаимодействии человека с окружающей средой. Контакты с окружающей средой, вне всяких сомнений, являются источником знаний как и о самом человеке, так и об окружающей его среде. Причем этот процесс длится на протяжении всей жизни человека. Взаимодействия с окружающей средой дают полную информацию о связи причин и следствий, о последовательности действий, которые нужно выполнить, чтобы добиться определенных целей. Именно обучение через взаимодействие является той основополагающей идеей, на которой базируются почти все теории обучения и интеллекта.

Изучение того, как сопоставить ситуации с действиями, чтобы получить максимальную выгоду, называется обучением

с подкреплением. Обучающийся не знает заранее, какие действия необходимо предпринять, чтобы максимизировать вознаграждение. Поэтому он должен самостоятельно выяснить, какие действия нужно для этого сделать. К тому же, нужно учитывать, что выбор действия влияет не только на вознаграждение на данном этапе, но и на то, какую выгоду агент обучения может получить в дальнейшем. Эти две характеристики – поиск методом проб и ошибок и отложенное вознаграждение – являются двумя наиболее важными отличительными чертами обучения с подкреплением.

Главное при обучении с подкреплением – это зафиксировать основные аспекты реальной проблемы при взаимодействии обучающегося с окружающим миром для достижения своей цели. То есть агент обучения должен иметь цель, связанную с состоянием окружающей среды. Также

учащийся должен иметь возможность ощущать среду и совершать действия, влияющие на нее.

Формулировка задачи обучения с подкреплением должна учитывать все три аспекта – ощущение, действие и цель – в их наиболее простых формах.

Агент обучения, чтобы получить наибольшее вознаграждение, как правило, отдаст предпочтение уже проверенным действиям, то есть тем, которые оказались эффективными для него в прошлом и дали ему лучшее вознаграждение. Но он не сможет найти такие действия, если бы ранее он не пробовал действия, которые раньше не выбирал. Из-за этого, одной из проблем, возникающей при обучении с подкреплением, является компромисс между изучением и применением. Получается, что обучающийся должен использовать то, что он уже испытал, чтобы получить вознаграждение, но он также должен изучить новое, чтобы сделать лучший выбор действий в будущем [1]. Агент обучения всегда должен пробовать различные действия, в дальнейшем отдавая предпочтения тем, которые оказались лучшими. Нельзя только использовать проверенные действия или только искать новые – в этом и состоит проблема.

В стохастической задаче каждое действие должно быть опробовано много раз, чтобы получить надежную оценку ожидаемого вознаграждения. Дилемма «изучение – применение» интенсивно изучается математиками на протяжении многих десятилетий, но до сих пор остается нерешенной.

Выделяют еще четыре основных элемента, помимо агента обучения и среды, входящих в систему обучения с подкреплением:

1. Стратегия.
2. Вознаграждение.
3. Ценность состояния.
4. Модель среды (необязательно).

Сопоставление состояний окружающей среды с действиями, которые должны быть выполнены в этих состояниях, называется стратегией. Проще говоря, именно стратегия определяет то, какой способ поведения выберет агент обучения в конкретный момент времени. В общем случае стратегия может быть случайной. В некоторых случаях она выражается функцией или таблицей, в более сложных вариантах – может даже включать какие-либо вычисления. Как правило, стратегии вполне достаточно для определения поведения агента обучения, является ядром агента.

Целью задачи обучения с подкреплением является вознаграждение – число, которое получает обучающийся на каждом шаге среды. Вознаграждение несет как по-

ложительный смысл, так и отрицательный для агента обучения, так как его основная цель – это максимизация общего вознаграждения, которое агент планирует получить в долгосрочной перспективе. Если провести аналогию с биологией, то вознаграждение можно сравнить с опытом боли или удовольствия. Если агент обучения получил низкое вознаграждение после действия, выбранного какой-либо стратегией, то это может являться основанием для смены стратегии в будущем. В общем случае вознаграждения могут быть случайными.

Ценность состояния – это общая сумма вознаграждения, которую агент обучения может получить в будущем, начиная с этого состояния. Как видно из определения, основным отличием ценности от вознаграждения является то, что она определяет то, что хорошо в дальнейшей перспективе. То есть конкретное состояние может давать небольшое моментальное вознаграждение, но за ним могут идти состояния, приносящие высокую выгоду, а следовательно, такое состояние будет иметь высокую ценность. Обратная ситуация так же может быть правдой. Если провести человеческую аналогию, то ценности соответствуют тому, насколько обучающийся доволен или недоволен тем, что его окружение находится в конкретном состоянии, тогда как вознаграждения в чем-то схожи с удовольствием, если значение выгоды высокое, или же с болью, если выгода низкая.

Цель оценки ценностей – получение большего вознаграждения. Без вознаграждения не может быть и ценностей, поэтому они в некотором роде первичны, а ценности – вторичны. Но, несмотря на это, именно ценности более интересны для принятия и оценки решений. Так как ценность рассматривает выгоду именно в долгосрочной перспективе, то выбор действий осуществляется именно на основе оценочных суждений, то есть таких действий, которые приводят к состояниям наивысшей ценности, а не наивысшего вознаграждения. К сожалению, определить вознаграждения намного легче, чем ценности. Ценности необходимо вычислять снова и снова из всей последовательности наблюдений, тогда вознаграждения в основном можно получить непосредственно из самой среды. Наиболее важный компонент практически всех алгоритмов обучения с подкреплением – это метод эффективной оценки значений функции ценностей.

И, наконец, последним и необязательным элементом является модель окружающей среды. С помощью модели можно делать выводы о том, как поведет себя среда, то есть

задача модели – имитировать поведение самой окружающей среды. При помощи модели можно рассматривать возможные будущие ситуации и, в зависимости от этого, принимать решения о курсе действий. То есть модели используются для планирования: учитывая состояние и действие, она может предсказать то, в каком состоянии окажется окружающая среда, и соответствующее ему вознаграждение.

Таким образом, если для решения задач обучения с подкреплением используются модели и планирование, то методы решения называются методами, основанными на моделях. Противоположностью этих методов являются более простые методы, которые не используют модели, а учатся методом проб и ошибок.

Методы Монте-Карло – общее название группы численных методов. Они базируются на получении как можно большего количества реализаций случайного процесса, который формируется так, чтобы его вероятностные характеристики совпадали с аналогичными величинами решаемой задачи [2].

Методы Монте-Карло способны решить проблемы обучения с подкреплением, основываясь на усреднении результатов выборки. Чтобы обеспечить доступность четко определенных результатов, определим методы Монте-Карло только для эпизодических задач. Предполагается, что данные делятся на эпизоды, которые, так или иначе, будут завершены, независимо от того, какие действия выбраны. Только после того, как эпизод завершится, может произойти оценивание ценности или изменение стратегии. Таким образом, методы Монте-Карло могут быть инкрементными на уровне эпизодов.

Начнем с рассмотрения методов Монте-Карло для изучения функции значения состояния для заданной стратегии. Вспомним, что ценность состояния – это будущее накопленное вознаграждение, начиная с данного состояния. Таким образом, можно оценить выгоду, усреднив результаты полученной выгоды после пройденного состояния. Если число наблюдений будет расти, то, следовательно, количество значений выгоды также будет увеличиваться. А значит, среднее значение выгоды будет стремиться к ожидаемой величине. Данная идея лежит в основе методов Монте-Карло.

Введем следующие обозначения: пусть s – состояние, π – стратегия. Учитывая набор эпизодов, которые получились с помощью применения стратегии и прохождения через состояние, оценим ценность состояния s при стратегии π . Эта величина будет обозначаться следующим образом – $v_{\pi}(s)$.

Посещением s называется каждое появление состояния s в эпизоде. Конечно, s может посетить один и тот же эпизод несколько раз. Назовем первое посещение в эпизоде первым посещением s . Метод Монте-Карло первого посещения оценивает $v_{\pi}(s)$ как усреднение значения выгоды, которые соответствуют первым посещениям s , тогда как метод Монте-Карло всех посещений оценивает величину как среднее значение после всех посещений s в эпизодах. Эти два метода Монте-Карло очень похожи, но имеют разные теоретические свойства.

В случае если число посещений стремится к бесконечности, то результат, который получается при использовании любого из вышеуказанных методов, сходится к $v_{\pi}(s)$.

В методе Монте-Карло оценка одного состояния никоим образом не базируется на оценке какого-либо другого состояния. Таким образом, эти оценки являются независимыми друг от друга. Этот факт является важной особенностью методов Монте-Карло.

Также интересной особенностью данного метода является то, что вычислительные затраты на оценку значения одного состояния не зависят от количества состояний. То есть можно создать большую выборку только нужных для работы эпизодов, не обращая внимания на остальные. И считать среднее значение выгоды только для этой выборки. Такая особенность делает методы Монте-Карло очень полезными, если необходимо оценить ценность только одного или некоторого подмножества состояний.

Если модель присутствует, то для определения стратегий достаточно ценностей состояния. Нужно просто сделать шаг и выбрать такое действие, которое приведет к наилучшему вознаграждению. Но если модель отсутствует, то таких данных будет недостаточно. И в таком случае лучше оценивать значения пар состояние – действие. Чтобы значения были полезны при выборе стратегии, нужно явно оценивать ценность каждого действия.

Таким образом, одной из основных целей для применения методов Монте-Карло является оценка q_{π} . Чтобы достичь этого, сначала рассмотрим задачу оценки стратегии.

Задача оценки стратегии на основе значений действий состоит в том, чтобы оценить $q_{\pi}(s,a)$ – ожидаемую выгоду при начале в состоянии s , выполнении действия a и последующем следовании стратегии π . Метод Монте-Карло для этого случая такой же, как и для рассмотренного ранее случая для значений состояния, за исключением того, что теперь оценивается пара состояние – действие, а не состояния. Ме-

тод Монте-Карло всех посещений оценивает ценность пары состояние – действие как среднее значение выгоды, полученной после всех посещений. Метод Монте-Карло первого посещения усредняет значения выгод после первого посещения в каждом эпизоде состояния s и выбора в нем действия a . Значения, получаемые при использовании этих методов, сходятся квадратично к истинным значениям ожидаемых ценностей, поскольку количество посещений каждой пары состояние – действие приближается к бесконечности.

Единственная сложность заключается в том, что многие пары состояние – действие могут никогда не быть посещены. В случае, когда π – детерминированная стратегия, выгода будет учитываться только для одного из действий каждого состояния. И тогда оценки для других действий не будут улучшаться с опытом, так как значения выгоды не будут усредняться. Вспомним, что целью изучения значений ценности действий является следующее: помощь в выборе действий, доступных в каждом состоянии. Но тогда вышеперечисленное становится серьезной проблемой, так как невозможно сравнить действия между собой, чтобы выбрать наилучшее, поскольку нужно оценить ценность всех действий из каждого состояния, а не только того состояния, которое в данный момент предпочтительно.

Необходимо обеспечить постоянное изучение для того, чтобы оценить через ценность действия стратегию. Чтобы гарантировать, что все пары состояние – действие будут посещены бесконечное число раз при бесконечном числе эпизодов, нужно указать, что первый шаг каждого эпизода

начинается в паре состояние – действие и что каждая пара имеет отличную от нуля вероятность быть выбранной в качестве начала. Это называется предположением об изучающих стартах.

К сожалению, на предположение об изучающих стартах нельзя полагаться в целом, так как стартовые условия не всегда могут быть полезны. Например, при обучении непосредственно на основе фактического взаимодействия с окружающей средой. Другим вариантом для обеспечения появления всех пар состояние – действие может быть подход, который заключается в рассмотрении только стохастических стратегий с ненулевой вероятностью выбора всех действий в каждом состоянии.

Рассмотрим, как можно использовать оценку методом Монте-Карло для формирования управлением, то есть для аппроксимации оптимальных стратегий.

Стратегия улучшается в несколько раз, чтобы приблизиться к функции ценности. Но и функция ценности, в свою очередь, постоянно меняется, чтобы наиболее точно приблизиться к текущей стратегии. Каждый из этих двух видов создает постоянно меняющуюся цель друг для друга, то есть в некоторой степени работают друг против друга. Но, несмотря на это, они приближаются к оптимальности и функцию ценности, и стратегию.

Сначала рассмотрим метод классической итерации по стратегиям. Будем выполнять чередующиеся шаги: сначала полную ее оценку, затем – полное улучшение стратегии. Начнем с произвольной стратегии π_0 , а закончим оптимальной стратегией и оптимальной функцией ценности (рис. 1).

$$\pi_0 \xrightarrow{E} q_{\pi_0} \xrightarrow{I} \pi_1 \xrightarrow{E} q_{\pi_1} \xrightarrow{I} \pi_2 \xrightarrow{E} \dots \xrightarrow{I} \pi_* \xrightarrow{E} q_*$$

Рис. 1. Схема метода

$\xrightarrow{E}$ обозначает полную оценку стратегии, а $\xrightarrow{I}$ – полное улучшение стратегии. Реализуется много эпизодов, где приближительная функция ценности действия асимптотически приближается к истинной функции. Предположим, что будем наблюдать бесконечное количество эпизодов и что, кроме того, они будут генерироваться с помощью изучающих стартов. При этих предположениях методы Монте-Карло будут точно вычислять каждое q_{π_k} для произвольного π_k .

Стратегию можно улучшить, сделав ее «жадной» по отношению к текущей функции ценности. Тогда в данном случае будем иметь функцию действие-ценность, следовательно, чтобы построить «жадную» стратегию, модель не потребуется.

Для любой функции ценности действия q , соответствующей «жадной» стратегии, это такая стратегия, которая для каждого $s \in S$ выбирает действие с максимальной ценностью:

$$\pi(s) = \arg \max_a q(s, a).$$

Затем можно улучшить стратегию, построив каждое π_{k+1} как жадную по отношению к q_{π_k} . Для всех $s \in S$

$$q_{\pi_k}(s, \pi_{k+1}(s)) = q_{\pi_k}(s, \arg \max_a q_{\pi_k}(s, a)) = \\ = \max_a q_{\pi_k}(s, a) \geq q_{\pi_k}(s, \pi_k(s)) \geq v_{\pi_k}(s).$$

Тогда каждая стратегия π_{k+1} лучше, чем π_k , или равна ей в том случае, если они обе являются оптимальными стратегиями. Это, в свою очередь, гарантирует, что весь процесс сходится к оптимальной стратегии и оптимальной функции ценности.

Таким образом, методы Монте-Карло можно использовать для нахождения оптимальных стратегий, учитывая только выборку эпизодов, при отсутствии других знаний о динамике окружающей среды.

Чтобы понять, как работает метод Монте-Карло на практике, рассмотрим следующий пример. Сыграем в карточную игру блэкджек и вычислим функцию ценности.

Суть игры блэкджек состоит в следующем: необходимо собрать карты таким образом, чтобы сумма была максимальной, но при этом не превышала 21. Король, дама, валет имеют значение 10, туз принимает значения 1 либо 11, тогда он называется играющим. Остальные карты имеют значения согласно своему номиналу. Игрок играет со сдающим, независимо от других участников. В начале игры им обоим раздаются по две карты, одна из карт раздающего открывается. Игрок может взять себе еще одну карту, или же он может остановиться. Если он останавливается, то сдающий берет себе карты из колоды до тех пор, пока их сумма не окажется больше или равна 17. Если игрок или сдающий получает в сумме больше 21, то он проигрывает. В остальных случаях выигрывает тот, у кого сумма карт окажется больше, чем у другого. В случае равной суммы – ничья.

Для имитации окружающей среды воспользуемся средой Blackjack библиотеки Gym [3]. Она описывается следующим образом:

1. Каждый эпизод представляет собой марковский процесс принятия решений, в начале которого оба участника получают свои две карты, при этом одна карта сдающего открыта.

2. Эпизод заканчивается в случае, если кто-то выигрывает или игра завершается ничьей. Вознаграждение начисляется в конце эпизода: 1, если игрок выиграл; 0 – ничья; -1, если игрок проиграл.

3. В каждом раунде у игрока есть два действия: получить еще одну карту (1) или больше не брать карт (0).

Посмотрим, как работает данная среда.

Для начала подключим библиотеки PyTorch и Gym и создадим экземпляр окружающей среды Blackjack [4]:

```
import torch
import gym
env = gym.make('Blackjack-v0')
```

Затем приведем среду в исходное состояние командой `env.reset()` и получим следующий результат (рис. 2).

(10, 2, False)

Рис. 2. Исходное состояние

Возвращаются три переменные, которые определяют:

1. Количество очков у игрока – в данном случае 10.

2. Количество очков у сдающего – в данном случае 2.

3. Наличие играющего туза у игрока – в данном случае отсутствует.

Можно попросить еще одну карту командой `env.step(1)`. Получим результат (рис. 3).

((19, 2, False), 0.0, False, {})

Рис. 3. Результат работы команды

После выполнения данной команды возвращаются три переменные состояния (19, 2, False), вознаграждение, равное нулю в данном случае, и признак завершения эпизода – False. После этого игрок перестает брать карты с помощью команды `env.step(0)`.

После этого к действиям приступает сдающий, и в данном случае игрок проигрывает (рис. 4).

((19, 2, False), -1.0, True, {})

Рис. 4. Завершение игры

Теперь перейдем к предсказанию ценности для простой стратегии, когда игрок перестает брать карты, если он набрал 19 очков.

Для начала напомним функцию, которая имитирует эпизод Blackjack при следовании простой стратегии:

```
def run_episode(env, hold_score):
    state = env.reset()
    rewards = []
    states = [state]
    is_done = False
    while not is_done:
        action = 1 if state[0] < hold_score else 0
        state, reward, is_done, info = env.step(action)
        states.append(state)
        rewards.append(reward)
    if is_done:
        break
    return states, rewards
```

Теперь определим функцию, которая оценивает простую стратегию методом Монте-Карло первого посещения:

```
from collections import defaultdict
def mc_prediction_first_visit(env, hold_score, gamma, n_episode):
    V = defaultdict(float)
    N = defaultdict(int)
    for episode in range(n_episode):
        states_t, rewards_t = run_episode(env, hold_score)
        return_t = 0
        G = {}
        for state_t, reward_t in zip(states_t[1::-1], rewards_t[1::-1]):
            return_t = gamma * return_t + reward_t
            G[state_t] = return_t
        for state, return_t in G.items():
            if state[0] <= 21:
                V[state] += return_t
                N[state] += 1
    for state in V:
        V[state] = V[state] / N[state]
    return V
```

Данная функция выполняет следующие действия [5]:

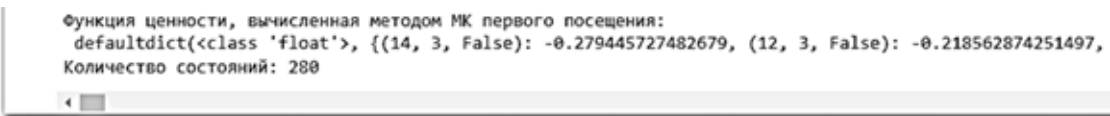
1. Прогоняет `n_episode` эпизодов, следуя простой стратегии.
2. Вычисляет доходы при первом посещении каждого состояния в каждом эпизоде.
3. Усредняет доходы, полученные при первом посещении каждого состояния по всем эпизодам, вычисляя тем самым ценность. Состояния, в которых игрок набрал больше 21 очка, игнорируются, так как вознаграждение в них равно -1.

