

ISSN 2713-3192
DOI 10.15622/ia.2023.22.2
<http://ia.spcras.ru>

ТОМ 22 № 2

ИНФОРМАТИКА И АВТОМАТИЗАЦИЯ

INFORMATICS AND AUTOMATION



СПб ФИЦ РАН

Санкт-Петербург
2023



INFORMATICS AND AUTOMATION

Volume 22 № 2, 2023

Scientific and educational journal primarily specialized in computer science, automation, robotics, applied mathematics, interdisciplinary research

Founded in 2002

Founder and Publisher

St. Petersburg Federal Research Center of the Russian Academy of Sciences (SPC RAS)

Editor-in-Chief

R. M. Yusupov, Prof., Dr. Sci., Corr. Member of RAS, St. Petersburg, Russia

Editorial Council

A. A. Ashimov	Prof., Dr. Sci., Academician of the National Academy of Sciences of the Republic of Kazakhstan, Almaty, Kazakhstan
N. P. Veselkin	Prof., Dr. Sci., Academician of RAS, St. Petersburg, Russia
I. A. Kalyaev	Prof., Dr. Sci., Academician of RAS, Taganrog, Russia
Yu. A. Merkuriev	Prof., Dr. Sci., Academician of the Latvian Academy of Sciences, Riga, Latvia
A. I. Rudskoi	Prof., Dr. Sci., Academician of RAS, St. Petersburg, Russia
V. Sgurev	Prof., Dr. Sci., Academician of the Bulgarian Academy of Sciences, Sofia, Bulgaria
B. Ya. Sovetov	Prof., Dr. Sci., Academician of RAE, St. Petersburg, Russia
V. A. Soyfer	Prof., Dr. Sci., Academician of RAS, Samara, Russia

Editorial Board

O. Yu. Gusikhin	Ph. D., Dearborn, USA
V. Delic	Prof., Dr. Sci., Novi Sad, Serbia
A. Dolgui	Prof., Dr. Sci., St. Etienne, France
M. N. Favorskaya	Prof., Dr. Sci., Krasnoyarsk, Russia
M. Zelezny	Assoc. Prof., Ph.D., Plzen, Czech Republic
H. Kaya	Assoc. Prof., Ph.D., Utrecht, Netherlands
A. A. Karpov	Assoc. Prof., Dr. Sci., St. Petersburg, Russia
S. V. Kuleshov	Dr. Sci., St. Petersburg, Russia
A. D. Khomonenko	Prof., Dr. Sci., St. Petersburg, Russia
D. A. Ivanov	Prof., Dr. Habil., Berlin, Germany
K. P. Markov	Assoc. Prof., Ph.D., Aizu, Japan
R. V. Meshcheryakov	Prof., Dr. Sci., Moscow, Russia
N. A. Moldovian	Prof., Dr. Sci., St. Petersburg, Russia
V. V. Nikulin	Prof., Ph.D., New York, United States
V. Yu. Osipov	Prof., Dr. Sci., St. Petersburg, Russia
V. K. Pshikhopov	Prof., Dr. Sci., Taganrog, Russia
A. L. Ronzhin	Prof., Dr. Sci., Deputy Editor-in-Chief, St. Petersburg, Russia
H. Samani	Assoc. Prof., Ph.D., Plymouth, UK
A. V. Smirnov	Prof., Dr. Sci., St. Petersburg, Russia
B. V. Sokolov	Prof., Dr. Sci., St. Petersburg, Russia
L. V. Utkin	Prof., Dr. Sci., St. Petersburg, Russia

Editor: A. S. Lopotova

Interpreter: Ya. N. Berezina

Art editor: N. A. Dormidontova

Editorial office address

SPC RAS, 39 litera A, 14-th line V.O., St. Petersburg, 199178, Russia

e-mail: ia@spcras.ru, web: <http://ia.spcras.ru>

The journal is indexed in Scopus

The journal is published under the scientific-methodological supervision of Department for Nanotechnologies and Information Technologies of the Russian Academy of Sciences

© St. Petersburg Federal Research Center of the Russian Academy of Sciences, 2023

ИНФОРМАТИКА И АВТОМАТИЗАЦИЯ

Том 22 № 2, 2023

Научный, научно-образовательный журнал с базовой специализацией
в области информатики, автоматизации, робототехники, прикладной математики
и междисциплинарных исследований.

Журнал основан в 2002 году

Учредитель и издатель

Федеральное государственное бюджетное учреждение науки
«Санкт-Петербургский Федеральный исследовательский центр Российской академии наук»
(СПб ФИЦ РАН)

Главный редактор

Р. М. Юсупов, чл.-корр. РАН, д-р техн. наук, проф., Санкт-Петербург, РФ

Редакционный совет

А. А. Ашимов	академик Национальной академии наук Республики Казахстан, д-р техн. наук, проф., Алматы, Казахстан
Н. П. Веселкин	академик РАН, д-р мед. наук, проф., Санкт-Петербург, РФ
И. А. Калыев	академик РАН, д-р техн. наук, проф., Таганрог, РФ
Ю. А. Меркурьев	академик Латвийской академии наук, д-р, проф., Рига, Латвия
А. И. Рудской	академик РАН, д-р техн. наук, проф., Санкт-Петербург, РФ
В. Сгурев	академик Болгарской академии наук, д-р техн. наук, проф., София, Болгария
Б. Я. Советов	академик РАО, д-р техн. наук, проф., Санкт-Петербург, РФ
В. А. Сойфер	академик РАН, д-р техн. наук, проф., Самара, РФ

Редакционная коллегия

О. Ю. Гусихин	д-р наук, Диаборн, США
В. Делич	д-р техн. наук, проф., Нови-Сад, Сербия
А. Б. Долгий	д-р наук, проф. Сент-Этьен, Франция
М. Железны	д-р наук, доцент, Пльзень, Чешская республика
Д. А. Иванов	д-р экон. наук, проф., Берлин, Германия
Х. Кайя	д-р наук, доцент, Утрехт, Нидерланды
А. А. Карпов	д-р техн. наук, доцент, Санкт-Петербург, РФ
С. В. Кулешов	д-р техн. наук, Санкт-Петербург, РФ
К. П. Марков	д-р наук, доцент, Аизу, Япония
Р. В. Мещеряков	д-р техн. наук, проф., Москва, РФ
Н. А. Молдовян	д-р техн. наук, проф., Санкт-Петербург, РФ
В. В. Никулин	д-р наук, проф., Нью-Йорк, США
В. Ю. Осипов	д-р техн. наук, проф., Санкт-Петербург, РФ
В. Х. Пшихолопов	д-р техн. наук, проф., Таганрог, РФ
А. Л. Ронжин	д-р техн. наук, проф., зам. главного редактора, Санкт-Петербург, РФ
Х. Самани	д-р наук, доцент, Плимут, Соединённое Королевство
А. В. Смирнов	д-р техн. наук, проф., Санкт-Петербург, РФ
Б. В. Соколов	д-р техн. наук, проф., Санкт-Петербург, РФ
Л. В. Уткин	д-р техн. наук, проф., Санкт-Петербург, РФ
М. Н. Фаворская	д-р техн. наук, проф., Красноярск, РФ
А. Д. Хомоненко	д-р техн. наук, проф., Санкт-Петербург, РФ
Л. Б. Шереметов	д-р техн. наук, Мехико, Мексика

Выпускающий редактор: А. С. Лопотова

Переводчик: Я. Н. Березина

Художественный редактор: Н. А. Дормидонтова

Адрес редакции

14-я линия В.О., д. 39, лит. А, г. Санкт-Петербург, 199178, Россия

e-mail: ia@spcras.ru, сайт: <http://ia.spcras.ru>

Журнал индексируется в международной базе данных Scopus

Журнал входит в «Перечень ведущих рецензируемых научных журналов и изданий,
в которых должны быть опубликованы основные научные результаты диссертации
на соискание ученой степени доктора и кандидата наук»

Журнал выпускается при научно-методическом руководстве Отделения нанотехнологий
и информационных технологий Российской академии наук

© Федеральное государственное бюджетное учреждение науки

«Санкт-Петербургский Федеральный исследовательский центр Российской академии наук», 2023
Разрешается воспроизведение в прессе, а также сообщение в эфир или по кабелю опубликованных
в составе печатного периодического издания - журнала «ИНФОРМАТИКА И АВТОМАТИЗАЦИЯ»
статей по текущим экономическим, политическим, социальным и религиозным вопросам
с обязательным указанием имени автора статьи и печатного периодического издания
журнала «ИНФОРМАТИКА И АВТОМАТИЗАЦИЯ»

CONTENTS

Digital Information Telecommunication Technologies

M. Gofman, A. Kornienko

LIMIT BIPOLAR SEQUENCES FOR PATCHWORK-BASED ROBUST
DIGITAL AUDIO WATERMARKING 221

S.V. Dvornikov, S.S. Dvornikov, K. Zheglov

NOISE IMMUNITY OF SINGLE-SIDEBAND MODULATION SIGNALS
WITH A CONTROLLED CARRIER LEVEL 261

H.V. Nguyen, N. Tan, N.H. Quan, T.T. Huong, N.H. Phat

BUILDING A CHATBOT SYSTEM TO ANALYZE OPINIONS OF ENGLISH
COMMENTS 289

A. Trofimov, F. Taubin

PERFORMANCE ANALYSIS OF CONCATENATED CODING TO
INCREASE THE ENDURANCE OF MULTILEVEL NAND FLASH MEMORY 316

Mathematical Modeling, Numerical Methods

D. Efanov, T. Pogodina

PROPERTIES INVESTIGATION OF SELF-DUAL COMBINATIONAL
DEVICES WITH CALCULATION CONTROL BASED ON HAMMING
CODES 349

Yu. Zelenkov

OPTIMIZATION OF THE REGRESSION ENSEMBLE SIZE 393

E. Semenenko, A. Belolipetskaya, R. Yuriev, A. Alodjants, I. Bessmertny,
I. Surov

DISCOVERY OF ECONOMIC COLLUSION BY METRICS OF QUANTUM
ENTANGLEMENT 416

E. Kulakov, A. Mikhalev, A. Sarenkov, A. Shutalev, A. Fedoreev

POSITION CORRECTION ALGORITHM OF WELL PADS WHEN SOLVING
THE PROBLEM OF DEVELOPING OIL FIELDS 447

СОДЕРЖАНИЕ

Цифровые информационно-телекоммуникационные технологии

М.В. Гофман, А.А. Корниенко.

ПРЕДЕЛЬНЫЕ БИПОЛЯРНЫЕ ПОСЛЕДОВАТЕЛЬНОСТИ ДЛЯ
РОБАСТНОГО МАРКИРОВАНИЯ ЦИФРОВЫХ АУДИОСИГНАЛОВ ПО
МЕТОДУ ЛОСКУТА 221

С.В. Дворников, С.С. Дворников, К.Д. Жеглов

ПОМЕХОУСТОЙЧИВОСТЬ СИГНАЛОВ ОДНОПОЛОСНОЙ
МОДУЛЯЦИИ С УПРАВЛЯЕМЫМ УРОВНЕМ НЕСУЩЕГО КОЛЕБАНИЯ 261

Х.В. Нгуен, Н. Тан, Н.Х. Куан, Ч.Т. Хыонг, Н.Х. Пхат

СОЗДАНИЕ СИСТЕМЫ ЧАТ-БОТОВ ДЛЯ АНАЛИЗА МНЕНИЙ
АНГЛОЯЗЫЧНЫХ КОММЕНТАРИЕВ 289

А.Н. Трофимов, Ф.А. Таубин

АНАЛИЗ ЭФФЕКТИВНОСТИ КАСКАДНОГО КОДИРОВАНИЯ ДЛЯ
ПОВЫШЕНИЯ ВЫНОСЛИВОСТИ МНОГОУРОВНЕВОЙ NAND ФЛЕШ-
ПАМЯТИ 316

Математическое моделирование и прикладная математика

Д.В. Ефанов, Т.С. Погодина

ИССЛЕДОВАНИЕ СВОЙСТВ САМОДВОЙСТВЕННЫХ
КОМБИНАЦИОННЫХ УСТРОЙСТВ С КОНТРОЛЕМ ВЫЧИСЛЕНИЙ НА
ОСНОВЕ КОДОВ ХЭММИНГА 349

Ю.А. Зеленков

ОПТИМИЗАЦИЯ РАЗМЕРА АНСАМБЛЯ РЕГРЕССОРОВ 393

Е.К. Семененко, А.Г. Белолипецкая, Р.Н. Юрьев, А.П. Алоджанц,

И.А. Бессмертный, И.А. Суров

ВЫЯВЛЕНИЕ ЭКОНОМИЧЕСКИХ СГОВОРОВ МЕТРИКАМИ
КВАНТОВОЙ ЗАПУТАННОСТИ 416

Е.Д. Кулаков, А.С. Михалев, А.В. Саренков, А.Д. Шуталев, А.Е. Федорев

АЛГОРИТМ КОРРЕКТИРОВКИ ПОЛОЖЕНИЯ КУСТОВЫХ ПЛОЩАДОК
ПРИ РЕШЕНИИ ЗАДАЧИ РАЗРАБОТКИ НЕФТЯНЫХ
МЕСТОРОЖДЕНИЙ 447

М.В. ГОФМАН, А.А. КОРНИЕНКО
**ПРЕДЕЛЬНЫЕ БИПОЛЯРНЫЕ ПОСЛЕДОВАТЕЛЬНОСТИ
ДЛЯ РОБАСТНОГО МАРКИРОВАНИЯ ЦИФРОВЫХ
АУДИОСИГНАЛОВ ПО МЕТОДУ ЛОСКУТА**

Гофман М.В., Корниенко А.А. Пределные биполярные последовательности для робастного маркирования цифровых аудиосигналов по методу лоскута.

Аннотация. Обеспечение устойчивости маркирования цифровых аудиосигналов в условиях действия помех, различных преобразований и возможных атак является актуальной проблемой. Одним из наиболее используемых и достаточно устойчивых методов маркирования является метод лоскута. Его робастность обеспечивается применением расширяющих биполярных числовых последовательностей при формировании и внедрении маркера в цифровой аудиосигнал и корреляционного детектирования при обнаружении и извлечении маркерной последовательности. Анализ свойств биполярных последовательностей, реализуемых в методе лоскута, показал, что абсолютные значения величины отношения максимума автокорреляционной функции (АКФ) к её минимуму для расширяющих биполярных последовательностей и расширенных маркерных последовательностей, используемых при традиционном маркировании, с высокой точностью приближаются к 2. Это позволило сформулировать критерии для поиска специальных расширяющих биполярных последовательностей, обладающих улучшенными корреляционными свойствами и большей устойчивостью. В статье разработан математический аппарат для поиска и построения предельных расширяющих биполярных последовательностей, используемых при решении задачи робастного маркирования цифровых аудиосигналов по методу лоскута. Пределные биполярные последовательности определены как последовательности, у которых автокорреляционные функции обладают максимально возможными по абсолютному значению отношениями максимума к минимуму. Сформулированы и доказаны теоремы и следствия из них: о существовании верхней границы минимальных значений автокорреляционных функций предельных биполярных последовательностей и о значениях первого и второго лепестков АКФ. На этой основе дано строгое математическое определение предельных биполярных последовательностей. Разработаны метод поиска полного множества предельных биполярных последовательностей на основе рационального перебора и метод построения предельных биполярных последовательностей произвольной длины с использованием порождающих функций. Представлены результаты компьютерного моделирования по оценке значений абсолютной величины отношения максимума к минимуму автокорреляционной и взаимной корреляционных функций исследуемых биполярных последовательностей для слепого приема. Показано, что предложенные предельные биполярные последовательности характеризуются лучшими корреляционными свойствами в сравнении с традиционно используемыми биполярными последовательностями и обладают большей устойчивостью.

Ключевые слова: стеганография, аудиосигнал, метод лоскута, маркирование цифровых объектов, биполярная последовательность, корреляционная функция.

1. Введение. Широкое распространение и использование цифрового аудиоконтента остро ставит вопрос о необходимости его стеганографической защиты для доказательства авторского права,

права собственности, возможности отслеживания сделок и оплаченной рекламы, контроля целостности и решения других задач [1, 2]. Это, в свою очередь, требует развития робастных методов цифрового маркирования и построения аудиостегосистем, устойчивых к атакам, действию помех и преобразованиям. Целью возможных атак может быть удаление, разрушение или нелегитимное извлечение внедренного цифрового маркера (цифрового водяного знака).

Существует достаточно большое количество робастных методов маркирования цифровых аудиосигналов, связанных с погружением цифровых водяных знаков или скрытием передаваемой аудиоинформации [1, 2]. Одним из наиболее используемых и достаточно устойчивых является метод лоскута (от англ. patchwork). Этот метод представляет собой статистический метод маркирования цифровых объектов [3, 4, 5]. Принцип маркирования методом лоскута состоит в следующем. В цифровом объекте в соответствии со стегоключом случайным образом выбирается пара одинаково распределенных элементов. Затем амплитуда одного элемента из каждой пары увеличивается на определенную величину, тогда как амплитуда другого элемента этой пары уменьшается на такую же величину. Путем такой модификации амплитуд можно закодировать, например, один бит внедряемого цифрового маркера (цифрового водяного знака). Этот процесс случайного выбора пар с последующим изменением амплитуд элементов повторяется достаточное количество раз для того, чтобы обеспечивалась статистическая устойчивость внедренного маркера к помехам и предполагаемым атакам на маркированный цифровой объект. Реализация метода лоскута основана на применении биполярных последовательностей.

Метод лоскута и принципы, лежащие в его основе, широко используются для маркирования цифровых аудиосигналов (речь и музыка), цифровых изображений и видео, а также в кодах компьютерных программ [6 – 16]. При этом робастность маркирования методом лоскута обеспечивается применением специальных биполярных последовательностей при формировании и внедрении маркера в покрывающий объект – цифровой аудиосигнал, и корреляционного детектирования при обнаружении и извлечении маркера (выделении информационных битов маркерной последовательности).

В последние годы получило широкое развитие маркирование цифровых аудиосигналов методом лоскута с использованием специальных расширяющих биполярных последовательностей для решения задач, связанных с использованием цифровых водяных

знаков, скрываем передаваемой аудиоинформации, позиционированием маркированными аудиосигналами нелегитимного приемного устройства. В статьях [17–21] изложен традиционный подход к маркированию цифровых аудиосигналов по методу лоскута в частотной области Фурье, когда маркированный аудиосигнал предполагается передавать через воздушный аудиоканал или радиоканал. В этом случае маркер формируется из псевдослучайной биполярной числовой последовательности с почти равномерным распределением элементов $(1 \text{ и } -1)$ и из расширяющей ее последовательности определенного вида. При этом корреляционные свойства расширенной маркерной последовательности практически не зависят от соответствующих свойств исходной последовательности, а определяются корреляционными свойствами расширяющей биполярной последовательности (РБП). Одно из свойств автокорреляционной функции (АКФ) расширяющей биполярной последовательности – абсолютное значение величины отношения максимума АКФ к её минимуму при этих условиях с высокой точностью приближается к 2. Это позволяет использовать данное свойство как критерий при выборе специальных РБП и построении на их основе специальных маркерных последовательностей с улучшенными корреляционными свойствами, повышающими качество детектирования и устойчивость маркирования цифровых аудиосигналов.

Другие особенности и свойства АКФ расширенной маркерной последовательности накладывают ограничения на выбор расширяющей последовательности при решении задачи позиционирования маркированными аудиосигналами приемного устройства. В работах [21, 22] показано, что для этой задачи наилучшая длина расширяющей последовательности равна двум разрядам. Увеличение длины расширяющей последовательности негативно сказывается на точности определения местоположения нелегитимного приемника с помощью приема и передачи маркированных аудиосигналов.

В статьях [23, 24] предлагается метод маркирования аудиосигналов с использованием выбранных из эвристических соображений расширяющих биполярных последовательностей с улучшенными корреляционными свойствами, повышающими качество детектирования, которым свойственно большее, чем два, по абсолютному значению отношение максимума ее АКФ к минимуму этой функции. Повышение качества детектирования происходит вследствие того, что увеличение абсолютного значения этой величины

обычно приводит к увеличению устойчивости к атакам, и, в частности, к атаке аддитивным белым гауссовским шумом при использовании корреляционного детектора.

В связи с этим представляется актуальной задача разработки математического аппарата построения таких специальных биполярных последовательностей с улучшенными корреляционными свойствами, у которых по абсолютному значению отношение максимума автокорреляционной функции к её минимуму будет максимально возможным.

Целью данной работы является разработка методов поиска и построения нового класса последовательностей – предельных расширяющих биполярных последовательностей, у которых автокорреляционные функции обладают максимально возможными по абсолютному значению отношениями максимума к минимуму, и применение которых повышает устойчивость цифрового маркирования.

Для достижения этой цели:

- Исследованы особенности АКФ расширяющих биполярных последовательностей и сформированных на их основе маркерных последовательностей, используемых при традиционном маркировании цифровых аудиосигналов. Предложены специальные функции для построения РБП с улучшенными корреляционными свойствами для робастного маркирования и предельных биполярных последовательностей, АКФ которых обладают максимально возможными по абсолютному значению отношениями максимума к минимуму. Сформулированы и доказаны теоремы и следствия из них о существовании верхней границы минимальных значений автокорреляционных функций таких последовательностей и о значениях первого и второго лепестков АКФ. На этой основе дано строгое математическое определение предельных биполярных последовательностей.

- Разработаны два метода поиска и построения предельных биполярных последовательностей: с использованием рационального перебора и на основе введенных порождающих функций, функционально связанных с такими последовательностями.

- Проведено компьютерное моделирование по оценке значений абсолютной величины отношения максимума к минимуму автокорреляционной и взаимной корреляционных функций исследуемых биполярных последовательностей для наихудших условий извлечения маркерных последовательностей (слепой прием).

В этом случае искажающим воздействием (шумом) становился сам покрывающий цифровой аудиосигнал.

Теоретический вклад настоящей работы в решение задачи робастного маркирования цифровых аудиосигналов по методу лоскута заключается в разработке математического аппарата для построения предельных расширяющих биполярных последовательностей с улучшенными корреляционными свойствами, более устойчивых к действию помех и искажающих воздействий при обнаружении и извлечении цифровой маркерной последовательности.

2. Традиционное маркирование с использованием расширяющих биполярных последовательностей. Традиционно маркер формируется из псевдослучайной биполярной числовой последовательности:

$$\mathbf{a} = (\alpha_1, \alpha_2, \dots, \alpha_{N_a}), \quad \alpha_i \in \{1, -1\}, \quad (1)$$

где N_a – длина последовательности. После формирования последовательности $\mathbf{a}$ её элементы расширяются путем подстановки вместо них последовательности следующего вида:

$$\gamma = \left(\underbrace{1, \dots, 1}_{N_n}, \underbrace{-1, \dots, -1}_{N_n} \right), \quad (2)$$

где N_n – положительное целое число, определяющее половину длины расширяющей последовательности.

Затем формируется маркерная последовательность:

$$\mathbf{m} = (m_1, m_2, \dots, m_{2N_n N_a}) = \mathbf{a} \otimes \gamma, \quad (3)$$

где $\otimes$ – оператор кронекерова произведения,

$$m_i = (-1)^{\left\lfloor \frac{i-1}{N_n} \right\rfloor \bmod 2} \alpha_{\left\lfloor \frac{i-1}{2N_n} \right\rfloor + 1}, \quad (4)$$

где $\lfloor x \rfloor$ – ближайшее целое, меньшее или равное числу x .

Прежде чем, перейти к анализу корреляционных свойств маркерной последовательности, определим АКФ произвольной вещественной последовательности:

$$\mathbf{p} = (p_1, p_2, \dots, p_L), \quad (5)$$

как дискретную функцию от целочисленной переменной k , удовлетворяющую следующему равенству:

$$f_{\text{АКФ}}(\mathbf{p}, k) = \sum_{i=1}^L p_i p_{i-k}, \quad (6)$$

где числа p_l являются элементами вещественной последовательности $\mathbf{p}$, когда индекс $l \in \{1, 2, \dots, L\}$, но для сокращения обозначений будем считать $p_l = 0$, когда индекс $l \notin \{1, 2, \dots, L\}$.

Использование расширяющих последовательностей γ вида, описываемого равенством (2), приводит к тому, что у АКФ сформированной маркерной последовательности $\mathbf{m}$ в соответствии с (3) появляются большие по амплитуде отрицательные значения. Это уменьшает абсолютное значение отношения максимума АКФ к её минимуму по сравнению с таким же значением АКФ для биполярной последовательности $\mathbf{a}$. При этом абсолютное значение этого отношения оказывается приблизительно равным два в тех случаях, когда исходная псевдослучайная последовательность $\mathbf{a}$ генерируется таким образом, что вероятности появления 1 и -1 оказываются почти равными.

Это свойство АКФ биполярной последовательности, используемой при традиционном маркировании, можно обосновать с помощью следующих рассуждений. Когда в псевдослучайных биполярных последовательностях $\mathbf{a}$ вероятности появления 1 и -1 приблизительно одинаковые, тогда можно получить оценку максимального падения АКФ маркерной последовательности $\mathbf{m}$, опираясь на среднее значение АКФ маркерной последовательности $\mathbf{m}$, вычисленное для единичного сдвига от центрального значения АКФ маркерной последовательности $\mathbf{m}$. Действительно, значение АКФ маркерной последовательности $\mathbf{m}$, получаемое при единичном сдвиге маркерной последовательности $\mathbf{m}$ относительно центрального значения АКФ, включает в себя все подстановочные (расширяющие)

последовательности $(1, -1)$, значит, падение значения АКФ окажется максимальным. Среднее значение на отдельную подстановку (в зависимости от результатов соседних подстановок) можно вычислить, исходя из значений для каждого из возможных вариантов, составляющих значения АКФ. С учетом симметрии АКФ возможны два варианта составляющих значения, как показано на рисунке 1. Без учета – четыре варианта составляющих, два из которых будут давать такой же вклад в значение АКФ, какой дают элементы другой пары составляющих.

$$\begin{array}{cc}
 \begin{array}{c} \times \\ \hline \end{array} \begin{array}{c} (\dots, 1, -1, \dots) \\ (\dots, -1, 1, \dots) \end{array} & \begin{array}{c} \times \\ \hline \end{array} \begin{array}{c} (\dots, 1, -1, \dots) \\ (\dots, 1, 1, \dots) \end{array} \\
 \begin{array}{c} \dots, \underbrace{-1, -1}, \dots \\ \text{вклад в} \\ \text{значение} \\ \text{АКФ} \\ \text{равен } -2 \end{array} & \begin{array}{c} \dots, \underbrace{1, -1}, \dots \\ \text{вклад в} \\ \text{значение} \\ \text{АКФ} \\ \text{равен } 0 \end{array}
 \end{array}$$

Рис. 1. Составляющие значения АКФ

Так как таких составляющих будет в среднем одинаковое количество, то, исходя из них, среднее арифметическое значение оказывается равным единице со знаком минус, что указывает на падение центрального значения АКФ в 2 раза. То есть, если центральное значение АКФ псевдослучайной последовательности после подстановки расширяющей последовательности $(1, -1)$, равно X , то минимальное значение в среднем будет равно $-X/2$.