Далее задаются начальные параметры игры: количество очков, при которых игра останавливается, равное 19; коэффициент обесценивания 1; количество эпизодов 500000:

```
hold_score = 19,
gamma = 1,
n_episode = 500000.
```

Выполним предсказание методом Монте-Карло с данными параметрами, распечатаем получившуюся функцию ценности, выведем количество состояний (рис. 5):

```
value = mc_prediction_first_visit(env, hold_score, gamma, n_episode)
print('Функция ценности, вычисленная методом МК первого посещения:\n',
      value)
print('Количество состояний:', len(value))
```



```
Функция ценности, вычисленная методом МК первого посещения:
defaultdict(<class 'float'>, {(14, 3, False): -0.279445727482679, (12, 3, False): -0.218562874251497,
Количество состояний: 280
```

Рис. 5. Результат работы программы

Заключение

Проведенное исследование дало возможность показать, насколько эффективно можно вычислить функцию ценности 280 состояний в среде BlackJack с помощью предсказания методом Монте-Карло. При этом был применен нестандартный подход к обучению с заранее неизвестными обучающими примерами, которые подбирались автоматически, в процессе оптимизации.

Приведены возможные пути улучшения стратегии:

- Увеличение числа оптимизируемых параметров.
- Применение других способов вознаграждения агента.

– Создание нескольких конкурирующих между собой агентов для увеличения пространства вариантов.

Список литературы

1. Sutton R., Barto A. Reinforcement Learning: An Introduction. MIT Press; second edition, 2018. 552 p. P. 115–124.
2. Кашникова А.П., Беляева М.Б. Метод Монте-Карло в задачах моделирования процессов и систем // Modernscience. 2021. № 1–2. С. 358–362.
3. Шолле Ф. Глубокое обучение на Python. СПб.: Питер, 2018. 400 с. С. 95–118.
4. Лю Ю. (Х.) Обучение с подкреплением на PyTorch: сборник рецептов. М.: ДМК Пресс, 2020. 282 с. С. 122–124.
5. Акимов А.А., Мустафина С.И., Барабанов В.Ф., Морозкин Н.Д., Мустафина С.А. Алгоритмы распознавания девиантного поведения на основе методов искусственного интеллекта // Актуальные проблемы прикладной математики, информатики и механики: сборник трудов Международной научной конференции. Воронеж, 2022. С. 141–149.

МЕТОД ВЫБОРА КОЛИЧЕСТВА УРОВНЕЙ СИГНАЛЬНОГО СОЗВЕЗДИЯ СИГНАЛОВ С КВАДРАТУРНОЙ АМПЛИТУДНОЙ МОДУЛЯЦИЕЙ С УЧЕТОМ ВЗАИМНОГО ВЛИЯНИЯ СИГНАЛОВ В МНОГОКАНАЛЬНОЙ ТЕЛЕКОММУНИКАЦИОННОЙ СИСТЕМЕ

Власов С.В.

ФГБОУ ВО «Пензенский государственный технологический университет», Пенза,
e-mail: vlasov_s.v@mail.ru

Для проведения исследования взаимного влияния сигналов в телекоммуникационных системах с использованием многомерных метрических пространств, предлагается использовать как аналитические, так и графические модели. В данной статье разработан метод выбора количества уровней сигнального созвездия квадратурной многоуровневой многофазовой модуляции с учетом взаимного влияния сигналов по соседним каналам многоканальной телекоммуникационной системы с использованием многомерных метрических пространств. Взаимное влияние оценивается как расстояние между сферами в многопараметрическом пространстве. Результат исследования заключается в том, что при использовании полученных моделей обнаруживается эффект, когда различные передатчики, работающие на различных частотах с использованием квадратурной модуляции, способны создавать помехи (влиять друг на друга) при определенном значении амплитуд и начальных фаз как несущих частот, так и значениях частот, и амплитуд модулирующих колебаний. Научный вывод проведенных исследований: передача информации на различных несущих частотах не исключает взаимного негативного влияния сигналов друг на друга при использовании сигналов с квадратурной модуляцией, и необходимо при выборе количества уровней сигнального созвездия квадратурной модуляции учитывать взаимное влияние соседних каналов, создающее в используемом канале энергетические и фазовые шумы. Как следствие, определяя влияние друг на друга сигналов различных информационных систем, можно вырабатывать комплекс мероприятий по повышению надежности и помехоустойчивости телекоммуникационных систем.

Ключевые слова: сигналы, метрическое пространство, телекоммуникационная система, сигнальное созвездие

METHOD FOR SELECTING THE NUMBER OF LEVELS OF A SIGNAL CONSTELLATION OF SIGNALS WITH QUADRATIVE AMPLITUDE MODULATION TAKING INTO ACCOUNT THE MUTUAL INFLUENCE OF SIGNALS IN A MULTI-CHANNEL TELECOMMUNICATION SYSTEM

Vlasov S.V.

Penza State Technological University, Penza, e-mail: vlasov_s.v@mail.ru

To conduct a study of the mutual influence of signals (VVS) in telecommunication systems using multidimensional metric spaces, it is proposed to use both analytical and graphical models. In this article, a method has been developed for choosing the number of levels of a quadrature multilevel multiphase modulation (QAM) signal constellation, taking into account the mutual influence of signals on adjacent channels of a multichannel TCS using multidimensional metric spaces. Mutual influence is estimated as the distance between the spheres in a multi-parameter space. The result of the study is that when using the obtained models, an effect is detected when various transmitters operating at different frequencies using QAM are able to interfere (influence each other) at a certain value of the amplitudes and initial phases of both carrier frequencies and frequency values, and amplitudes of modulating oscillations. The scientific conclusion of the conducted research is: the transmission of information at different carrier frequencies does not exclude the mutual negative influence of signals on each other when using QAM signals, and when choosing the number of levels of the QAM signal constellation, it is necessary to take into account the mutual influence of adjacent channels, which creates energy and phase noise in the channel used. As a result, by determining the influence of the signals of various information systems on each other, it is possible to develop a set of measures to improve the reliability and noise immunity of telecommunication systems.

Keywords: signals, metric space, telecommunication system, signal constellation

Для современных исследований процессов, происходящих в различных сферах человеческой жизнедеятельности, использование чисто эмпирического подхода выбора оптимального количества уровней сигнального созвездия квадратурной амплитудной модуляции (КАМ) неприемлемо в связи со сложностью и параметрической емкостью используемых телекоммуникационных систем.

Современные технологии контроля и диагностирования каналов передачи данных телекоммуникационных систем предусматривают использовать комплексные показатели качества при многопараметрическом анализе, однако они не обладают наглядностью, требуют значительных затрат в практической реализации. Результат корреляционного анализа и анализа коэффициентов взаимного различия сигналов не пред-

усматривает оценку их взаимного влияния в среде распространения [1]. Необходимо сформировать научный базис выбора оптимального количества уровней сигнального созвездия КАМ не только с учетом дестабилизирующих факторов работы телекоммуникационных систем (мультипликативных и аддитивных помех, дефектов передающих систем), но и взаимного влияния КАМ сигналов друг на друга.

Современные методы оценки помеховой обстановки в линиях связи каналов передачи данных позволяют измерять уровень помех псофометрами, но приборы и существующие методы не позволяют определять значения помех, обусловленных взаимным влиянием сигналов.

Цель исследования – разработать метод оперативного выбора сигнального созвездия сигналов с КАМ для обеспечения заданной скорости передачи цифровой информации с учетом взаимного влияния сигналов в телекоммуникационных системах.

Материалы и методы исследования

Сигнальное созвездие – это квадратная матрица, в которой уровни амплитуды I и Q компонент КАМ сигнала отображены в виде значащих точек в квадратной системе координат I x Q [2].

Целочисленное значение каждой полученной точки определяется ячейкой матрицы, в которую она попадает. Ошибка определяется как выпадение измеренной точки из ячейки.

16-QAM (КАМ) диаграмма – это 4x4 матрица, в которой каждая из 16 ячеек представляет одну из 16 возможных бинарных комбинаций (рис. 1). Вертикальное и горизонтальное положение каждой точки соответствует I и Q уровням амплитуды сигнала переданного в течение одного цикла (рис. 2).

На рис. 3 представлена осциллограмма сигнала с КАМ модуляцией, где прослеживается изменение амплитуды и фазы сигнала

ла в зависимости от кодовой комбинации, поступающей на вход модулятора.

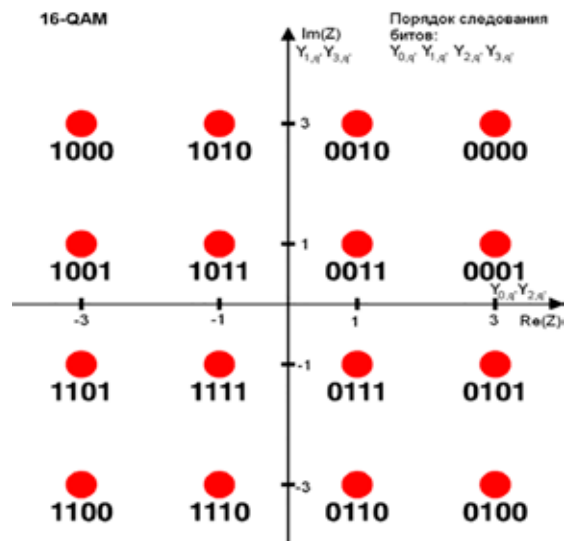


Рис. 1. Сигнальное созвездие 16-QAM (КАМ)

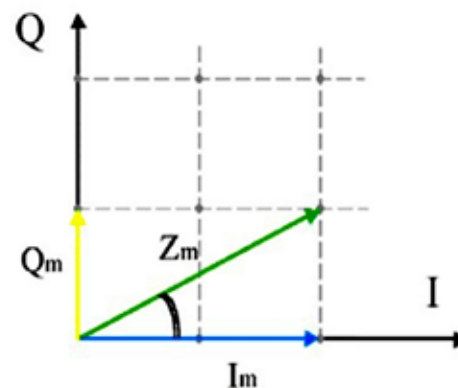


Рис. 2. Сигнальное созвездие 16-QAM (КАМ)

Для оптимизации сигнального созвездия необходимо провести анализ причин искажения сигнального созвездия КАМ.

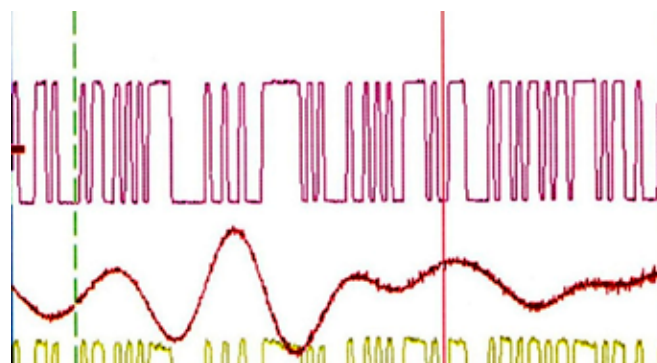


Рис. 3. Осциллограмма сигнала с КАМ модуляцией

Внешний вид значащих точек в ячейках сигнального созвездия может дать ключевую информацию о том, что происходит при передаче сигнала.

Значительный уровень внешних шумов, обусловленных в том числе взаимным влиянием сигналов соседних каналов (отношение энергии сигнала к спектральной плотности шума) – может привести к полной потере информации, так как расплывчатый образ точки занимает практически все пространство ячейки (рис. 4).

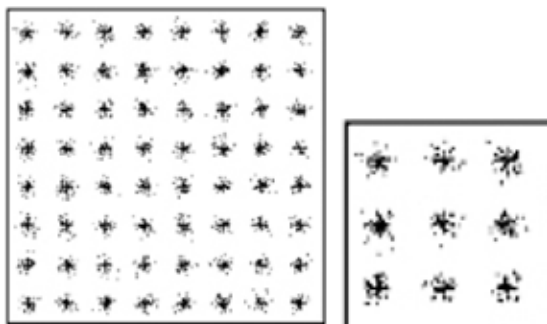


Рис. 4. Внешний вид сигнального созвездия с большим соотношением сигнал/шум в канале

Для проведения исследования взаимного влияния сигналов (ВВС) в телекоммуникационных системах (ТКС), предлагается использовать как геометрические объемные фигуры, так и их геометрическое взаимодействие в многомерных метрических пространствах [3]. Близость сфер, размеры и положение которых в виртуальном пространстве определяется параметрами модулированных сигналов [3–5], можно использовать как комплексный показатель оценки взаимного влияния сигналов друг на друга.

Результаты исследования и их обсуждение

Особенностью сигналов с КАМ является количество уровней сигнального созвездия, которое используется при передаче информации [2, 6]. Безусловно, чем больше уровней сигнального созвездия с КАМ, тем больше скорость передачи информации, в соответствии с теоремой Найквиста [7].

$$V_{max} = 2H \log_2 2M, \quad (1)$$

где V_{max} – максимальная скорость передачи; H – ширина полосы пропускания канала, выраженная в Гц; M – количество уровней сигнала, которые используются при передаче.

Однако при этом наблюдается снижение достоверности демодулированного цифрового сигнала, обусловленной как шумами в линии связи канала передачи данных, так

и взаимным влиянием сигналов как в многоканальной ТКС с частотным разделением каналов (OFDM) и влиянием энергий сигналов, расположенных на соседних частотах, но излучаемых другими станциями. Взаимное влияние сигналов наблюдается как в проводных, так и в радиоканалах. ВВС в проводных каналах обусловлено влиянием электрических и магнитных полей в проводах магистральных кабелей [8], а в радиоканалах взаимным влиянием электромагнитных волн. Неидеальность реализации модулятора, радиотракта, нелинейность амплитудно-частотных характеристик входных селективных фильтров является причиной возникновения помехи по соседнему каналу [9]. В OFDM-системе связи передача данных осуществляется блоками из N отсчетов, образующих один OFDM-символ. Набор из N комплексных отсчетов $X = [X_0, \dots, X_{N-1}]$ используется для параллельной модуляции N ортогональных поднесущих в частотной области. Вектор X включает в себя поднесущие, используемые для передачи данных, пилотные поднесущие, позволяющие осуществлять подстройку частот и фаз генераторов несущей и тактирующей частот во время приема пакета данных, а также нулевые поднесущие, где сигнал не передается и которые используются для формирования требуемого спектра сигнала, путем организации защитных частотных интервалов на краях используемой частотной полосы [10]. Для передачи данных на соответствующих поднесущих, как правило, используется квадратурная амплитудная модуляция (КАМ).

Изменение фаз модулированных сигналов с КАМ влечет возникновение фазовых шумов по соседнему каналу, что в совокупности с энергетическим влиянием приведет к искажению и потере полезной информации при демодуляции.

Несмотря на зависимость скорости передачи цифровой информации от количества уровней сигнала с КАМ, их количество ограничивается теоремой Шеннона, в соответствии с которой максимальная скорость передачи данных по каналу с шумом равна

$$V_{max} = H \log_2 \left(1 + \frac{S}{N} \right), \quad (2)$$

где V_{max} – максимальная скорость передачи; H – ширина полосы пропускания канала, выраженная в Гц; S/N – соотношение энергии сигнала к спектральной плотности шума в канале.

Следует отметить зависимость влияния размерности сигнального созвездия и погрешности амплитуд на чувствительность

системы связи к самим погрешностям амплитуд. То есть чем обширнее само сигнальное созвездие, тем ближе располагаются амплитуды аналогового сигнала, которые определяют различный цифровой код, следствием этого является повышенная чувствительность системы связи к погрешностям и снижение достоверности цифрового сигнала на выходе компьютера.

Если приравнять выражения (1) и (2), путем несложных математических преобразований получается выражение

$$M = \sqrt{1 + \frac{S}{N}}. \quad (3)$$

Из выражения следует, что количество уровней многоуровневой многофазовой квадратурной модуляции зависит от отношения энергии сигнала к спектральной плотности шума.

В современных ТКС спектральная плотность шума определяется эмпирически, путем измерения значения соответствующими контрольно-измерительными приборами (псифометрами). Комплексная оценка качества передаваемых сигналов не позволяет учитывать взаимное влияние сигналов по соседним каналам [1]. Перед передачей сигнала можно измерить значения искусственных и естественных шумов на данной несущей частоте. Но данные методы не позволяют определить значения энергетических и фазовых шумов по соседним

каналам, пока не осуществлена попытка передачи сигнала с вариантом сигнального созвездия. Для обеспечения требуемой достоверности передачи цифрового информационного сигнала осуществляется передача сигнала в линию связи канала передачи данных, а затем, в зависимости от помеховой обстановки и информационной загруженности соседних каналов, необходимо увеличивать или уменьшать количество уровней сигнального созвездия. Это влечет за собой значительные временные и аппаратные затраты.

Предлагается для выбора сигнального созвездия предварительно использовать коэффициент взаимного влияния (КВВ) в многомерном метрическом пространстве.

В общем случае мощность помехи Рп (АСП), проникающей в соседний канал, определяется выражением [8]:

$$ACI = A(W) = \frac{\int_{-\infty}^{\infty} G(f) |H(f - \omega)|^2 df}{\int_{-\infty}^{\infty} G(f) |H(f)|^2 df}, \quad (4)$$

где ω – нормированный частотный разнос каналов, $G(f)$ – спектральная плотность мощности модулированного или линейно усиленного сигнала; $H(f)$ – частотная характеристика селективного входного фильтра приемника.

Тогда спектральная плотность шума выражения (3) будет иметь вид

$$N = ASI + J = k(A_{m1}, A_{m2}, \omega_1, \omega_2, \varphi_1, \varphi_2) G(f), \quad (5)$$

где

$$k(A_{m1}, A_{m2}, \omega_1, \omega_2, \varphi_1, \varphi_2) = \frac{1}{d(A_{m1}, A_{m2}, \omega_1, \omega_2, \varphi_1, \varphi_2)} - \text{коэффициент взаимного влияния сигналов в многоканальной телекоммуникационной системе};$$

$$J_i = \frac{1}{N} \sum_{n=0}^{N-1} \exp(j\varphi_n) \exp \frac{-j2\pi ni}{N} - \text{фазовый шум по соседнему каналу};$$

$d(A_{m1}, A_{m2}, \omega_1, \omega_2, \varphi_1, \varphi_2)$ – расстояние между сферами в виртуальном многомерном метрическом пространстве.

Из выражения (5) следует, что, зная значения амплитуд A_{m1}, A_{m2} , несущих частот ω_1, ω_2 , мгновенных фаз сегментов сигналов φ_1, φ_2 сигнальных созвездий КАМ соседних каналов, можно предварительно произвести расчет расстояния между точками сфер в виртуальном многомерном метрическом пространстве для оценки взаимного энергетического и фазового влияния соседних каналов для выбора

соответствующего количества уровней сигнального созвездия для обеспечения требуемой скорости передачи дискретных сообщений в соответствии с выражением (3).

На рис. 5 представлен вариант значения расстояния между точками сфер, определяемыми параметрами сигналов, представленных в таблице в частотно-фазовом метрическом пространстве.

Параметры исследуемых сигналов

Номер сигнала	Частота (кГц)	Фаза	Амплитуда (В)
1 сигнал	62,5	$\pi/8$	0,2
2 сигнал	166	$\pi/3$	0,4

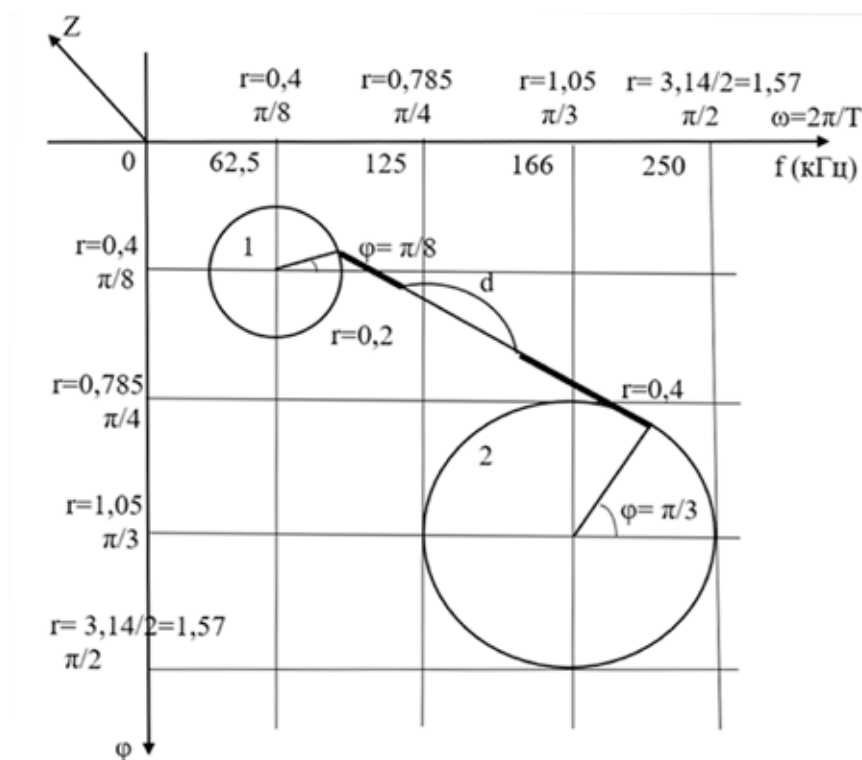


Рис. 5. Вариант значения расстояния между точками сфер в частотно-фазовом пространстве

Заключение

Таким образом, использование коэффициента взаимного влияния, полученного в результате расчета расстояния между точками сфер в многомерном метрическом пространстве, позволит значительно сократить временные и аппаратные затраты при выборе сигнального созвездия, должного обеспечить заданную скорость передачи информации сигналов с КАМ в телекоммуникационной системе с частотным разделением каналов.

Список литературы

1. Будко П.А., Федоренко В.В. Управление в сетях связи. Математические модели и методы оптимизации: монография. М.: Издательство физико-математической литературы, 2007. 228 с.
2. Муравьев В.В., Корневский С.А., Печень Т.М. Полосовая модуляция в системах телекоммуникаций: учеб.-метод. пособие. Минск: БГУИР, 2019. 79 с.
3. Власов С.В., Власов В.И. Использование геометрических сфер для контроля качества информационных си-

стем // Universum: Технические науки. № 3 (36). М.: Изд-во «МЦНО», 2017. С. 9–12.

4. Власов С.В., Власов В.И. Модель контроля безопасности информационных систем // Современные наукоемкие технологии. 2020. № 6. С. 31–37.

5. Власов С.В., Власов В.И. Использование тел вращения для контроля качества информационных сигналов каналов передачи данных // Colloquium journal, Physics and Mathematics. 2020. № 6 (58). С. 12–14.

6. Пушкарев В.П. Аналоговые и цифровые радиоприемные устройства: учебное пособие. Томск: РТФ, ТУСУР, 2018. 237 с.

7. Коберниченко В.Г. Основы цифровой обработки сигналов: учебное пособие. Министерство науки и высшего образования Российской Федерации. Екатеринбург: Издательство Уральского университета, 2018. 150 с.

8. Семенов А.Б., Стрижаков С.К., Сунчелей И.Р. Структурированные кабельные системы. 5-е изд. М.: Компания АйТи; ДМК Пресс, 2019. 640 с.

9. Феер К. Беспроводная цифровая связь. Методы модуляции и расширения спектра / Пер. с англ.; Под ред. В.И. Журавлева. М.: Радио и связь, 2000. 520 с.

10. Мальцев А.А., Масленников Р.О., Хоряев А.В. Влияние фазового шума на OFDM-системы передачи данных // Известия вузов. Радиофизика. 2010. Т. LIII. № 8. С. 528–542.

УДК 004.428.4

ИСПОЛНЕНИЕ И АВТОРИЗАЦИЯ EMV-ТРАНЗАКЦИЙ

Ермаков К.С., Сидякин И.М.

*ФГБОУ ВО «Московский государственный технический университет имени Н.Э. Баумана
(национальный исследовательский университет)», Москва,
e-mail: cosmozz@yandex.ru, ivan.sidyakin@bmstu.ru*

В статье рассматривается практическая реализация стандарта EMV, разработанного для микропроцессорных карт, которые являются важной частью современной инфраструктуры электронных платежей. Приводятся основные положения стандарта EMV, определяющие информационное взаимодействие между банковской картой и платежным терминалом, необходимое для исполнения финансовых операций. Представлена общая архитектура системы электронных платежей. В статье подробно описан цикл EMV-транзакции, исполняемый при совершении платежной операции. Рассмотрены отдельные этапы обработки транзакции, включая выбор приложения, верификацию владельца карты, проверку подлинности данных карты, онлайн- и офлайн-авторизацию. Приведены некоторые особенности обработки операций приложениями разных платежных систем, описан процесс аутентификации EMV-карт, приведено сравнение способов авторизации транзакций оплаты картой с магнитной полосой и смарт-чипом. Рассмотрены элементы интерфейса EMV-приложения, установленного на карте. Приведены форматы команд и ответов, пересылаемых между картой и терминалом в процессе исполнения транзакций. Практическим результатом работы является разработанное и протестированное на соответствие требованиям стандарта EMV-приложение для платежного терминала. Проведенное функциональное тестирование включает более 1,5 тыс. тестов, собранных в плане тестирования контактного интерфейса EMV.

Ключевые слова: EMV, смарт-карты, платежные карты, платежное приложение, транзакция

EMV-TRANSACTION PROCESSING AND AUTHORIZATION

Ermakov K.S., Sidyakin I.M.

*Bauman Moscow State Technical University (National Research University), Москва,
e-mail: cosmozz@yandex.ru, ivan.sidyakin@bmstu.ru*

The article discusses the practical implementation of the EMV-standard developed for microprocessor cards, which are an important part of the modern electronic payment infrastructure. Basic statements of the EMV-standard are given, which determine the information interaction between a bank card and a Point-of-Sale terminal required for the execution of financial transactions. The general architecture of the electronic payment system is presented. The article describes in detail the EMV-transaction process. The individual steps of transaction processing are considered, including application selection, cardholder verification, card data authentication, online and offline authorization. Some proprietary features of the transaction processing introduced by different payment systems are discussed. The process of authentication of EMV-cards is described and comparison of methods for authorizing payment transactions with a magnetic stripe data and a smart chip data is given as well. The interface of the smart-card EMV-application are depicted. The data formats of application data unit and corresponding response used for data exchange between the card and the terminal applications are given. The practical result of the work is the software application designed for the payment terminal developed and tested for compliance with the requirements of the EMV-standard. The functional testing carried out includes more than one and a half thousand tests declared in the EMV-contact interface test plan.

Keywords: EMV, smart-card, payment cards, payment applications, transaction

За последнее десятилетие платежные карты прошли два важных этапа развития. Контактные карты, разработанные по стандартам EMV, сменили карты с магнитной полосой. В настоящее время идет процесс замены контактных карт на бесконтактные средства оплаты, к которым кроме карт относятся мобильные телефоны и другие электронные устройства. По состоянию на начало 2022 г. большинство дебетовых и кредитных карт используют стандарт EMV [1], но при этом имеют магнитную полосу для совместимости со старым терминальным оборудованием.

Требования безопасности являются ключевым фактором смены технологий. Контактная карта EMV имеет существенно бо-

лее высокую степень защищенности, чем карта с магнитной полосой. Стандарты EMV включают надежные способы аутентификации данных карты и защиты от копирования.

Стандарт EMV – это совместная разработка международных платежных систем VISA Inc и MasterCard Worldwide. Контактные карты EMV разных платежных систем в целом соответствуют этому стандарту, но могут иметь некоторые особенности реализации, не входящие в противоречие с требованиями спецификации.

Целью исследования является разработка и подготовка к сертификационному тестированию программного модуля ядра EMV Level 2, входящего в состав программного обеспечения платежного терминала.

Материалы и методы исследования

Для достижения поставленной цели использовались различные методы исследования, включая теоретический анализ спецификаций EMV и сравнительный анализ реализаций EMV-приложений различных платежных систем. Материалом исследования является опубликованная в открытом доступе документация EMVCO [1], а также несколько стандартов, применяемых в области электронных платежей.

Авторизация EMV-транзакции

Авторизация транзакции по карте – это процесс проверки данных карты и принятия решения об отказе или одобрении транзакции, которое принимается совместно несколькими участниками системы электронных платежей. В процессе исполнения банковской транзакции задействованы:

- платежная система;
- банк-эмитент – банк, который выпустил платежную карту и обслуживает связанный с ней счет;
- банк-эквайер – банк, который управляет платежным терминалом торговой точки или транспортным терминалом;
- платежный терминал – устройство, которое обеспечивает работу с платежной картой.

Архитектура системы электронных платежей показана на рис. 1.



Рис. 1. Архитектура системы электронных платежей

Микропроцессорная карта стандарта EMV – это программно-аппаратный комплекс, в состав которого входят:

- процессор;
- оперативная память;
- подсистема хранения данных;
- операционная система.

В смарт-карте должно быть установлено платежное EMV-приложение. Требования к этому приложению указаны в серии стандартов EMV Level 2. Общие требования к интерфейсу между терминалом и EMV-приложением карты приводятся в [2]. Вопросы безопасности, включая проверку подлинности данных карты, рассмотрены в [3]. Требования к алгоритму исполнения транзакции указаны в [4]. Требования к интерфейсам с держателем карты, продавцом и эквайером перечислены в [5].

Интерфейс EMV-приложения

Интерфейс приложения включает набор APDU (Application Protocol Data Unit) команд, используемых для проведения транзакций и управления EMV-приложением.

Последовательность команд и APDU определяет протокол обмена между картой и устройствами чтения карт в составе Point-of-Sale (POS) терминалов, банкоматов и пр. Устройство передает карте команду APDU-C, а карта возвращает ответ APDU-R. Инициатором обмена всегда является устройство чтения. Структура команды APDU-C [6] приведена в табл. 1.

Таблица 1

Формат APDU-C

Элемент	Размер (байт)	Описание
Заголовок		
CLA	1	Класс команды
INS	1	Код инструкции
P1	1	Параметр № 1
P2	1	Параметр № 2
Тело команды		
Lc	1	Длина отправляемых данных
Data	Lc	Данные
Le	1	Длина ожидаемого ответа

Ответ APDU-R состоит из заголовка, тела сообщения и трейлера с байтами статуса Status Word 1 (SW1) и Status Word 2 (SW2). Байты статуса составляют код, по которому определяется результат обмена. Структура команды APDU-R [6] приведена в табл. 2.

Байт SW1 содержит основную информацию, а SW2 дополнительную. В табл. 3 приведены примеры кодов статуса.

Для исполнения операции «оплата» по карте платежный терминал запускает цикл EMV-транзакции, который включает обмен APDU командами и ответами на них в заданной последовательности, а также

проверку подлинности данных карты и некоторые другие действия, описание которых приводится ниже.

Таблица 2

Формат APDU-R

Элемент	Размер (байт)	Описание
Заголовок		
Data	Le	Данные
Трейлер		
SW1	1	Слово статуса 1
SW2	1	Слово статуса 2

Таблица 3

Примеры комбинаций SW1SW2

SW1SW2	Значение
9000	Ок
66XX	Команда не была исполнена по причинам безопасности
6700	Неправильная длина команды
6E00	Неизвестный CLA

EMV-транзакция выполняется в несколько этапов. Транзакция начинается с того, что карта возвращает двоичный блок данных Answer to reset, ATR, подтверждающий готовность к дальнейшему обмену данными.

Затем производится выбор платежного приложения. Карта и терминал совместно выбирают приложение, которое будет использовано для проведения операции. В настройках терминала указывается список поддерживаемых EMV-приложений. Кроме приложений EMV, включая приложения НСПК МИР, Visa International VSDC, MasterCard International M/Chip, на карте могут присутствовать программы лояльности, бонусные или идентификационные приложения, которые используют собственные, отличные от EMV, стандарты. В специально выделенной области памяти карты размещен список всех установленных в карте приложений.

Существует два метода построения списка приложений: Payment System Environment, PSE и Payment System Application, PSA [2]. Большинство современных смарт-карт поддерживают оба этих метода.

Метод PSE основан на данных каталога приложений, который считывается с карты первой командой APDU-C SELECT. Заданное в спецификации имя каталога передается в параметрах этой команды. В случае, если каталог PSE присутствует, карта возвращает статус SW1/SW2 = 9000 (OK)

и параметры, необходимые для выбора приложения. Терминал выбирает одно из приложений каталога и запускает его с помощью команды SELECT.

Метод PSA для поиска приложения использует идентификатор Application Identifier, AID. Идентификатор AID состоит из двух частей:

– Registered Application Provider Identifier, RID – зарегистрированный идентификатор провайдера приложения. RID – это уникальный идентификатор платежной системы или другого эмитента, выпустившего данный карточный продукт.

– Proprietary Application Identifier Extension, PIX – частный идентификатор приложения или бизнес-код карточного продукта.

Данные карты хранятся в формате DER [7]. Номера тегов всех параметров карты указаны в спецификации EMV Level 2.

Идентификатор AID в карте хранится в теге 4F или 84. В конфигурации терминала имеется список идентификаторов AID всех поддерживаемых терминалом приложений. Терминал последовательно отправляет карте команды SELECT с разными идентификаторами AID из своего списка. Если приложение с указанным в команде AID имеется в карте, карта возвращает код статуса SW12 = 9000 и переходит к исполнению приложения. Иначе карта возвращает ошибку «файл не найден» в поле статуса SW12 = 6A82 и терминал переходит к следующему идентификатору. Если в конфигурации устройства отсутствует AID приложения карты, транзакция завершается с ошибкой.

Команда выбора приложения SELECT запускает приложение карты и возвращает информацию об этом приложении. Различные платежные системы используют для передачи информации о приложении как общие теги, так и собственные частные теги и форматы. Например, Mastercard указывает региональную принадлежность карты в теге Third Party Data, а Visa – в теге Application Program. Эти теги несут сходную информацию, но имеют разный формат.

Карты многих платежных систем в ответе на команду SELECT возвращают тег Processing Options Data Object List, PDOL, содержащий список параметров, которые терминал должен передать карте для продолжения обработки транзакции. В списке указывается номер тега каждого параметра и ожидаемая длина данных. Например:

– PDOL Amex 9F3501 запрашивает тип терминала.

– PDOL Visa 9F66049F02069F37045F2A-029F1A02 запрашивает сумму операции, валюту, код страны терминала и некоторые другие элементы.

– PDOL.MIP9F7A015F2A029F02069F35019F40059F1A029F3303, запрашивает подробную информацию о терминале, его возможностях и местонахождении.

Терминал в параметрах команды Get Processing Options, GPO, должен передать всю запрошенную картой информацию, иначе транзакция будет прервана.

На следующем шаге цикла транзакции терминал выполняет проверку держателя карты. Если сумма транзакции не превышает заданный в настройках терминала предел, то разрешается провести операцию без проверки. Если сумма превышает предел, то также, в зависимости от настроек карты, выбирается один из методов проверки держателя карты. Для пластиковых карт обычно проверка держателя карты производится с помощью онлайн- или офлайн-проверки ПИН-кода.

Эмитент карты может делегировать проверку ПИН-кода платежному приложению карты. Значение ПИН-кода, введенное пользователем на терминальном устройстве, сравнивается с образцом, безопасно хранящимся в карте. Этот метод называется офлайн-проверкой ПИН-кода.

Для онлайн-проверки ПИН-код шифруется в терминале и передается на проверку к эмитенту.

После окончания проверки держателя карты выполняется проверка ограничений, накладываемых на использование карты, например:

- проверка срока действия карты,
- проверка лимитов по списаниям,

– проверка совместимости приложения и терминала,

– проверка кода банка-эмитента.