Например, последовательность Баркера длиной 7:

$$\alpha = (1, 1, 1, -1, -1, 1, -1), \quad (7)$$

после подстановки последовательности $\gamma = (1, -1)$, станет маркерной последовательностью:

$$\mathbf{m} = (1, -1, 1, -1, 1, -1, -1, 1, -1, 1, 1, -1, -1, 1), \quad (8)$$

у которой максимум АКФ равен 14, а минимум -7 и, следовательно, по абсолютной величине их отношение равно $|14 / (-7)| = 2$.

Исходя из этих рассуждений и результатов проведенного моделирования, можно сделать следующие выводы относительно особенностей и свойств АКФ маркерной (расширенной) последовательности, сформированной в соответствии с равенствами (1)–(3). При традиционном маркировании значение этой АКФ оказывается минимальным на «расстоянии» одного сдвига относительно центрального значения, что связано с длиной расширяющей последовательности $(1, -1)$. Максимальное значение АКФ маркерной последовательности m зависит от длины расширяемой псевдослучайной последовательности α и от длины расширяющей последовательности $(1, -1)$. Тогда как минимальное значение АКФ маркерной последовательности m тем ближе по абсолютному значению к половине центрального значения, чем больше длина псевдослучайной последовательности α и чем более равновероятным является появление 1 и -1 в ней. При выполнении последнего условия, когда в качестве расширяемой последовательности α используется псевдослучайная биполярная последовательность, у которой вероятности появления 1 и -1 почти равны, АКФ расширенной последовательности m будет определяться корреляционными свойствами расширяющей последовательности γ , а не корреляционными свойствами расширяемой последовательности α . И еще одно важнейшее свойство АКФ – абсолютное значение величины отношения максимума АКФ к её минимуму с высокой точностью приближается к 2. Это позволяет сформулировать критерии для поиска расширяющих биполярных последовательностей γ , обладающих улучшенными корреляционными свойствами и большей устойчивостью.

С учетом выявленных выше свойств АКФ расширенной биполярной последовательности и в развитие статей [23, 24] далее предлагается метод построения специальных расширяющих биполярных последовательностей для робастного маркирования, позволяющий обеспечить большее, чем 2 по абсолютному значению отношение максимума АКФ маркерной последовательности m к минимуму этой функции.

3. Предельные биполярные последовательности для робастного маркирования. Если вместо традиционной расширяющей последовательности $\gamma = (1, -1)$ в равенстве (3) использовать расширяющую последовательность вида:

$$\gamma = \Phi \otimes \gamma_{\text{трад}}, \quad (9)$$

где:

$$\gamma_{\text{трад}} = (1, -1). \quad (10)$$

При этом, когда для формирования γ применяется специально выбираемая биполярная последовательность:

$$\Phi = (\varphi_1, \varphi_2, \dots, \varphi_{N_\Phi}), \quad \varphi_i \in \{1, -1\}, \quad (11)$$

то можно обеспечить превышающее 2 значение абсолютной величины отношения максимума АКФ к её минимуму у маркерной последовательности $\mathbf{m}$.

Использование биполярной последовательности γ вида (9), вместо биполярной последовательности вида (2), приводит к тому, что маркерная последовательность $\mathbf{m}$ уже не будет определяться уравнением (3), а будет с учетом (1) удовлетворять следующему равенству:

$$\mathbf{m} = (m_1, m_2, \dots, m_{2N_\Phi N_\alpha}) = \alpha \otimes \gamma, \quad (12)$$

где:

$$m_i = (-1)^{(i-1) \bmod 2} \varphi_{\left(\left\lfloor \frac{i-1}{2} \right\rfloor \bmod N_\Phi\right)+1} \alpha_{\left\lfloor \frac{i-1}{2N_\Phi} \right\rfloor+1}. \quad (13)$$

Например, если:

$$\Phi = (-1, -1, 1, -1), \quad (14)$$

то получится расширяющая последовательность:

$$\gamma = (-1, 1, -1, 1, 1, -1, -1, 1), \quad (15)$$

у которой максимум АКФ равен 8, а минимум -3 , то есть, в этом случае у последовательности γ значение абсолютной величины отношения максимума АКФ к её минимуму равно $|-8/3| = 2.667$.

Выполнив подстановку этой последовательности γ вместо элементов последовательности Баркера длины 7:

$$\alpha = (1, 1, 1, -1, -1, 1, -1), \quad (16)$$

будет получена биполярная последовательность длиной 56:

$$\begin{aligned} \mathbf{m} = & (-1, 1, -1, 1, 1, -1, -1, 1, -1, 1, -1, 1, 1, -1, -1, 1, \\ & -1, 1, -1, 1, 1, -1, -1, 1, -1, 1, -1, 1, 1, -1, \\ & 1, -1, 1, -1, -1, 1, 1, -1, -1, 1, -1, 1, 1, -1, \\ & 1, 1, -1, 1, -1, -1, 1, 1, -1), \end{aligned} \quad (17)$$

у которой максимум АКФ равен 56, а минимум -21 (рисунок 2). Следовательно, у последовательности $\mathbf{m}$ абсолютное значение отношения максимума АКФ к минимуму этой функции равно $|-56/21| = 2.667$, что больше на 33%, чем в случае, если бы вместо последовательности (15) использовалась последовательность $\gamma_{\text{трад}}$, применяемая в традиционном маркировании (рисунок 3). Напомним, что при традиционном подходе значение этого отношения с высокой точностью приближается к 2.

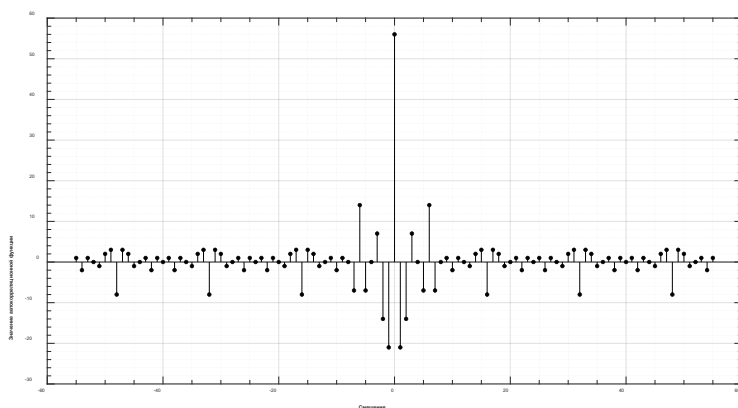


Рис. 2. График автокорреляционной функции маркерной последовательности (17), основанной на специальной расширяющей последовательности

$$\gamma = (1, -1, -1, 1, 1, -1, -1, 1).$$

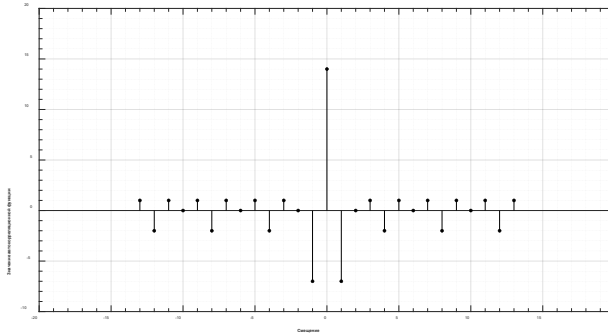


Рис. 3. График автокорреляционной функции маркерной последовательности (8), основанной на традиционной расширяющей последовательности

$$\gamma_{\text{трад}} = (1, -1).$$

Таким образом, с учетом того, что расширяющая последовательность γ в основном определяет автокорреляционные свойства расширенной последовательности $\mathbf{m}$, когда биполярная последовательность $\mathbf{a}$ является псевдослучайной, то требуется метод получения таких последовательности γ , у которых автокорреляционные функции имеют наибольшие минимумы среди минимумов автокорреляционных функций всех прочих удовлетворяющих равенству (9) биполярных последовательностей такой же длины.

По определению значение АКФ $f_{\text{АКФ}}(\gamma, k)$ последовательности γ , описываемой формулой (9), когда смещение k принадлежит множеству:

$$\{-2N_{\varphi} + 1, -2N_{\varphi} + 2, \dots, -1, 0, 1, 2, \dots, 2N_{\varphi} - 1\}, \quad (18)$$

удовлетворяет следующему равенству:

$$\begin{aligned} f_{\text{АКФ}}(\gamma, k) &= \sum_{i=1}^{2N_{\varphi}-|k|} (-1)^{i-1} \varphi_{\left[\frac{i-1}{2}\right]+1} (-1)^{i+|k|-1} \varphi_{\left[\frac{i+|k|-1}{2}\right]+1} = \\ &= \sum_{i=1}^{2N_{\varphi}-|k|} (-1)^{2i+|k|-2} \varphi_{\left[\frac{i-1}{2}\right]+1} \varphi_{\left[\frac{i+|k|-1}{2}\right]+1}, \end{aligned} \quad (19)$$

где $|x|$ – абсолютное значение числа x , а $\lfloor x \rfloor$ – ближайшее целое, меньшее или равное числу x .

Полезно эту АКФ представить в следующем виде:

$$f_{\text{АКФ}}(\gamma, k) = \begin{cases} 2z_{\frac{k}{2}}, & \text{если } k \bmod 2 = 0, \\ -z_{\frac{k-1}{2}} - z_{\left[\frac{k-2}{2}\right]+1}, & \text{если } k \bmod 2 = 1 \text{ и } k \neq 2N_{\varphi} - 1, \\ -z_{\frac{k-1}{2}}, & \text{если } k = 2N_{\varphi} - 1, \end{cases} \quad (20)$$

где:

$$z_j = z_{-j} = \sum_{i=1}^{N_{\varphi}-j} \varphi_i \varphi_{i+j}, \quad j \in \{0, 1, \dots, N_{\varphi} - 1\}. \quad (21)$$

При этом, учитывая, что $\varphi_i \in \{1, -1\}$, будет выполняться следующее равенство:

$$z_0 = N_{\varphi}. \quad (22)$$

Так как конкретные значения z_j и z_{-j} (при $j \in \{1, 2, \dots, N_{\varphi} - 1\}$) будут определяться элементами последовательности φ , использованной при формировании последовательности γ , то для z_j и z_{-j} , когда $j \in \{1, 2, \dots, N_{\varphi} - 1\}$, можно указать лишь множества потенциальных значений:

$$z_j, z_{-j} \in \{\pm(a + j - 1), \pm(a + j + 1), \pm(a + j + 3), \dots, \pm(N_{\varphi} - j)\}, \quad (23)$$

где:

$$a = (N_{\varphi} - 1) \bmod 2. \quad (24)$$

Следовательно, количество различных автокорреляционных функций, соответствующих одному и тому же N_{φ} , меньше количества последовательностей γ . Поэтому всё множество последовательностей γ , соответствующих одному и тому же N_{φ} , можно поделить

на непересекающиеся подмножества, в которые попадут последовательности, имеющие одинаковые автокорреляционные функции.

Теорема 1 (О существовании верхней границы для минимального значения АКФ биполярной последовательности γ).

Среди всех биполярных последовательностей Φ , определенных равенством (11) и имеющих одинаковую длину N_Φ , нет такой, которая позволяет получить в соответствии с правилом (9) биполярную последовательность γ , у которой АКФ имеет минимум, который больше величины:

$$\left\lfloor \frac{N_\Phi - 1}{3} \right\rfloor - N_\Phi, \quad (25)$$

то есть для всех таких биполярных последовательностей γ одинаковой длины справедливо неравенство:

$$\min_{\forall k} f_{\text{АКФ}}(\gamma, k) \leq \left\lfloor \frac{N_\Phi - 1}{3} \right\rfloor - N_\Phi. \quad (26)$$

Доказательство. Так как для $N_\Phi = 1$ условие этой теоремы выполняется, поэтому ниже предлагается доказательство для $N_\Phi > 1$.

Предположим обратное – а именно, что среди биполярных последовательностей Φ , определяемых равенством (11), существует такая последовательность, обозначим ее $\hat{\Phi}$, из которой в соответствии с правилом (9) получается биполярная последовательность $\hat{\gamma}$, у которой АКФ имеет минимальное значение, которое больше величины (25). В таком случае автокорреляционная функция $f_{\text{АКФ}}(\hat{\gamma}, k)$ (19), (20) последовательности $\hat{\gamma}$, имея максимум, равный $2z_0 = 2N_\Phi$, имеет минимум, который больше¹:

¹ Так, например, среди последовательностей Φ , длиной $N_\Phi = 4$, должны были бы существовать такие $\hat{\Phi}$, из которых получались бы последовательности $\hat{\gamma}$, у которых минимальное значение автокорреляционной функции было бы больше -3, то есть равнялось бы -2 или -1.

$$\left\lfloor \frac{N_{\varphi} - 1}{3} \right\rfloor - N_{\varphi}. \quad (27)$$

Значит, для этой АКФ обязательно выполняются следующие два неравенства:

$$f_{\text{АКФ}}(\hat{\gamma}, 1) = -z_0 - z_1 > \left\lfloor \frac{N_{\varphi} - 1}{3} \right\rfloor - N_{\varphi}, \quad (28)$$

$$f_{\text{АКФ}}(\hat{\gamma}, 2) = 2z_1 > \left\lfloor \frac{N_{\varphi} - 1}{3} \right\rfloor - N_{\varphi}. \quad (29)$$

Исходя из того, что $z_0 = N_{\varphi}$, неравенство (28) позволяет заключить, что z_1 должно удовлетворять следующему неравенству:

$$z_1 < -\left\lfloor \frac{N_{\varphi} - 1}{3} \right\rfloor. \quad (30)$$

Тогда как на основании неравенства (29) получается, что:

$$z_1 > \frac{\left\lfloor \frac{N_{\varphi} - 1}{3} \right\rfloor - N_{\varphi}}{2}. \quad (31)$$

Следовательно, величина z_1 должна принадлежать диапазону:

$$\frac{\left\lfloor \frac{N_{\varphi} - 1}{3} \right\rfloor - N_{\varphi}}{2} < z_1 < -\left\lfloor \frac{N_{\varphi} - 1}{3} \right\rfloor. \quad (32)$$

Нетрудно показать, что расстояние между правой и левой границами этого диапазона удовлетворяет равенству:

$$-\left\lfloor \frac{N_{\varphi}-1}{3} \right\rfloor - \frac{\left\lfloor \frac{N_{\varphi}-1}{3} \right\rfloor - N_{\varphi}}{2} = \frac{1+3\left\{ \frac{N_{\varphi}-1}{3} \right\}}{2},$$

где $\{a\}$ – это дробная часть вещественного числа a . Так как в данном случае выполняется равенство:

$$3\left\{ \frac{N_{\varphi}-1}{3} \right\} = (N_{\varphi} - 1) \bmod 3, \quad (33)$$

то,

$$\frac{1}{2} \leq \frac{1+3\left\{ \frac{N_{\varphi}-1}{3} \right\}}{2} \leq \frac{3}{2}. \quad (34)$$

Таким образом, возможны три ситуации:

$$\frac{1+3\left\{ \frac{N_{\varphi}-1}{3} \right\}}{2} = \begin{cases} 1/2, \text{ если } N_{\varphi} \in \{1, 4, 7, \dots\}, \\ 1, \text{ если } N_{\varphi} \in \{2, 5, 8, \dots\}, \\ 3/2, \text{ если } N_{\varphi} \in \{3, 6, 9, \dots\}. \end{cases} \quad (35)$$

В первом и во втором случае, когда расстояние между границами неравенства (32) равно $1/2$ и 1 , соответственно, не существует подходящих целочисленных значений z_1 . Однако, в третьем случае, расстояние между границами неравенства (32) становится наибольшим и оказывается равным $3/2$. Например, когда

$$N_{\varphi} = 9, \text{ тогда } -\frac{7}{2} < z_1 < -2, \text{ или, когда } N_{\varphi} = 12, \text{ тогда } -\frac{9}{2} < z_1 < -3.$$

Можно предположить, что, когда $N_{\varphi} \in \{3, 6, 9, \dots\}$, подходящим значением для z_1 является:

$$\left\lceil \frac{\left\lfloor \frac{N_{\varphi} - 1}{3} \right\rfloor - N_{\varphi}}{2} \right\rceil, \quad (36)$$

где $\lceil a \rceil$ – ближайшее целое, большее или равное вещественному числу a . Но, в тех случаях, когда число N_{φ} четно, тогда и величина (36) является четной, однако, в таком случае z_1 не может принимать четные значения (23). И наоборот, когда величина (36) нечетная, тогда и число N_{φ} также является нечетным, но в таком случае z_1 не может принимать нечетные значения. Следовательно, во всех этих случаях z_1 не может принять значение равное величине (36).

Таким образом, среди биполярных последовательностей φ , определяемых равенством (11), нет такой последовательности $\hat{\varphi}$, из которой в соответствии с правилом (9) получается биполярная последовательность $\hat{\gamma}$, у которой АКФ имеет минимум, больше величины (25).

Следствие из Теоремы 1 (О предельной величине верхней границы отношения по абсолютному значению максимума АКФ биполярной последовательности γ к минимуму этой функции). В пределе с увеличением длины N_{φ} последовательности φ , абсолютное значение отношения максимума к минимуму АКФ последовательности γ равно 3.

Действительно, проведя преобразования, получим:

$$\lim_{N_{\varphi} \rightarrow \infty} \left| \frac{2N_{\varphi}}{\left\lfloor \frac{N_{\varphi} - 1}{3} \right\rfloor - N_{\varphi}} \right| = \lim_{N_{\varphi} \rightarrow \infty} \left| \frac{6N_{\varphi}}{-2N_{\varphi} - 1 - 3 \left\{ \frac{N_{\varphi} - 1}{3} \right\}} \right| = 3, \quad (37)$$

где $\{a\}$ – это дробная часть числа a .

Определение. Предельной биполярной последовательностью γ , удовлетворяющей равенству (9), называется такая (обозначим

ее γ_∞), у которой АКФ имеет минимальное численное значение, равное величине (25):

$$\min_{\forall k} f_{\text{АКФ}}(\gamma_\infty, k) = \left\lfloor \frac{N_\varphi - 1}{3} \right\rfloor - N_\varphi. \quad (38)$$

Исходя из равенства (38) можно заключить, что у предельных последовательностей γ_∞ автокорреляционные функции такие, что для всех k из множества (18) выполняются неравенства:

$$f_{\text{АКФ}}(\gamma_\infty, k) \geq \left\lfloor \frac{N_\varphi - 1}{3} \right\rfloor - N_\varphi, \quad (39)$$

которые, с учетом равенства (20) и $z_0 = N_\varphi$, можно свести к следующим (когда $N_\varphi > 1$):

$$\frac{\left\lfloor \frac{N_\varphi - 1}{3} \right\rfloor - N_\varphi}{2} \leq z_j \leq N_\varphi - \left\lfloor \frac{N_\varphi - 1}{3} \right\rfloor - z_{j-1}, \quad (40)$$

где $j \in \{1, \dots, N_\varphi - 1\}$.

Автокорреляционные функции всех предельных биполярных последовательностей γ_∞ обладают общим свойством, доказательство которого приведено далее в Теореме 2.

Теорема 2 (о значениях первого и второго лепестков АКФ предельной биполярной последовательности). АКФ предельной последовательности γ_∞ является такой, что для неё выполняется либо равенство:

$$f_{\text{АКФ}}(\gamma_\infty, 1) = \left\lfloor \frac{N_\varphi - 1}{3} \right\rfloor - N_\varphi, \quad (41)$$

либо равенство:

$$f_{\text{АКФ}}(\gamma_{\infty}, 2) = \left\lfloor \frac{N_{\varphi} - 1}{3} \right\rfloor - N_{\varphi}, \quad (42)$$

но не оба сразу.

Доказательство. Так как для $N_{\varphi} = 1$ условие этой теоремы выполняется, поэтому ниже предлагается доказательство для $N_{\varphi} > 1$.

Предположим обратное, что для любой предельной последовательности γ_{∞} при выполнении:

$$\min_{\forall k} f_{\text{АКФ}}(\gamma_{\infty}, k) = \left\lfloor \frac{N_{\varphi} - 1}{3} \right\rfloor - N_{\varphi}, \quad (43)$$

для всех $k \in \{3, 4, \dots, 2N_{\varphi} - 1\}$ выполняется неравенство:

$$f_{\text{АКФ}}(\gamma_{\infty}, k) \geq \left\lfloor \frac{N_{\varphi} - 1}{3} \right\rfloor - N_{\varphi}, \quad (44)$$

но для всех $k \in \{1, 2\}$ выполняется либо неравенство:

$$f_{\text{АКФ}}(\gamma_{\infty}, k) > \left\lfloor \frac{N_{\varphi} - 1}{3} \right\rfloor - N_{\varphi}, \quad (45)$$

либо равенство:

$$f_{\text{АКФ}}(\gamma_{\infty}, k) = \left\lfloor \frac{N_{\varphi} - 1}{3} \right\rfloor - N_{\varphi}. \quad (46)$$

Если для всех $k \in \{1, 2\}$ выполняется неравенство (45), то в таком случае z_1 должно удовлетворять двойному неравенству (32), для которого при доказательстве Теоремы 1 было показано, что не существует такого z_1 , который удовлетворял бы ему.

Теперь, если будет опровергнуто равенство:

$$f_{\text{АКФ}}(\gamma_{\infty}, 1) = f_{\text{АКФ}}(\gamma_{\infty}, 2), \quad (47)$$

то теорема будет доказана. Предположим обратное, что равенство (47) выполняется, тогда выполняется и равенство:

$$-z_0 - z_1 = 2z_1. \quad (48)$$

В таком случае, с учетом равенства $z_0 = N_{\varphi}$, будет:

$$z_1 = -\frac{N_{\varphi}}{3}. \quad (49)$$

Так как z_1 является целым числом, то это равенство не может выполняться для значений N_{φ} не кратных 3. Тогда, рассмотрим только такие значения N_{φ} , которые кратны 3 – в таких случаях отношение справа от знака равенства (49) будет целым числом. Во-первых, это отношение будет равно нечетному числу, когда N_{φ} равно нечетному числу, кратному 3. Во-вторых, это отношение будет равно четному числу, когда N_{φ} равно четному числу, кратному 3. Но при нечетном N_{φ} число z_1 может быть только четным (23), и, наоборот, – при четном N_{φ} число z_1 может быть только нечетным. Значит, равенство $f_{\text{АКФ}}(\gamma_{\infty}, 1) = f_{\text{АКФ}}(\gamma_{\infty}, 2)$ также не выполняется.

Следствие из Теоремы 2. У предельных последовательностей γ_{∞} , когда $N_{\varphi} > 1$, автокорреляционные функции такие, что:

$$z_1 = \begin{cases} -\left\lfloor \frac{N_{\varphi}-1}{3} \right\rfloor - 1, & \text{если } (N_{\varphi} + 1) \bmod 3 = 0, \\ -\left\lfloor \frac{N_{\varphi}-1}{3} \right\rfloor, & \text{иначе.} \end{cases} \quad (50)$$

Доказательство. Когда для АКФ предельной последовательности γ_{∞} выполняется равенство:

$$f_{\text{АКФ}}(\gamma_{\infty}, 1) = \left\lfloor \frac{N_{\varphi} - 1}{3} \right\rfloor - N_{\varphi}, \quad (51)$$

тогда:

$$-z_0 - z_1 = \left\lfloor \frac{N_{\varphi} - 1}{3} \right\rfloor - N_{\varphi}. \quad (52)$$

В таком случае, с учетом равенства $z_0 = N_{\varphi}$, будет выполняться равенство:

$$z_1 = - \left\lfloor \frac{N_{\varphi} - 1}{3} \right\rfloor. \quad (53)$$

Однако это равенство, учитывая свойство (23), может выполняться только, когда четность числа $N_{\varphi} - 1$ совпадает с четностью числа

$\left\lfloor \frac{N_{\varphi} - 1}{3} \right\rfloor$, то есть равенство (53) выполняется только для тех чисел

N_{φ} , для которых выполняется равенство:

$$(N_{\varphi} - 1) \bmod 2 = \left\lfloor \frac{N_{\varphi} - 1}{3} \right\rfloor \bmod 2. \quad (54)$$

Таковыми являются числа $N_{\varphi} \in \{1, 3, 4, 6, 7, 9, 10, 12, 13, \dots\}$. Этим доказана нижняя половина кусочно-определенного равенства (50).

Теперь докажем верхнюю половину (50). Когда для АКФ предельной последовательности γ_{∞} выполняется равенство:

$$f_{\text{АКФ}}(\gamma_{\infty}, 2) = \left\lfloor \frac{N_{\varphi} - 1}{3} \right\rfloor - N_{\varphi}, \quad (55)$$

тогда:

$$z_1 = \frac{\left\lfloor \frac{N_\varphi - 1}{3} \right\rfloor - N_\varphi}{2}. \quad (56)$$

С одной стороны, правая часть равенства (56) принимает целые значения, только когда $N_\varphi \in \{2, 5, 8, 11, 14, 17, 20, \dots\}$, то есть когда:

$$(N_\varphi + 1) \bmod 3 = 0. \quad (57)$$

С другой стороны, учитывая свойство (23), равенство (56) может выполняться, только когда четность числа $N_\varphi - 1$ совпадает с четностью правой части равенства (56) числа, то есть для таких положительных целых чисел N_φ , для которых выполняется равенство:

$$(N_\varphi - 1) \bmod 2 = \left| \frac{\left\lfloor \frac{N_\varphi - 1}{3} \right\rfloor - N_\varphi}{2} \right| \bmod 2. \quad (58)$$

Но это равенство как раз и выполняется, когда N_φ удовлетворяет равенству (57). При этом, когда N_φ удовлетворяет равенству (57), тогда выполняется равенство:

$$\frac{\left\lfloor \frac{N_\varphi - 1}{3} \right\rfloor - N_\varphi}{2} = - \left\lfloor \frac{N_\varphi - 1}{3} \right\rfloor - 1. \quad (59)$$

Действительно, если вычесть из левой части этого равенства правую, то получается следующее равенство:

$$\frac{\left\lfloor \frac{N_\varphi - 1}{3} \right\rfloor - N_\varphi}{2} + \left\lfloor \frac{N_\varphi - 1}{3} \right\rfloor + 1 = \frac{1 - 3 \left\lfloor \frac{N_\varphi - 1}{3} \right\rfloor}{2}, \quad (60)$$

где $\{a\}$ – это дробная часть числа a . Но,

$$3 \left\{ \frac{N_\varphi - 1}{3} \right\} = 1, \quad (61)$$

когда N_φ удовлетворяет равенству (57). Значит, когда N_φ удовлетворяет равенству (57), будет выполняться равенство:

$$\frac{1 - 3 \left\{ \frac{N_\varphi - 1}{3} \right\}}{2} = 0. \quad (62)$$

Следовательно, когда N_φ удовлетворяет равенству (57), тогда выполняется равенство (59). Тем самым доказана верхняя половина кусочно-определенного равенства (50).