По результатам этих проверок принимается одно из трех возможных решений о завершении операции:

- отклонение операции в режиме офлайн;
- отправка запроса на онлайн-авторизацию;
- одобрение операции в режиме офлайн.

Режим офлайн означает принятие решения терминалом без участия банка. Режим онлайн означает принятие решения терминалом совместно с банком.

Аутентификация карты – это обязательный шаг цикла транзакции. Аутентификация подтверждает достоверность данных, считанных с карты, а также то, что банк – эмитент карты авторизован платежной системой, выпустившей приложение карты.

Для аутентификации карты с магнитной полосой используются статические данные. Данные магнитной полосы каждый раз при авторизации карты передаются в банк-эмитент. Платежный терминал не имеет возможности проверить достоверность данных магнитной полосы самостоятельно, без участия банка. Банк-эмитент не может отличить карту от ее полной копии, поэтому имеется высокая вероятность проведения мошеннических операций копиями таких карт. Схема авторизации транзакции картой с магнитной полосой показана на рис. 2.



Рис. 2. Схема авторизации транзакции картой с магнитной полосой

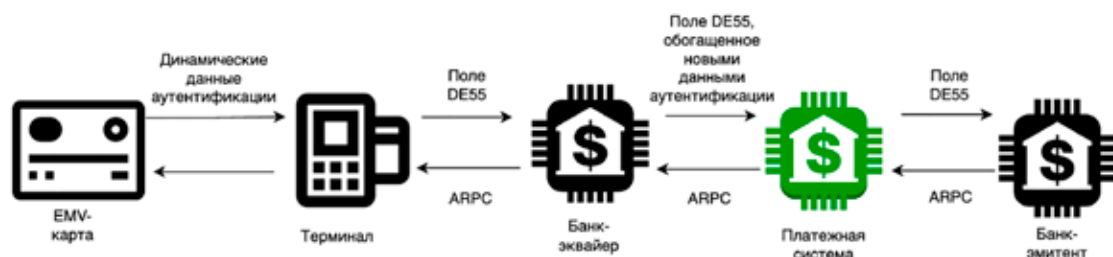


Рис. 3. Схема авторизации EMV-транзакции

EMV-карты для повышения уровня безопасности используют динамически вычисляемую цифровую подпись статических данных карты и динамических данных транзакции, которые передаются банку-эмитенту и подтверждают достоверность данных. Значение цифровой подписи уникально для каждой EMV-транзакции, поэтому транзакцию нельзя повторить, отправив банку копию перехваченной транзакции. Карта хранит в зашифрованном виде криптографический ключ, использующийся для вычисления цифровой подписи. Этот ключ недоступен для чтения. Изготовление копии карты без информации о ключе невозможно.

На рис. 3 показана процедура авторизации EMV-транзакции.

Транзакция начинается, когда контактная карта вставляется в слот устройства чтения или бесконтактная карта обнаруживается в поле терминала. Терминал передает карте данные транзакции, включая сумму, код валюты, код страны и прочее. Затем карта и терминал производят проверку рисков транзакции. Если проверка завершается успешно, карта возвращает терминалу данные транзакции и их цифровую подпись, которая называется криптограмма, а терминал, используя полученные от карты данные, создает и отправляет банку-эквайеру сообщение авторизации. Банк-эквайер передает сообщение банку-эмитенту.

Банк-эмитент проверяет подлинность криптограммы, которая вычисляется по динамическим данным текущей транзакции. Одновременно эта проверка позволяет удостовериться в подлинности карты.

Приведенное выше упрощенное описание цикла EMV-транзакции раскрывает главный аспект, обеспечивающий безопасность финансовых операций по карте – использование для аутентификации карты динамических данных.

EMV-карты добавляют важный функционал, который может использовать банк-эмитент, включая:

- проверку достоверности данных карты;

- идентификацию каждой транзакции;
- взаимную аутентификацию банка и карты. Банк возвращает собственную криптограмму терминалу в ответе на запрос авторизации;

- изменение данных карты после авторизации. Банк, например, может заблокировать карту или изменить предельные суммы операций.

Заключение

В статье приводится описание основных функций приложения, разработанного авторами в соответствии с требованиями стандарта EMV и предназначенного для использования в составе программного обеспечения платежного терминала. Рассмотрены основные этапы цикла исполнения платежной транзакции. Приведена схема авторизации платежной транзакции, а также интерфейс с EMV-приложением, установленным на смарт-карте. Выполнено функциональное тестирование приложения в соответствии с планом тестирования EMV [8], который содержит более 1,5 тыс. тестов.

Список литературы

1. EMVCO. [Электронный ресурс]. URL: <https://www.emvco.com/> (дата обращения: 24.10. 2022).
2. EMVCO. EMV Integrated Circuit Card Specifications for Payment Systems Book 1 Application Independent ICC to Terminal Interface Requirements. Version 4.4. 2022. 81 с.
3. EMVCO. EMV Integrated Circuit Card Specifications for Payment Systems Book 2. Security and Key Management. Version 4.4. 2022. 192 с.
4. EMVCO. EMV Integrated Circuit Card Specifications for Payment Systems. Book 3. Version 4.4. 2022. 230 с.
5. EMVCO. EMV Integrated Circuit Card Specifications for Payment Systems Book 4 Cardholder, Attendant, and Acquirer Interface Requirements. Version 4.4. 2022. 133 с.
6. ISO/IEC 7816-4:2020. Identification cards — Integrated circuit cards, 2020. 176 с.
7. ISO X.690 Series X: Data networks and open system communications OSI networking and system aspects – abstract syntax notation one (ASN.1). Information technology – ASN.1 encoding rules: specification of basic encoding rules (BER), canonical encoding rules (CER) and distinguished encoding rules (DER). 2003. 39 с.
8. EMVCO. Terminal Type Approval Level 2 Test Cases. Version 43k. 2021. 1808 с.

УДК 004.056.5

ГОСУДАРСТВЕННЫЕ ИНФОРМАЦИОННЫЕ СИСТЕМЫ. ВОПРОСЫ БЕЗОПАСНОСТИ

Кечкина Н.И., Наумова Е.Г.*Дзержинский политехнический институт (филиал)**ФГБОУ ВО «Нижегородский государственный технический университет им. Р.Е. Алексеева»,**Дзержинск, e-mail: nataliyabaranova@yandex.ru, soboll007@yandex.ru*

В статье рассматриваются вопросы безопасности информации в информационных системах органов государственного и муниципального управления. Для достижения поставленной цели решен ряд задач. Для начала рассмотрено понятие информационных систем и представлена их классификация. Особое внимание уделяется группе государственных информационных систем. Поскольку данные системы предполагают обращение большого объема разнородной информации, то вопросы информационной безопасности в данной области являются наиболее актуальными и требуют неотложного решения. Основная цель работы – рассмотреть вопросы безопасности государственных информационных систем, обозначить возможные способы защиты информации. Для достижения поставленной цели проведен анализ информации, циркулирующей в государственных информационных системах. Предлагается классифицировать информацию на общедоступную и информацию с ограниченным доступом. Направления и класс защиты выбирают исходя из уровня значимости информации, а также размера информационной системы. Представлена классификация возможных угроз. Основной принята классификация по аспектам, на которые угрозы направлены: угрозы доступности, угрозы целостности и угрозы конфиденциальности. Для каждого типа приведена характеристика угроз и предложены мероприятия и средства по обеспечению информационной безопасности. Отмечается, что для обеспечения требуемого уровня информационной безопасности стоит рассматривать совокупность организационно-правовых и программно-технических средств.

Ключевые слова: государственные информационные системы, класс защиты, уровень значимости информации, масштабы информационной системы, угрозы безопасности информации

STATE INFORMATION SYSTEMS. SECURITY QUESTIONS

Kechkina N.I., Naumova E.G.*Dzerzhinsk Polytechnic Institute, R.E. Alekseev Nizhny Novgorod State Technical University,**Dzerzhinsk, e-mail: nataliyabaranova@yandex.ru, soboll007@yandex.ru*

The article deals with the issues of information security in information systems of state and municipal authorities. To achieve this goal, a number of tasks have been solved. To begin with, the concept of information systems is considered and their classification is presented. Particular attention is paid to the group of state information systems. Since these systems involve the circulation of a large amount of heterogeneous information, the issues of information security in this area are the most urgent and require an urgent solution. The main purpose of the work is to consider the security issues of state information systems, to outline possible ways to protect information. To achieve this goal, the analysis of information circulating in state information systems was carried out. It is proposed to classify information into public and restricted information. The level of significance of information and the scale of the information system make it possible to determine the class of protection, and therefore to choose the direction of protection. The classification of possible threats is presented. The main classification is adopted according to the aspects to which the threats are directed: threats to availability, threats to integrity and threats to confidentiality. For each type, the characteristics of threats are given and measures and means for ensuring information security are proposed. It is noted that to ensure the required level of information security, it is worth considering a set of organizational and legal and software and hardware tools.

Keywords: state information systems, protection class, information significance level, information system scale, information security threats

В соответствии с ГОСТ Р 52653-2006 информационные технологии (ИТ) представляют собой совокупность методов и способов сбора, передачи, накопления, обработки, хранения, представления и использования информации.

В настоящее время обязательным элементом компьютерной ИТ является наличие вычислительной техники, а также телекоммуникационных средств, интуитивно-понятного пользовательского интерфейса.

Структура компьютерной ИТ включает:

- технические средства, в том числе средства вычислительной, коммуникационной и организационной техники;

- программные средства: системное (общее) и прикладное программное обеспечение;

- инструкции, нормативные документы, методические материалы по организации работы, образующие организационно-методическое обеспечение [1].

Информационная система (ИС) является средой реализации информационной технологии. ИС организации предполагает наличие и взаимосвязь средств накопления, обработки, передачи, хранения и контроля информации, формирующие информационную инфраструктуру, а также персонала, выполняющего эти действия с информацией.

Согласно определению, приведенному в ГОСТ Р 52653-2006, информационная система – это совокупность содержащейся в базах данных информации, а также обеспечивающих ее обработку информационных технологий и технических средств.

Среди многочисленных целей функционирования ИС организации основной является производство информации, необходимой для функционирования организации. Для поддержания возможности управления ИС организацией необходимой является реализация информационной и технической сред. ИТ позволяют осуществить в ИС все необходимые процессы преобразования информации.

Цель исследования – выявить основные угрозы информационной безопасности и на основе их анализа определить средства, необходимые для формирования системы информационной безопасности государственных информационных систем (ГИС).

Материалы и методы исследования

С использованием методов анализа и синтеза в исследовании проводится выявление основных направлений и методов обеспечения информационной безопасности. На основе современных публикаций, законов РФ в области информационных технологий и систем, методических документов, отражающих меры защиты информации в государственных информационных системах, рассмотрена классификация информационных систем, способы оценки класса защиты, определены направления защиты.

Результаты исследования и их обсуждение

Далее (рисунок) приводится классификация ИС, сформированная в соответствии с ч. 1 ст. 13 Федерального закона от 27 июля 2006 г. № 149-ФЗ [1].

Если коснуться более подробно государственных и муниципальных ИС, то первые создаются на основании федеральных зако-

нов, законов субъектов Российской Федерации, а также на основании правовых актов государственных органов [2].

Муниципальные ИС разрабатываются на основании решения органа местного самоуправления [2].

И государственные, и муниципальные ИС необходимы для организации обмена информацией между органами государственного и муниципального управления (ГМУ), между органами управления и юридическими и физическими лицами, а также в информировании общества [2].

При рассмотрении информационных систем с точки зрения безопасности выделяют субъекты (активные компоненты) и объекты (пассивные компоненты).

К субъектам государственных ИС относят госслужащих, обслуживающий ИС персонал, представителей юридических лиц и физических лиц, которые могут обращаться к ИС по различным вопросам.

Объектами ИС являются информация и информационные процессы. В информационных системах ГМУ циркулирует самая разная информация, в первом приближении она делится на общедоступную информацию и информацию ограниченного доступа.

К категории общедоступной информации относят как общеизвестные сведения, так и иные данные, доступ к которым не должен быть ограничен. Это позволяет использовать и размещать ее свободно любыми лицами, в том числе в сети Интернет в открытом доступе, если это не противоречит установленным федеральными законами ограничениям в отношении распространения подобной информации. Так, для государственных ИС согласно ст. 13 Федерального закона от 09.02.2009 № 8-ФЗ в сети Интернет может быть размещена общая информация о деятельности государственных органов и органов местного самоуправления [3]. В ст. 14 этого же закона дополнительно указан состав общедоступной информации, размещаемой в сети «Интернет», в том числе в форме открытых данных [3].



Классификация ИС

Разновидности ИС ГМУ РФ в зависимости от масштаба

Масштаб	Область функционирования	Сегменты
ИС федерального масштаба	В пределах федерального округа РФ [5]	Сегменты находятся в субъектах РФ, муниципальных образованиях и организациях [5]
ИС регионального масштаба	На территории отдельного субъекта РФ [5]	Сегменты находятся в одном или нескольких муниципальных образованиях и подведомственных и других организациях [5]
ИС объектового масштаба	На территории одного федерального органа государственной власти (ОГВ), ОГВ субъекта РФ, муниципального образования и организации [5]	Не имеют сегментов в территориальных органах, представительствах, филиалах, подведомственных и иных организациях [5]

Информацией ограниченного доступа являются государственная тайна и конфиденциальная информация. Последняя в свою очередь разбивается на следующие типы: коммерческая тайна, персональные данные, служебная информация и т.п.

Приведенная выше информация задействована в таких информационных процессах, характерных для сферы ГМУ, как ведение документооборота, сбор и хранение информации, поиск и обработка информации, анализ данных, планирование и принятие управленческих решений, информирование населения [4] и др.

Вышеизложенное определяет строгие требования к обеспечению должного уровня защиты подобного типа информации. Направление защиты определяется характером информации и формой ее представления. Устанавливаются четыре класса защищенности ИС: от самого высокого (первого) к самому низкому (четвертому). Класс защиты (K) определяется в зависимости от уровня значимости информации ($УЗ$), обрабатываемой в ИС, и от масштаба ИС ($M_{ИС}$) [3]:

$$K = [УЗ; M_{ИС}]. \quad (1)$$

Если со временем происходит изменение масштаба ИС или в ней появляется информация с другим уровнем значимости, то требуется пересмотреть класс защищенности ИС.

Для определения масштаба ИС необходимо рассмотреть ее назначение, из каких сегментов она состоит, а также их распределенность. Масштабы ИС ГМУ РФ представлены в таблице.

Основной целью защиты информации можно считать предотвращение возможного ущерба для владельца и/или пользователя информации в результате совершения угрозы. Размер ущерба будет зависеть от формы самой угрозы и от того, какая информация подвергается угрозам. Поэтому одним из вариантов определения уровня

значимости информации является оценка возможного ущерба для владельца и/или пользователя информации, возникшего в результате нарушения конфиденциальности, целостности и/или доступности информации [5]:

$$УЗ = [(конфиденциальность, степень ущерба) (целостность, степень ущерба) (доступность, степень ущерба)]. \quad (2)$$

Здесь (2) оценивается вероятность совершения той или иной угрозы, а также степень возможного ущерба, которая определяется владельцем/пользователем одним из доступных методов. Можно использовать следующую градацию ущерба в зависимости от последствий в социальной, политической, международной, экономической, финансовой или иных областях деятельности организации, а также в зависимости от степени влияния реализованной угрозы на работоспособность ИС:

- высокий ущерб, если возможны существенные негативные последствия в какой-либо области деятельности, а также если ИС и/или оператор ИС (владелец информации) не могут выполнять возложенные на них функции;

- средний ущерб, если возможны умеренные негативные последствия в какой-либо области деятельности и/или ИС и/или оператор ИС (владелец информации) не могут выполнять хотя бы одну из возложенных на них функций;

- низкий ущерб, если возможны незначительные негативные последствия в какой-либо области деятельности и/или ИС и/или оператор ИС (владелец информации) могут выполнять возложенные на них функции с недостаточной эффективностью, или выполнение функций возможно только с привлечением дополнительных сил и средств [5].

Окончательно уровень значимости информации устанавливается по наивысшим значениям степени возможного ущерба,

определенным для конфиденциальности, целостности, доступности информации. Кроме того, ИС чаще всего работают с информацией разного вида, следовательно, и оценка уровня значимости выполняется отдельно для каждого вида информации.

Оценка угроз безопасности информации (*УБИ*) осуществляется по ряду факторов: оценке возможностей (потенциала, оснащенности и мотивации) внешних и внутренних нарушителей, результату анализа возможных уязвимостей ИС, возможных способов реализации угроз безопасности информации и последствий от нарушения свойств безопасности информации (конфиденциальности, целостности, доступности).

УБИ = [(возможности нарушителя;
уязвимости информационной системы;
способ реализации угрозы;
последствия от реализации угрозы)] [5]. (3)

Существует большое количество классификаций угроз в зависимости от используемых критериев, например:

- по аспекту информационной безопасности (основной направленности угрозы) выделяют угрозы нарушения доступности, нарушения целостности, нарушения конфиденциальности;
- по целевым объектам ИС, на которые воздействуют, различают угрозы данным, информационным процессам, программному или аппаратному обеспечению, поддерживающей инфраструктуре, носителям информации, персоналу и др.;
- по типу и характеру различают случайные/преднамеренные действия природного/техногенного характера;
- по расположению источника угроз и ее потенциалу: внешние и внутренние угрозы с низким, средним или высоким потенциалом [5].

В соответствии с моделью (2) за основу выбрана первая классификация, согласно которой выделяют: угрозы доступности, угрозы целостности и угрозы конфиденциальности.

1. Под **доступностью** понимают возможность за приемлемое время получить требуемую информационную услугу. В большей степени это свойство информации касается общедоступной информации, которая циркулирует в ИС ГМУ и к которой могут обращаться представители юридических лиц и физические лица [2].

К угрозам доступности можно отнести случайные ошибки операторов, штатных пользователей, системных администраторов и других лиц, обслуживающих ИС. Последствиями от реализации таких

угроз могут быть внутренний отказ ИС, отказ поддерживающей инфраструктуры. Для снижения вероятности исполнения угроз доступности необходимо своевременно информировать сотрудников об изменениях в функционировании информационной системы, проводить обучение персонала, выполнять диагностику программно-технического комплекса и т.п.

2. Если рассматривать **целостность** как свойство информации, то это актуальность, полнота информации, отсутствие непротиворечивости, а также ее защищенность от уничтожения и несанкционированного изменения [6]. Это касается информации как общедоступной, так и с ограниченным доступом. Угрозам нарушения целостности подвержены не только данные, но и сама ИС.