Следует отметить, что с формальной точки зрения, величина z_1 не определена, когда $N_\varphi = 1$ (21). Но если принять, что $z_1 = 0$, когда $N_\varphi = 1$, то (50) будет выполняться и для случая $N_\varphi = 1$.

4. Метод поиска полного множества предельных последовательностей на основе рационального выбора. Равенство (50) с учетом равенства (21) можно представить в следующем виде:

$$\sum_{i=1}^{N_\varphi-1} \varphi_i \varphi_{i+1} = \begin{cases} - \left\lfloor \frac{N_\varphi - 1}{3} \right\rfloor - 1, & \text{если } (N_\varphi + 1) \bmod 3 = 0, \\ - \left\lfloor \frac{N_\varphi - 1}{3} \right\rfloor, & \text{иначе.} \end{cases} \quad (63)$$

Поэтому, опираясь на Следствие из Теоремы 2, можно сказать, что при поиске предельных последовательностей γ_∞ можно ограничиться только множеством Φ из таких биполярных последовательностей φ длиной N_φ , для которых выполняется кусочно-определенное равенство (63).

С учетом равенства (21) числа z_j можно определить как сумму элементов последовательности следующего вида:

$$\mathbf{v}_j = (v_{1,j}, v_{2,j}, \dots, v_{N_\varphi-j,j}), \quad (64)$$

где:

$$v_{i,j} = \varphi_i \varphi_{i+j}. \quad (65)$$

Так как $\varphi_i \in \{1, -1\}$, то $v_{i,j} \in \{1, -1\}$.

Важно отметить, что элементы последовательности $\mathbf{v}_j$, когда $j \geq 2$, можно вычислить по элементам двух предыдущих последовательностей $\mathbf{v}_{j-1}$ и $\mathbf{v}_{j-2}$ следующим образом:

$$v_{i,j} = v_{i,j-1} v_{i+1,j-2} v_{i+1,j-1}. \quad (66)$$

Действительно, произведение справа от знака этого равенства совпадает с правой частью равенства (65):

$$v_{i,j-1} v_{i+1,j-2} v_{i+1,j-1} = \varphi_i \varphi_{i+j-1} \varphi_{i+1} \varphi_{i+1+j-2} \varphi_{i+1} \varphi_{i+1+j-1} = \varphi_i \varphi_{i+j-1}^2 \varphi_{i+1}^2 \varphi_{i+j} = \varphi_i \varphi_{i+j}.$$

Любой вектор $\mathbf{v}_1$ связан с двумя биполярными последовательностями $\boldsymbol{\varphi}$. При этом эти две последовательности определяют две последовательности $\boldsymbol{\gamma}$, у которых совпадают автокорреляционные функции. Так, когда известен вектор $\mathbf{v}_1$, тогда восстановить элементы упомянутых двух биполярных последовательностей $\boldsymbol{\varphi}$, с которыми он связан, можно по следующему правилу:

- 1) Прямая последовательность $\boldsymbol{\varphi}$: пусть $\varphi_1 = -1$, тогда:

$$\varphi_i = -\prod_{l=1}^{i-1} v_{l,1}, \quad i \in \{2, 3, \dots, N_\varphi\}. \quad (67)$$

- 2) Инвертированная прямая последовательность $\boldsymbol{\varphi}$: пусть $\varphi_1 = 1$, тогда:

$$\varphi_i = \prod_{l=1}^{i-1} v_{l,1}, \quad i \in \{2, 3, \dots, N_\varphi\}. \quad (68)$$

Таким образом, множество Φ соответствует множеству V_1 из таких последовательностей $\mathbf{v}_1$, у которых сумма элементов равна числу z_1 , удовлетворяющему равенству (50). А так как $v_{i,j} \in \{1, -1\}$, то мощность множества V_1 равна биномиальному коэффициенту:

$$|V_1| = C_{N_\Phi - 1}^{\frac{N_\Phi - 1 + z_1}{2}}, \quad (69)$$

где:

$$C_{N_\Phi - 1}^{\frac{N_\Phi - 1 + z_1}{2}} = \frac{(N_\Phi - 1)!}{\left(\frac{N_\Phi - 1 + z_1}{2}\right)! \left(\frac{N_\Phi - 1 - z_1}{2}\right)!}. \quad (70)$$

Если некоторая последовательность $\mathbf{v}_1 \in V_1$ связана с последовательностью Φ , которая в свою очередь определяет предельную последовательность γ_∞ , тогда сумма элементов каждой из последовательностей множества $\{\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_{N_\Phi - 1}\}$ удовлетворяет двойному неравенству (40).

А теперь, для примера выполним поиск всех предельных последовательностей γ_∞ для $N_\Phi = 4$. Поиск представлен в виде следующей последовательности действий:

1) Так как $N_\Phi = 4$, то, в соответствии с равенством (50), выполняется следующее равенство: $z_1 = -1$.

2) Так как $N_\Phi = 4$ и $z_1 = -1$, то поиск подходящих последовательностей $\mathbf{v}_1$ длиной $N_\Phi - 1 = 3$, следует вести среди всех таких биполярных последовательностей, у которых сумма элементов равна -1 .

3) В соответствии с равенством (69), множество V_1 проверяемых последовательностей $\mathbf{v}_1$, длиной $N_\Phi - 1 = 3$, имеет мощность $|V_1| = 3$.

С учетом того, что в это множество должны войти все последовательности $\mathbf{v}_1$ длиной 3, у которых сумма элементов равна -1 , то оно будет следующим:

$$\mathbf{V}_1 = \{(-1, -1, 1), (-1, 1, -1), (1, -1, -1)\}.$$

4) Определим по формуле (66) для каждой из этих трех последовательностей последующие за ними последовательности $\mathbf{v}_2$ и $\mathbf{v}_3$: для последовательности $\mathbf{v}_1 = (-1, -1, 1)$, это будут: $\mathbf{v}_2 = (1, -1)$ и $\mathbf{v}_3 = 1$; для последовательности $\mathbf{v}_1 = (-1, 1, -1)$: $\mathbf{v}_2 = (-1, -1)$ и $\mathbf{v}_3 = 1$; для последовательности $\mathbf{v}_1 = (1, -1, -1)$, последующими будут: $\mathbf{v}_2 = (-1, 1)$ и $\mathbf{v}_3 = 1$.

5) Так как сумма элементов последовательности $\mathbf{v}_j$ равна величине z_j , то можно проверить, определяет ли какая-то из последовательности множества $\mathbf{V}_1$ предельную последовательность γ_∞ . Учитывая, что $z_1 = -1$ и, что у предельных последовательностей γ_∞ выполняются двойные неравенства (40). Следует проверить выполняются ли два двойных неравенства:

$$-\frac{3}{2} \leq z_2 \leq 4, \quad -\frac{3}{2} \leq z_3 \leq 3 - z_2.$$

– Выполним проверку для вектора $\mathbf{v}_1 = (-1, -1, 1)$.

Суммы элементов векторов $\mathbf{v}_2$ и $\mathbf{v}_3$ равны $z_2 = 0$ и $z_3 = 1$, соответственно. Оба двойных неравенства удовлетворяются. Следовательно, вектор $\mathbf{v}_1 = (-1, -1, 1)$ успешно проходит проверку, так как определяет предельные последовательности γ_∞ .

– Выполним проверку для вектора $\mathbf{v}_1 = (-1, 1, -1)$.

Суммы элементов векторов $\mathbf{v}_2$ и $\mathbf{v}_3$ равны $z_2 = -2$ и $z_3 = 1$, соответственно. Первое двойное неравенство не удовлетворяются. Следовательно, вектор $\mathbf{v}_1 = (-1, 1, -1)$ не проходит проверку, так как не определяет предельные последовательности γ_∞ .

– Выполним проверку для вектора $\mathbf{v}_1 = (1, -1, -1)$.

Суммы элементов векторов $\mathbf{v}_2$ и $\mathbf{v}_3$ равны $z_2 = 0$ и $z_3 = 1$, соответственно. Оба двойных неравенства удовлетворяются. Следовательно, вектор $\mathbf{v}_1 = (1, -1, -1)$ успешно прошел проверку, так как определяет предельные последовательности γ_∞ .

6) Используя вектора, которые прошли проверку, восстановим по правилам (67)–(68) биполярные последовательности Φ :

- Для вектора $\mathbf{v}_1 = (-1, -1, 1)$:
 - прямая последовательность $\Phi_{1,1} = (-1, 1, -1, -1)$;
 - инвертированная прямая последовательность $\Phi_{1,2} = (1, -1, 1, 1)$.
- Для вектора $\mathbf{v}_1 = (1, -1, -1)$:
 - прямая последовательность $\Phi_{2,1} = (-1, -1, 1, -1)$;
 - инвертированная прямая последовательность $\Phi_{2,2} = (1, 1, -1, 1)$.

Видно, что последовательность $\Phi_{2,1}$ получается из последовательности $\Phi_{1,1}$ путем обращения порядка всех элементов, тогда как последовательность $\Phi_{2,2}$ получается из $\Phi_{1,1}$ путем обращения порядка элементов с последующим инвертированием их значений. А так как последовательность $\Phi_{1,2}$ получается из последовательности $\Phi_{1,1}$ путем инвертирования значений всех элементов, то, получается, что все три последовательности $\Phi_{1,2}$, $\Phi_{2,1}$ и $\Phi_{2,2}$ можно получить из $\Phi_{1,1}$.

7) Теперь по последовательностям $\Phi_{1,1}$, $\Phi_{1,2}$, $\Phi_{2,1}$ и $\Phi_{2,2}$, в соответствии с правилом (9), вычислим биполярные последовательности, которые являются предельными:

$$\gamma_{\infty,1,1} = (-1, 1, 1, -1, -1, 1, -1, 1),$$

$$\gamma_{\infty,1,2} = (1, -1, -1, 1, 1, -1, 1, -1),$$

$$\gamma_{\infty,2,1} = (-1, 1, -1, 1, 1, -1, -1, 1),$$

$$\gamma_{\infty,2,2} = (1, -1, 1, -1, -1, 1, 1, -1).$$

Все эти четыре последовательности имеют одну и ту же автокорреляционную функцию $f_{\text{АКФ}}(\gamma_{\infty}, k)$:

$$-1, 2, -1, 0, 1, -2, -3, 8, -3, -2, 1, 0, -1, 2, -1.$$

5. Метод построения предельных биполярных последовательностей с использованием порождающих функций.

С ростом величины N_{φ} возрастает и мощность множества V_1 (69) последовательностей-кандидатов, среди которых находятся те, которые связаны с предельными последовательностями γ_{∞} . Поэтому если найти функциональную зависимость между предельными последовательностями и значениями специальных функций, которые далее называются порождающими, то можно будет сразу определять некоторое подмножество предельных последовательностей для любого заданного значения N_{φ} (без необходимости выполнения поиска предельных последовательностей на основе рационального перебора).

Каждой биполярной последовательности φ (11) можно сопоставить целое неотрицательное число a , по следующему правилу:

$$a = \sum_{i=1}^{N_{\varphi}} 2^{N_{\varphi}-i} \left(\frac{\varphi_i + 1}{2} \right). \quad (71)$$

Например, биполярной последовательности $\varphi = (1, -1, -1, -1, 1)$, будет сопоставлено число $a = 17$.

Среди целых неотрицательных чисел a , в контексте решения задачи нахождения предельных последовательностей γ_{∞} , особый интерес представляют такие целые неотрицательные числа, которые соответствуют последовательностям φ , лежащим по правилу (9) в основе предельных последовательностей γ_{∞} . Например, когда $N_{\varphi} = 5$, тогда такие целые числа a образуют последовательность: 5, 9, 11, 13, 18, 20, 22, 26. Видно, что минимальным среди них является число 5, значит, числа с 1 по 4 не связаны с предельными последовательностями.

Целое неотрицательное число a , соответствующее биполярной последовательности Φ , определенной равенством (11) и имеющей длину N_Φ , которая позволяет получить в соответствии с правилом (9) биполярную последовательность, которая является предельной последовательностью γ_∞ , можно вычислить с помощью следующей рекуррентной функции:

$$f_1(N_\Phi) = \begin{cases} f_1(N_\Phi - 1), & \text{если } N_\Phi \in \{3, 6, 9, \dots\}, \\ 2f_1(N_\Phi - 1), & \text{если } N_\Phi \in \{4, 7, 10, \dots\}, \\ 2f_1(N_\Phi - 1) + 1, & \text{если } N_\Phi \in \{2, 5, 8, \dots\}, \\ 0, & \text{если } N_\Phi = 1. \end{cases} \quad (72)$$

Например, пусть $N_\Phi = 19$, тогда:

$$\begin{aligned} a = f(19) &= 2f(18) = 2(f(17)) = 2(2f(16) + 1) = \\ &= 2(2(2f(15)) + 1) = \dots = 2730. \end{aligned} \quad (73)$$

В двоичном представлении в девятнадцатиразрядной двоичной сетке (так как $N_\Phi = 19$) это число выглядит так:

$$2730_{10} = 00000001010101010_2. \quad (74)$$

Значит, искомая биполярная последовательность будет:

$$\Phi = (-1, -1, -1, -1, -1, -1, -1, 1, 1, 1, 1, 1, 1, 1, 1, 1, 1, 1, 1). \quad (75)$$

Этой последовательности Φ соответствует биполярная последовательность:

$$\begin{aligned} \gamma &= (-1, 1, -1, 1, -1, 1, -1, 1, -1, 1, -1, 1, -1, 1, -1, 1, -1, 1, 1, \\ &1, -1, -1, 1, 1, -1, -1, 1, 1, -1, -1, 1, \\ &1, -1, -1, 1, 1, -1, -1, 1, 1, -1, -1, 1), \end{aligned} \quad (76)$$

у которой абсолютное значение отношения максимума к минимуму АКФ равно:

$$\left| \frac{\frac{2N_{\varphi}}{N_{\varphi}-1} - N_{\varphi}}{3} \right| = \frac{38}{13} \cong 2,92, \quad (77)$$

что является максимально достижимым значением, когда $N_{\varphi} = 19$ в формуле (26). Значит последовательность (76) является предельной последовательностью.

Важно отметить, что числа, получаемые из значения $f_1(N_{\varphi})$, путем инвертирования всех двоичных разрядов, путем обращения всех двоичных разрядов, а также путем обращения и инвертирования всех двоичных разрядов также соответствуют последовательностям φ , определяющим предельные последовательности γ_{∞} . Например, если такие операции выполнить с числом $f_1(19) = 2730$, представленном в 19-разрядной сетке, соответственно получатся следующие числа: 521557, 174720, 349567. Этим трем числам, также как и числу 2730 (73)–(77), соответствуют предельные последовательности γ_{∞} . И все они будут иметь одинаковые автокорреляционные функции.

Когда $N_{\varphi} \in \{1, 2, 3, 4\}$, тогда с помощью порождающей функции $f_1(N_{\varphi})$, определенной равенством (72), можно получить все возможные предельные последовательности γ_{∞} . Однако, когда $N_{\varphi} \geq 5$, среди предельных последовательностей γ_{∞} появляются такие, которые уже не получаются с помощью порождающей функции $f_1(N_{\varphi})$. Так, например, при $N_{\varphi} = 5$ существует предельные последовательности γ_{∞} , связанные с числами 9, 13, 18, 22, тогда как с помощью порождающей функции $f_1(N_{\varphi})$ можно получить числа: 5, 11, 20, 26. Число 9 является вторым по порядку в этой последовательности из 8 чисел. Если проанализировать вторые по порядку числа для $N_{\varphi} \geq 5$, можно получить вторую порождающую функцию:

$$f_2(N_\varphi) = \begin{cases} f_2(N_\varphi - 1), & \text{если } N_\varphi \in \{6, 9, 12, \dots\}, \\ 2f_2(N_\varphi - 1), & \text{если } N_\varphi \in \{7, 10, 13, \dots\}, \\ 2f_2(N_\varphi - 1) + 1, & \text{если } N_\varphi \in \{8, 11, 14, \dots\}, \\ 9, & \text{если } N_\varphi = 5. \end{cases} \quad (78)$$

Функции $f_1(N_\varphi)$ (для значений $N_\varphi \in \{1, \dots, 10000\}$) и $f_2(N_\varphi)$ (для значений $N_\varphi \in \{5, \dots, 10000\}$) успешно прошли проверку на то, что они действительно вычисляют подходящие числа a , с помощью которых можно получить предельные последовательности γ_∞ . Таким образом, с помощью порождающих функций f_1 и f_2 можно получать предельные последовательности в тех случаях, когда вычислительная сложность поиска предельных последовательностей в соответствии с предложенным выше методом неприемлемо велика.

6. Результаты компьютерного моделирования. Выше было доказано, что абсолютное значение отношения максимума АКФ маркерной последовательности $\mathbf{m}$ к минимуму АКФ, когда используется предельная последовательность γ_∞ , превышает число 2 и в пределе равно 3. В то время как абсолютное значение этого отношения, когда используется традиционная расширяющая последовательность $\gamma_{\text{трад}}$ при построении маркерной последовательности $\mathbf{m}$, с достаточно высокой точностью равно 2. Применение предложенных предельных последовательностей с улучшенными корреляционными свойствами позволяет повысить качество детектирования и, как следствие, устойчивость маркирования цифровых аудиосигналов к действию помех, преобразованиям и деструктивным воздействиям.

Важно выполнить экспериментальную оценку того, насколько применение предельных биполярных последовательностей отражается на корреляционных свойствах извлекаемого из стегоаудиосигнала маркера для наихудших (и часто реальных) условий приема, характеризующих прежде всего слепой прием. Напомним, что при слепом приеме аудиосигнал, в который внедряется маркер, неизвестен авторизованному получателю. В этом случае в процессе обнаружения и извлечения маркера из стегоаудиосигнала исходный (покрывающий) аудиосигнал становится шумом, влияющим на качество корреляционного детектирования.

Для этого необходимо, прежде всего, оценить то, насколько при слепом приеме изменяется взаимная корреляция между извлеченной из маркированного аудиосигнала последовательностью и внедренной маркерной последовательностью $\mathbf{m}$ при варьировании тех или иных параметров, а также провести сравнительную оценку абсолютных значений отношений максимумов к минимумам взаимных корреляционных функций (ВКФ), характеризующих эту корреляцию, для традиционно используемой последовательности $\gamma_{\text{трад}}$ и предложенной предельной последовательности γ_{∞} .

В процессе проведенного компьютерного моделирования в качестве покрывающих аудиосигналов использовались 100 популярных музыкально-песенных композиций из различных направлений музыки. Маркированные аудиосигналы не подвергались никаким дополнительным цифровым преобразованиям. Поэтому исходный цифровой аудиосигнал в условиях слепого приема выступал единственным искажающим воздействием (шумом). Варьируемым параметром была сила внедрения маркера, определяющая коэффициент масштабирования амплитуды частотной составляющей спектра Фурье блока отсчетов при внедрении элементов маркера. Оценивались два показателя – отношение сигнал-маркер и абсолютное значение отношения максимума взаимной корреляционной функции к ее минимуму. Оба этих показателя позволяли делать выводы об устойчивости (в условиях слепого приема) маркерных последовательностей, сформированных на основе предельных последовательностей γ_{∞} , и традиционно используемых маркерных последовательностей, основанных на $\gamma_{\text{трад}}$.

Для внедрения, обнаружения и извлечения маркерной последовательности методом слепого приема использовались процедуры, предложенные в работе [24].

При формировании маркерных последовательностей $\mathbf{m}$ для компьютерного моделирования последовательность $\mathbf{a}$ представляла собой биполярную последовательность длиной 64449 элементов, получаемую в результате кронекерова произведения двух псевдослучайных последовательностей:

$$\mathbf{a} = \mathbf{a}_1 \otimes \mathbf{a}_2,$$

где $\mathbf{a}_1$ – это биполярная последовательность Касами (длиной 1023 элемента), $\mathbf{a}_2$ – это биполярная последовательность Голда (длиной 63 элемента). В качестве расширяющей последовательности

использовалась либо биполярная последовательность $\gamma_{\text{трад}}$, определяемая равенством (10), соответствующая традиционному методу формирования маркерной последовательности по правилу (3), либо предельная биполярная последовательность:

$$\gamma_{\infty} = (1, -1, -1, 1, -1, 1, -1, 1, -1, 1, 1, -1, -1, 1, -1, 1, 1, -1, -1, 1, 1, -1, -1, 1, 1, -1),$$

соответствующая предлагаемому робастному методу формирования маркерной последовательности по правилу (12).

В результате, при использовании традиционного подхода к формированию маркерной последовательности вычисленное абсолютное значение отношения максимума АКФ к её минимуму составило:

$$\left| -\frac{128898}{64417} \right| \cong 2.001,$$

а при использовании предлагаемого подхода с использованием предельных биполярных последовательностей, это отношение равнялось:

$$\left| -\frac{1675674}{580009} \right| \cong 2.889.$$

Поскольку при слепом приеме покрывающий аудиосигнал становится шумом для процесса извлечения внедренных маркеров и вычисления ВКФ, то абсолютные значения этих отношений для ВКФ будут определяться отношением сигнал-маркер, которое в свою очередь, зависит от силы внедрения элементов маркера.

Результаты компьютерного моделирования по оценке абсолютной величины отношения максимума ВКФ к ее минимуму в зависимости от основных параметров маркирования представлены в Таблицах 1 и 2. Приняты следующие обозначения столбцов таблицы: [А] – сила внедрения маркера; [Б] – среднее значение отношения сигнал-маркер (единица измерения – дБ); [В] – среднее значение абсолютной величины отношения максимума ВКФ к минимуму этой функции, характеризующей корреляцию между извлеченной последовательностью и внедренной маркерной последовательностью.

Традиционно считается [25, 26, 27, 28], что, когда отношение сигнал-маркер больше 20 дБ, тогда акустические артефакты, появляющиеся в результате внедрения маркера в аудиосигнал, будут

не слышимы человеческим ухом. Поэтому при моделировании максимальная сила внедрения равнялась 0.1, что обеспечивает отношение сигнал-маркер больше 20 дБ.

Анализ полученных результатов показал, что при силе внедрения 0.1 абсолютное значение величины отношения максимума ВКФ к ее минимуму уменьшается по сравнению с этим же показателем, полученным для АКФ (2.001), почти на 0.12 (традиционное маркирование), тогда как для робастного маркирования, основанного на использовании предельной последовательности, это отличие составляет всего лишь около 0.05 (при значении для АКФ, равном 2.889). Традиционный маркер становится не отличим от белого шума, когда сила внедрения будет меньше 0.01, тогда как маркер, построенный на основе предельной последовательности, будет не отличим от белого шума только тогда, когда сила внедрения будет меньше 0.003.

Таким образом, при слепом приеме маркирование цифровых аудиосигналов с использованием предельных биполярных последовательностей, по сравнению с традиционным маркированием, не только обеспечивает большее по абсолютной величине отношение максимума ВКФ к ее минимуму, но и является на порядок более устойчивым к искажающему воздействию покрывающего аудиосигнала.

Таблица 1. Результаты моделирования для случая, когда использовалась традиционная биполярная последовательность $\gamma = (1, -1)$

[A]	[Б]	[B]
0.1	22.6537914407729	1.89400133362661
0.09	23.5689087161114	1.88940377072308
0.08	24.5919137027463	1.88378546192023
0.07	25.7516863726654	1.87676523118745
0.06	27.0905201404766	1.86774872443684
0.05	28.6739760385219	1.85576203481233
0.04	30.6118656605248	1.83912315246444
0.03	33.1099713242667	1.80857006872274
0.02	36.6298991976145	1.66904242496330
0.01	42.6405797512073	1.22115540454256
0.009	43.5527289528179	1.17197219241071
0.008	44.5716571653026	1.12916956984856
0.007	45.7256236763519	1.09483301588215
0.006	47.0558026726118	1.06745248183093
0.005	48.6255673720736	1.04974406741366
0.004	50.5399521171489	1.03956744534308
0.003	52.9924466534771	1.03586059278761
0.002	56.4032438747814	1.03275100298008
0.001	62.0084277052975	1.03056260054767

последовательностей произвольной длины с использованием порождающих функций.

Представлены результаты компьютерного моделирования по оценке значений абсолютной величины отношения максимума к минимуму автокорреляционной и взаимной корреляционных функций исследуемых биполярных последовательностей для наихудших условий приема – слепого приема. Их анализ показал, что предложенные предельные биполярные последовательности характеризуются лучшими корреляционными свойствами в сравнении с традиционно используемыми биполярными последовательностями и обладают большей устойчивостью.