Как правило, угрозы целостности связаны:

- с нарушением достоверности информации в результате нарушения правил и/или инструкций по работе с системой вследствие незнания или халатности;
- с непредумышленным или намеренным предоставлением или вводом в систему ошибочной информации.

Защита в этом случае должна быть комплексной и включать организационно-правовые и программно-аппаратные средства защиты. Например, для обеспечения достоверности и полноты информации необходимо выполнять контроль ввода новых данных, отслеживать изменения внутренней информации о ГМУ и изменений, касающихся законодательства РФ, своевременно обновляя данные. Для поддержания целостности информации и информационной системы необходимо использовать лицензионное программное обеспечение, сертифицированное аппаратное обеспечение, средства защиты от вредоносных программ, брандмауэры, средства резервного копирования, архивирования информации и т.д.

3. Для информации с ограниченным доступом важным направлением защиты является обеспечение конфиденциальности данных. **Конфиденциальность** – это свойство информации, характеризующее ее защиту от несанкционированного доступа. Источниками угроз конфиденциальности выступают лица, работающие в органах ГМУ или обслуживающие информационные системы ГМУ, пресса, иностранные агенты, спецслужбы, различные организации, физические лица, действующие в своих целях или целях заинтересованных лиц. Наибольшую опасность представляют собой злоумышленники, действующие внутри ГМУ и злоупотребляющие своими полномочиями [7].

Основные угрозы конфиденциальности – хищение или раскрытие информации с ограниченным доступом.

Порядок предоставления доступа к информации в ИС ГМУ определяется Правительством РФ с учетом законодательства РФ, в частности с учетом таких федеральных законов, как «О государственной тайне», «Об информации, информационных технологиях и о защите информации», «О персональных данных».

Средствами защиты информации будут:

– организационно-правовые средства, в которых определяются права, обязанности и ответственность лиц, которые могут участвовать в любых информационных процессах, характерных для деятельности органов ГМУ;

– программно-технические средства, реализующие управление доступом к конфиденциальной или закрытой информации, хранение данных в зашифрованном виде, использование электронной подписи, защиту от вредоносных программ и т.д.

Заключение

Таким образом, в статье рассмотрены информационные системы. Представленная классификация позволяет выделить крупную область ИС – государственные информационные системы. Вопросы информационной безопасности в данной области являются наиболее актуальными и требуют неотложного решения. Для государственных информационных систем определен состав с точки зрения безопасности. Пред-

ставлена методика классификации информационной системы по требованиям защиты информации, что позволило выявить основные угрозы информационной безопасности и их источники. Анализ угроз позволил определить направления и способы защиты информации.

Список литературы

1. Лебедкин А.П., Капелюшный Э.Д. Применение информационных технологий в экономике и управлении // Экономика и социум. Саратов, 2017. С. 851–854.
2. Закон Российской Федерации «Об информации, информационных технологиях и о защите информации» от 08.07.2006 (с изменениями на 14 июля 2022 года) № 149-ФЗ. [Электронный ресурс]. URL: http://www.consultant.ru/document/cons_doc_LAW_61798/ (дата обращения: 23.08.2022).
3. Закон Российской Федерации «Об обеспечении доступа к информации о деятельности государственных органов и органов местного самоуправления» от 21.01.2009 (с изменениями на 30 апреля 2021 года) № 8-ФЗ. [Электронный ресурс]. URL: http://www.consultant.ru/document/cons_doc_LAW_84602/ (дата обращения: 10.05.2021).
4. Белов Е.Б., Лось В.П., Мещеряков Р.В., Шелупанов А.А. Основы информационной безопасности: учебное пособие для вузов. М.: Горячая линия – Телеком, 2011. С. 69–71.
5. Методический документ. Меры защиты информации в государственных информационных системах [утвержден ФСТЭК России 11 февраля 2014 года]. [Электронный ресурс]. URL: <https://normativ.kontur.ru/document?moduleId=1&documentId=252834> (дата обращения: 10.05.2021).
6. Проза.ру Информационная безопасность. [Электронный ресурс]. URL: <https://proza.ru/2019/08/10/1345> (дата обращения: 01.04.2021).
7. Шептура С.В. Информационная безопасность. М., 2010. [Электронный ресурс]. URL: http://www.e-biblio.ru/book/bib/01_informatika/Inform_bezopast/sg/sg.html (дата обращения: 01.04.2021).

НАУЧНЫЙ ОБЗОР

УДК 004.6:330

**АНАЛИЗ, ПРОБЛЕМЫ И ВОЗМОЖНОСТИ ИСПОЛЬЗОВАНИЯ
БОЛЬШИХ ДАННЫХ В ГОРОДСКОМ ХОЗЯЙСТВЕ****Массеров Д.Д., Массеров Д.А.***ФГБОУ ВО «Национальный исследовательский Мордовский государственный университет
им. Н.П. Огарева», Саранск, e-mail: resurs2003@bk.ru*

Аналитика больших данных может быть использована во многих областях экономики, в том числе для анализа больших объемов данных в городском хозяйстве. В данной работе мы анализируем применение больших данных в таких областях, как производство, энергетика, транспорт и даже здравоохранение, везде, где интеллектуальные машины вовлечены в бизнес-процесс. Назрела необходимость в согласовании аппаратных технологий (т.е. машин и датчиков) с программными технологиями (т.е. представление данных, связь, хранение, анализ и контроль из машинного оборудования). Авторы рассмотрели будущие разработки встраиваемых систем, которые превращаются в «киберфизические системы». Перед наукой и специалистами стоят задачи по синхронизации совместной разработки аппаратного обеспечения (вычислений, датчиков и сетей) и программного обеспечения (форматов данных, операционных систем и систем анализа и управления). С одной стороны, это добавляет новую зависимость от использования больших данных, а именно зависимость от аппаратных систем, их разработки и ограничений. С другой стороны, это открывает новые возможности для решения более интегрированных систем с приложениями для использования больших объемов данных в качестве основы поддержки бизнес-решений в городском хозяйстве.

Ключевые слова: интеллектуальная сеть, хранение данных, анализ данных, обработка данных, сбор данных, безопасность данных

**ANALYSIS, CHALLENGES AND OPPORTUNITIES
FOR THE USE OF BIG DATA IN URBAN MANAGEMENT****Masserov D.D., Masserov D.A.***Mordovia State University, Saransk, e-mail: resurs2003@bk.ru*

Big data analytics can be used in many areas of the economy, including big data analytics in the urban economy. In this paper, we analyze the application of big data in areas such as manufacturing, energy, transportation, and even healthcare, wherever smart machines are involved in the business process. The need to align hardware technology (i.e., machines and sensors) with software technology (i.e., data presentation, communication, storage, analysis, and control from the machinery) is urgent. The authors reviewed future developments in embedded systems that are evolving into “cyber-physical systems.” Science and specialists face the challenge of synchronizing the joint development of hardware (computing, sensors, and networks) and software (data formats, operating systems, and analysis and control systems). On the one hand, this adds a new dependency on the use of big data, namely the dependency on hardware systems, their development and limitations. On the other hand, it opens up new opportunities to solve more integrated systems with applications to use big data as the basis for supporting business decisions in urban areas.

Keywords: phase measurement devices, smart grid, data storage, data analysis, data processing, data collection, data security

Огромный потенциал для инноваций имеет применение аналитики больших данных в государственном секторе [1], банковском деле [2], здравоохранении [3], сфере услуг [4], агропромышленном комплексе [5], Интернете вещей (IoT) [6, 7], связи [8], умных городах [9, 10] и транспорте [11]. Соответственно существует большой потенциал для применения больших данных в электросетях городского хозяйства, как в настоящее время, так и в будущем [12, 13]. Для эффективной работы электроэнергетических систем, обеспечивающих доступную, надежную, устойчивую и качественную энергию для конечных потребителей, в электросетях используются всевозможные инновации в области управления, связи, измерений и информатики. Для сбора общесистемных электрических измерений

и развертывания массивной инфраструктуры в электрических сетях начал широко использоваться расширенный учет (АМУ) в умных счетчиках и фазоизмерительных устройствах (PMU) [14]. Эффективное использование больших данных повышает надежность за электросетями. Для надежной и экономичной работы электросетей в условиях города необходим мониторинг общесистемных условий сети, поведения конечных потребителей и доступность возобновляемых источников энергии.

Масштаб увеличившихся измерительных устройств, а также данных на основе симуляторов и данных из неэлектрических источников приводит к колоссальному объему широко варьирующихся данных в электросетях городов, измеряемых тысячами терабайт информации каждый год [15]. Эти данные

поступают из различных источников, включая, программируемые термостаты, умные штекеры, умные приборы и счетчики, PMU, географическую информационную систему (ГИС), информацию о погоде, систему диспетчерского контроля и сбора данных (SCADA), информацию о трафике и сетей коммуникаций, как показано на рисунке [16].

Наличие и доступ к данным будут основой любой городской системы, ориентированной на данные. Здоровая городская система (назовем ее «экосистемой») данных будет состоять из широкого спектра различных типов данных: структурированных, неструктурированных, многоязычных, генерируемых машинами и датчиками, статических данных и данных в реальном времени. Полученный объем достоверной информации позволяет использовать новые алгоритмы управления городским хозяйством, которые могут привести к революционным изменениям в способах планирования и эксплуатации городских сетей.

Следовательно, аналитика больших данных будет играть все возрастающую роль не только для эффективной работы будущих электрических сетей, но и для разработки эффективных бизнес-моделей в городском хозяйстве [17].

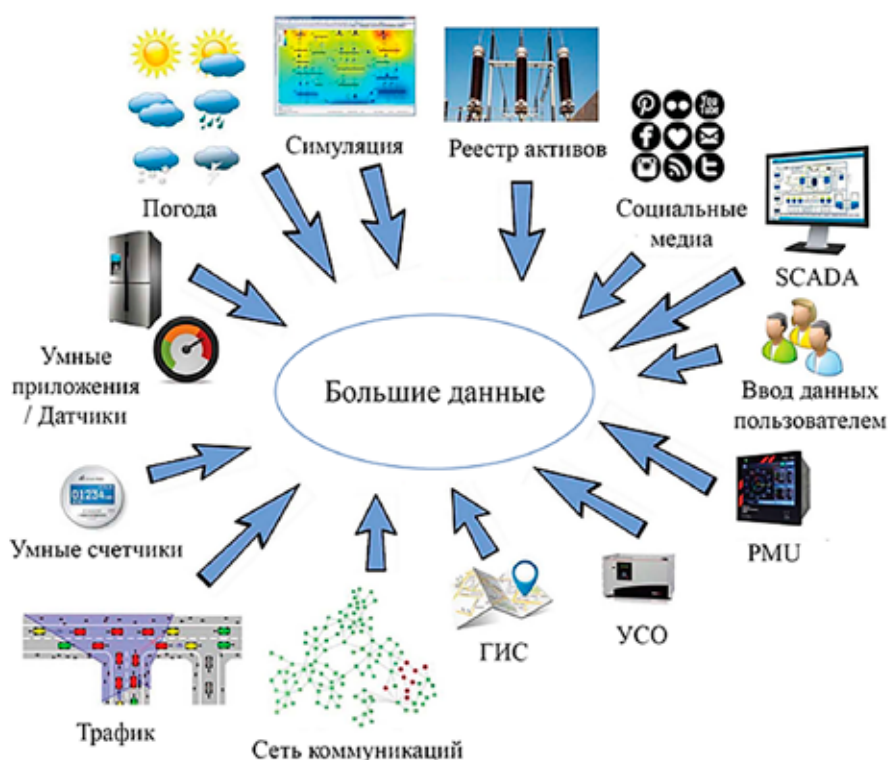
Цель исследования – дать анализ использования больших данных для текущего

применения в управлении городским хозяйством, а также определить проблемы и потенциал их применения.

Основные проблемы, связанные с развертыванием аналитики больших данных в будущих электросетях городов, заключаются в том, что, когда энергия поступает в распределительную сеть, она должна использоваться в это время. Поставщики энергии экспериментируют с устройствами хранения данных, чтобы помочь решить эту проблему, но они находятся на стадии становления и очень дороги. Поэтому проблема решается с помощью интеллектуальных измерительных приборов.

В настоящее время и в будущем объем данных (хранение данных, обработка данных, запрос данных и индексация данных) как у крупных, так и у мелких потребителей растет с экспоненциальной скоростью. Необходимы новые инновационные решения, заключающиеся в создании распределенной и масштабируемой вычислительной архитектуры [18].

Недостоверные данные как характеристика представляются отличительной особенностью достоверных данных умной сети. Данные не могут быть получены со 100% правдоподобностью, поскольку реальные данные восприимчивы к ошибкам из-за недостаточной приходящей информации [19].



Источники больших массивов данных в умных сетях города

Большие данные уже оказали влияние на многие предприятия и потенциально могут повлиять на все секторы бизнеса. Несмотря на наличие ряда технических проблем, влияние на управление и принятие решений и даже корпоративную культуру будет не менее значительным.

Однако все еще существует несколько границ. Именно проблемы конфиденциальности и безопасности должны решаться с помощью этих систем и технологий. Многие системы уже генерируют и собирают большие объемы данных, но лишь небольшой фрагмент активно используется в бизнес-процессах. Кроме того, многие из этих систем не имеют требований к работе в режиме реального времени [20]. Однако управление данными – это нечто большее, чем просто технические задачи по обработке данных.

Сбор данных связан со сбором данных из множества разнородных источников. Особенностью их сбора является конфиденциальность и безопасность при получении и передаче данных. Для этого используются подходы шифрования-дешифрования данных и агрегации-деагрегации [21].

Хранение данных представляет собой управление данными в специальных хранилищах. При сборе данных с интеллектуальных измерительных устройств первой проблемой является хранение большого объема данных. Например, если предположить, что 1 миллион устройств сбора данных извлекает 5 Кб данных за один сбор, потенциальный рост объема данных за год может составить до 2920 ТБ [22].

Анализ данных предназначен для извлечения ценной информации из набора данных. Вытекающие из этого задачи заключаются в анализе этого огромного объема данных, сопоставлении этих данных с информацией о клиентах, распределении сети и пропускной способности по сегментам, информацией о местной погоде и данными о стоимости энергии на спотовом рынке [23].

Использование этих данных позволит коммунальным службам лучше понять структуру затрат и стратегические варианты в рамках своей сети.

Одним из таких подходов от компании Lavastorm является проект, который исследует проблемы аналитики с такими инновационными компаниями, как FalbygdensEnergi AB (FEAB) и Sweco. Чтобы ответить на ключевые вопросы, используется аналитическая платформа Lavastorm. Аналитический движок Lavastorm – это решение для бизнес-аналитики с самообслуживанием, которое позволяет аналитикам быстро

получать, преобразовывать, анализировать и визуализировать данные, а также делиться ключевыми идеями и надежными ответами на бизнес-вопросы с нетехническими менеджерами и руководителями [24].

Поступающие из нескольких источников большие данные в умных сетях приносят ценную информацию для всех заинтересованных сторон, планирования и принятия оперативных решений.

Технологии хранения больших данных являются ключевым фактором для продвинутой аналитики, которая потенциально может трансформировать общество и способ принятия ключевых бизнес-решений. Это имеет особое значение в традиционно не связанных с ИТ секторах, таких как энергетика. В то время как эти секторы сталкиваются с нетехническими проблемами, такими как нехватка квалифицированных экспертов по большим данным и регуляторные барьеры, новые технологии хранения данных обладают потенциалом для создания новых аналитических возможностей, генерирующих ценность, в различных отраслях промышленности [25].

В дополнение к преимуществам, описанным выше, существуют также угрозы технологиям хранения больших объемов данных, которые необходимо устранить, чтобы избежать любого негативного воздействия. Это относится, например, к задаче защиты данных отдельных лиц и снижения энергопотребления центров обработки данных. Киберфизическая уязвимость является критической инфраструктурой в умной сети, которая может привести к широкомасштабным последствиям. Поэтому использование аналитики больших данных в кибербезопасности позволяет своевременно выявлять кибератаки и предотвращать огромный ущерб городской системе [26].

На сегодняшних сложных рынках электроэнергии или интермодальной мобильности практически не существует сценария, в котором все необходимые данные для ответа на бизнес- или инженерный вопрос поступали бы из баз данных одного департамента. Тем не менее большинство установленных в настоящее время передовых измерительных инфраструктур обеспечивают привязку полученных данных об энергопотреблении к платежным системам коммунальных предприятий. Блокировка затрудняет использование энергетических данных для другой ценной аналитики. Кроме того, объем данных, подлежащих обмену в прошлом, был намного меньше, так что интерфейсы, протоколы и процессы для обмена данными были довольно рудиментарными.

Открытыми остаются следующие вопросы: как представлять конкретные данные об энергии и мобильности, возможно, в нескольких измерениях – и как разработать алгоритмы, которые изучают ответы на конкретные вопросы в областях энергетики и мобильности лучше, чем это могут сделать операторы – люди, и делают это проверяемым образом. Основными вопросами машинного обучения являются экономически эффективное хранение и вычисления для огромных объемов данных с высокой выборкой, разработка новых эффективных структур данных и таких алгоритмов, как тензорное моделирование и сверточные нейронные сети. Встроенная аналитика и анализ распределенных данных, облегчающие внутрисетевую и полевую аналитику (иногда называемую периферийными вычислениями) в сочетании с аналитикой, проводимой на уровне предприятия, станут катализатором инноваций в энергетике и транспорте.

Энергетический и транспортный секторы с точки зрения инфраструктуры, а также эффективности использования ресурсов, глобальной конкурентоспособности и качества жизни очень важны для России.

Анализ доступных источников данных в энергетике, а также примеров их использования в различных категориях ценности больших данных, операционной эффективности, клиентского опыта и новых бизнес-моделей поможет определить промышленные потребности и требования к технологиям больших данных. При изучении этих требований становится ясно, что простого использования существующих технологий больших данных будет недостаточно. Необходима адаптация к конкретным областям и устройствам для использования в киберфизических энергетических и транспортных системах. Инновации, касающиеся управления и анализа данных с сохранением конфиденциальности и секретности, являются главной заботой заинтересованных сторон энергетического и транспортного секторов.