Литература

1. Алексеев В.И., Коржик В.И. Погружение цифровых водяных знаков в аудиосигналы с помощью использования частотно селективного изменения фазы // Научные технологии в космических исследованиях Земли. 2019. Т. 11. № 6. С. 22–29.
2. Шелухин О.И., Рыбаков С.Ю., Магомедова Д.И. Скрытие информации в аудиосигналах с использованием детерминированного хаоса // Научные технологии в космических исследованиях Земли. 2021. Т. 13. № 1. С. 80–91.
3. Bender W., Gruhl D., Morimoto N., Lu A. Techniques for data hiding // IBM systems journal. 1996. vol. 35. № 3.4. pp. 313–336.
4. Wendzel S., Caviglione L., Mazurczyk W., Mileva A., Dittmann J., Krätzer C., Kevin L., Claus V., Laura H., Jorg K., Tom N., Sebastian Z. A Generic Taxonomy for Steganography Methods. TechRxiv. Preprint. IEEE. 2022. DOI: 10.36227/techrxiv.20215373.v2.
5. Makhdoom I., Abolhasan M., Lipman J. A Comprehensive Survey of Covert Communication Techniques, Limitations and Future Challenges // Computers & Security. 2022. vol. 120. DOI: 10.1016/j.cose.2022.102784.
6. Yeo I., Kim H. Modified patchwork algorithm: A novel audio watermarking scheme // IEEE Transactions on speech and audio processing. 2003. vol. 11. № 4. pp. 381–386.
7. Liu Z., Huang Y., Huang J. Patchwork-based audio watermarking robust against de-synchronization and recapturing attacks // IEEE transactions on information forensics and security. 2018. vol. 14. № 5. pp. 1171–1180.
8. Chincholkar Y., Ganorkar S. Audio watermarking algorithm implementation using patchwork technique // 2019 IEEE 5th International Conference for Convergence in Technology (I2CT). IEEE. 2019. pp. 1–5.
9. Chincholkar Y., Kude S. Effective robust patchwork method to the vulnerable attack for digital audio watermarking // ICTACT Journal on Image & Video Processing. 2018. vol. 8. № 4. pp. 1753–1758.
10. Li Q., Wang X., Pei Q. Compression Domain Reversible Robust Watermarking Based on Multilayer Embedding // Security and Communication Networks. 2022. vol. 2022. Article ID 4542705. 13 p. DOI: 10.1155/2022/4542705.
11. Saritas O., Ozturk S. A color channel multiplexing approach for robust discrete wavelet transform based image watermarking // Concurrency and Computation: Practice and Experience. 2022. e7255. DOI: 10.1002/cpe.7255.
12. Zhu X., Lai Z., Zhou N., Wu J. Steganography with High Reconstruction Robustness: Hiding of Encrypted Secret Images. Mathematics. 2022. vol. 10. № 16: 2934. DOI: 10.3390/math10162934.

13. Paul S., Mishra D. Hiding images within audio using deep generative model // *Multimedia Tools and Applications*. 2022. pp. 1–24. DOI: doi.org/10.1007/s11042-022-13034-4.
14. Luo X., Goebel M., Barshan E., Yang F. LECA: A Learned Approach for Efficient Cover-agnostic Watermarking // *arXiv preprint arXiv:2206.10813*. 2022. DOI: 10.48550/arXiv.2206.10813.
15. Ghosal S., Roy S., Basak R. LSB Steganography Using Three Level Arnold Scrambling and Pseudo-random Generator // Eds.: Giri D., Mandal J., Sakurai K., De D. In: *Proceedings of International Conference on Network Security and Blockchain Technology. ICNSBT 2021. Lecture Notes in Networks and Systems*. 2022. vol. 481. pp. 105–116. DOI: 10.1007/978-981-19-3182-6_9.
16. Rai D.K. Invisible Unicode Programming // *International Journal of Research Publication and Reviews*. vol. 3. № 4. pp. 1–19. DOI: 10.55248/gengpi.2022.3.4.1.
17. Nakashima Y., Tachibana R., Babaguchi N. Watermarked movie soundtrack finds the position of the camcorder in a theater // *IEEE Transactions on Multimedia*. 2009. vol. 11. № 3. pp. 443–454.
18. Tachibana R. Sonic watermarking // *EURASIP Journal on Advances in Signal Processing*. 2004. vol. 2004. № 13. DOI: 10.1155/S1110865704403138.
19. Tachibana R. Audio watermarking for live performance // *Security and Watermarking of Multimedia Contents V. SPIE*, 2003. vol. 5020. pp. 32–43. DOI: 10.1117/12.476832.
20. Zhang Z., Wu X. An audio covert communication system for analog channels // *International Conference on Electrical and Control Engineering. IEEE*, 2010. pp. 3279–3282. DOI: 10.1109/ICECE.2010.800.
21. Kaneto R., Nakashima Y., Babaguchi N. Real-time user position estimation in indoor environments using digital watermarking for audio signals // *20th International Conference on Pattern Recognition. IEEE*, 2010. pp. 97–100.
22. Гофман М.В., Корниенко А.А., Глухов А.П. Методика позиционирования маркированными аудиосигналами // *Проблемы информационной безопасности. Компьютерные системы*. 2018. № 4. С. 120–129.
23. Гофман М.В. Методика скрытой передачи данных при связи через воздушный аудиоканал // *Труды СПИИРАН*. 2017. Вып. 51. С. 97–122.
24. Гофман М.В. Помехоустойчивое маркирование цифровых аудиосигналов в аудиостегосистемах с множественным входом и множественным выходом // *Проблемы информационной безопасности. Компьютерные системы*. 2021. Вып. 3. С. 83–95.
25. Özer H., Sankur B., Memon N., Avcıbaşı İ. Detection of audio covert channels using statistical footprints of hidden messages // *Digital Signal Processing*. 2006. vol. 16. № 4. pp. 389–401.
26. Su Z., Zhang G., Yue F., Chang L., Jiang J., Yao X. SNR-constrained heuristics for optimizing the scaling parameter of robust audio watermarking // *IEEE Transactions on Multimedia*. 2018. vol. 20. № 10. pp. 2631–2644.
27. Appadurai E., Bhatt M., Geetha D. Semi Fragile Audio Crypto-Watermarking based on Sparse Sampling with Partially Decomposed Haar Matrix Structure // *Acta Cybernetica*. 2020. vol. 24. № 4. pp. 679–697.
28. Khillare A., Malviya A. Reversible Digital Audio Watermarking Scheme Using Wavelet Transformation // *Journal of Engineering Research and Application*. 2018. vol. 8(7). pp. 62–72.

Гофман Максим Викторович — канд. техн. наук, доцент, кафедра «информатика и информационная безопасность», Петербургский государственный университет путей сообщения Императора Александра I. Область научных интересов: информационная

безопасность. Число научных публикаций — 45. maxgof@gmail.com; Московский проспект, 9, 190031, Санкт-Петербург, Россия; р.т.: +7(812)310-3472.

Корниенко Анатолий Адамович — д-р техн. наук, профессор кафедры, кафедра «информатика и информационная безопасность», Петербургский государственный университет путей сообщения Императора Александра I. Область научных интересов: информационная безопасность. Число научных публикаций — 421. kaa.pgups@yandex.ru; Московский проспект, 9, 190031, Санкт-Петербург, Россия; р.т.: +7(812)310-3472.

M. GOFMAN, A. KORNIENKO
**LIMIT BIPOLAR SEQUENCES FOR PATCHWORK-BASED
ROBUST DIGITAL AUDIO WATERMARKING**

Gofman M., Kornienko A. Limit Bipolar Sequences for Patchwork-Based Robust Digital Audio Watermarking.

Abstract. Ensuring the robustness of digital audio watermarking under the influence of interference, various transformations and possible attacks is an urgent problem. One of the most used and fairly stable marking methods is the patchwork method. Its robustness is ensured by the use of expanding bipolar numerical sequences in the formation and embedding of a watermark in a digital audio and correlation detection in the detection and extraction of a watermark. An analysis of the patchwork method showed that the absolute values of the ratio of the maximum of the autocorrelation function (ACF) to its minimum for expanding bipolar sequences and extended marker sequences used in traditional digital watermarking approach 2 with high accuracy. This made it possible to formulate criteria for searching for special expanding bipolar sequences, which have improved correlation properties and greater robustness. The article developed a mathematical apparatus for searching and constructing limit-expanding bipolar sequences used in solving the problem of robust digital audio watermarking using the patchwork method. Limit bipolar sequences are defined as sequences whose autocorrelation functions have the maximum possible ratios of maximum to minimum in absolute value. Theorems and corollaries from them are formulated and proved: on the existence of an upper bound on the minimum values of autocorrelation functions of limit bipolar sequences and on the values of the first and second petals of the ACF. On this basis, a rigorous mathematical definition of limit bipolar sequences is given. A method for searching for the complete set of limit bipolar sequences based on rational search and a method for constructing limit bipolar sequences of arbitrary length using generating functions are developed. The results of the computer simulation of the assessment of the values of the absolute value of the ratio of the maximum to the minimum of the autocorrelation and cross-correlation functions of the studied bipolar sequences for blind reception are presented. It is shown that the proposed limit bipolar sequences are characterized by better correlation properties in comparison with the traditionally used bipolar sequences and are more robust.

Keywords: steganography, digital audio, patchwork method, digital watermarking, bipolar sequences, correlation function.

References

1. Alekseev V.G., Korzhik V.I. [Embedding of digital watermark signals with the use of a frequency selective phase changing]. *Naukoemkie tehnologii v kosmicheskikh issledovaniyah Zemli – H&ES Research*. 2019. vol. 11. № 6. pp. 22–29. DOI: 10.24411/2409-5419-2018-10292. (In Russ.).
2. Sheluhin O.I., Rybakov S.Y., Magomedova D.I. [Audio steganography method using determined chaos]. *Naukoemkie tehnologii v kosmicheskikh issledovaniyah Zemli – H&ES Research*. 2021. vol. 13. № 1. pp. 80–91. DOI: 10.36724/2409-5419-2021-13-1-80-91. (In Russ.).
3. Bender W., Gruhl D., Morimoto N., Lu A. Techniques for data hiding. *IBM systems journal*. 1996. vol. 35. № 3.4. pp. 313–336.
4. Wendzel S., Caviglione L., Mazurczyk W., Mileva A., Dittmann J., Krätzer C., Kevin L., Claus V., Laura H., Jorg K., Tom N., Sebastian Z. A Generic Taxonomy for

- Steganography Methods. TechRxiv. Preprint. IEEE. 2022. DOI: 10.36227/techrxiv.20215373.v2.
5. Makhdoom I., Abolhasan M., Lipman J. A Comprehensive Survey of Covert Communication Techniques, Limitations and Future Challenges. *Computers & Security*. 2022. vol. 120. DOI: 10.1016/j.cose.2022.102784.
6. Yeo I., Kim H. Modified patchwork algorithm: A novel audio watermarking scheme. *IEEE Transactions on speech and audio processing*. 2003. vol. 11. № 4. pp. 381–386.
7. Liu Z., Huang Y., Huang J. Patchwork-based audio watermarking robust against desynchronization and recapturing attacks. *IEEE transactions on information forensics and security*. 2018. vol. 14. № 5. pp. 1171–1180.
8. Chincholkar Y., Ganorkar S. Audio watermarking algorithm implementation using patchwork technique. *IEEE 5th International Conference for Convergence in Technology (I2CT)*. IEEE. 2019. pp. 1–5.
9. Chincholkar Y., Kude S. Effective robust patchwork method to the vulnerable attack for digital audio watermarking. *ICTACT Journal on Image & Video Processing*. 2018. vol. 8. № 4. pp. 1753–1758.
10. Li Q., Wang X., Pei Q. Compression Domain Reversible Robust Watermarking Based on Multilayer Embedding. *Security and Communication Networks*. 2022. vol. 2022, Article ID 4542705. 13 p. DOI: doi.org/10.1155/2022/4542705.
11. Saritas O., Ozturk S. A color channel multiplexing approach for robust discrete wavelet transform based image watermarking. *Concurrency and Computation: Practice and Experience*. 2022. e7255. DOI: doi.org/10.1002/cpe.7255.
12. Zhu X., Lai Z., Zhou N., Wu J. Steganography with High Reconstruction Robustness: Hiding of Encrypted Secret Images. *Mathematics*. 2022. vol. 10. № 16:2934. DOI: doi.org/10.3390/math10162934.
13. Paul S., Mishra D. Hiding images within audio using deep generative model. *Multimedia Tools and Applications*. 2022. pp. 1–24. DOI: 10.1007/s11042-022-13034-4.
14. Luo X., Goebel M., Barshan E., Yang F. LECA: A Learned Approach for Efficient Cover-agnostic Watermarking. *arXiv preprint arXiv:2206.10813*. 2022. DOI: 10.48550/arXiv.2206.10813.
15. Ghosal S., Roy S., Basak R. LSB Steganography Using Three Level Arnold Scrambling and Pseudo-random Generator. Eds.: Giri D., Mandal J., Sakurai K., De D. In: *Proceedings of International Conference on Network Security and Blockchain Technology. ICNSBT 2021. Lecture Notes in Networks and Systems*. 2022. vol. 481. pp. 105–116. DOI: 10.1007/978-981-19-3182-6_9.
16. Rai D. Invisible Unicode Programming. *International Journal of Research Publication and Reviews*. vol. 3. № 4. pp. 1–19. DOI: 10.55248/gengpi.2022.3.4.1.
17. Nakashima Y., Tachibana R., Babaguchi N. Watermarked movie soundtrack finds the position of the camcorder in a theater. *IEEE Transactions on Multimedia*. 2009. vol. 11. № 3. pp. 443–454.
18. Tachibana R. Sonic watermarking. *EURASIP Journal on Advances in Signal Processing*. 2004. vol. 2004. № 13. DOI: 10.1155/S1110865704403138.
19. Tachibana R. Audio watermarking for live performance. *Security and Watermarking of Multimedia Contents V*. SPIE. 2003. vol. 5020. pp. 32–43. DOI: 10.1117/12.476832.
20. Zhang Z., Wu X. An audio covert communication system for analog channels. *International Conference on Electrical and Control Engineering*. IEEE. 2010. pp. 3279–3282. DOI: 10.1109/ICECE.2010.800.
21. Kaneto R., Nakashima Y., Babaguchi N. Real-time user position estimation in indoor environments using digital watermarking for audio signals. *20th International Conference on Pattern Recognition*. IEEE. 2010. pp. 97–100.

22. Gofman M.V., Kornienko A.A., Glukhov A.P. [A Positioning Method Based On Digital Audio Watermarking]. Problemy informacionnoj bezopasnosti. Komp'juternye sistemy – Problems of information security. Computer systems. 2018. № 4. pp. 120–129. (In Russ.).
23. Gofman M.V. [A Method of Hidden Data in Communication via Air Audio Channel] Trudy SPIIRAN – SPIIRAS Proceedings. 2017. № 2(51). pp. 97–122. (In Russ.)
24. Gofman M.V. [Noisproof Digital Audio Watermarking In MIMO Audio Stegosystems]. Problemy informacionnoj bezopasnosti. Komp'juternye sistemy – Problems of information security. 2021. № 3. pp. 83–95. (In Russ.).
25. Özer H., Sankur B., Memon N., Avcıbaş İ. Detection of audio covert channels using statistical footprints of hidden messages. Digital Signal Processing. 2006. vol. 16. № 4. pp. 389–401.
26. Su Z., Zhang G., Yue F., Chang L., Jiang J., Yao X. SNR-constrained heuristics for optimizing the scaling parameter of robust audio watermarking. IEEE Transactions on Multimedia. 2018. vol. 20. № 10. pp. 2631–2644.
27. Appadurai E., Bhatt M., Geetha D. Semi Fragile Audio Crypto-Watermarking based on Sparse Sampling with Partially Decomposed Haar Matrix Structure. Acta Cybernetica. 2020. vol. 24. № 4. pp. 679–697.
28. Khillare A., Malviya A. Reversible Digital Audio Watermarking Scheme Using Wavelet Transformation. Journal of Engineering Research and Application. 2018. vol. 8(7). pp 62–72.

Gofman Maksim — Ph.D., Associate Professor of the Department, Department «informatics and information security», Emperor Alexander I St. Petersburg State Transport University. Research interests: information security. The number of publications — 45. maxgof@gmail.com; 9, Moskovsky Av., 190031, St. Petersburg, Russia; office phone: +7(812)310-3472.

Kornienko Anatolij — Ph.D., Dr.Sci., Professor of the department, Department «informatics and information security», Emperor Alexander I St. Petersburg State Transport University. Research interests: information security. The number of publications — 421. kaa.pgups@yandex.ru; 9, Moskovsky Av., 190031, St. Petersburg, Russia; office phone: +7(812)310-3472.

С.В. Дворников, С.С. Дворников, К.Д. Жеглов
**ПОМЕХОУСТОЙЧИВОСТЬ СИГНАЛОВ ОДНОПОЛОСНОЙ
МОДУЛЯЦИИ С УПРАВЛЯЕМЫМ УРОВНЕМ НЕСУЩЕГО
КОЛЕБАНИЯ**

Дворников С.В., Дворников С.С., Жеглов К.Д. Помехоустойчивость сигналов однополосной модуляции с управляемым уровнем несущего колебания.

Аннотация. Однополосная модуляция активно используется при организации связи посредством ионосферного канала в дециметровом диапазоне радиоволн. Это обусловлено, тем, что передачи с однополосной модуляции позволяют минимизировать полосу частот при сохранении скорости передачи информации и при этом повысить помехоустойчивость приема по отношению к передачам с амплитудной и частотной аналоговой модуляцией. Вместе с тем широкое применение технологий квадратурного синтеза открыли новые возможности по формированию передач с однополосной модуляцией без непосредственного применения процедур фильтрации. Анализ особенностей реализации метода квадратурного синтеза сигналов с однополосной модуляцией показал, что введение в состав его процедур дополнительного параметра позволит регулировать остаточный уровень несущего колебания и тем самым управлять помехоустойчивостью приема. Открывшиеся возможности позволили разработать способ и реализующее его устройство формирования сигнала однополосной модуляции с регулируемым уровнем несущего колебания. Рассмотрены технологии квадратурного синтеза сигналов амплитудной модуляции и однополосной модуляции с подавленной несущей как на уровне аналитического моделирования, так и с применением стандартного квадратурного модулятора. Обоснована необходимость перехода к аналитической форме представления модулирующего сигнала. Показана роль и место преобразователя Гильберта при формировании сигналов с однополосной модуляцией. Рассмотрены известные технологии формирования сигналов однополосной модуляции с сохраненным пилот-сигналом. Обоснована возможность управления величиной сохраненного пилот-сигнала на уровне процедур квадратурного синтеза. Разработана аналитическая модель и на ее основе структурная схема, позволяющая формировать сигналы однополосной модуляции с регулируемым уровнем пилот-сигнала. Демонстрируются результаты аналитического моделирования. Рассчитана величина обеспечиваемого энергетического выигрыша в результате регулирования остаточным уровнем несущего колебания. Проанализированы подходы к оценке помехоустойчивости передач с однополосной модуляцией. Предложен подход к расчету вероятности битовой ошибки передач с однополосной модуляцией, манипулированных дискретными колебаниями по результатам перераспределения энергии между несущим колебанием и боковой полосой, определяемого остаточным уровнем пилот-сигнала. Сформулированы выводы и предложения по практической реализации полученных результатов.

Ключевые слова: однополосная модуляция, управление уровнем пилот-сигнала, синтез сигналов однополосной модуляции, помехоустойчивость передач с однополосной модуляцией.

1. Введение. Однополосная модуляция (ОМ), в английском варианте *single-sideband modulation (SSB)*, получила свое развитие благодаря исследовательской работе Дж. Р. Карсона (патент на способ

передачи сигналов с эффективным использованием канального спектра) [1]. С того момента передачи с ОМ активно используются при организации связи посредством ионосферного канала в декаметровом диапазоне радиоволн [2 – 4].

Основное преимущество сигналов ОМ в каналах декаметровой радиосвязи, перед другими видами модуляции состоит в том, что их применение позволяет минимизировать полосу частот при сохранении скорости передачи информации [5, 6]. Теоретические основы формирования и обработки сигналов ОМ достаточно хорошо проработаны, что позволило получить им широкую практическую апробацию не только в декаметровом диапазоне радиоволн, но и в оптике [7, 8].

Однако основное применение передачи с ОМ преимущественно находят на линиях коротковолновой связи [9, 10], а также в аппаратуре, использующей технологии мультиплексирования с частотным разделением каналов (ЧРК) [11, 12], в английском варианте *frequency-division multiplexing (FDM)* [8]. Начиная с середины 50 гг. XX века однополосная модуляция применяется в качестве основного стандарта для связи с воздушными судами в воздухе.

Вместе с тем, несмотря на глубокую проработку теоретических аспектов, связанных с передачами ОМ по различным каналам связи, проведенный анализ публикационной активности показал, что в настоящее время наблюдается определенный ренессанс технологий *SSB*, вызванный переходом к цифровым методам квадратурной обработки сигналов [13 – 15]. Указанные обстоятельства определяют актуальность данного направления исследования.

В связи с этим, в настоящей статье представлены результаты исследования по разработке способа и реализующего его устройства, позволяющего формировать сигналы ОМ с регулируемым уровнем несущего колебания. Данное направление представляет собой дальнейшее развитие технологии синтеза сигналов ОМ, получившей в англоязычной литературе название *single-sideband suppressed-carrier modulation (SSB-SC)* [16].

2. Теоретические основы синтеза сигналов ОМ. В аналоговых модуляторах формирование сигналов ОМ осуществлялось путем соответствующей фильтрации предварительно формируемых сигналов амплитудной модуляции (АМ). В результате указанных процедур осуществлялся выбор верхней или нижней боковой полосы спектра сигнала АМ, при необходимости, с частичным сохранением пилот-сигнала [6]. Очевидно, что такая технология достаточно сложна в своей реализации, поскольку предполагает наличие системы

высокооборотных фильтров с достаточно узкой полосой пропускания и низким уровнем боковых лепестков [17, 18].

Развитие цифровых технологий открыло возможность синтеза сигналов ОМ на основе их квадратурной обработки [19], предполагающей переход к аналитической форме представления обрабатываемого сигнала [20, 21].

С учетом того, что процедуры синтеза сигналов ОМ базируются на технологии формирования сигналов амплитудной модуляции, первоначально рассмотрим сигнал АМ, который в терминах квадратурной обработки представим в следующем виде:

$$s_{AM}(t) = \frac{1}{\sqrt{2}}[1 + m_{AM}s(t)]\cos(\omega_0 t) + \frac{1}{\sqrt{2}}[1 + m_{AM}s(t)]\sin(\omega_0 t), \quad (1)$$

где m_{AM} – индекс амплитудной модуляции; $s(t)$ – модулирующий сигнал; $\omega_0 = 2\pi f_0$, f_0 – несущая частота.

Заметим, что формула (1) может быть упрощена, но данный вариант представлен в терминах квадратурного синтеза, необходимого для дальнейшего исследования.

На рисунке 1 демонстрируется структурная схема модулятора, позволяющего осуществлять квадратурный синтез сигналов АМ, в соответствии с формулой (1).

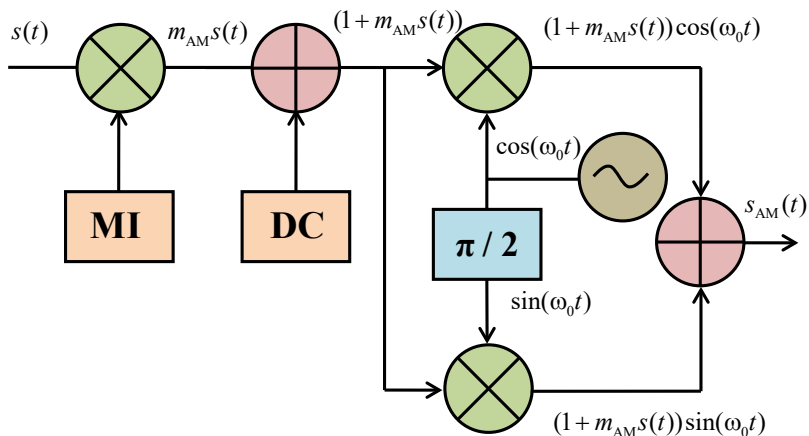


Рис. 1. Структурная схема модулятора сигналов АМ $s_{AM}(t)$ (схема предложена авторами)

На рисунке 1 введены следующие обозначения: генератор формирования значения индекса модуляции (MI – *modulation index*); генератор формирования единичного уровня напряжения постоянного тока (DC – *direct current*); $\pi/2$ – фазовращатель.

Уникальность предложенного модулятора в том, что он позволяет путем изменения напряжения генератора MI формировать сигналы АМ с заданным уровнем глубины модуляции.

В качестве примера на рисунке 2 показаны сигналы АМ: сигнал $s1_{AM}(t)$, сформированный при $m_{AM} = 1$, и сигнал $s2_{AM}(t)$, сформированный при $m_{AM} = 0,5$.

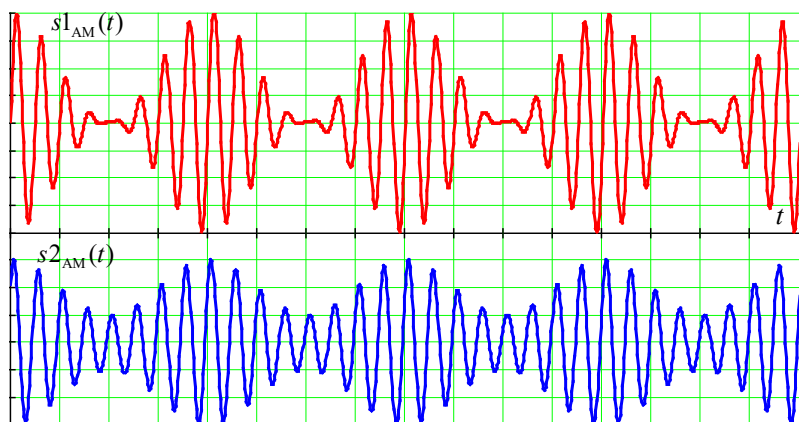


Рис. 2. Сигналы АМ при разных индексах модуляции

На рисунке 3 представлены спектры сигналов $s1_{AM}(t)$ и $s2_{AM}(t)$.

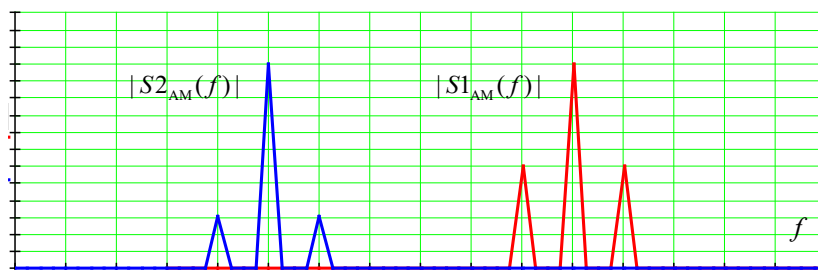


Рис. 3. Спектры сигналов АМ при разных индексах модуляции

Анализ спектров, представленных на рисунке 3, указывает на то, что изменение индекса модуляции для сигналов АМ, сформированных методом квадратурного синтеза, не ведет к перераспределению энергии между пилот-сигналом и боковыми полосами, содержащими информацию, а лишь снижает энергию последних. Это подтверждается результатами исследования, представленного в [22].