Среди заинтересованных сторон энергетического и транспортного секторов есть ощущение, что «больших данных» будет недостаточно. Растущий интеллект, встроенный в инфраструктуру, сможет анализировать данные для получения «умных данных». Это представляется необходимым, поскольку для аналитики потребуются гораздо более сложные алгоритмы, чем для других секторов. Кроме того, ставки в сценариях использования больших данных в энергетике и транспорте очень высоки, поскольку возможности оптимизации затронут критически важные инфраструктуры.

В энергетическом и транспортном секторах есть несколько примеров, когда технология сбора данных, т.е. интеллектуальное устройство, существует уже много лет, или заинтересованные стороны уже измеряют и собирают значительный объем данных. Благодаря последним достижениям появилась возможность передавать, хранить и обрабатывать данные с минимальными затратами.

Многие из современных технологий больших данных только ожидают адаптации и использования в этих традиционных секторах. Технологическая дорожная карта определяет и разрабатывает приоритетные требования и технологии, которые выведут энергетический и транспортный секторы за рамки современного уровня, чтобы они могли сосредоточиться на создании увеличения стоимости путем адаптации и применения этих технологий в своих конкретных областях применения городского хозяйства.

Такие инициативы, как «умные города», показывают, как различные секторы (например, энергетика и транспорт) могут сотрудничать, чтобы максимизировать потенциал оптимизации и отдачи от стоимости. Взаимное обогащение заинтересованных сторон и наборов данных из разных секторов является ключевым элементом для продвижения экономики больших данных. За последние годы сам объем данных, которые постоянно обрабатываются, увеличился. 90% данных в современном мире было получено за последние два года. Источник и характер этих данных разнообразны. Она варьируется от данных, собираемых датчиками, до данных, отражающих транзакции (онлайн). Все большая часть производится в социальных сетях и с помощью мобильных устройств. Тип данных (структурированные или неструктурированные) и семантика также различны. Тем не менее все эти данные должны быть агрегированы, чтобы помочь ответить на бизнес-вопросы и сформировать общую картину рынка.

Для бизнеса эта тенденция открывает ряд возможностей и проблем как для создания новых бизнес-моделей, так и для улучшения текущих операций, тем самым создавая рыночные преимущества. Например, интеллектуальные системы учета тестируются в энергетическом секторе. Кроме того, в сочетании с новыми системами выставления счетов эти системы также могут быть полезны в других секторах, таких как телекоммуникации и транспорт.

Анализ доступных источников данных в области энергетики, а также примеров их использования в различных категориях

ценности больших данных, операционной эффективности, клиентского опыта и новых бизнес-моделей помог определить промышленные потребности и требования к технологиям больших данных. При изучении этих требований становится ясно, что простого использования существующих технологий больших данных, используемых компаниями, занимающимися онлайн-обработкой данных, будет недостаточно. Необходима адаптация к конкретным предметным областям и устройствам для использования в киберфизических энергетических и транспортных системах. Инновации в области обеспечения конфиденциальности и управления данными и их анализа с сохранением конфиденциальности являются первоочередной задачей заинтересованных сторон энергетического и транспортного секторов. Без удовлетворения потребности в конфиденциальности всегда будет существовать неопределенность в регулировании и неопределенность в отношении принятия пользователями нового предложения, основанного на данных.

Список литературы

1. Ключева Е.Г., Ошакбай Д.М. Использование технологии «Больших данных» в государственном секторе // Автоматика. Информатика. 2019. № 2 (45). С. 35–38.
2. Trelewicz, J.Q. Big data and big money: the role of data in the financial sector. IT Prof. 2017. No. 19 (3). P. 8–10.
3. Islam M.M., Razzaque M.A., Hassan M.M. et al. Mobile cloud-based big healthcare data processing in smart cities. IEEE Access. 2017. No. 5. P. 11887–11899.
4. Боков И.С. Анализ больших данных как эффективное средство управления клиентами // Молодой ученый. 2019. № 22 (260). С. 605–609.
5. Отмахова Ю.С., Девяткин Д.А., Усенко Н.И. Анализ цифровых технологий в агропродовольственной сфере с использованием методов обработки больших данных // Информационное общество. 2021. № 4–5. С. 334–344. DOI: 10.52605/16059921_2021_04_334.
6. Marjani M., Nasaruddin F., Gani A. et al. Big IoT data analytics: architecture, opportunities, and open research challenges, IEEE Access, 2017. P. 5247–5261.
7. Массеров Д.А., Массеров Д.Д. Применение искусственного интеллекта в концепции Интернета вещей // Отходы и ресурсы. 2022. Т. 9. № 2. DOI: 10.15862/051TOR222.
8. He X., Ai Q., Qiu R.C. et al. A big data architecture design for smart grids based on random matrix theory, IEEE Trans. Smart Grid. 2017. No. 8 (2). P. 674–686.
9. Массеров Д.А., Массеров Д.Д. Обеспечение кибербезопасности умного города в процессе цифровизации городской среды // Информационные технологии и математическое моделирование в управлении сложными системами. 2022. № 1 (13). С. 32–42. DOI: 10.26731/2658-3704.2022.1(13).32-41.
10. Массеров Д.Д. Безопасность «умного города» в процессе цифровизации городской среды // XXIV Всероссийская студенческая научно-практическая конференция Нижневартковского государственного университета: материалы конференции (Нижневартковск, 05–06 апреля 2022 г.) / Под общ. ред. Д.А. Погонишева. Нижневартковск: Нижневартковский государственный университет, 2022. С. 135–141.
11. Zhu L., Yu F.R., Wang Y., Ning B. and Tang T. Big Data Analytics in Intelligent Transportation Systems: A Survey, in IEEE Transactions on Intelligent Transportation Systems. 2018. Vol. 20. No. 1. P. 383–398. DOI: 10.1109/TITS.2018.2815678.
12. Hong T. Big data analytics: making the smart grid smarter [Guest Editorial], IEEE Power Energy Mag. 2018. No. 16 (3). P. 12–16.
13. Кузнецов А. Новая эпоха в энергетике и умные сети // Наука и инновации. 2017. № 8 (174). С. 22–27.
14. Meliopoulos A.P.S., Cokkinides G., Huang R. et al. Smart grid technologies for autonomous operation and control, IEEE Trans. Smart Grid. 2011. No. 2 (1). P. 1–10.
15. Михеев А.В. Анализ больших данных для обоснования решений по научно-технологическому развитию в энергетике // Информационные и математические технологии в науке и управлении. 2020. № 4 (20). С. 158–167. DOI: 10.38028/ESI.2020.20.4.014.
16. Пейсахович В.Я. «Умные сети» и их место в инженерной инфраструктуре поселений // Региональная энергетика и энергосбережение. 2017. № 2. С. 82–83.
17. Zinaman O., Miller M., Adil A. et al. Power systems of the future. Electr. J. 2015. No. 28 (2). P. 113–126.
18. Денисова О.Ю., Мухутдинов Э.А. Большие данные – это не только размер данных // Вестник технологического университета. 2015. Т. 8. № 4. С. 226–230.
19. Рыговский И. Анализ эффективности методов обработки больших массивов данных с использованием вычислительных систем // Проблемы информатики. 2014. № 2 (23). С. 54–58.
20. Копылов А.А. Киберфизические атаки на умные сети // Инновации. Наука. Образование. 2021. № 48. С. 1710–1720.
21. Бочков А.П., Проурзин В.А., Проурзин О.В. Анализ больших данных надежности невосстанавливаемых многоканальных систем // Научные исследования в космических исследованиях Земли. 2021. Т. 13. № 4. С. 49–55. DOI: 10.36724/2409-5419-2021-13-4-49-55.
22. Аль-Раммахи А.А.Х., Громов Ю.Ю., Кошелев Е.В., Минин Ю.В. Обзор и анализ алгоритмов кластеризации больших данных // Приборы и системы. Управление, контроль, диагностика. 2021. № 9. С. 30–38. DOI: 10.25791/pribor.9.2021.1292.
23. Шайтура С.В., Галкин Д.А. Геомаркетинговый анализ больших данных // Информационные технологии. 2021. Т. 27. № 4. С. 180–187. DOI: 10.17587/it.27.180-187.
24. Олейник А.Н. Контент-анализ больших качественных данных. International Journal of Open Information Technologies. 2019. Т. 7. № 10. С. 36–49.
25. Касперская Н.И. Анализ больших данных в ИБ предприятий. Перспективы развития // Защита информации. Инсайд. 2019. № 3 (87). С. 34–43.
26. Березкин Д.В., Грошев В.Ю. Анализ современных средств 3D-визуализации больших данных для решения аналитических задач // Динамика сложных систем – XXI век. 2022. Т. 16. № 1. С. 22–34. DOI: 10.18127/j19997493-202201-03.

СТАТЬИ

УДК 004.89

**ИСПОЛЬЗОВАНИЕ ГРАФА СОАВТОРСТВА
ДЛЯ ОПРЕДЕЛЕНИЯ АВТОРСТВА ТЕКСТА
С НЕСКОЛЬКИМИ АВТОРАМИ****Батурин М.М., Белов Ю.С.***ФГБОУ ВО «Московский государственный технический университет имени Н.Э. Баумана»,
Калужский филиал, Калуга, e-mail: k4dys@yandex.ru*

Стилометрия успешно применяется для идентификации авторства документов с одним автором (authorship identification of single-author documents – AISD). Задача AISD связана с идентификацией первоначального автора анонимного документа из группы авторов-кандидатов. Однако методы AISD неприменимы к идентификации авторства документов с несколькими авторами (authorship identification of multi-author documents – AIMD). Из-за комбинаторного характера документов в AIMD отсутствует основная информация об истинных авторах, то есть информация о пишущих и непишущих авторах в документе с несколькими авторами, что усложняет решение этой проблемы. Помимо этого, из-за своей комбинаторной природы один и тот же список авторов не может повторяться в корпусе документов, что усложняет моделирование этой проблемы. В этой статье предлагается структура AIMD, называемая графом соавторства, которую можно использовать для сбора стилистической информации каждого автора в корпусе документов с несколькими авторами. Предлагаемая структура AIMD основана на наблюдении, что стилистически похожие фрагменты, вероятно, были написаны аналогичной группой авторов. Кроме того, предлагается итеративный алгоритм для идентификации оригинального автора каждого фрагмента документа.

Ключевые слова: определение авторства текста, граф соавторства, AIMD**USING THE CO-AUTHORSHIP GRAPH TO DETERMINE
THE AUTHORSHIP OF A TEXT WITH MULTIPLE AUTHORS****Baturin M.M., Belov Yu.S.***Bauman Moscow State Technical University, Kaluga branch, Kaluga, e-mail: k4dys@yandex.ru*

Stylometry is successfully used to identify the authorship of documents with one author (authorship identification of single-author documents – AISD). The aim of AISD is to identify the original author of an anonymous document from a group of candidate authors. However, AISD methods are not applicable to the identification of authorship of documents with multiple authors (authorship identification of multi-author documents – AIMD). Due to the combinatorial nature of the documents, AIMD lacks basic information about the true authors, that is, information about writing and non-writing authors in a document with multiple authors, which complicates the solution of this problem. In addition, due to its combinatorial nature, the same list of authors cannot be repeated in the corpus of documents, which complicates the modeling of this problem. This article proposes an AIMD structure called a co-authorship graph, which can be used to collect stylistic information of each author in a corpus of documents with multiple authors. The proposed AIMD structure is based on the observation that stylistically similar fragments were probably written by a similar group of authors. In addition, an iterative algorithm is proposed to identify the original author of each fragment of the document.

Keywords: determination of authorship of the text, graph of co-authorship, AIMD

В отличие от AISD, где каждый документ пишется одним автором, AIMD фокусируется на работе с документами, написанными несколькими авторами. Решения AIMD имеют ряд ограничений: лучшие решения AIMD на основе стилометрии имеют низкую точность; увеличение числа соавторов текстов негативно влияет на производительность решений AIMD; и решения AIMD не были предназначены для работы с непишущими авторами (non-writing authors – NWA). Тем не менее NWA существуют в реальных случаях, то есть существуют тексты, для которых не каждый из перечисленных соавторов внес свой вклад в качестве автора.

Цель исследования – изучить способы определения авторства текста с несколькими авторами.

Предложенное решение

Часть предварительной обработки фреймворка отвечает за три основных процесса: извлечение признаков, построение графа соавторства (Co-Authorship Graph – CAG) и обучение CAG. Для процесса извлечения признаков каждый документ с несколькими авторами представляется в виде набора фрагментов (т.е. коллекции наборов точек). Каждая точка данных рассчитывается из 500 токенов, с использованием 56 стилометрических признаков [1].

После завершения процесса выделения признаков CAG строится таким образом, чтобы каждая вершина представляла фрагмент, а наличие ребра между двумя вершинами означало, что они стилистически похожи. После завершения построения CAG обучение производится таким образом, что-

бы каждый фрагмент отражал только своего истинного автора (авторов). Структура CAG представлена на рис. 1.

После того как часть предварительной обработки была завершена, полученные в ходе обучения данные используются для создания прогноза нескольких авторов для любого заданного документа [2]. Чтобы получить вероятностный прогноз для данной выборки, используются вероятностные метки стилистически похожих выборок. Таким образом, эффективно фиксируется совместный характер документа как в тестовом, так и в обучающем наборе текстов.

Как объяснялось ранее, каждый документ в корпусе представляется в виде набора фрагментов, где каждый фрагмент представлен как набор точек [3]. Есть две основные причины для представления каждого документа в виде коллекции наборов точек. Это позволяет нам приписывать различные фрагменты одного и того же документа их первоначальным авторам, а также применять установленные меры подобия, связанные с методами обработки выбросов, такими как модифицированное расстояние Хаусдорфа (modified Hausdorff distance –

MHD), что может помочь улучшить общую производительность.

Чтобы получить достоверную стилистическую информацию из каждой точки данных, размер устанавливается равным 500 токенам. Однако использование 500 токенов на фрагмент дает только 24 точки данных для документа с 12000 токенов, чего недостаточно для стилометрического анализа. Чтобы преодолеть эту проблему, применяется концепция скользящего окна для создания точек данных с перекрывающимися последовательностями токенов. Этот процесс проиллюстрирован на рис. 2. Например, используя 500 токенов в качестве размера скользящего окна и 50 токенов в качестве приращения скользящего окна, мы можем создать 231 точку данных для каждого документа с 12 000 токенов. Этот же принцип применяется на уровне фрагментов, чтобы получить достаточное количество фрагментов для проведения анализа. В частности, установив размер фрагмента на 6 точек данных и значение приращения скользящего окна на 2, мы можем сгенерировать 113 фрагментов для каждого документа с 12 000 токенов.

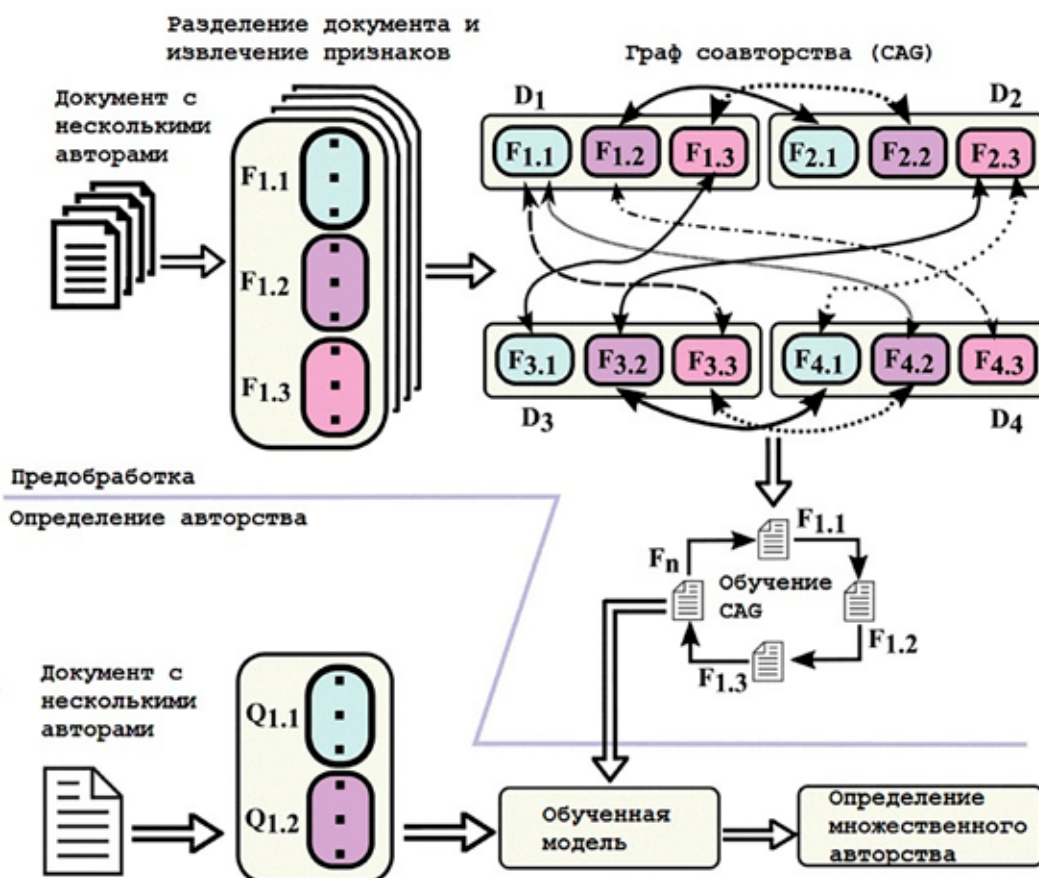


Рис. 1. Структура предложенной модели CAG

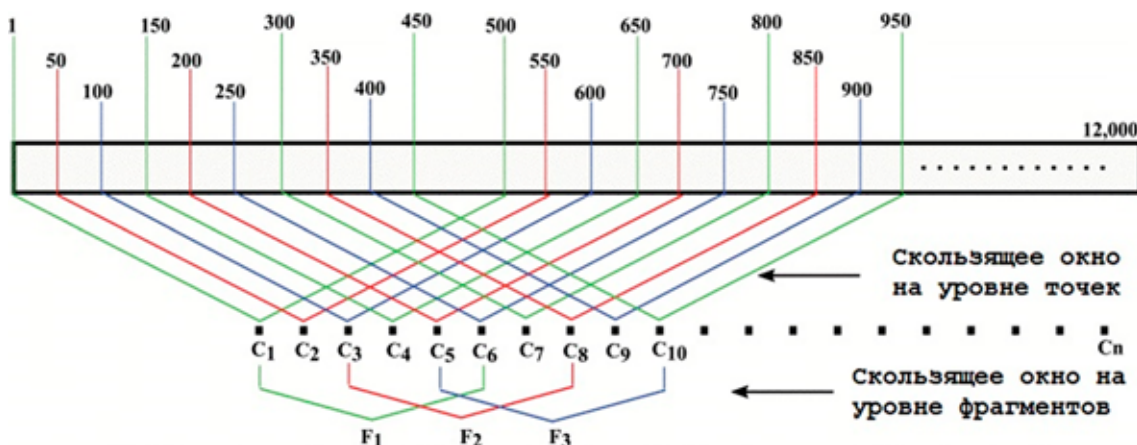


Рис. 2. Получение стилистической информации

В данной работе для обработки документов с несколькими авторами используются символьные n -граммы переменной длины [4]. После того как измеряются их показатели $tf-idf$, n -граммы ранжируются в порядке убывания в соответствии с их показателями $tf-idf$, чтобы выбрать 500 лучших n -грамм для использования в качестве признаков вместе с предыдущим пространством признаков.