Действительно, согласно выражению (1), в каждом из квадратурных каналов используется один и тот же модулирующий сигнал $s(t)$, поэтому в результирующем спектре амплитудной модуляции левая и правая боковые полосы имеют одно и то же информационное наполнение. В результате энергия, приходящая на пилот-сигнал, который не участвует в передаче информации, в четыре раза выше, чем энергия каждой из боковых полос спектра. Именно поэтому амплитудная модуляция относится к низкоэнергетическим модуляционным форматам.

Но сама идея квадратурного синтеза открывает новые возможности по формированию сигналов ОМ, используя аналогичный подход.

В интересах раскрытия сущности квадратурного синтеза сигналов ОМ рассмотрим формирование сигнала ОМ, модулированного низкочастотной гармоникой. С позиций выражения (1) – это частный случай, при котором вместо сложения происходит вычитание квадратурных составляющих, при условии, что модулирующий сигнал приведен к аналитическому виду:

$$s_{\text{ОМ}}(t) = \frac{1}{\sqrt{2}} s(t) \cos(\omega_0 t) \pm \frac{1}{\sqrt{2}} s^*(t) \sin(\omega_0 t), \quad (2)$$

где $s(t)$ – модулирующий сигнал; $*$ – знак комплексного сопряжения по Гильберту; $\omega_0 = 2\pi f_0$; f_0 – несущая частота.

Именно использование процедуры комплексного сопряжения в выражении (2) обеспечивает переход к аналитической форме представления сигнала $s_a(t)$, согласно которой в результирующем спектре $S_a(f)$ наблюдаются только лишь положительные составляющие:

$$S_a(f) = \begin{cases} 2S(f), & f > 0; \\ S(f), & f = 0; \\ 0, & f < 0, \end{cases} \quad (3)$$

где $S(f)$ и $S_a(f)$ – соответственно, спектральные представления исходного сигнала $s_a(t)$ и его аналитической формы $s(t)$, полученные в результате преобразования Фурье.

Учитывая, что спектр аналитического сигнала содержит только положительные составляющие, то соответственно преобразование Фурье функции сдвига $S_a(f - f_0)$ будет содержать только одну из частотных полос спектра $S(f)$, при условии, что в качестве модулирующего сигнала использовано гармоническое колебание [23, 24]:

$$s_{\text{ОМ}}(t) + js_{\text{ОМ}}^*(t) = \Phi^{-1} \{ S_a(f - f_0) \} = s_a(t) \exp(j2\pi f_0 t). \quad (4)$$

В формуле (4) Φ^{-1} – процедура обратного преобразования Фурье; j – знак мнимой единицы [24].

Следует отметить, что в формуле (2) используется двойной знак. Использование знака минус позволит получить сигнал ОМ с верхней боковой полосой, а знак плюс – с нижней боковой.

Структурная схема модулятора сигналов ОМ, в соответствии с формулой (2), будет иметь вид, представленный на рисунке 4.

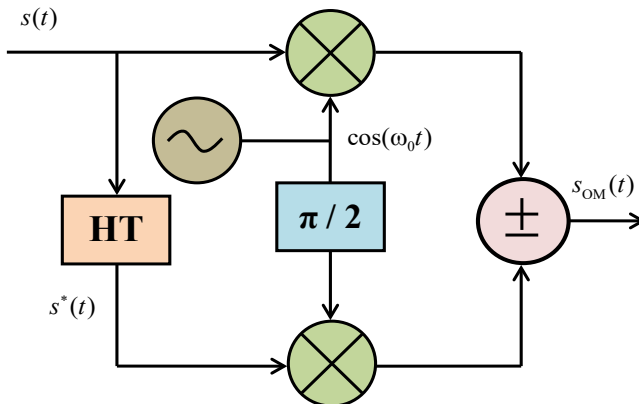


Рис. 4. Структурная схема модулятора сигналов ОМ $s_{\text{ОМ}}(t)$

Основное отличие модулятора ОМ от модулятора АМ в наличии преобразователя Гильберта, в английском варианте (*Hilbert Transformer* – НТ), который как раз и обеспечивает формирование комплексно-сопряженной формы $s^*(t)$ модулирующего сигнала $s(t)$. Сам модулятор ОМ построен по классической технологии квадратурного синтеза [10].

Для более детального раскрытия сущности работы модулятора на рисунке 5 представлено временное представление сигналов ОМ с нижней боковой полосой $s1_{OM}(t)$ и верхней боковой $s2_{OM}(t)$.

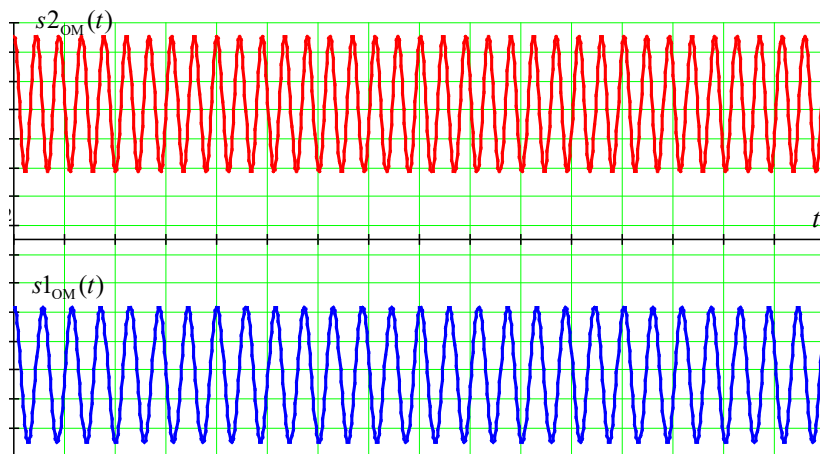


Рис. 5. Сигналы ОМ с нижней и верхней боковыми полосами

Заметим, что синтез сигналов на основе модулятора ОМ (рисунок 4), приводит к полному подавлению пилот сигнала, которое обеспечивается за счет квадратурной компенсации несущего колебания на сумматоре.

Правомерность рассмотренных процедур можно доказать аналитически, используя формулу (4), путем проведения следующих преобразований:

$$\begin{aligned}
 s_{OM}(t) &= \operatorname{Re}\{s_a(t) \exp(j2\pi f_0 t)\} = \\
 &= \operatorname{Re}\{[s(t) + js^*(t)][\cos(2\pi f_0 t) + j \sin(2\pi f_0 t)]\} = \\
 &= s(t) \cos(2\pi f_0 t) - s^*(t) \sin(2\pi f_0 t).
 \end{aligned}
 \tag{5}$$

Для более детального раскрытия сущности квадратурной компенсации несущего колебания рассмотрим синтез сигнала ОМ при использовании в качестве модулирующего сигнала $s(t)$ низкочастотной гармоники вида $s(t) = A \cos(\Omega t)$ с нулевым значением начальной фазы, где A – амплитуда колебания.

Если подставить это значение в формулу (2) и предположить, что амплитуда несущего колебания равна U_0 , а сопряженная по Гильберту форма $s^*(t) = A \sin(\Omega t)$, то с учетом известного тригонометрического преобразования получим – $\cos(\alpha + \beta) = \cos \alpha \cos \beta - \sin \alpha \sin \beta$:

$$\begin{aligned} s_{\text{ОМ}}(t) &= s(t) \cos(\omega_0 t) - s^*(t) \sin(\omega_0 t) = \\ &= A \cos(\Omega t) U_0 \cos(\omega_0 t) - A \sin(\Omega t) U_0 \sin(\omega_0 t) = \\ &= A U_0 \cos([\omega_0 + \Omega]t). \end{aligned} \quad (6)$$

Заметим, что в соответствии с выражением (6) результирующий сигнал ОМ содержит только лишь одну косинусоидальную составляющую.

Полученный результат соответствует синтезу сигнала ОМ с верхней боковой полосой.

Синтез сигнала ОМ с нижней боковой полосой возможен при переходе к модулирующему колебанию вида $s(t) = A \sin(\Omega t)$ и, соответственно, $s^*(t) = A \cos(\Omega t)$. Искомый результат будет определяться следующим тригонометрическим выражением:

$$\cos(\alpha - \beta + \pi/2) = \sin \alpha \cos \beta - \cos \alpha \sin \beta.$$

Следует отметить, что для рассмотренного случая, когда в качестве модулирующего сигнала использована низкочастотная гармоника, спектры сигналов ОМ независимо от способа формирования (нижняя боковая полоса или верхняя боковая полоса) будут иметь одинаковую структуру. При использовании более сложных модулирующих сигналов формируемые спектры будут иметь зеркальную структуру относительно несущего колебания.

Так, на рисунке 6 представлены спектры сигналов ОМ с нижней боковой $s_{3\text{ОМ}}(t)$ и верхней боковой $s_{4\text{ОМ}}(t)$ полосой, модулированных сложным многокомпонентным колебанием.

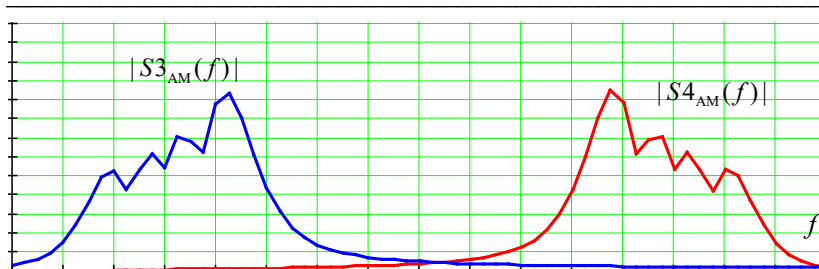


Рис. 6. Спектры сигналов ОМ сложной структуры

Результирующие спектры имеют явно зеркальную структуру.

По отношению к сигналам АМ, энергия боковой информационной составляющей у сигналов ОМ возрастает в два раза.

Несмотря на очевидную простоту реализации рассмотренного подхода, сигналы ОМ не получили широкого применения ввиду того, что прием таких сигналов связан с определенными сложностями.

Так, для того, чтобы обеспечить передачу сообщения без искажения, приемник должен быть точно настроен на частоту передатчика. Однако в силу нестабильности опорных генераторов и канальных искажений это достаточно сложно решаемая задача [25]. В результате передачи на основе ОМ без несущего колебания после демодуляции на приеме могут звучать на приеме очень неестественно с плохой разборчивостью речи.

3. Сигналы ОМ с сохраненной несущей. Для снижения негативных последствий данного эффекта на практике используют сигналы ОМ с частично сохраненной (подавленной) несущей (ОМ-ПН). Такие сигналы, как отмечалось ранее, получили название *SSB-SC* [26]. Наличие у таких сигналов несущего колебания обеспечивает на приеме возможность частотной подстройки опорного генератора приемника.

По своей сути передачи *SSB-SC* аналогичны передачам с АМ, но при этом для передачи информации используется полоса частот, которая в два раза уже, чем требуется для передач с амплитудной модуляцией. Поэтому режим работы с полной или частично подавленной несущей получил название режима, эквивалентного амплитудной модуляции, в английском варианте *amplitude modulation equivalent (AME)* [27].

Технологически (в рамках аналогового синтеза) режим *AME* не является эффективным, хотя его применение как раз и позволяет сохранить допустимое качество и требуемую разборчивость речи. Так, при реализации режима *AME* гармонические искажения могут

достигать величины порядка 25% [27], а возникающие попутно интермодуляционные искажения по своей величине намного выше, чем в традиционных режимах с АМ, но в целом словесная разборчивость остается на уровне порядка 95% и даже выше.

В теории и практике широкое применение получили два аналоговых способа реализации режима АМЕ.

Первый способ основан на совместном использовании амплитудной и фазовой модуляций, в английском варианте *compatible single sideband (CSSB)* [27]. Но такой подход предусматривает наличие высокостабильного фазового конвертора с постоянной фазовой характеристикой на всех частотах в пределах полосы пропускания приемного тракта.

Вместе с тем у такой системы имеется серьезный недостаток, заключающийся в возникновении высокого уровня интермодуляционных компонентов второго порядка при индексе модуляции $m_{AM} < 100\%$.

Стремление компенсации этого негативного явления, даже с учетом современных технологий, приводит к асимметрии структуры боковых полос. В результате возникает смещение спектральных компонент сигнала, и искомый спектр проявляется только в пределах некоторой части выделенной полосы частот приемного тракта. Из-за такого смещения происходит подавление высокочастотных составляющих спектра в стандартной полосе канала 3,1 кГц почти на 20 дБ по отношению к его низкочастотным составляющим.

Другая особенность аналоговой реализации данного режима связана с тем, что для дополнительной фазовой модуляции используется логарифмическая функция, характер поведения которой существенно зависит от уровня несущей. Поэтому при очень малом индексе модуляции сигнал CSSB становится по своей структуре близок к сигналам ОМ, что приводит к потере синхронизации на приеме.

Второй способ разработан Леонардом Р. Каном [28], который предложил в интересах снижения уровня интермодуляционных компонентов второго порядка использовать процедуры предварительного искажения. Для этого он разработал модулятор на основе функций $\arcsin$.

Но такая реализация достаточно сложна в техническом плане, так как для генерации точной формы сигнала $\arcsin$ необходимо использовать обратную связь с несколькими контурами как в модуляторе при формировании сигнала, так и в демодуляторе при его приеме. Несмотря на это, данная технология получила развитие как метод STR-84 [27].

Вместе с тем проведенные аналитические исследования показали, что смена знака в выражении (1) (как в выражении (2)) приведет к синтезу сигнала *SSB-SC*:

$$s_{\text{OM-H}}(t) = [1 + m_{\text{AM}}s(t)]\cos(\omega_0 t) \pm [1 + m_{\text{AM}}s^*(t)]\sin(\omega_0 t). \quad (7)$$

На рисунке 7 представлены эпюры временных фрагментов сигнала АМ и сигнала ОМ-Н, который синтезирован в соответствии с формулой (7).

Для дальнейшего исследования определим сигнал $s_{\text{OM-H}}(t)$ как сигнал однополосной модуляции с сохраненной несущей (пилот-сигналом) (ОМ-Н).

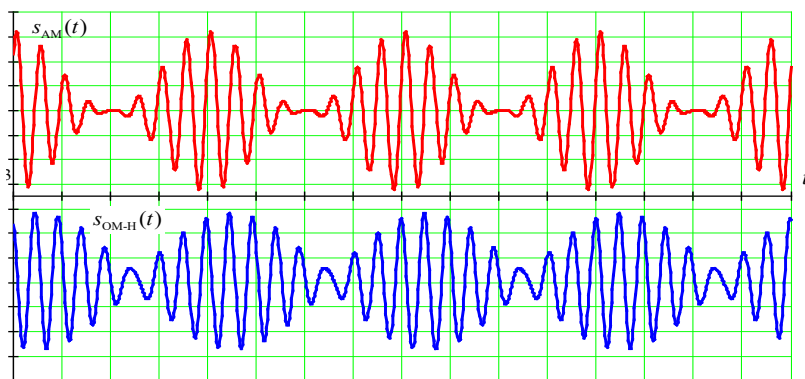


Рис. 7. Временное представление сигналов АМ $s_{\text{AM}}(t)$ и ОМ-Н $s_{\text{OM-H}}(t)$

Отметим, что сигналы АМ и ОМ-Н имеют близкую структуру. Но если у сигнала АМ при смене модулирующей посылки происходит инверсия фазы, то у сигналов ОМ-Н фаза остается непрерывной по всей длительности сигнала. Данный эффект объясняется наличием лишь одной боковой полосы у $s_{\text{OM-H}}(t)$.

На рисунке 8 показаны модули спектров сигналов $s_{\text{OM-H}}(t)$ и $s_{\text{AM}}(t)$ $|S_{\text{OM-H}}(f)|$ и $|S_{\text{AM}}(f)|$.

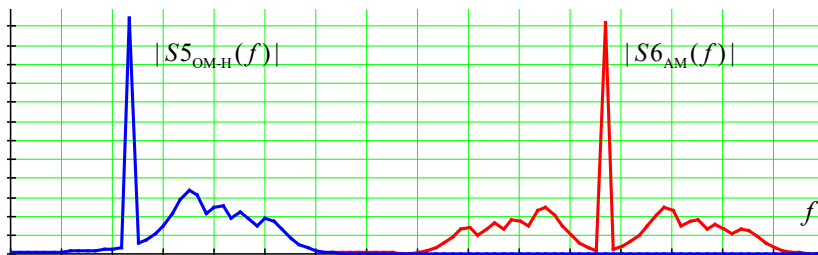
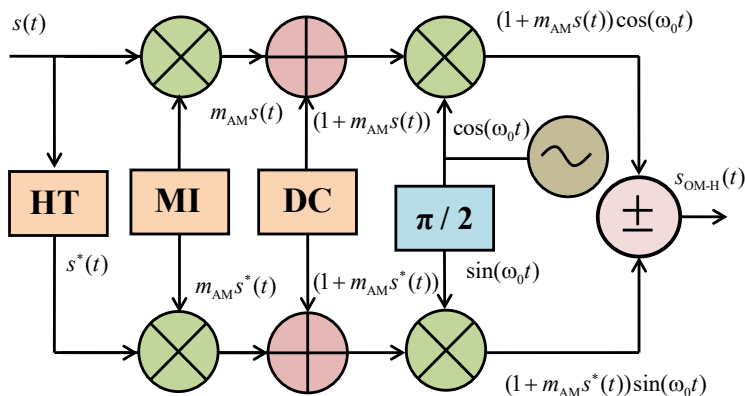


Рис. 8. Спектры сигналов ОМ-Н и АМ сложной структуры

В формуле (7) знак плюс или минус определяет выбор верхней или нижней боковой полосы. Другой важной особенностью формулы (7) является то, что в отличие от формулы (1) модулирующий сигнал в синфазной части представляет собой комплексно-сопряженную по Гильберту копию исходного сигнала. Поэтому синтез сигнала $s_{\text{OM-H}}(t)$ предполагает наличие преобразователя Гильберта. Схема такого устройства представлена на рисунке 9.

Рис. 9. Структурная схема модулятора сигналов ОМ-Н $s_{\text{OM-H}}(t)$

Элементы структурной схемы модулятора сигналов $s_{\text{OM-H}}(t)$, по своим функциям аналогичны элементам, представленным на рисунках 1 и 4.

Согласно результатам рисунка 8, сигналы $s_{\text{OM-H}}(t)$ и $s_{\text{AM}}(t)$ имеют одинаковую структуру несущего колебания (пилот-сигнала), при том, что энергия боковых полос у сигнала амплитудной модуляции в два раза меньше. При этом следует отметить, что спектр $|S5_{\text{OM-H}}(f)|$

не имеет ни искажений, ни смещений, возникающих при формировании аналоговых сигналов на основе аналоговых технологий в режиме АМЕ.

4. Сигналы ОМ с управляемой несущей. Очевидно, что энергетика, приходящаяся на информационные составляющие спектра у сигналов $s_{\text{ОМ-Н}}(t)$ существенно выше, чем у $s_{\text{АМ}}(t)$. Однако наличие такой мощной несущей не является положительным моментом. Поэтому необходим поиск возможности управления ее уровнем в зависимости от качества канала.

Проведенные исследования показали, что один из подходов решения этой задачи, обеспечивающий достижение желаемого эффекта, заключается во введение в формулу (7) дополнительного параметра $m_{\text{ОМ}}$, который как раз и позволяет регулировать остаточный уровень несущего колебания.

Полученный указанным образом сигнал определим как сигнал однополосной модуляции с управляемым уровнем несущей (пилот-сигналом) (ОМ-У).

Тогда искомое выражение для синтеза сигнала $s_{\text{ОМ-У}}(t)$ представим в следующем виде:

$$s_{\text{ОМ-У}}(t) = [m_{\text{ОМ}} + m_{\text{АМ}}s(t)]\cos(\omega_0 t) \pm [m_{\text{ОМ}} + m_{\text{АМ}}s^*(t)]\sin(\omega_0 t). \quad (8)$$

Следует отметить, что регулирование параметром $m_{\text{ОМ}}$ в пределах от нуля до единицы не ведет к перераспределению энергии, а лишь уменьшает уровень несущего колебания. Для синтеза сигналов $s_{\text{ОМ-У}}(t)$ в модулятор (рисунок 9) необходимо генератор формирования единичного уровня напряжения постоянного тока заменить на генератор формирования уровня пилот-сигнала (PS – *pilot signal*) (рисунок 10). На рисунке 11 демонстрируются спектры сигналов $s_{\text{ОМ-У}}(t)$ при значении $m_{\text{ОМ}}$ равном 1, 0,7 и 0,3.

Изменение значения $m_{\text{ОМ}}$ ведет не только к уменьшению уровня пилот-сигнала, но и изменяет структуру самого сигнала, характеризуемую снижением величины пик-фактора:

$$\Pi^2 = \frac{E_{\text{max}}}{E_0}, \quad (9)$$

где E_{max} – максимальное (пиковое) значение энергии; E_0 – среднее значение энергии.

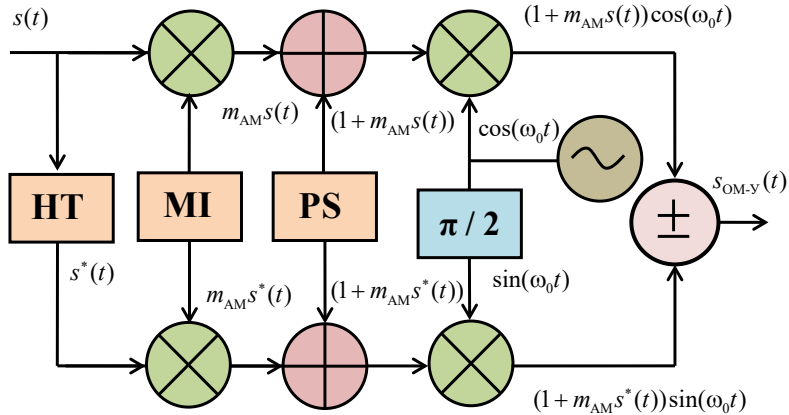


Рис. 10. Структурная схема модулятора сигналов ОМ-Н $s_{\text{ОМ-У}}(t)$

В таблице 1 представлены результаты исследования зависимости изменения значения пик-фактора от изменения величины $m_{\text{ОМ}}$. Результирующим показателем здесь рассматривается относительная величина изменения пик фактора, рассчитываемая как:

$$\delta_{\Pi} = \frac{\Pi_1^2}{\Pi_m^2}. \quad (10)$$

В формуле (10) Π_1^2 – значение пик-фактора при $m_{\text{ОМ}} = 1$; Π_m^2 – значение пик-фактора при вариативной величине $m_{\text{ОМ}}$.

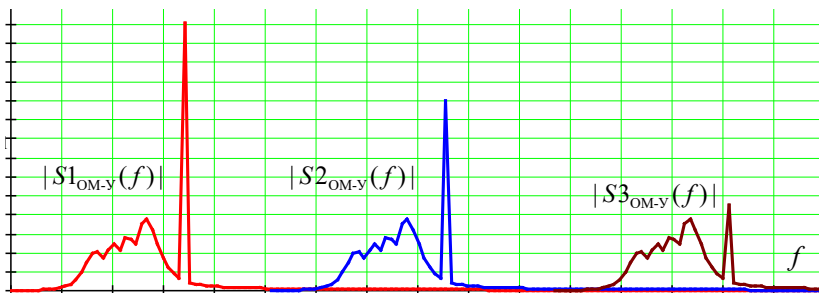


Рис. 11. Спектры сигналов ОМ-У при различном значении $m_{\text{ОМ}}$

Таблица 1. Зависимость относительной величины пик-фактора от величины коэффициента управления величиной уровня пилот-сигнала

m_{OM}	1	0,9	0,8	0,7	0,6	0,5	0,4	0,3	0,2	0,1
δ_{Π}	1	0,99	0,98	0,97	0,98	1	1,05	1,13	1,28	1,53

Результаты, представленные в таблице 1, получены при использовании в качестве модулирующего сигнала гармонического колебания.

Так, до величины $m_{OM} > 0,5$ сигналы s_{OM-y} имеют более высокий показатель пик-фактора по отношению к сигналам s_{OM-n} , несмотря на то, что они обладают более низким уровнем пилот-сигнала. А, начиная с $m_{OM} < 0,5$, у сигналов s_{OM-y} происходит стремительное снижение величины пик-фактора.

В подтверждение полученных расчетных значений на рисунке 12 представлены фрагменты сигналов $s1_{OM-y}(t)$ при значении $m_{OM} = 1$, и $s2_{OM-y}(t)$ при значении $m_{OM} = 0,3$.

Анализ эпюр временного и спектрального представления указанных сигналов показывает, что снижение величины m_{OM} не ведет к перераспределению общей энергии сигнала между боковой составляющей и несущим колебанием. Уменьшается только уровень пилот-сигнала. Однако если указанным образом формировать сигнал до его подачи на усилитель мощности, то рассмотренный подход обеспечит повышение энергетического потенциала информационных составляющих на выходе усилителя по отношению к сигналу без измененного уровня пилот-сигнала.

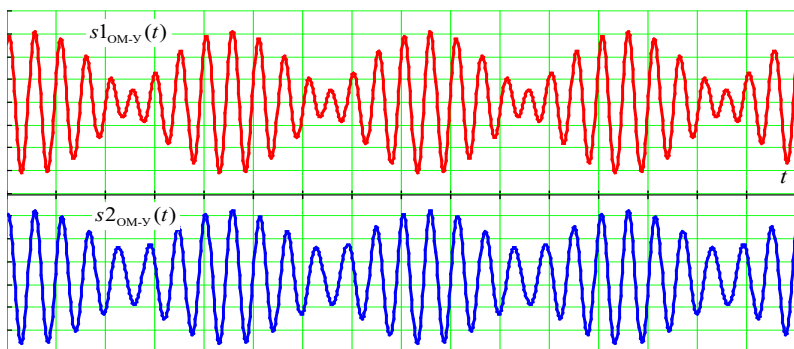


Рис. 12. Временное представление сигналов $s_{OM-y}(t)$ при различных значениях m_{OM}

Наблюдаемый эффект достигается за счет того, что при квадратурном синтезе происходит взаимная компенсация энергии синфазной и квадратурной составляющих. При $m_{OM} = 0$ сигнал $s_{OM-y}(t)$ вырождается в обычный сигнал $s_{OM}(t)$. Возможность изменения уровня пилот-сигнала позволяет управлять энергетическим балансом боковой полосы, содержащей информационное наполнение, и открывает новые возможности по организации радиосвязи.