Символьные n -граммы представляют собой непрерывную последовательность из n символов из образца текста. Доказано, что признаки на основе n -грамм хорошо работают при решении проблем идентификации авторства независимо от длины текстовых образцов. В частности, n -граммы символов могут эффективно фиксировать стилистическую информацию авторов из небольших образцов текста (например, 500 токенов) [5] по сравнению с функциями, основанными на словарном запасе. Характеристики символьной n -граммы обеспечивают наилучшие результаты, когда значение n равно 5, 4 или 3 [6]. Особенности n -граммы символов могут фиксировать сложную стилистическую информацию об авторах на синтаксическом, структурном и лексическом уровнях. Символьные n -граммы могут допускать шум в текстовых образцах. Извлечение символьных n -грамм не требует токенизаторов, теггировщиков, синтаксических анализаторов или каких-либо языково-зависимых и нетривиальных инструментов обработки текста, что делает их пригодными для выполнения многоязычных задач атрибуции авторства.

Создание CAG

Одна из основных проблем, связанных с проблемой AIMD, заключается в том, что каждый документ связан с несколькими авторами. Из-за своей комбинаторной

природы один и тот же список авторов не может повторяться в корпусе документов с несколькими авторами, что усложняет моделирование этой проблемы. Более того, в документе с несколькими авторами (например, в научной статье) некоторые из авторов могут не участвовать в написании самой статьи (NWA). То есть в задачах AIMD отсутствует базовая истинная информация, что усложняет решение этой проблемы. Таким образом, метод прогнозирования AIMD должен быть способен делать выводы об авторстве документа с несколькими авторами без абсолютной достоверной информации.

Предлагаемая структура AIMD основана на наблюдении, что стилистически похожие фрагменты, вероятно, были написаны аналогичной группой авторов [7]. Чтобы зафиксировать стилистическое сходство фрагментов документов, предлагается структура данных, называемая графом соавторства (CAG). Кроме того, предлагается итеративный алгоритм для идентификации оригинального автора каждого фрагмента документа.

После завершения процесса извлечения признаков строится граф соавторства (CAG), в котором каждая вершина представляет собой фрагмент, а ребро между двумя вершинами показывает, что они стилистически похожи. Для построения ребер CAG для каждого фрагмента идентифицируются k стилистически похожих фрагментов, где в качестве расстояния используется модифицированное расстояние Хаусдорфа. Эти ближайшие соседи – вершины графа, а MHD-расстояния – веса ребер. Предполагается, что каждый фрагмент документа связан со списком авторов из документа, который может включать одного или нескольких авторов. Список инициализируется путем назначения равной вероятности каждому автору.

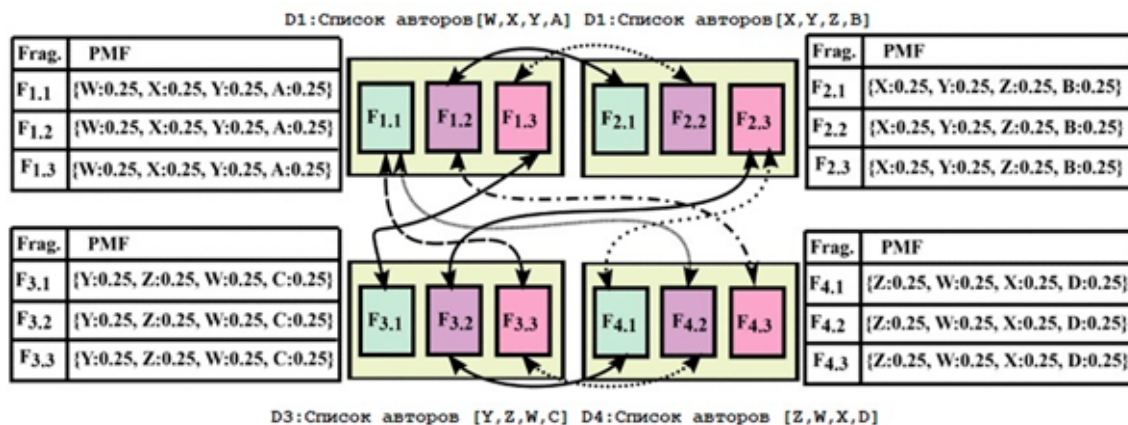


Рис. 3. Визуализация графа соавторства четырех текстов с четырьмя авторами

Обучение CAG

Теперь объясним процесс обучения CAG на примере, показанном на рис. 3. Имеется четыре документа, каждый связан с четырьмя авторами. Мы устанавливаем информацию об истинности данных следующим образом. Во-первых, предположим, что только первые три автора каждого документа внесли свой вклад в написание текста. Например, для документа D1 только авторы W, X и Y написали части документа D1, в то время как автор A был непишущим автором. Точно так же B, C и D являются NWA для D2, D3 и D4 соответственно, эта основная истинная информация скрыта от модели.

Поскольку основная информация о непишущих авторах скрыта от модели, каждый фрагмент документа [8] связывается со всеми перечисленными соавторами с равной распределенной вероятностью. Например, авторские PMF F1.1, F1.2, F1.3 являются однородными {W:0.25, X:0.25, Y:0.25, Z:0.25}. Таким же образом выводим начальные PMF остальных фрагментов, показанных на рис. 3.

Алгоритм построения CAG возвращает ребра, соединяющие стилистически похожие фрагменты. Например, по выявленным ребрам графа видно, что фрагмент F1.1 стилистически похож на F3.3 и F4.2. Точно так же F1.2 стилистически похож на F2.1 и F4.3. Несмотря на то, что каждый фрагмент в документе инициализируется с равной вероятностью, распределенной между всеми соавторами, они связаны с разными наборами стилистически сходных фрагментов, связанных с разными соавторами.

Алгоритм обучения CAG преследует две основные цели: изменить PMF каждого фрагмента таким образом, чтобы он отра-

жал истинного автора (авторов) этого фрагмента, и исключить непишущих авторов из списка. Один и тот же алгоритм выполняется в каждой вершине корпуса с несколькими итерациями, называемыми супершагами. В этом алгоритме каждая вершина отмечает k наиболее похожих фрагментов в качестве соседей. Этот алгоритм состоит из трех основных частей: получение, вычисление и отправка.

На каждом супершаге применяется тот же процесс, супершаги повторяются до тех пор, пока все PMF не сойдутся или количество итераций не достигнет заданного значения.

Пример

Рассмотрим теперь, как работает обучение в контексте примера, приведенного на рис. 3. Фрагмент F1.1 получает два PMF от двух своих соседей F3.3 и F4.2 как {Y:0.25, Z:0.25, W:0.25, C:0.25} и {Z:0.25, W:0.25, X:0.25, D:0.25} соответственно. Затем сравниваются PMF каждого фрагмента (соседей) со списком соавторов [W, X, Y, A], чтобы удалить авторов, которые не указаны в списке F1.1. В этом примере авторы Z и C исключены из фрагмента F3.3. Точно так же авторы Z и D исключаются из фрагмента F4.2. После исключения авторов, которых нет в списке соавторов F1.1, повторно нормализуем PMF. Повторная нормализация приводит к {Y:0.5, W:0.5} как PMF для F3.3 и {W:0.5, X:0.5} как PMF для F4.2. Для простоты изложения предположим, что два ближайших соседа находятся на одинаковом расстоянии от соответствующего фрагмента и вносят одинаковый вклад в PMF фрагмента. Следовательно, средневзвешенное значение двух PMF равно {W:0.5, X:0.25, Y:0.25} после первого супершага.

Следуя тому же процессу, получаем

1. {W:0,25, X:0,5, Y:0,25} для F1.2,
{W:0,25, X:0,25, Y:0,5} для F1.3.
2. {X:0,5, Y:0,25, Z:0,25} для F2.1,
{X:0,25, Y:0,5, Z:0,25} для F2.2,
{X:0,25, Y:0,25, Z:0,5} для F2.3.
3. {Y:0,5, Z:0,25, W:0,25} для F3.1,
{Y:0,25, Z:0,5, W:0,25} для F3.2,
{Y:0,25, Z:0,25, W:0,5} для F3.3.
4. {Z:0,5, W:0,25, X:0,25} для F4.1,
{Z:0,25, W:0,5, X:0,25} для F4.2,
{Z:0,25, W:0,25, X:0,5} для F4.3.

Как видно, все PMF становятся менее однородными только после первого супершага. Для каждого документа PMF сходятся к следующим значениям.

1. Документ D1: {W:1} для F1.1,
{X:1} для F1.2, и {Y:1} для F1.3.
2. Документ D2: {X:1} для F2.1,
{Y:1} для F2.2, и {Z:1} для F2.3.
3. Документ D3: {Y:1} для F3.1,
{Z:1} для F3.2, и {W:1} для F3.3.
4. Документ D4: {Z:1} для F4.1,
{W:1} для F4.2, и {X:1} для F4.3.

NWA каждого документа не включены в PMF, а списки авторов D1, D2 и D3 правильно определены как [W, X, Y], [X, Y, Z], [Y, Z, W] и [Z, W, X] соответственно.

Закключение

В этом исследовании предложен эффективный и масштабируемый подход для идентификации авторства документов с несколькими авторами. Основное преимущество предлагаемой структуры заключается в ее способности вероятно приписывать различные фрагменты (части)

одного и того же документа с несколькими авторами разным авторам. В частности, структура может фиксировать стилистическое сходство между парами фрагментов во всем корпусе документов. Кроме того, алгоритм обучения графа эффективен при изучении истинного автора (авторов) каждого фрагмента и определении NWA документов с несколькими авторами.

Список литературы

1. Soler J., Wanner L. On the relevance of syntactic and discourse features for author profiling and identification. Proceedings of the 15th Conference of the European Chapter of the Association for Computational Linguistics. 2017. Vol. 2. P. 681–687.
2. Sundararajan K., Woodard D. What represents 'style' authorship attribution? Proceedings of the 27th International Conference on Computational Linguistics, 2018. P. 2814–2822.
3. Zhang R., Hu Z., Guo H. Syntax encoding with application in authorship attribution. Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing. 2018. P. 2742–2753.
4. Батурин М.М., Белов Ю.С. Применение многозадачного обучения для определения авторства текста на основе механизма конкурентного внимания // Научное обозрение. Технические науки. 2022. № 3. С. 5–9.
5. Shrestha P., Sierra S., González F. Convolutional Neural Networks for Authorship Attribution of Short Texts. Proceedings of the 15th Conference of the European Chapter of the Association for Computational Linguistics. 2017. Vol. 2. P. 669–674.
6. Gómez-Adorno H., Posadas-Durán J.P., Sidorov G. Document embeddings learned on various types of n-grams for cross-topic authorship attribution. Computing. 2018. P. 1–16.
7. Батурин М.М., Белов Ю.С. Использование сверточных, рекуррентно-сверточных нейронных сетей и метода опорных векторов для определения авторства текста // Научные исследования в современном мире. Теория и практика: сборник избранных статей Всероссийской (национальной) научно-практической конференции. СПб., 2022. С. 47–49.
8. Stamatatos E. Authorship Attribution Using Text Distortion. Proceedings of the 15th Conference of the European Chapter of the Association for Computational Linguistics. 2017. Vol. 1. P. 1138–1149.

УДК 004.89

МОДИФИКАЦИИ АРХИТЕКТУРЫ WAVENET ДЛЯ РЕАЛИЗАЦИИ ВОКОДЕРА В ГЕНЕРАТИВНОЙ МОДЕЛИ ПРЕОБРАЗОВАНИЯ ТЕКСТА В РЕЧЬ

Белонозко П.Е., Белов Ю.С.

*ФГБОУ ВО «Московский государственный технический университет имени Н.Э. Баумана»,
Калужский филиал, Калуга, e-mail: iu4-kf@mail.ru*

Механизм преобразования текста в речь – это часть программного обеспечения, которая преобразует текст в речь (аудио). На сегодняшний день существует множество моделей, реализующих такой механизм. Среди них основное место занимают параметрическая, последовательная и генеративная модели. Генеративная модель является современной и эффективной, она не является полноценной системой преобразования текста в речь, каждая ее часть – это большой набор моделей и эвристик. Такая модель обычно разделена на конвейер, основными элементами которого являются синтезатор спектрограмм и вокодер, задача которого – построение формы волны по заданной мел-спектрограмме и акустическим характеристикам. WaveNet и Tacotron – это модели нейронных сетей, которые учитывают один шаг конвейера генеративной модели. В частности, WaveNet является нейронным вокодером и отвечает за этап «синтеза формы сигнала» конвейера. Tacotron – это последовательная модель для синтеза спектрограмм, предназначенная для этапа «высокоуровневого синтеза звука». Оригинальная модель WaveNet имеет ряд недостатков, влияющих на качество синтезируемой речи. Поэтому существуют модификации этой модели, улучшающие ее работу. Среди них встречаются подходы линейного предсказания (LP-WaveNet), авторегрессии и кондиционирования (WaveRNN), симбиоз с параметрической моделью (WaveGlow).

Ключевые слова: преобразование текста в речь, TTS, генеративная модель, вокодер, WaveNet, LP, WaveRNN, WaveGlow

MODIFICATIONS OF THE WAVENET ARCHITECTURE FOR THE IMPLEMENTATION OF A VOCODER IN A GENERATIVE MODEL OF TEXT-TO-SPEECH CONVERSION

Belonozhko P.E., Belov Yu.S.

Bauman Moscow State Technical University, Kaluga branch, Kaluga, e-mail: iu4-kf@mail.ru

The text-to-speech engine is a piece of software that converts text to speech (audio). To date, there are many models that implement such a mechanism. Among them, the main place is occupied by the given, subsequent and generative models. The generating model is modern and voluminous, it does not have a large amount of text-to-speech, each part of it is a set of models and heuristics. Such a model, as a rule, is used on a conveyor, in particular, for which there are a spectrogram synthesizer and a vocoder, the task of which is to build waveforms according to a given chalk spectrogram and acoustic characteristics. WaveNet and Tacotron are neural network models that take into account one step of the generating model pipeline. Specifically, WaveNet is a neural encoder and is responsible for the “shape synthesis” step of the pipeline. Tacotron is a sequential spectrogram synthesis model designed for “high-level audio synthesis” steps. The original WaveNet model has a number of shortcomings that affect the quality of the synthesized speech. As a result, modifications to this model improve its performance. Among them there are approaches of linear prediction (LP-WaveNet), autoregression and dependence (WaveRNN), symbiosis with a parametric model (WaveGlow).

Keywords: text-to-speech, TTS, generative model, vocoder, WaveNet, LP, WaveRNN, WaveGlow

Технологии глубокого обучения становятся все более популярными в области искусственного интеллекта. Одной из таких технологий является преобразование текста в речь (text-to-speech, TTS), задачей которой является преобразование заданного текста в разговорный человеческий голос.

Синтез речи – это область, тесно связанная с искусственным воспроизведением человеческих голосов, к которой относится технология TTS, где входными данными является текст на естественном языке, а выходными данными – аудио/речь, имитирующая человеческий голос. Данные системы начали свое развитие в 2016 г. и до сих пор большая часть исследований сосредоточена на том, чтобы сделать их более эффектив-

ными, звучащими более естественно, с правильным произношением, живой интонацией и минимумом фонового шума [1].

В последние пару десятилетий доминирующими методами для TTS были последовательный синтез и статистический параметрический синтез речи. Позже были представлены генеративные модели, такие как WaveNet, которые являются полностью авторегрессионными и вероятностными моделями. Каждый отдельный аудиосэмпл прогнозируется из условного распределения вероятностей (обусловленного всеми ранее предсказанными аудиосэмплими), которое вычисляется для соответствующего прогнозируемого аудиосэмпла. Эта модель может использоваться для TTS, когда

распределение также обусловлено лингвистическими особенностями, такими как фонетика, полученная из текста или его предсказанной спектрограммы.

Основной проблемой генерации качественной речи все еще является дефицит наборов данных. Для обучения обычной модели TTS, такой как Tacotron, требуются сотни часов профессионально записанной речи. Однако решением этой проблемы является генеративная модель синтеза речи из текста, которая клонирует голос, а не преобразует его. Если преобразование голоса представляет собой форму передачи стиля в сегменте речи от одного голоса к другому, то клонирование состоит в захвате голоса говорящего для преобразования текста в речь на основании произвольных данных. Системы генерации сигналов с использованием WaveNet значительно улучшили качество синтеза систем преобразования текста в речь (TTS) на основе глубокого обучения. Поскольку вокодер WaveNet может генерировать речевые образцы в единой унифицированной нейронной сети, он не требует какого-либо ручного конвейера обработки сигналов. Таким образом, он обеспечивает гораздо более высокое синтетическое качество, чем традиционные параметрические вокодеры. Существует множество моделей, построенных на архитектуре WaveNet: LP-WaveNet, WaveRNN, WaveGlow [2].

Архитектура LP-WaveNet

Чтобы еще больше улучшить качество восприятия синтезированной речи, в более поздних вокодерах с нейронным возбуждением используются преимущества как вокодера с линейным предсказанием (LP), так и структуры WaveNet. Спектральная структура речевого сигнала, связанная с формантами, развязывается фильтром анализа LP, и WaveNet только оценивает распределение своего остаточного сигнала (т.е. возбуждение). Поскольку физическое поведение сигнала возбуждения проще, чем речевого сигнала, процессы обучения и генерации становятся более эффективными.