Так, на первоначальном этапе целесообразно использовать сигнал $s_{OM-y}(t)$ с $m_{OM} = 1$. А затем, при установлении жесткой синхронизации, уменьшать уровень пилот-сигнала, тем самым повышая энергетический баланс боковой полосы.

В таблице 2 представлена зависимость относительного изменения энергии в сигнале $s_{OM-y}(t)$ δ_E :

$$\delta_E = \frac{E_B}{E_{PC}}, \quad (11)$$

где E_B – энергия, приходящейся на информационные составляющие; E_{PC} – энергия, приходящейся на пилот-сигнал.

В таблице 2 представлены результаты исследования зависимости относительного изменения энергии при различных вариациях величины m_{OM} , рассчитанной согласно формуле (11).

Таблица 2. Зависимость относительного изменения энергии, приходящейся на информационные составляющие к энергии пилот-сигнала, при изменении величины коэффициента управления величиной уровня пилот-сигнала

m_{OM}	1	0,9	0,8	0,7	0,6	0,5	0,4	0,3	0,2	0,1
δ_E	0,5	0,62	0,78	1	1,4	2,0	3,1	5,6	12,5	5,0

Тогда, с учетом данных в таблице 2, можно оценить результирующий энергетический выигрыш, приходящийся на информационные составляющие в сигнале $s_{OM-y}(t)$ при различной величине m_{OM} (таблица 3) без учета пик-фактора, в соответствии с выражением:

$$\delta_E = \frac{E_B + E_{PC}(1 - m_{OM})}{E_B}. \quad (12)$$

Таблица 3. Результирующий энергетический выигрыш, приходящейся на информационные составляющие, при изменении величины коэффициента управления величиной уровня пилот-сигнала

m_{OM}	1	0,9	0,8	0,7	0,6	0,5	0,4	0,3	0,2	0,1
δ_{Σ}	1	1,2	1,4	1,6	1,8	2	2,2	2,4	2,6	2,8

Окончательный энергетический выигрыш с учетом изменения пик-фактора (таблица 1) представлен в таблице 4.

Таблица 4. Окончательный энергетический выигрыш, приходящийся на информационные составляющие, при изменении величины коэффициента управления величиной уровня пилот-сигнала

m_{OM}	1	0,9	0,8	0,7	0,6	0,5	0,4	0,3	0,2	0,1
$\Xi_{П}$	1	1,19	1,37	1,55	1,76	2	2,31	2,71	3,33	4,28

Данные таблицы 4 можно рассматривать как результат, характеризующий энергетический выигрыш, обеспечиваемый переходом от сигналов $s_{OM-H}(t)$ к сигналам $s_{OM-Y}(t)$.

5. Помехоустойчивость передач с однополосной модуляцией.

Традиционно помехоустойчивость передач ОМ рассматривается с позиций выигрыша, обеспечиваемого в помехоустойчивости приема до $W_{OM}^{вх}$ и после $W_{OM}^{вых}$ демодулятора в сравнении с другими видами модуляции [29].

Так, средняя мощность сигнала ОМ на входе приемника будет определяться средней мощностью модулирующего сигнала $s_M^2(t)$ и средней мощности несущего колебания U_0^2 :

$$P_{OM}^2 = \frac{s_M^2(t)U_0^2}{2}. \quad (13)$$

Выражение (13) соответствует нагрузке величиной 1 Ом.

С учетом того, что основное значение пик-фактора P^2 будет определяться характером модулирующего колебания, а значение пиковой мощности – несущим колебанием, то P_{OM}^2 можно записать в виде:

$$P_{\text{OM}}^2 = \frac{U_0^2}{\Pi^2}. \quad (14)$$

С учетом этого отношение сигнал/шум (ОШ) $W_{\text{OM}}^{\text{BX}}$ на входе демодулятора будет определяться как:

$$W_{\text{OM}}^{\text{BX}} = \frac{P_{\text{OM}}^2}{v_{\text{BX}}^2 \Delta F_{\text{BX}}} = \frac{U_0^2}{v_{\text{BX}}^2 \Delta F_{\text{BX}} \Pi^2}, \quad (15)$$

где ΔF_{BX} – занимаемая полоса частот; v_{BX}^2 – спектральная плотность мощности шума на входе тракта обработки.

На выходе демодулятора, соответственно, значение ОСШ:

$$W_{\text{OM}}^{\text{ВЫХ}} = \frac{P_{\text{OM}}^2}{v_{\text{ВЫХ}}^2 \Delta F_{\text{ВЫХ}}} = \frac{U_0^2}{v_{\text{ВЫХ}}^2 \Delta F_{\text{ВЫХ}} \Pi^2}. \quad (16)$$

Тогда, показатель обобщенного энергетического выигрыша, обеспечиваемого в результате обработки передач ОМ, будет равен единице [29]:

$$B_{\text{OM}} = \frac{W_{\text{OM}}^{\text{BX}}}{W_{\text{OM}}^{\text{ВЫХ}}} = 1. \quad (17)$$

Искомый результат получен с учетом того, что характер шумов на выходе $v_{\text{ВЫХ}}^2$ будет аналогичен входным шумам v_{BX}^2 . А занимаемая полоса частот не изменится, т.е. $\Delta F_{\text{ВЫХ}} = \Delta F_{\text{ВЫХ}}$.

Заметим, что значение, равное единице в формуле (17), указывает на то, что передачи ОМ обладают потенциально возможной (достижимой) помехоустойчивостью приема.

Например, у передач с амплитудной модуляцией, такой показатель будет существенно зависеть от величины пик-фактора [29]:

$$B_{\text{AM}} = \frac{1}{(1 + \Pi^2)}. \quad (18)$$

И тогда, учитывая, что для речи значение пик-фактора достигает значения $\Pi^2 \approx 3,3$, то можно показать, относительную величину

энергетического выигрыша, обеспечиваемого однополосной модуляцией по отношению к амплитудной модуляции.

$$B_{OM} \approx 11,8 B_{AM}. \quad (19)$$

Очевидно, что такой способ оценки имеет место и используется при характеристике аналоговых систем связи [28]. Однако такая оценка является относительной и только в общем характеризует помехоустойчивость передач.

Вместе с тем в [30] представлено выражение для оценки вероятности битовой ошибки для приема сигналов бинарной амплитудной манипуляции (*binary amplitude shift keying* – *BASK*), представляющих собой разновидной амплитудной модуляции:

$$p_b(h) = Q\left(\sqrt{h^2/2}\right). \quad (20)$$

В формуле (14) h^2 – отношение энергии, приходящейся на бит, к спектральной плотности мощности шума (далее рассматриваем как ОСШ); $Q(*)$ – гауссов интеграл ошибок.

Учитывая, что переход от АМ к ОМ-Н позволит в два раза повысить энергию, приходящуюся на информационные составляющие, а применение сигналов ОМ-У при $m_{OM} = 0,3$ и $m_{OM} = 0,1$ соответственно повысит энергию в 2,71 и 4,28 раза (таблица 4), по отношению к сигналам ОМ-Н, можно на основе формулы (14) построить следующие графические зависимости, представленные на рисунке 13.

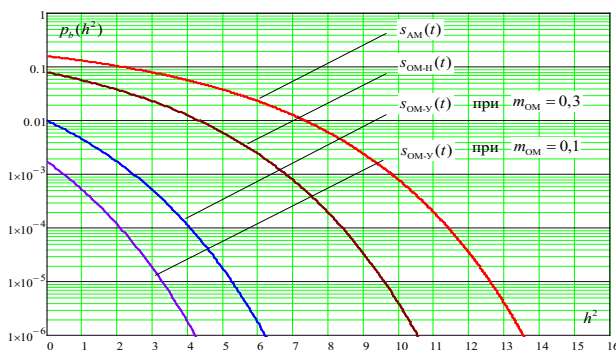


Рис. 13. Вероятностная оценка помехоустойчивости приема сигналов $s_{AM}(t)$, $s_{OM-H}(t)$ и $s_{OM-Y}(t)$ при $m_{OM} = 0,3$ и $m_{OM} = 0,1$

Полученные результаты следует рассматривать со следующих позиций. Так, переход к технологии однополосной модуляции приводит к перераспределению сигнальной энергии в более узкой полосе частот. В результате происходит фактическое изменение величины ОСШ в канале для обрабатываемого символа сигнала ОМ-У по отношению к сигналу ОМ-Н. То есть предложенный подход фактически изменяет помехоустойчивость приема не за счет изменения структуры символа как такового, а за счет дополнительного повышения энергии, приходящейся на информационные составляющие, в результате ее перераспределения между несущей и боковой полосой, обеспечиваемой путем изменения индекса $m_{\text{ОМ}}$. Поэтому шкала абсцисс на рисунке 13 отображает эквивалентное значение ОСШ, характерное только для символа сигнала АМ.

В качестве примера на рисунке 14 демонстрируются фрагменты сигналов $BASK$ $s_{\text{АМ}}(t)$ и ОМ-У $s_{\text{ОМ-У}}(t)$ при $m_{\text{ОМ}} = 0$, после перераспределения мощности (т.е. после усилителя).

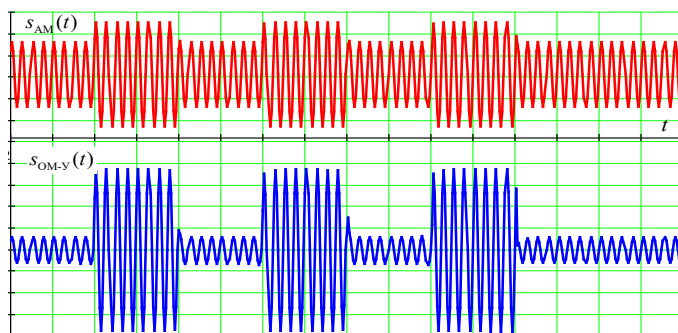


Рис. 14. Фрагменты сигналов $BASK$, сформированных на основе амплитудной модуляции $s_{\text{АМ}}(t)$ и на основе однополосной модуляции с управляемым уровнем несущей $s_{\text{ОМ-У}}(t)$ при $m_{\text{ОМ}} = 0,3$

Очевидно, что после перераспределения мощности энергетика, приходящаяся на информационные составляющие, существенно возросла. В результате при одинаковой интенсивности шумов, текущее значение ОСШ для сигналов ОМ-У существенно будет выше, чем для сигналов АМ.

Полученные результаты в помехоустойчивости для $s_{\text{ОМ-У}}(t)$ и $s_{\text{АМ}}(t)$ сопоставимы с результатами, определяемые формулой (13) для передач ОМ и АМ, модулированных аналоговой речью.

Данный факт указывает на состоятельность предложенного подхода к оценке помехоустойчивости приема сигналов ОМ различных видов.

7. Заключение. Полученные результаты исследования показали, что применение технологий квадратурного синтеза сигналов передач с однополосной модуляцией открывает новые возможности по их применению в дециметровых каналах радиосвязи. Возможность управления напряжением несущего колебания обеспечивает устойчивое вхождение в связь на уровне шумов, при которых не возможна без искажения передача сообщений при использовании сигналов ОМ-Н. И после установления синхронизации, снижением значения параметра регулирования величины пилот-сигнала, повысить энергетику боковой составляющей для устойчивой работы радиолинии. Такой подход обеспечивает энергетический выигрыш на уровне порядка 7 дБ.

Обоснованная в работе аналитическая модель сигнала ОМ с управляемым уровнем пилот-сигнала и разработанная на его основе структурная схема модулятора позволяют реализовать предложенный способ как на программном, так и аппаратном уровнях.

Дальнейшие исследования авторы связывают с применением методов совместной частотно-временной обработки сигналов, предложенных в [31], при решении задач демодуляции передач с однополосной модуляцией.

Литература

1. Carson J. Method and Means for Signaling with High Frequency Waves № US 1449382. AT&T, 1923.
2. Дворников С.В., Овчинников Г.Р., Балыков А.А. Программный симулятор ионосферного радиоканала дециметрового диапазона // Информация и космос. 2019. № 3. С. 6–12.
3. Леушин А.В. Потенциальная помехоустойчивость командной радиолинии управления // Успехи современной радиоэлектроники. 2022. Т. 76. № 7. С. 20–29.
4. Kokhanov A.B., Yemeljanov S.V., Derevyagin Y.V. Single Sideband Hartley Amplitude Modulation // Radioelectronics and Communications Systems. 2020. vol. 63. no. 11. pp. 574–585. DOI 10.3103/S0735272720110023.
5. Коханов А.Б. Однополосная квадратурная модуляция // Известия высших учебных заведений. Радиоэлектроника. 2017. Т. 60. № 3 (657). С. 123–131.
6. What is SSB: Single Sideband Modulation. Electronics Notes. Available at: <https://www.electronics-notes.com/articles/radio/modulation/single-sideband-ssb-basics.php>. (accessed 29.01.2022).
7. Ye Z., Jun M., Leilei W., Dongyan W., Li Zh., Jiangnan X. Optical Polarization Division Multiplexing Transmission System Based on Simplified Twin-SSB Modulation. Sensors. 2022. vol. 22(20). pp. 7700. DOI: 10.3390/s22207700.

8. Hualong Ye., Leihong Zh., Kaimin W., Dawei Zh. Study on the key technology of ghost imaging based on orthogonal frequency division multiplexing // *Opticheskiy Zhurnal*. 2021. vol. 88. no. 8. pp. 20–31.
9. Рахлин В.П., Сак П.В. Повышение энергетических показателей КВ-передатчика с однополосной модуляцией речевой информации при применении автоматической регулировки режима // *Техника радиосвязи*. 2021. № 4(51). С. 37–43.
10. A-Imam Al-S., Ahmed J., Ayman A. Bahrain Polytechnic. Graphical Analysis of Single Sideband Modulation. *International Journal of Computing and Digital Systems*. IJCDs Journal, 2021. vol. 10(1). DOI: 10.12785/ijcds/1001111.
11. Дворников С.В. Демодуляция сигналов на основе обработки их модифицированных частотно-временных распределений // *Цифровая обработка сигналов*. 2009. № 2. С. 7–11.
12. Никишкин П.Б., Витязев В.В. Методы широкополосной передачи данных на основе сигналов с частотным разделением каналов // *Цифровая обработка сигналов*. 2020. № 3. С. 45–49.
13. Щербаков В.В., Солодков А.Ф., Задерновский А.А. Расширенные форматы однополосной модуляции двухэлектродного модулятора Маха-Цендера // *Электроника и микроэлектроника СВЧ*. 2019. Т. 1. С. 395–399.
14. Алексеев А.А., Железняк В.К., Комарович В.Ф., Дворников С.В. Автоматизированная система контроля интенсивности физических полей рассеивания сигналов // *Научное приборостроение*. 2000. Т. 10. № 3. С. 77–87.
15. Sinicya E., Davydov V., Galichina A., Lukiyanov A., Shishkov A., Podstrigaev A. A study of temperature dependence of phase shift in optoelectronic path of direction finder channels // *Journal of Physics: Conference Series*. 2019. pp. 012075.
16. Rosepreet K., Manoj S. Generation of Single Sideband-Suppressed carrier (SSB-SC) Signal Based on Stimulated Brillouin Scattering. *Journal of Physics: Conference Series*. vol. 2327. IOP Publishing, 2022. pp. 012025. DOI:10.1088/1742-6596/2327/1/012025.
17. Дворников С.В., Бородин Е.Ю., Маджар Х., Махлуф Ю.Х. Частотно-временное оценивание параметров сигналов на основе функций огибающих плотности распределения их энергии // *Информация и космос*. 2007. № 4. С. 41–45.
18. Нефедов Е.А. Исследование двухкаскадной системы формирования сигналов с однополосной амплитудной модуляцией // *Сборник тезисов работ XLVII Международной молодёжной научной конференции Гагаринские чтения*. Москва: Издательство "Перо", 2021. С. 551.
19. Бойсунов Б.П., Короткова Л.А. Формирование радиосигналов с использованием преобразования частоты с одной боковой полосой // *Достижения науки и образования*. 2021. № 3(75). С. 21–23.
20. Дворников С.В. Теоретические основы синтеза билинейных распределений. СПб.: Изд-во Политехн. ун-та, 2007. 268 с.
21. Егоров С.Б., Горбачев Р.И. Аналитическая модель шумового сигнала с медленной спектрально-амплитудной модуляцией для пассивного шумолокатора // *Труды Санкт-Петербургского государственного морского технического университета*. 2020. Т. 1. № 52. С. 6.
22. Kulkarni A., Kothavade S., Patel D. Assessment of SSB, Modified-SSB and VSB Modulation Techniques based on Modulation Index, Extinction Ratio, Chromatic Dispersion and Received RF power. *International Conference on Communication information and Computing Technology (ICCICT)*. 2021. pp. 1–7. DOI:10.1109/ICCICT50803.2021.9509947.7.
23. Дворников С.В., Пшеничников А.В. Формирование спектрально-эффективных сигнальных конструкций в радиоканалах передачи данных контрольно-

- измерительных комплексов // Известия высших учебных заведений. Приборостроение. 2017. Т. 60. № 3. С. 221–228. DOI 10.17586/0021-3454-2017-60-3-221-228.
24. Бобков В.И., Снытко Ю.Н. Алгоритм формирования аналитического сигнала инфракрасного газоанализатора устойчивого к качке // Энергетика, информатика, инновации: Сборник трудов XI Международной научно-технической конференции. Смоленск: Универсум; филиал ФГБОУ ВО «НИУ «МЭИ», 2021. Т. 1. С. 63–67.
 25. Кобяков Р.С., Новожилов Р.Н., Писарев И.А., Жеглов А.В., Медведев С.Ю. Некоторые методы повышения точности компенсации фазовой нестабильности при передаче сигналов частоты и времени // Труды Института прикладной астрономии РАН. 2021. № 58. С. 36–40. DOI: 10.32876/AplAstron.58.36-40.
 26. Huang C., Chan E. Photonic techniques for generating a single RF sideband with no second order sidebands. IEEE Photonics Journal. 2022. vol. 14. no 1. DOI: 10.1109/JPHOT.2021.3123168.
 27. Weber P. The History of Single Sideband Modulation Archived 2004-01-03. Wayback Machine, 2004.
 28. Kahn L. Single-sideband transmission by envelope elimination and restoration. Proceedings of the IRE. 1952. vol. 40(7). pp. 803–806. DOI: 10.1109/JRPROC.1952.273844.
 29. Дворников А.С., Гудков М.А., Аюков Б.А., Федосов А.Ю., Подгорный А.В., Заседателев А.Н., Дворников С.В., Крячко А.Ф., Пшеничников А.В. Анализ помехоустойчивости передач с однополосной модуляцией в каналах с флуктуационными помехами // Вопросы радиоэлектроники. Серия: Техника телевидения. 2022. № 4. С. 58–64.
 30. Суржигов В.Ф., Компанийцев А.В. Физическое моделирование цифровых каналов СВЧ-связи с бинарным амплитудно-манипулированным сигналом в среде MATLAB // Мир науки, культуры, образования. 2022. № 1(92). С. 119–122. DOI: 10.24412/1991-5497-2022-192-119-122.
 31. Дворников С.В. Демодуляция сигналов на основе обработки их модифицированных частотно-временных распределений // Цифровая обработка сигналов. 2009. № 2. С. 7–11.

Дворников Сергей Викторович — д-р техн. наук, профессор, кафедра радиотехнических и оптоэлектронных комплексов, Федеральное государственное автономное образовательное учреждение высшего образования «Санкт-Петербургский государственный университет аэрокосмического приборостроения»; профессор кафедры, кафедра радиосвязи, Федеральное государственное казенное военное образовательное учреждение высшего образования «Военная академия связи имени Маршала Советского Союза С.М. Буденного». Область научных интересов: построение помехозащищенных систем радиосвязи, формирование и обработка сигналов сложных структур. Число научных публикаций — 448. practicdsv@yandex.ru; Тихорецкий проспект, 3, 194064, Санкт-Петербург, Россия; р.т.: +7(812)247-9400.

Дворников Сергей Сергеевич — канд. техн. наук, доцент, кафедра конструирования и технологий электронных и лазерных средств, Федеральное государственное автономное образовательное учреждение высшего образования «Санкт-Петербургский государственный университет аэрокосмического приборостроения»; научный сотрудник, научно-исследовательский отдел, Федеральное государственное казенное военное образовательное учреждение высшего образования «Военная академия связи имени Маршала Советского Союза С.М. Буденного». Область научных интересов: теория передачи сигналов, спектральная эффективность сигналов, помехозащищенность

каналов управления и связи радиотехнических систем. Число научных публикаций — 114. dvornik.92@mail.ru; Тихорецкий проспект, 3, 194064, Санкт-Петербург, Россия; р.т.: +7(812)247-9811.

Жеглов Кирилл Дмитриевич — аспирант, Федеральное государственное автономное образовательное учреждение высшего образования «Санкт-Петербургский государственный университет аэрокосмического приборостроения». Область научных интересов: системы радиосвязи дециметрового диапазона, управление частотно-энергетическим ресурсом радиолиний. Число научных публикаций — 2. zhiglov.k@gmail.com; улица Большая Морская, 67А, 190000, Санкт-Петербург, Россия; р.т.: +7(812)247-9400.

S.V. DVORNIKOV, S.S. DVORNIKOV, K. ZHEGLOV
**NOISE IMMUNITY OF SINGLE-SIDEBAND MODULATION
SIGNALS WITH A CONTROLLED CARRIER LEVEL**

Dvornikov S.V., Dvornikov S.S., Zheglov K. Noise Immunity of Single-Sideband Modulation Signals with a Controlled Carrier Level.

Abstract. Single-sideband modulation is actively used in the organization of communication through the ionospheric channel in the decameter range of radio waves. This is due to the fact that transmissions with single-sideband modulation make it possible to minimize the frequency band while maintaining the information transfer rate and at the same time increase the noise immunity of reception in relation to transmissions with amplitude and frequency analog modulation. At the same time, the widespread use of quadrature synthesis technologies has opened up new possibilities for the formation of transmissions with single-sideband modulation without the direct use of filtering procedures. An analysis of the implementation features of the method of quadrature synthesis of signals with single-sideband modulation showed that the introduction of an additional parameter into its procedures will allow you to control the residual level of the carrier wave, and thereby control the noise immunity of the reception. The opened opportunities made it possible to develop a method and a device for generating a single-sideband modulation signal with an adjustable level of the carrier wave that implements it. The technologies of quadrature synthesis of signals of amplitude modulation and single-sideband modulation with the suppressed carrier are considered both at the level of analytical modeling and using a standard quadrature modulator. The necessity of transition to the analytical form of representation of the modulating signal is substantiated. The role and place of the Hilbert converter in the formation of signals with single-sideband modulation are shown. Known technologies for generating single-sideband modulation signals with a stored pilot signal are considered. The possibility of controlling the value of the stored pilot signal at the level of quadrature synthesis procedures is substantiated. An analytical model and, based on it, a structural diagram have been developed that allow one to generate single-sideband modulation signals with an adjustable pilot signal level. The results of analytical modeling are demonstrated. The value of the provided energy gain as a result of regulation by the residual level of the carrier wave is calculated. Approaches to assessing the noise immunity of transmissions with single-sideband modulation are analyzed. An approach is proposed for calculating the bit error probability of SSB transmissions manipulated by discrete oscillations based on the results of energy redistribution between the carrier oscillation and the sideband, determined by the residual pilot signal level. Conclusions and proposals for the practical implementation of the results obtained are formulated.

Keywords: single-sideband modulation, pilot signal level control, single-sideband modulation signal synthesis, noise immunity of single-sideband modulation transmissions.

References

1. Carson J. Method and Means for Signaling with High Frequency Waves № US 1449382. AT&T, 1923.
2. Dvornikov S.V., Ovchinnikov G.R., Balykov A.A. [Software simulator of decameter ionospheric radio channel]. *Informacija i kosmos – Information and space*. 2019. No. 3. pp. 6–12. (In Russ.).
3. Leushin A.V. [Potential noise immunity of the command radio control line]. *Uspеhi sovremennoj radioelektroniki – Successes of modern radio electronics*. 2022. vol. 76. No. 7. pp. 20-29. (In Russ.).

4. Kokhanov A.B., Yemelianov S.V., Derevyagin Y.V. Single Sideband Hartley Amplitude Modulation. *Radioelectronics and Communications Systems*. 2020. vol. 63. no. 11. pp. 574–585. DOI 10.3103/S0735272720110023.
5. Kokhanov A.B. Single-sideband quadrature modulation. *Radioelektronika – Radioelectronics*. 2017. vol. 60. No. 3(657). pp. 123–131. (In Russ.).
6. What is SSB: Single Sideband Modulation. *Electronics Notes*. Available at: <https://www.electronics-notes.com/articles/radio/modulation/single-sideband-ssb-basics.php>. (accessed 29.01.2022).
7. Ye Z., Jun M., Leilei W., Dongyan W., Li Zh., Jiangnan X. Optical Polarization Division Multiplexing Transmission System Based on Simplified Twin-SSB Modulation. *Sensors*. 2022. vol. 22(20). pp. 7700. DOI: 10.3390/s22207700.
8. Hualong Ye., Leihong Zh., Kaimin W., Dawei Zh. Study on the key technology of ghost imaging based on orthogonal frequency division multiplexing. *Opticheskii Zhurnal*. 2021. vol. 88. No. 8. pp. 20–31.
9. Rakhlin V.P., Sak P.V. [Improving the energy performance of a HF transmitter with single-sideband modulation of speech information when using automatic mode control]. *Tehnika radiosvjazi – Radio communication technology*. 2021. No. 4(51). pp. 37–43. (In Russ.).
10. A-Imam Al-S., Ahmed J., Ayman A. Bahrain Polytechnic. Graphical Analysis of Single Sideband Modulation. *International Journal of Computing and Digital Systems. IJCDS Journal*, 2021. vol. 10(1). DOI: 10.12785/ijcds/1001111.
11. Dvornikov S.V. [Demodulation of signals based on the processing of their modified time-frequency distributions]. *Cifrovaja obrabotka signalov – Digital Signal Processing*. 2009. No. 2. pp. 7–11. (In Russ.).
12. Nikishkin P.B., Vityazev V.V. [Methods of broadband data transmission based on signals with frequency division of channels]. *Cifrovaja obrabotka signalov – Digital signal processing*. 2020. No. 3. pp. 45–49. (In Russ.).
13. Shcherbakov V.V., Solodkov A.F., Zadernovsky A.A. [Extended formats of single-sideband modulation of a two-electrode Mach-Zehnder modulator]. *Jelektronika i mikroelektronika SVCh – Electronics and microwave microelectronics*. 2019. vol. 1. pp. 395–399. (In Russ.).
14. Alekseev A.A., Zheleznyak V.K., Komarov V.F., Dvornikov S.V. [Automated control system for the intensity of physical signal scattering fields]. *Nauchnoe priborostroenie – Scientific Instrumentation*. 2000. vol. 10. No. 3. pp. 77–87. (In Russ.).
15. Sinicyna E., Davydov V., Galichina A., Lukyanov A., Shishkov A., Podstrigaev A. A study of temperature dependence of phase shift in optoelectronic path of direction finder channels. *Journal of Physics: Conference Series*. 2019. pp. 012075.
16. Rosepreet K., Manoj S. Generation of Single Sideband-Suppressed carrier (SSB-SC) Signal Based on Stimulated Brillouin Scattering. *Journal of Physics: Conference Series*. vol. 2327. IOP Publishing, 2022. pp. 012025. DOI:10.1088/1742-6596/2327/1/012025.
17. Dvornikov S.V., Borodin E.Yu., Madzhar Kh., Makhluif Yu.Kh. [Time-frequency estimation of signal parameters based on their energy distribution density envelope functions]. *Informacija i kosmos – Information and space*. 2007. No. 4. pp. 41–45. (In Russ.).
18. Nefedov E.A. [Study of a two-stage signal generation system with single-sideband amplitude modulation]. *Sbornik tezisev rabot XLVII Mezhdunarodnoj molodjozhnoj nauchnoj konferencii Gagarinskie chtenija [Collection of abstracts of the XLVII International Youth Scientific Conference Gagarin Readings]*. Moscow: Izdatel'stvo "Pero", 2021. pp. 551. (In Russ.).

19. Boysunov B.P., Korotkova L.A. [Formation of radio signals using frequency conversion with one sideband]. *Dostizheniya nauki i obrazovaniya – Achievements of science and education*. 2021. No. 3 (75). pp. 21–23. (In Russ.).
20. Dvornikov S.V. *Teoreticheskie osnovy sinteza bilinejnyh raspredelenij* [Theoretical foundations for the synthesis of bilinear distributions]. St. Petersburg: Izd-vo Politehn. un-ta, 2007. 268 p. (In Russ.).
21. Egorov S.B., Gorbachev R.I. [Analytical model of a noise signal with slow spectral-amplitude modulation for a passive noise radar]. *Trudy Sankt-Petersburgskogo gosudarstvennogo morskogo tehnickeskogo universiteta* [Proceedings of the St. Petersburg State Marine Technical University]. 2020. vol. 1. No. 52. pp. 6. (In Russ.).
22. Kulkarni A., Kothavade S., Patel D. Assessment of SSB, Modified-SSB and VSB Modulation Techniques based on Modulation Index, Extinction Ratio, Chromatic Dispersion and Received RF power. *International Conference on Communication information and Computing Technology (ICCICT)*. 2021. pp. 1–7. DOI: 10.1109/ICCICT50803.2021.9509947.7.
23. Dvornikov S.V., Pshenichnikov A.V. [Formation of spectrally efficient signal structures in radio channels of data transmission of control and measuring complexes]. *Izvestiya vysshih uchebnyh zavedenij. Priborostroenie – Instrumentation*. 2017. vol. 60. No. 3. pp. 221–228. DOI 10.17586/0021-3454-2017-60-3-221-228. (In Russ.).
24. Bobkov V.I., Snytko Ju.N. [Algorithm for generating an analytical signal of an infrared gas analyzer resistant to rolling]. *Jenergetika, informatika, innovacii: Sbornik trudov XI Mezhdunarodnoj nauchno-tehnicheskoy konferencii* [Energy, Informatics, Innovations: Proceedings of the XI International Scientific and Technical Conference]. Smolensk: Universum; filial FGBOU VO «NIU «MJeI», 2021. vol. 1. pp. 63–67. (In Russ.).
25. Kobayakov R.S., Novozhilov R.N., Pisarev I.A., Zhiglov A.V., Medvedev S.Ju. [Some methods for improving the accuracy of phase instability compensation in the transmission of frequency and time signals]. *Trudy Instituta prikladnoj astronomii RAN* [Proceedings of the Institute of Applied Astronomy RAS]. 2021. no. 58. pp. 36–40. DOI 10.32876/ApplAstron.58.36-40. (In Russ.).
26. Huang C., Chan E. Photonic techniques for generating a single RF sideband with no second order sidebands. *IEEE Photonics Journal*. 2022. vol. 14. no 1. DOI: 10.1109/JPHOT.2021.3123168.
27. Weber P. The History of Single Sideband Modulation Archived 2004-01-03. Wayback Machine, 2004.
28. Kahn L. Single-sideband transmission by envelope elimination and restoration. *Proceedings of the IRE*. 1952. vol. 40(7). pp. 803–806. DOI: 10.1109/JRPROC.1952.273844.
29. Dvornikov A.S., Gudkov M.A., Ayukov B.A., Fedosov A.Ju., Podgornij A.V., Zasedatelev A.N., Dvornikov S.V., Krjachko A.F., Pshenichnikov A.V. [Analysis of the noise immunity of transmissions with single-sideband modulation in channels with fluctuation noise] *Voprosy radioelektronics. Serija: Tehnika televidenija. – Questions of radio electronics. Series: TV Technique*. 2022. no. 4. pp. 58–64. (In Russ.).
30. Surzhikov V.F., Kompaniitsev A.V. [Physical modeling of digital channels of microwave communication with a binary amplitude-shift keyed signal in the MATLAB environment]. *Mir nauki, kul'tury, obrazovaniya – World of science, culture, education*. 2022. no. 1(92). pp. 119–122. DOI 10.24412/1991-5497-2022-192-119-122. (In Russ.).
31. Dvornikov S.V. [Demodulation of signals based on the processing of their modified time-frequency distributions]. *Cifrovaja obrabotka signalov – Digital signal processing*. 2009. no. 2. pp. 7–11. (In Russ.).

Dvornikov Sergey V. — Ph.D., Dr.Sci., Professor, Department of radio engineering and optoelectronic complexes, Federal State Autonomous Educational Institution of Higher Education «St. Petersburg State University of Aerospace Instrumentation»; Professor of the department, Department of radio communications, Federal State Military Educational Institution of Higher Education «Military Telecommunications Academy named after the Soviet Union Marshal Budienny S.M.». Research interests: construction of noise-protected radio communication systems, formation and processing of signals of complex structures. The number of publications — 448. practicedsv@yandex.ru; 3, Tikhoretsky Av., 194064, St. Petersburg, Russia; office phone: +7(812)247-9400.

Dvornikov Sergey S. — Ph.D., Associate professor, Department of design and technologies of electronic and laser means, Federal State Autonomous Educational Institution of Higher Education «St. Petersburg State University of Aerospace Instrumentation»; Researcher, Research department, Federal State Military Educational Institution of Higher Education «Military Telecommunications Academy named after the Soviet Union Marshal Budienny S.M.». Research interests: theory of signal transmission, spectral efficiency of signals, noise immunity of control and communication channels of radio engineering systems. The number of publications — 114. dvornik.92@mail.ru; 3, Tikhoretsky Av., 194064, St. Petersburg, Russia; office phone: +7(812)247-9811.

Zheglov Kirill — Post-graduate student, Federal State Autonomous Educational Institution of Higher Education «St. Petersburg State University of Aerospace Instrumentation». Research interests: decimeter radio communication systems, management of the frequency-energy resource of radio links. The number of publications — 2. zheglov.k@gmail.com; 67A, Bolshaya Morskaya St., 190000, St. Petersburg, Russia; office phone: +7(812)247-9400.

H.V. NGUYEN, N. TAN, N.H. QUAN, T.T. HUONG, N.H. PHAT
**BUILDING A CHATBOT SYSTEM TO ANALYZE OPINIONS OF
ENGLISH COMMENTS**

Nguyen H.V., Tan N., Quan N.H., Huong T.T., Phat N.H. **Building a Chatbot System to Analyze Opinions of English Comments.**

Abstract. Chatbot research has advanced significantly over the years. Enterprises have been investigating how to improve these tools' performance, adoption, and implementation to communicate with customers or internal teams through social media. Besides, businesses also want to pay attention to quality reviews from customers via social networks about products available in the market. From there, please select a new method to improve the service quality of their products and then send it to publishing agencies to publish based on the needs and evaluation of society. Although there have been numerous recent studies, not all of them address the issue of opinion evaluation on the chatbot system. The primary goal of this paper's research is to evaluate human comments in English via the chatbot system. The system's documents are preprocessed and opinion-matched to provide opinion judgments based on English comments. Based on practical needs and social conditions, this methodology aims to evolve chatbot content based on user interactions, allowing for a cyclic and human-supervised process with the following steps to evaluate comments in English. First, we preprocess the input data by collecting social media comments, and then our system parses those comments according to the rating views for each topic covered. Finally, our system will give a rating and comment result for each comment entered into the system. Experiments show that our method can improve accuracy better than the referenced methods by 78.53%.

Keywords: chatbot, offensive comments, behavioral culture, online, ontology, opinion mining, sentiment analysis.

1. Introduction. Many people use the Internet these days to communicate information. Information disseminates often, and many people's thoughts and comments are expressed. Therefore, taking into account and comprehending these remarks are beneficial. Therefore, there are numerous researches on user opinions in online journals [1] and numerous programs to study people's psychology and ideas can use with social networks like Twitter and Facebook [2]. For instance, keeping an eye on a particular brand's reputation on social media might give candidates an insight into the ambitions of voters, allowing them to adapt their speeches and actions. Financial market analysis can also use the comment analysis system, so similar to the stock market.

Nowadays, a ChatBot is a computer program that conducts an instant messaging conversation [3]. It can automatically answer questions and handle situations. In a ChatBot, the creators' algorithm determines the scope and complexity of the chatbot. It uses in various applications such as e-commerce, customer service, healthcare, banking and finance, and entertainment.

From another perspective, artificial intelligence and Natural Language Processing (NLP) integrated with machine learning algorithms play a significant role in today's technology. The survey on artificial intelligence in chatbots is based on a computer program that uses artificial intelligence to mimic human decision-making while providing various services [4]. In [4], the paper provides a survey based on multiple platforms used to build a chatbot to deliver various services to various users. The designed techniques used to create the chatbot are determined by the services that will be provided to the users. The chatbot will gain experience by learning from previous experiences and employing various algorithms. The data can be trained to the chatbot, allowing it to check with the knowledge base and to provide accurate answers to the user's query via client-side applications.

The research [5] on the structure of a chatbot using artificial intelligence (AI) aims to assess, diagnose, and recommend immediate safety and prevention measures for patients who have been exposed to nCOV-19 and acts as a virtual assistant to assist in measuring the severity of the infection through symptom analysis and connecting with competent medical facilities as it moves into the severe stage. In addition, some researches use Natural Language Processing (NLP) and Deep Learning (DL) techniques to develop a chatbot that can engage in interactive conversations with visitors during the MPU opening day [6], or AI-based chatbots to engage customers [7]. The authors recommend Restaurant Chatbots [8], Chatbot Using AI in the Healthcare Service Market [9], and Conversational AI [1] in the restaurant, medical, or communication industries. Furthermore, in the art exhibition, a group of authors presented an NLP-based solution for conversational agents [11]. In general, the authors have mentioned a lot of words, all of which have produced significant results, but using NLP to analyze and evaluate opinions has yet to be mentioned, and the evaluation analysis still needs to be improved in terms of comment evaluation.

With its clear benefits, Chatbot has become more and more popular. With the advent of technology 4.0, people have integrated more into the virtual and real world. Because of the rapidly increasing needs of society, it requires a large amount of manipulation, but choosing a large number of valuable comments for ourselves is essential. In this article, we propose a chatbot system that performs the following tasks:

- First: Collect and identify positive comments and remove negative ones;
- Second: Use many comments to evaluate the quality of the words;
- Third: Utilize an evaluation method to choose valuable comments for businesses.

The remaining paper is as follows. Section 2 gives an overview of the related work. The theory background is described in Section 3, followed by criteria for evaluating comments in Section 4. Next, the proposed model is presented in Section 5. Section 6 provides an evaluation. Finally, the paper is concluded in Section 7.

2. Related Work. According to the findings, chatbots are being expanded into almost every field, such as [12], an educational chatbot for the Facebook Messenger platform. Similarly, in a study [13], the software would also ask questions based on the candidate's previous responses, using a Natural Language Processing (NLP) model, which is very useful in this process. Following the interview, the software would analyze the data gathered to determine the best candidate for the offered position. As a result, the project JARO chatbot aims to simplify the hiring process. On the other hand [14], conversationally built with technology in mind, automated medical chatbots have the potential to reduce healthcare costs while improving access to medical services and knowledge. The method created a diagnosis bot that converses with patients about their medical questions and problems to provide an individualized diagnosis based on their diagnosed manifestation and profile. A practical method for determining discoverability and features such as language, subject matter, and developer platform [12]. Or for the patients, a framework that acts as a virtual assistant is created using ML algorithms. It can predict symptoms, recommend doctors, and investigate patient treatments by interacting with them – efficient patient health care with encouraging results [15].

On the one hand, the system creates a chatbot for academic purposes using NLP and ML that various educational institutions can use. There are two modes available: audio and text. Instead of being placed on the inquiry desk's waiting list, users can interact with the bot. The same question is asked in various forms to test accuracy [16]. As a result, the plan is to combine intent classification and natural language processing to create an interactive user interface and a chatbot. The model is intended to recognize user's queries and generate SPARQL queries [17]. Deep learning techniques were used to develop an online video lecture assistant that improves Q&A data quality by incorporating multiple chatbots from various perspectives for a single video [18].

On the other hand, another approach is building a Chatbot model for Vietnamese comment management [19]; the author has also mentioned some changes to the algorithm proposal and significant improvements. However, this method is limited to this study's scope – Vietnamese language research. In this work, our approach is more extensive, and we develop an assessment

based on the views of all the social network members. In addition, our method is based on the analysis of reviews and personal opinions about one or more products that are extended on social networks.

Furthermore, one report indicated that the chatbot design is primarily based on the device learning the rule. There are three steps to enforcing this chatbot. Raw records are pre-processed at the start. As a result, a data set is created. The splitting process occurs during the second step [20]. Some other techniques are used, such as [21]. The authors use an academic website, for example, how the system quickly calculates the bully word or no longer. In [21], various techniques are used here, including device mastery, fuzzy judgment, sample matching, and sentiment evaluation.

Recently, according to a survey by [22], there are 74 articles suggested chatbots. The studies mainly focused on application, methodology (methods used, sample size, sample type, and countries studies), and bibliometrics (publishing, citation, and spotlight agency). The main objective is to conduct a systematic review of high-quality journal research articles to summarize the current state of research on chatbots to identify their role in digital business transformation. On the contrary, we conducted a deep dive to assess the views of the English commentaries, from which to evaluate and give helpful information for businesses when promoting products on social networks.

As described in [23], a Chatbot application is a direct communication channel between the company and the end user in various fields, such as e-commerce or customer service. The authors used Xatkit in this paper. Xatkit solves these issues by providing a set of platform-independent Domain-specific Languages for defining chatbots. Xatkit also includes a runtime engine that deploys the chatbot application and manages the defined conversation logic on selected platforms. However, evaluating our opinion with this application is difficult because Xatkit depends on the accompanying tool, whereas our method is independent and has strong analytical and adaptive capabilities. The same experience-based research using Evatalk[24] is still limited when it only focuses on measuring people's satisfaction, not analyzing the human point of view carefully.

In addition to the articles on application development, there are a few articles on the exploratory potential of chatbots in providing online emotional support to people based on stress triggers. That is quite exciting; the authors have developed a social interaction agent based on empirical research to converse with stressed people seeking mental support. The author also addressed chatbot questions in helping users deal with stressful situations in [25]. In a similar study [26] the authors used queries developed on QA forums by Software Engineering practitioners. Both of the primary research methods

are very intriguing. Nonetheless, both investigate the Software Engineering practitioner's feelings or processing.

Chatbot communication research [27], in which robots communicate with humans in natural language in an open domain, has made significant progress. However, it still has several unsolved issues, such as a need for more diversity and contextual relevance. Based on the retrieved prototype, the author proposed a retrieval polishing model (RP) to generate feedback polishing. The method relies on the response receiver and is focused on selecting the prototype for contextual retrieval. However, while this method improves improve fluency, contextual relevance, and response diversity significantly based on the number of prototypes, it results in a too complex system when dealing with a large sample flow. However, the authors wanted to know how/if the usability of the significantly improved method [28], based on the SOCIO chatbot prototype model, had changed. The author also attempts an empirical evidence-based evaluation of the usability of SOCIO V1 to the updated version, which necessitates comprehensive verification of test results to change performance, efficiency, satisfaction, and quality. That is feasible when evaluating with a small margin for error for experimentation, but there are still many risks when measuring only one parameter without measuring another.

In general, with the parsing technique of English sentences and words, with a recommendation system, we have improved the accuracy up to 78.53%.

3. Theory background. Natural Language Processing (NLP) literature with a variety of language analysis techniques, which only a small subset has been observed with regularity in sentiment mining. The most common is part of speech tagging [29], but there are also papers detailing classifiers using resolution [30] and even using a full syntactic parse tree [31], or Nasukawa et al [31] proposes a method to apply parsing to sentiment analysis.

In this part, we conduct Natural Language Processing (NLP) as follows:

- First: We analyze sentences and divide the comment sentences into words;
- Second: We use The Penn Treebank POS tags are divided into three categories: adjectives, nouns, and verbs;
- Third: Using SentiWordNet can be tweaked by using these tags to look at the meaning of words that match their POS tags.

Finally, to evaluate a comment, we also do a grammatical analysis to assess the statement. From there, we analyze and, based on the collected words, consider other people's comments online in society to conclude. We analyze and rewrite the sentence based on the conditions listed in Table 1.

Table 1. Meaning of the symbols

Type	Meaningful	Type	Meaningful
<i>S</i>	Sentence	<i>RP</i>	Particle
<i>NP</i>	Noun Phrase	<i>LST</i>	List marker
<i>PP</i>	Prepositional Phrase	<i>PRT</i>	Particle
<i>VP</i>	Verb Phrase	<i>UCP</i>	Unlike Coordinated Phrase
<i>CC</i>	Coordinating conjunction	<i>SYM</i>	Symbol
<i>CD</i>	Cardinal number	<i>TO</i>	to
<i>DT</i>	Determiner	<i>UH</i>	Interjection
<i>EX</i>	Existential there	<i>VB</i>	Verb, base form
<i>FW</i>	Foreign word	<i>VBD</i>	Verb, past tense
<i>IN</i>	Preposition/subordinate conjunction	<i>VBG</i>	Verb, gerund/present participle
<i>JJ</i>	Adjective	<i>VBN</i>	Verb, past participle
<i>JJR</i>	Adjective, comparative	<i>VBP</i>	Verb, non-3rd ps. sing. present
<i>JJS</i>	Adjective, superlative	<i>VBZ</i>	Verb, 3rd ps. sing. present
<i>LS</i>	List item marker	<i>WDT</i>	wh-determiner
<i>MD</i>	Modal	<i>WP</i>	wh-pronoun
<i>NN</i>	Noun, singular or mass	<i>WP\$</i>	Possessive wh-pronoun
<i>NNP</i>	Proper noun, singular	<i>WRB</i>	wh-adverb
<i>NNPS</i>	Proper noun, plural	<i>PRP</i>	Personal pronoun
<i>NNS</i>	Noun, plural	<i>PRP\$</i>	Possessive pronoun
<i>PDT</i>	Predeterminer	<i>RB</i>	Adverb
<i>POS</i>	Possessive ending	<i>RBR</i>	Adverb, comparative
<i>RBS</i>	Adverb, superlative		

Our system is designed to use many comments to evaluate the quality of comments for better analysis and evaluation using the comments entered into the system. Furthermore, our system provides an evaluation method to select valuable opinions for businesses, thereby bringing quality products to users. For example, we can see that we will rewrite the sentence in a parsed form in Figure 1, and the sentence above is rewritten in the form of TreeBank in Figure 2.

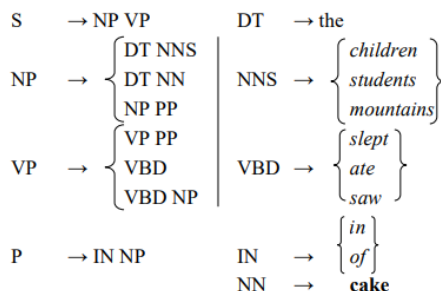


Fig. 1. Parsing by sentence

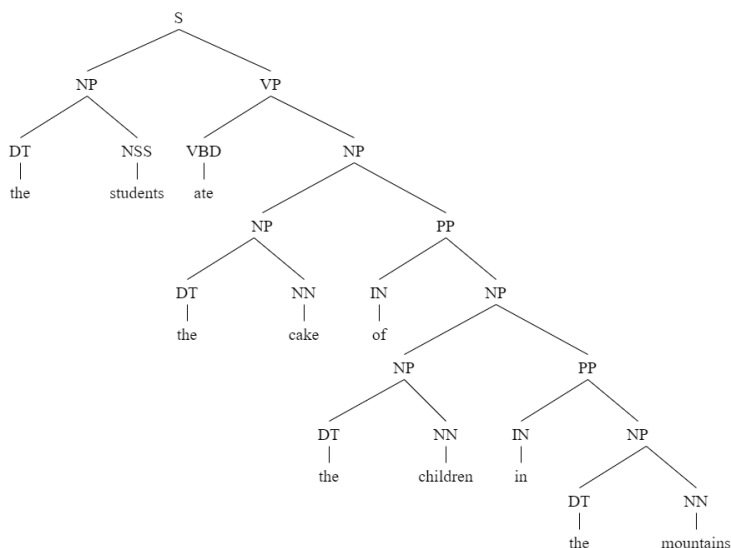


Fig. 2. Treebank

4. Design of Criteria for Evaluating Comments. The evaluation of emotions is quite essential to work, based on the general assessment of a specific topic is very reasonable when the opinions are more and more popular on social networking sites. The evaluation is always essential in the analysis for practical applications later. Based on this, we build a plan (Figure 3) to evaluate comments on social networks to find reasonable solutions for suitable products that society wants in the future.

Assessing the comments is a big challenge for Chatbots. Within the research scope of this paper, our system evaluates based on the parsing of the sentence and then evaluates the content and, finally, the commenter's evaluation conclusion. However, the problem encountered is that many comments need to be corrected; for example, short comments, acronyms, or grammatically incorrect, etc. Based on this content, we evaluate the comments to conclude whether they are negative or positive through previously trained vocabulary.

4.1. Negation. Negation is used to indicate that the comment is negative. In general, negative reviews will bring benefits to businesses.

In this paper, our Chatbot system will analyze comments based on the content and language processing analysis. The comments that our system evaluates can benefit businesses when the system can be a yardstick to assess products on social networking sites. We use categorical phrases to show negative comments through negative words. Such as "bad," "terrible," etc., in a sentence. For example, "This computer is terrible", the word "terrible" here indicates that the comment is negative. Here, we do not analyze whether the sentence is negative [32].

In natural terms, the evaluation is based on the proposed algorithm to determine the positive or the negative. The algorithm is based on a training dataset to make comment conclusions.

4.2. Positivity. Positivity is used to indicate that the comment is positive. We use categorical phrases to show positive comments through positive words. Such as "good," "excellent," etc., in a sentence. For example, "This computer is excellent," the word excellent here indicates that the comment is positive. Again, keep in mind that we do not analyze whether the sentence is positive [32].

Recently, many businesses often ignore responding to positive reviews because they think it is unnecessary, but in reality, you need to respond to all user experience reviews. That shows the brand's professionalism and effectively builds customer trust and loyalty. Using the following article [33] to learn how to respond to positive comments with Chatbot is to improve your business's customer service.

Moreover, through customer reviews, the enterprise can expect to increase or expand its production scale. On the other hand, mass-creating products without regarding to customers will directly impact the business and the company's interests.

5. Our proposed model. In this section, we propose a system structure in which the system is divided into several main parts, including Ontology and Preprocessing as described in Figure 3.

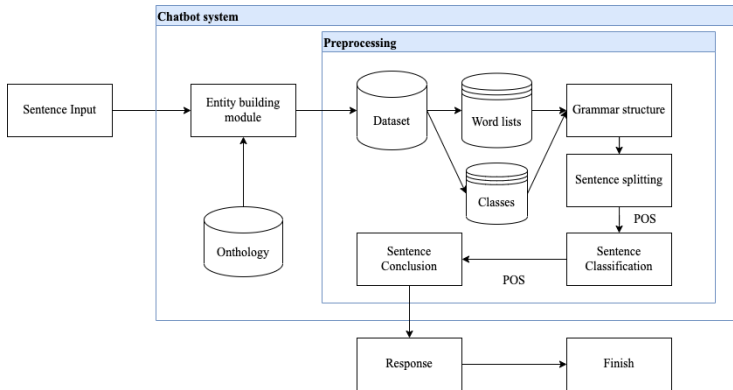


Fig. 3. Proposed System Structure

5.1. Ontology. In this part, we create a module (called: Entities building module) to take the product type as input. This module is in charge of retrieving the knowledge base from Ontology and extracting all entities from the corresponding branch of the product.

The entity is a critical component in the spam detection system, serving as the foundation for the search process and matching Ontology. A sentence can contain a single entity, multiple entities, or none. Preprocessing modules will perform and recognize this entity and save the entity identified in data preprocessing to assist the algorithm in identifying spam reviews. The entity here is not just the named entities but all entities in a sense defined by the researchers. We describe an entity as the word meaning, which brings specific knowledge in the reviews, to use these entities to find the product knowledge contained in the reviews. This definition states that adjectives and nouns are two types of words we have chosen to filter into the desired entity.

An ontology cannot cover all meaningful aspects of a field, so specific objectives are used for identifying spam reviews. Extracted entities are focused on product components or properties. As a result, the entity will be collected and distributed to the class groups based on their common characteristics following the statistics.

First, the system's input is a sentence, which is responsible for receiving a sentence to process and analyze whether that sentence is positive or negative. Next, the input data will pass through an Ontology used to create an entity (Entity building module), assisting the system in evaluating the words that the system has previously processed. This is the foundation for the search process. Ontologies search and match.