Однако синтезированная речь будет неестественной, когда ошибки предсказания при оценке возбуждения распространяются через процесс синтеза LP. Поскольку эффект синтеза LP не учитывается в процессе обучения, выходные данные синтеза уязвимы для изменения фильтра синтеза LP [3].

Чтобы облегчить эту проблему, существует LP-WaveNet, который позволяет совместно обучать сложные взаимодействия между фильтром возбуждения и синтеза

LP. Исходя из основного предположения, что прошлые речевые отсчеты и коэффициенты LP даны как условная информация, поэтому распределения речи и возбуждения лежат только на постоянной разности. Целевое распределение речи можно оценить путем суммирования средних параметров предсказанного результата и аппроксимации LP, которая определяется как линейная комбинация выборки прошлой речи, взвешенные по коэффициентам LP. LP-WaveNet легко обучать, потому что WaveNet нужно моделировать только компонент возбуждения, а сложная часть моделирования спектра встроена в приближение LP.

WaveNet представляет собой авторегрессивную генеративную модель на основе сверточной нейронной сети (CNN), которая предсказывает совместное распределение вероятностей выборок речи. Путем многократной укладки расширенных причинно-следственных сверточных слоев WaveNet эффективно расширяет свое восприимчивое поле до тысячи отсчетов.

WaveNet, также известная как WaveNet с μ -законом, определяет распределение выборки речи как 256 категориальных классов символов, получаемых с помощью 8-битных квантованных по закону μ -выборок речи. Чтобы смоделировать распределение выборки речи, категориальное распределение вычисляется путем применения операции softmax к выходным данным WaveNet. На этапе обучения веса WaveNet обновляются, чтобы минимизировать потери перекрестной энтропии. На этапе генерации образец речи генерируется авторегрессивно по образцу.

Поскольку WaveNet с μ -законом может генерировать речевой сигнал в единой унифицированной модели, он обеспечивает значительно лучшее синтетическое звучание, чем обычные параметрические вокодеры. Однако обучать сеть непросто, когда объем базы данных больше, а ее акустическая информация, такая как просодия, стиль или выразительность, шире. Кроме того, синтезированный звук WaveNet часто страдает от артефакта – фонового шума, поскольку целевой речевой сигнал слишком грубо квантован.

Подробная архитектура LP-WaveNet показана на рис. 1.

В предлагаемой системе распределение выборок речи определяется как распределение MoG (логнормальное распределение), следовательно, LP-WaveNet обучается генерировать параметры MoG, $[w_n, \mu_n, \sigma_n]$, обусловленные входными данными (акустическими признаками) [4].

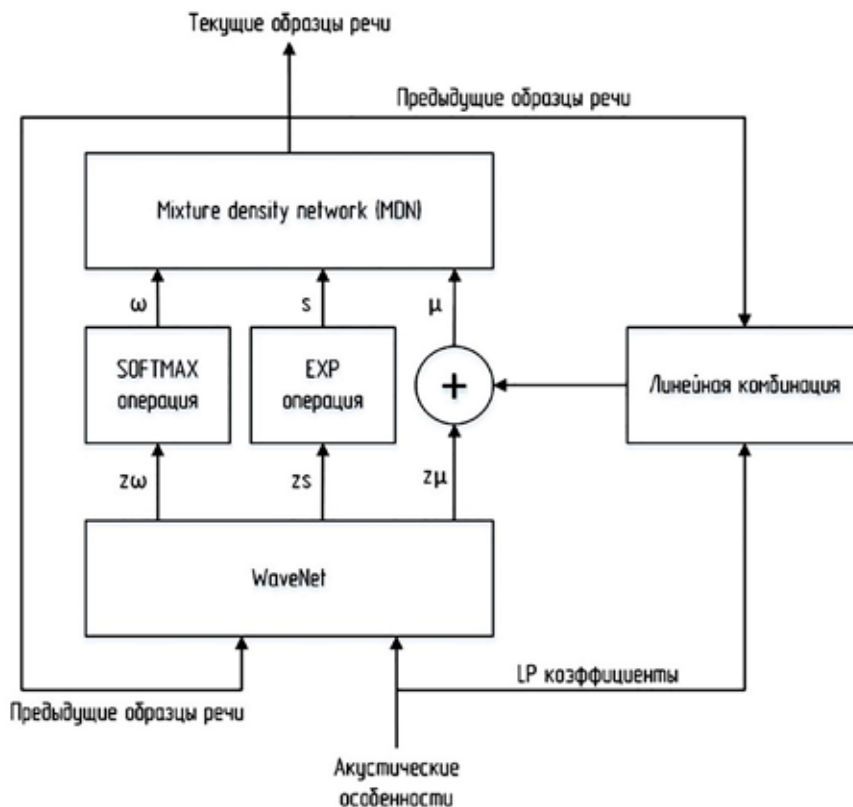


Рис. 1. Вокодер, основанный на архитектуре LP-WaveNet

В частности, акустические признаки проходят через два одномерных сверточных слоя с размером ядра 3 для явного наложения контекстной информации об изменении признаков. Затем применяется остаточное соединение по отношению к входному акустическому признаку, чтобы сделать сеть более сфокусированной на информации о текущем кадре. Наконец, транспонированная свертка применяется для повышения дискретизации временного разрешения этих признаков до разрешения речевого сигнала.

Сгенерированный сигнал часто может быть нестабильным, когда сгенерированные параметры логарифмического масштаба слишком высоки. Эту проблему можно предотвратить, обрезав верхнюю границу значения параметра масштаба. Если отсечение будет установлено слишком низким, то невокализованная область будет недостаточно смоделирована, что приводит к сухому синтетическому звуку, хотя волновая форма может стабильно генерироваться. Если отсечение будет установлено слишком высоким, то вероятность взрыва формы волны становится выше, но синтетический звук становится более живым, чем в случае более низкого значения отсечения.

Архитектура WaveGlow

WaveGlow – это генеративная модель, которая генерирует звук путем выборки из распределения. Чтобы использовать нейронную сеть в качестве генеративной модели, используются образцы из простого распределения, в данном случае сферического гауссова с нулевым средним значением с тем же числом измерений. Данные образцы проходят через серию слоев, которые преобразуют простое распределение к тому, которое требуется. В случае WaveGlow моделируется распределение аудиосэмплов на основе мел-спектрограммы.

Основной проблемой является минимизация отрицательного логарифмического правдоподобия данных. При использовании произвольной нейронной сети это неразрешимо. Поточковые сети решают эту проблему, обеспечивая обратимость отображения нейронной сети. Ограничив биективность каждого слоя, вероятность можно рассчитать напрямую, используя замену переменных [5].

Архитектура WaveGlow представлена на рис. 2. Первый член представляет собой логарифмическое правдоподобие сферического гауссиана. Этот член штрафует норму 12 преобразованной выборки.

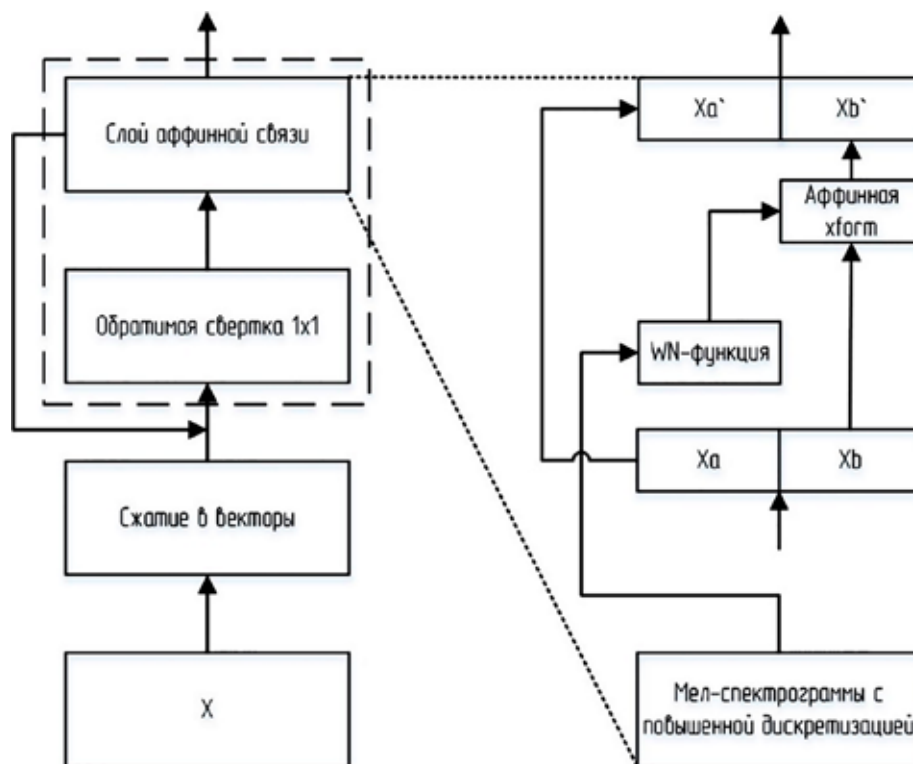


Рис. 2. Вокодер, основанный на архитектуре WaveGlow

Второй член возникает из-за замены переменных, J является якобианом. Лог-детерминант якобиана вознаграждает любой слой за увеличение объема пространства во время прямого прохода. Этот член также не позволяет слою просто умножать члены x на ноль, чтобы оптимизировать норму 12. Последовательность преобразований также называют нормализующим потоком. Затем эти векторы обрабатываются через несколько «этапов потока». Шаг потока здесь состоит из обратимой свертки 1×1 , за которой следует слой аффинной связи.

Обратимые нейронные сети обычно строятся с использованием связующих слоев (аффинная связь). Половина каналов служат входными данными, которые затем производят мультипликативные и аддитивные условия, которые используются для масштабирования и преобразования оставшихся каналов.

Здесь WN может быть любым преобразованием. Слой связи сохраняет обратимость всей сети, хотя WN необязательно должен быть обратимым. Это следует из того, что каналы, используемые в качестве входных данных для WN, в данном случае x_a , передаются без изменений на выход слоя. Соответственно, при инвертировании сети можно вычислить s и t из выходных данных

x_a , а затем инвертировать x_b для вычисления x_b , просто повторно вычислив WN (x_a , мел-спектрограмму). WN использует слои расширенных сверток с гейттановыми нелинейностями, а также остаточные соединения и пропуски соединений. Эта архитектура WN похожа на WaveNet и Parallel WaveNet, но свертки имеют три отвода и не являются причинно-следственными. Слой аффинной связи также включает мел-спектрограмму, чтобы обусловить сгенерированный результат на входе. Мел-спектрограммы с повышенной частотой дискретизации добавляются перед нелинейниками с гейттаном каждого слоя, как в WaveNet [6].

При использовании слоя аффинной связи только член s изменяет объем отображения и добавляет член изменения переменных к потерям. Этот термин также служит для наказания модели за необратимые аффинные отображения.

Архитектура WaveRNN

WaveRNN – это модель преобразования текста в речь (TTS), которая обеспечивает качество звука на уровне оригинальной WaveNet, но может быть сэмплирована в реальном времени на стандартном оборудовании. Она состоит из авторегрессионной и кондиционирующей сети.

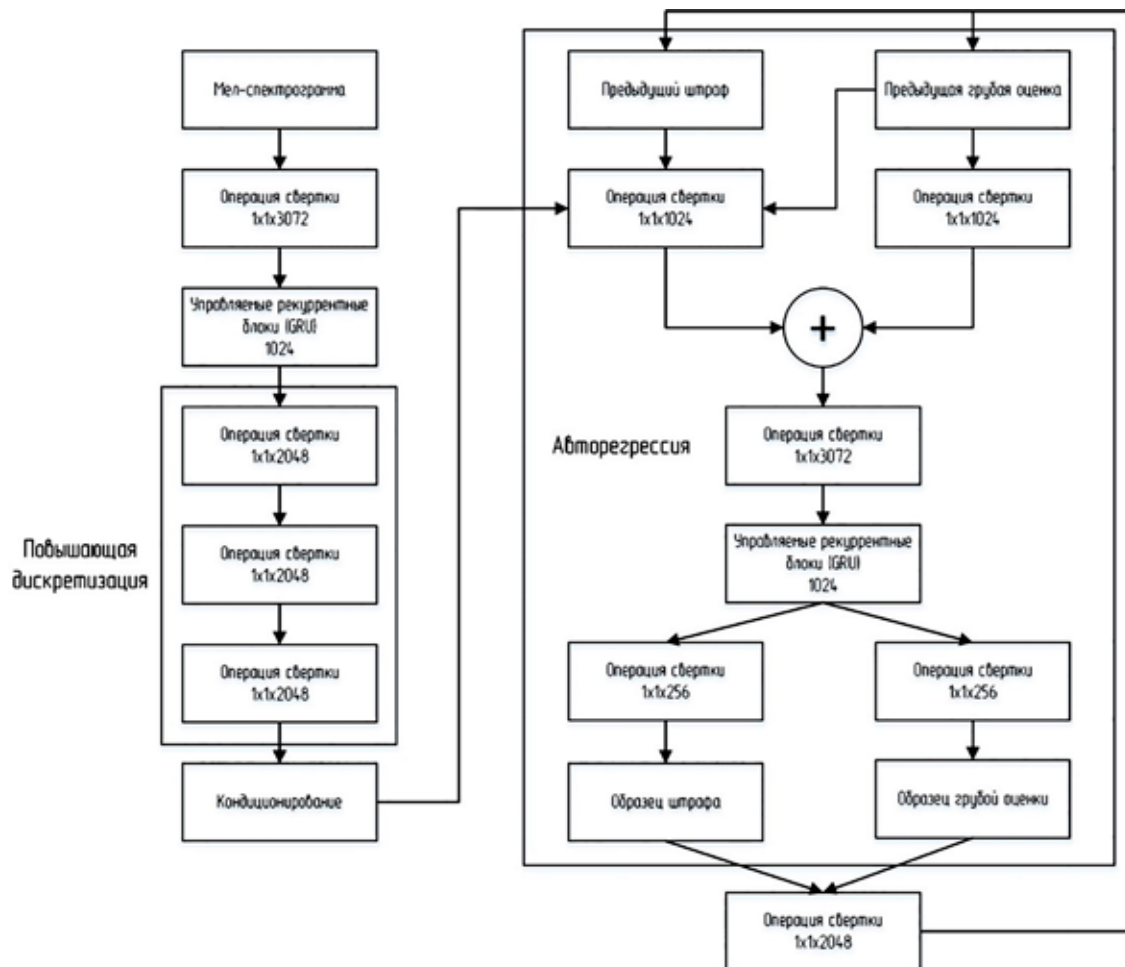


Рис. 3. Вокодер, основанный на архитектуре WaveRNN

Первая использует однослойную рекуррентную нейронную сеть с модифицированным вентилируемым рекуррентным блоком (GRU) в качестве ядра. В качестве входных данных авторегрессионный слой получает последнюю созданную им выборку и выходные данные сети кондиционирования, которые объединяются с помощью сверточных слоев. Выход представляет собой двойной слой softmax, который позволяет эффективно прогнозировать аудио-сэмплы с 16-битным разрешением. Кондиционирующая сеть состоит из свертки 1×1 , за которой следует RNN, использующая ту же ячейку GRU, что и авторегрессионная сеть, и три транспонированные свертки, которые повышают частоту дискретизации в 8 раз с частоты кондиционирования от 50 до 400 Гц. Активации из кондиционирующей сети затем транслируются на каждый временной шаг в авторегрессионной сети с ее исходной частотой дискретизации. Это означает, что существует отдельный новый вектор активации кондиционирования,

генерируемый сетью кондиционирования каждые 2,5 мс. Сэмплирование выполняется на частоте 8 кГц. На рис. 3 показана подробная архитектура WaveRNN.

Кондиционирующая сеть обычно имеет более широкое рецептивное поле, что позволяет ей моделировать медленно меняющиеся долгосрочные характеристики, в то время как авторегрессионный слой сам по себе способен моделировать краткосрочную динамику реалистичной речи.

Когда WaveRNN используется для TTS, кондиционирующий слой получает содержимое и просодию целевой речи в качестве входных данных и влияет на авторегрессионную сеть для создания правильных волновых форм. В конфигурации PLC нет доступа к этой информации. Вместо этого в кондиционирующую сеть подается окно логарифмической спектрограммы прошлого аудио, что позволяет ей извлекать необходимую информацию из прошлого аудио и направлять авторегрессионную сеть для формирования аудио таким образом,

чтобы оно соответствовало стилю и содержанию того, что было записано [7].

Заключение

В данной работе рассматриваются варианты реализации вокодера на основе архитектуры WaveNet. Оригинальная архитектура WaveNet имеет серьезный недостаток – качество восприятия синтезируемой речи. Для этого используются модифицированные архитектуры, такие как LP-WaveNet, WaveRNN, WaveGlow, которые не только улучшают качество синтезируемой речи, но и повышают эффективность работы системы, основанной на них. Каждый вариант реализации имеет свои преимущества и недостатки. Но несмотря на это они показывают равные результаты работы.

Список литературы

1. Shen J. Natural TTS Synthesis by Conditioning Wavenet on MEL Spectrogram Predictions. IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP). 2018. P. 4779–4783.
2. Cooper E. Zero-Shot Multi-Speaker Text-To-Speech with State-Of-The-Art Neural Speaker Embeddings. IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP). 2020. P. 6184–6188.
3. Weiss R.J., Skerry-Ryan R., Battenberg E. Wave-Tacotron: Spectrogram-Free End-to-End Text-to-Speech Synthesis. IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP). 2021. P. 5679–5683.
4. Gu Y. ByteSing: A Chinese Singing Voice Synthesis System Using Duration Allocated Encoder-Decoder Acoustic Models and WaveRNN Vocoders. 12th International Symposium on Chinese Spoken Language Processing (ISCSLP). 2021. P. 1–5.
5. Huang W.C. Refined WaveNet Vocoder for Variational Autoencoder Based Voice Conversion. 27th European Signal Processing Conference (EUSIPCO). 2019. P. 1–5.
6. Okamoto T., Toda T., Shiga Y. Tacotron-Based Acoustic Model Using Phoneme Alignment for Practical Neural Text-to-Speech Systems. IEEE Automatic Speech Recognition and Understanding Workshop (ASRU). 2019. P. 214–221.
7. Hwang M.J., Soong F., Song E. LP-WaveNet: Linear Prediction-based WaveNet Speech Synthesis. Asia-Pacific Signal and Information Processing Association Annual Summit and Conference (APSIPA ASC). 2020. P. 810–814.