Furthermore, a sentence may contain one or more entities, or none, which can be understood as critical keywords to support faster and more accurate processing and assessment of sentence properties. The preprocessing modules will recognize these entities, calculate and save the probability results determining the positive or negative of the entity to support the matching of entities in the sentence.

5.2. Preprocessing. The preprocessing module analyzes the content and title of the review and produces the data required by the classification model. Figure 3 depicts the division of preprocessing work.

The essential input to the model is the content of the review. As a result, normalizing must be completed before proceeding with the other processing steps to create standard data sources and avoid error analysis.

There are numerous methods for extracting words from a text. We chose the way n-gram models of unigram, combined with the POS tagging model, for this study. Stanford University's POS tagging tools (Stanford POS Tagger) have a relatively large database and have been widely used in the study of language processing. We chose this tool for the word-splitting module because of its high accuracy and processing performance. To do this, we did the following simulation:

First, the dataset here includes Comments collected from sources on the Internet such as forums, social networks, websites, etc. Then, through Grammar Structure, label POS for each element of the sentence through the POS labeling model (Stanford University's POS tagger has a sizeable available database and has been applied extensively in natural language processing research courses). We chose this tool in processing word decomposition in sentences and structural analysis of components in sentences because it has relatively high accuracy and has been applied in many places in analyzing nature processing language.

Second, the processing of sentences after they have been separated into words will match with WordLists and Classes used to classify what type of sentence the sentence belongs. Our system divides into five sentence types: simple, conditional, comparative, compound, and special sentence. Depending on the type of sentence, the system has different ways of processing words in the sentence to get accurate results. Specifically, after being decomposed into words and labeled through the POS model, the sentence will be sorted into an array to browse each word in the sentence, from which the system will evaluate whether each word is in it. How likely it is that the word has such a character depends on whether the word is positive or negative. After evaluating each word in the sentence, the system will classify positive or negative comments.

A specific generative grammar aims to provide a set of rules that generate (or, more abstractly, license, predict, etc.) all phrase structure trees that correspond to English grammatical sentences. This means that only the word sequences standardized by a linguist and a syntactic structure description (phrase structure) would be considered correct and complete sentences. The rules also include claims about the constituent structure of English.

5.3. Database design. Many languages are used to write social networks. As a result, it must apply to all languages and libraries. It must be constantly updated and expanded to ensure the chatbot's effectiveness. There are two factors required to ensure:

- The number of English words must be large enough to avoid confusion.
- The program can handle frequent updates without requiring cumbersome manipulations.

For example, English is challenging to evaluate due to its large area of vocabulary and grammar. Specifically, it contains 12 tenses [34] and complex structures. Moreover, with the number of comments increasing, we focus on assessing the opinions of the comments. Thus, a dictionary can break sentences into meaningful words and phrases. We calculate and draw opinion conclusions based on evaluation and analysis based on meaningful words or phrases.

5.4. System model. The opinion analysis system consists of 3 main components in Figure 3:

- Data set: Includes comments collected from sources on the Internet such as forums, social networks, websites, etc.
- Sentence classification in Figure 4: it is the function of classifying sentences commenting on their sentence types, such as simple sentences, compound sentences, negative sentences, comparative sentences, etc.

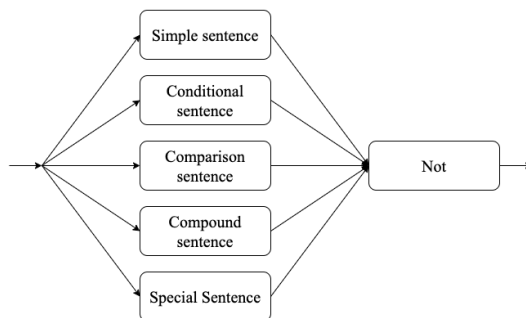


Fig. 4. Sentence classification model

– **Result:** The final feature will decide which group the comments belong to, positive, negative, or neutral.

The classification model has five classification features of basic sentence types, including simple sentences, conditional sentences, comparative sentences, compound sentences, and special sentences. In addition, for each of the above categories, there is a negative test in the sentence.

On the one hand, in the chatbot system, there is an equally important part of machine learning training.

Read data. The system reads data from the existing data file and adds data to lists. Word is a list containing POS-tagging for each word in a sentence, and classes is a list containing the classification type of the sentence here is positive. And the negative document is a list that includes both words and classes. In addition, there are ignore-words, a list of elements that can ignore in a sentence, and it does not affect the evaluation of sentence classification.

Word tokenization. In this step, we will look at each token (i.e., word by word of a sentence) and try to predict the phrase type of this token. It can be a noun, a verb, an adjective, and so on. Knowing each word's role in a sentence can help determine the meaning. The original part-of-speech model is trained by feeding millions of English sentences with each word of speech tagged, reproducing those behaviors. Note that, this model is purely based on statistics - it needs to understand what these words mean, as humans do. It only knows how to guess a part (i.e., a comment) of speech based on similar sentences and observations it has seen before.

Text lemmatization. Lowercase each word and remove duplicate comments. Lemmatization is to bring words to their original format by using a lookup table of the actual vocabulary of the words in the sentence. There may be rules for dealing with comments we can have never seen before.

Create training model. The prepared dataset is forwarded to the sequential build model to the respective weight of the process. Then the system will save the trained model of the system for deployment. Below is a detailed description of the proposed algorithm.

– Step 01: Layer 1, 128 neurons, dropout(0.01). Layer 2, 64 neurons, dropout(0.01).

– Step 02: Compile the network with parameters to train the network to evaluate the optimization.

– Step 03: Learning Rate = 0.001.

On the other hand, we define classify sentences in the Chatbot system as follows:

5.4.1. Simple sentences. Simple sentences express a straight point of view, in which the sentence can contain at most one verb (or possibly none, e.g., good or bad). Example sentences are as follows: Samsung is good.

After going through the POS tagging system, this sentence will give the following results: Samsung_NN is_VBZ good_JJ. The sentence contains one subject (Samsung), one verb (is), and one adjective (good).

The sentence above is positive because we see that there is the adjective good, which is a positive word.

Therefore, to determine the point of view on a single sentence, we need to identify the adjectives (or adverbs) in the sentence that indicate positive or negative. Currently, we use two dictionaries with more than 6,800 adjectives expressing opinions (1 set of positive adjectives and 1 set of negative adjectives).

For this type of sentence, we build a Chatbot system that handles the following steps in turn:

- Step 01: Label sentences (POS tagging).
- Step 02: Find adjectives.
- Step 03: Find out which order the adjective belongs to (positive/negative).
- Step 04: Conclude.

Negative on a simple sentence. Let's consider the example "The Dell's sound quality is not pretty good" in a sentence containing the adjective good. This is a positive word, so with the steps above, we would conclude that this is a positive opinion sentence, but it is clear that the above sentence has the exact opposite opinion. As we have seen in the above sentence containing the negative word (not), this is the key word to determine the point of view of the above sentence. Without it, the above sentence is entirely positive, so we need to check if it contains a negative word. We have to reverse the point of the sentence.

5.4.2. Conditional sentences. Conditional sentences are sentences that describe hypothetical effects or situations and their consequences. In English, various conditional joins can be used to form sentences. A conditional sentence consists of two clauses: A conditional clause and a consequential clause, which depend on each other. This relationship has important implications for describing the opinions of sentences. For the sake of simplicity, comment words (also known as opinion words) (e.g., great, beautiful, bad) alone cannot determine the opinion in a sentence. A conditional sentence can contain many comment words or phrases but may not express an opinion [35].

About here, we will describe a few examples to parse this sentence form:

Example 1: "If someone makes a beautiful and reliable car, I will buy it" expresses the opinion that there is no sympathy for any car, even though "beautiful" and "reliable" are two words that mean positive comments. However, this does not mean that a conditional cannot express opinions/comments.

Example 2: "If you are looking for a phone with good voice quality, don't buy this Nokia phone." It has a negative connotation with the "voice quality" of "Nokia phone," although here there is a positive comment "good" in the conditional on complementing "voice quality." However, according to the above sentence, the point of view is the opposite.

Furthermore, we noticed that most conditional sentences contain the word If. However, there are also many other dependent join words such as: even if, unless, in case, assumption/supposing, as long as, etc. Corresponding to each word (phrase) also gives different evaluation conditions.

For this type of sentence, we analyze some assessment skills as follows:

Firstly: Find the position of the commented words/phrases.

Secondly: POS tags comment words; comment words can use in some of the following cases, not all sentences containing an opinion, for example, I trust Motorola, and He has a trust fund both contain the word trust. But only the previous sentence has an idea.

Thirdly: Find words that do not indicate an opinion; similar to how to find comment words related to an idea, here are some words that imply the opposite. Words like wondering, thinking, and debating of the user ask questions or expresses doubt.

Fourthly: Tense patterns and basic tenses use to create a set of features. We identify the first word in both conditional and consequential clauses by searching for related words using POS tags. Fifthly: In sentences with or without the characters '?' and '!'.
Sixthly: Conditional associations used in sentences (if, even if, unless, only if, etc.) are considered a feature.

Seventhly: Conditional clause length and consequences. Using simple punctuation rules in the language, we automatically segment sentences into conditional and consequential clauses. Usually, the conditional is very short and does not affect the statements of opinion.

Eighthly: The use of negative words like not, don't, and never, etc., often change the opinion of a sentence; for example, adding negative words before comment words can change the idea of a sentence from positive to negative.

In addition, to handle more complex sentences, we propose to analyze and perform the following in Table 2:

If the condition contains VB/VBP/VBZ → 0 conditional;

If consequent contains VB/VBP/VBS → 0 conditional;

If the condition contains VBG → 1st conditional;

If the condition contains VBD → 2nd conditional;

If the condition contains VBN → 3rd conditional.

Table 2. Tenses to identify comparative sentence types

Type	Linguistic	Condition POS tags	Consequent POS tags
0	If + simple present → simple present	VB/VBP/VBZ	VB/VBP/VBZ
1	If + sim present → will + bare infinitive	VB/VBP/VBZ/VBG	MD+VB
2	If + past tense → would + infinitive	VBD	MD+VB
3	If + past perfect → present perfect	VBD/VBN	MD+VBD

5.4.3. Comparative sentences. A helpful note about comparative sentences is that in each such sentence, there is usually a comparative (e.g., “better,” “worse,” and –er) or superlative (e.g., “best,” “worst,” and the word –est). The objects to compare often appear on either side of the word comparison. An excellent sentence can have only one entity, e.g., “Camera X is the best.” For simplicity, the author uses comparative words (sentences) and superlative words (sentences).

The comparative words mainly identify the preferred entities in a comparative sentence in the sentence. Some comparative words explicitly indicate user preferences, for example, “better,” “worse,” and “best.” We call such words opinionated comparative words. For example, given the following sentence: “the picture quality of Camera X is better than that of Camera Y,” Camera X is a popular product due to its comparative point of view word “better.”

However, many comparators do not have a specific point of view, or their opinion (positive or negative) depends on the context or the application domain. For example, the word “longer” does not hold the conventional wisdom to show that the length of some feature of one entity is greater than that of another. However, it can represent a desired (positive) or undesirable (negative) state in a particular context. For example, given the following sentence: “the battery life of Camera X is longer than Camera Y,” “longer” clearly wants to represent the desired state of “battery life” (although the object in the sentence does not support a clear opinion). “Camera X” is also the favorite with “battery life” among the compared cameras. The opinions in the above sentence are called implicit opinions.

Sentences with opinion words (for example, "better" and "worse") are usually easy to deal with. The key to solving our's problem is to identify the opinions (positive or negative) of the context-dependent comparative words. To conclude, two questions arise: (1) what is context, and (2) how can context be used to help determine the opinion of a comparative word?

The simple answer to question (1) is the whole thing. However, the entire sentence and the context could be more straightforward because much irrelevant information, which can confuse the system. Intuitively, we want to use the most accurate context that can confirm the point of view of the word comparison. So, the context must have entity features to compare and find the comparison word. To answer the second question (2), we need more information or knowledge because there is no way a computer program can solve the problem by analyzing the sentences themselves. In this article, we propose to use customer evaluation information on the Chatbot system to help solve the problem.

5.4.4. Compound sentences. A compound sentence is a sentence that contains the following words: but, although, however, and nevertheless. For this type of sentence, the meaning will often be the opposite of the user's original desire.

Example: This is great. However, I hate it; although this is good, I won't buy it, etc.

For this type of sentence, we propose to evaluate according to the following steps:

- Step 01: Identify sentences containing but, although, however, and nevertheless, not.
- Step 02: Split the above sentence into simple sentences.
- Step 03: Processing simple sentences.
- Step 04: Conclusion.

5.4.5. Special sentences. Special sentences are sentences that contain special words or contain dichotomous words. For example:

- Dell has a fast processor.
- Dell has a fast battery.

We can see that fast here is fast, which is fine if you have a fast processor, but it is terrible for battery life. As for how to deal with this, we have established a dictionary of associated words and their views. Here, we can understand that "processor+fast->Positive", "Battery+fast->Negative".

6. Evaluation

6.1. Experimental Settings. In this section, we proceed to install and test. The experiment was written in Python and tested on a 64-bit Windows 11 Pro computer with the following configuration: 16384MB of RAM and

an 11th Gen Intel(R) Core(TM) i5-11400H 2.70GHz, 2688 Mhz, 6 Core(s), 12 Logical Processor(s). Besides, we use negative and positive words with 6800 words [32]. The data we collect are products, including computers (531 sentences), routers (879 sentences), and Speakers (689 sentences) [36] to test.

When working with text data, we must perform various preprocessing steps on the data before building a machine learning or deep learning model. We must apply various operations to preprocess the data based on the requirements. Tokenization is the most fundamental and first thing you can do with text data. Tokenizing is dividing a large text into small parts, such as words. We iterate through the patterns, tokenizing the sentence with the `nlTK.word.tokenize()` function and appending each word in the words list. We also make a class list for our tags. We will now lemmatize each word and remove any duplicates from the list. Lemmatizing converts a word to its lemma form and creates a pickle file to store the Python objects used during prediction.

Furthermore, in Table 3, we also test with ten comments that we give manually and evaluate and compare our system prices.

Table 3. Result of evaluation between the system Chatbot and Human

N.	Some English comments on social networks	Evaluates		Conclusion
		Our ChatBot	Human	
1	Hardware with very stable performance.	Positive	Positive	True
2	This program contains suspicious, malicious code.	Negative	Negative	True
3	More memory will run smoother.	Positive	Positive	True
4	If I leave this computer here for a month, it will malfunction.	Negative	Negative	True
5	In terms of performance, this year's gaming computer is better.	Positive	Positive	True
6	This program performed worse than expected compared to the previous run.	Negative	Negative	True
7	It's been a long time since he used his calculator, but it still works fine with minor repairs.	Positive	Positive	True
8	He tried to remove malicious programs from his computer, but it failed and even crashed his computer.	Negative	Negative	True
9	Fantastic, it works without problems!	Positive	Positive	True
10	Dammit, the laptop is so old!	Negative	Negative	True

The results show that the system we evaluated correctly with our manual evaluation reached 100%. To calculate the above result, we use the following formula [19]:

$$\text{Correct detection rate} = \frac{\text{The correct number of comments}}{\text{Total number of comments}}. \quad (1)$$

- **Correct detection rate:** Accuracy of comments on the system.
- **The correct number of comments:** The correctly evaluated sentences that the input already knows.
- **Total number of comments:** The total number of whole sentences tested.

In this part, the proposed is compared with two reference methods, A Chatbot for Changing Lifestyle in the Education method (called **ACCLE**) and Interactive Transport Enquiry with AI Chatbot (called **ITEAI**). A summary of those reference methods is described as follows.

– **ACCLE** [37]: The author proposes a Chatbot system to serve to learn between teachers and students. The system is implemented by having students ask questions in the Chatbot in the form of text. Then the system processes it through natural language processing and deep learning technology. Finally, the system processes to answer the student. However, this system only serves schools and has yet to analyze the respondents' emotions.

– **ITEAI** [38]: Similar to the **ACCLE** method, this method also builds a Chatbot system that confirms the user's current location and final destination by asking some questions. The design of this method checks the user's query and extracts the appropriate entries from the database. This approach aims at the receiver to get all the information about the bus name and number. Then so that the person can safely move to the desired location.

Although the methods used have their strengths, our approach has been evaluated based on training, information extraction, and evaluation based on human emotions to assess the overall and give good results for the desired user.

Based on the dataset used and having the correct sentence classification, we conducted a test with 230 (Negative: 67, Positive: 163) sentences with the data file Computer.txt, with Router.txt with 222 (Negative: 81, Positive: 141) sentences, and Speaker.txt with 284 (Negative: 63, Positive: 221) sentences rated as standard negative and positive (Figure 5).

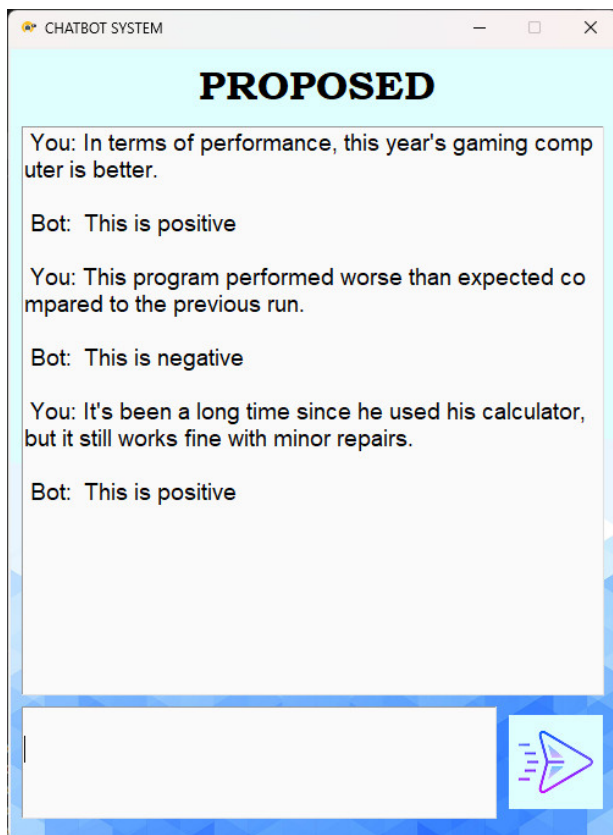


Fig. 5. GUI - The proposed

In this paper, we have referenced two related works. Unfortunately, our system gives better results than the two referenced methods in Table 4.

Table 4. Performance of the proposed and reference methods

Type	ACCLE	ITEAI	The proposed
The number of negative sentences	368	271	290
The number of positive sentences	2	290	288
Total number of comments	736	736	736
Rate of the positive sentence	50.27%	76.22%	78.53%
Build finished in 1 minute	10s	55s	46s
Build finished in	70.0679s	55.1462s	46.2025s

The results in Table 4 show that the proposed method always accounts for a higher percentage than the remaining methods when considering 736 sentences. The ACCLE approach shows low results when it comes to the average at 50.27%, while the ITEAI method gives results of 76.22% or more. The proposed method reached the lowest level at 78.53%.

7. Conclusion. We have built a Chatbot model to deal with some simple sentences, such as simple sentences, and comparison sentences with conditional and compound sentences, which are reliable, but for memorable sentences because there needs to be more time to solve the problem.

The paper has built an automatic evaluation model of opinion mining over a Chatbot system. This concept developed in response to the current issues that businesses face as social networks grow, but quality values remain limited. A series of document reviews to ensure consistency in all work, and the chatbot was determined to be the best model to meet the requirements. Chatbot research draws connections to learn more about emerging transient technologies and compatible algorithms such as Artificial Intelligence, Machine Learning, Python, and Natural Language Processing (NLP). The results show that our proposed method is up to 78.53%.

References

1. Scholz T., Conrad S. Opinion mining in newspaper articles by entropy-based word connections. Proceedings of the 2013 Conference on Empirical Methods in Natural Language Processing. Seattle, Washington, USA: Association for Computational Linguistics, 2013. pp. 1828–1839.
2. Feldman R. Techniques and applications for sentiment analysis. Communications of the ACM. 2013. vol. 56. pp. 82–89.
3. Mauldin M. Chatbot. Available at: <https://en.wikipedia.org/wiki/Chatbot>. (accessed: 10.11.2022).
4. Ganesan M., Deepika C., HarievashiniB., Krithikha A., Lokhratchana B. A survey on chatbots using artificial intelligence. International Conference on System, Computation, Automation and Networking (ICSCAN). 2020. pp. 1–5.
5. Singh S., Kaur M., Tanwar P., Sharma S. Design and development of conversational chatbot for covid-19 using nlp: an ai application. 6th International Conference on Computing Methodologies and Communication (ICCMC). 2022. pp. 1654–1658.
6. Liu F., Yang X., Wang Y. An interactive chatbot for university open day. IEEE 13th International Conference on Software Engineering and Service Science (ICSESS). 2022. pp. 214–217.
7. Rawat B., Bist A., Rahardja U., Aini Q., Ayu Sanjaya Y. Recent deep learning based nlp techniques for chatbot development: An exhaustive survey. 10th International Conference on Cyber and IT Service Management (CITSM). 2022. pp. 1–4.
8. Garg R., Riya R., Thakur S., Tyagi N., Basha K., Vij D., Sodhi G. Nlp based chatbot for multiple restaurants. 10th International Conference on System Modeling Advancement in Research Trends (SMART). 2021. pp. 439–443.
9. Wahal A., Aggarwal M., Poongodi T. Iot based chatbots using nlp and svm algorithms. 3rd International Conference on Intelligent Engineering and Management (ICIEM). 2022. pp. 484–489.

10. Adam M., Wessel M., Benlian A. AI-based chatbots in customer service and their effects on user compliance. *Electronic Markets*. 2021. vol. 31. no. 2. pp. 427–445.
11. Ceuca L., Rednic A., Chifu E. Safer museum guide interaction during a pandemic and further using nlp in human interactive museum visits: Museum guide chatbot. *IEEE 17th International Conference on Intelligent Computer Communication and Processing (ICCP)*. 2021. pp. 313–318.
12. Smutny P., Schreiberova P. Chatbots for learning: A review of educational chatbots for the facebook messenger. *Computers Education*. 2020. vol. 151. p. 103862.
13. Purohit J., Bagwe A., Mehta R., Mangaonkar O., George E. Natural language processing based jaro-the interviewing chatbot. *3rd International Conference on Computing Methodologies and Communication (ICCMC)*. 2019. pp. 134–136.
14. Srivastava P., Singh N. Automatized medical chatbot (medibot). *International Conference on Power Electronics IoT Applications in Renewable Energy and its Control (PARC)*. 2020. pp. 351–354.
15. Polignano M., Narducci F., Iovine A., Musto C., De Gemmis M., Semeraro G., Healthassistantbot: A personal health assistant for the italian language. *IEEE Access*. 2020. vol. 8. pp. 107479–107497.
16. Kumari S., Naikwadi Z., Akole A., Darshankar P. Enhancing college chat bot assistant with the help of richer human computer interaction and speech recognition. *International Conference on Electronics and Sustainable Communication Systems (ICESC)*. 2020. pp. 427–433.
17. Ait-Mlouk A., Jiang L. Kbot: A knowledge graph based chatbot for natural language understanding over linked data. *IEEE Access*. 2020. vol. 8. pp. 149220–149230.
18. Makhkamova O., Lee K.-H., Do K., Kim D. Deep learning-based multi-chatbot broker for qa improvement of video tutoring assistant. *IEEE International Conference on Big Data and Smart Computing (BigComp)*. 2020. pp. 221–224.
19. Nguyen Huu P., Do Manh C., Nguyen Trong H. Proposing chatbot model for managing comments in vietnam,” in *Industrial Networks and Intelligent Systems*. Eds. N.-S. Vo, V.-P. Hoang, Q.-T. Vien. Cham: Springer International Publishing, 2021. pp. 287–297.
20. Sinha S., Basak S., Dey Y., Mondal A. An educational chatbot for answering queries. *Emerging technology in modelling and graphics*. Springer, 2020. pp. 55–60.
21. Selvi V., Saranya S., Chidida K., Abarna R. Chatbot and bullyfree chat. *IEEE International Conference on System, Computation, Automation and Networking (ICSCAN)*. 2019. pp. 1–5.
22. Miklosik A., Evans N., Qureshi A. The use of chatbots in digital business transformation: A systematic literature review. *IEEE Access*. 2021. vol. 9. pp. 106530–106539.
23. Daniel G., Cabot J., Deruelle L., Derras M. Xatkit: A multimodal low-code chatbot development framework,” *IEEE Access*, 2020. vol. 8. pp. 15332–15346.
24. Santos G., De Andrade G., Silva G., Duarte F., Costa J., De Sousa R. A conversation-driven approach for chatbot management. *IEEE Access*. 2022. vol. 10. pp. 8474–8486.
25. Medeiros L., Bosse T., Gerritsen C. Can a chatbot comfort humans? studying the impact of a supportive chatbot on users’ self-perceived stress. *IEEE Transactions on Human-Machine Systems*. 2022. vol. 52. no. 3. pp. 343–353.
26. Abdellatif A., Badran K., Costa D., Shihab E. A comparison of natural language understanding platforms for chatbots in software engineering. *IEEE Transactions on Software Engineering*. 2022. vol. 48. no. 8. pp. 3087–3102.
27. Zhang L., Yang Y., Zhou J., Chen C., He L. Retrieval-polished response generation for chatbot,” *IEEE Access*. 2020. vol. 8. pp. 123882–123890.

28. Ren R., Perez-Soler S., Castro J., Dieste O., Acuna S. Using the socio chatbot for uml modeling: A second family of experiments on usability in academic settings. IEEE Access. 2022. vol. 10. pp. 130542–130562.
29. The stanford natural language processing (nlp). Available at: <https://nlp.stanford.edu/software/tagger.shtml>. (accessed: 02.11.2022).
30. Nicolov N., Salvetti F., Ivanova S. Sentiment analysis: Does coreference matter. AISB 2008 convention communication, interaction and social intelligence. 2008. vol. 1. p. 37.
31. Kanayama H., Nasukawa T., Watanabe H. Deeper sentiment analysis using machine translation technology. COLING 2004: Proceedings of the 20th International Conference on Computational Linguistics. 2004. pp. 494–500.
32. Hu M., Liu B. Mining and summarizing customer reviews. Proceedings of the tenth ACM SIGKDD international conference on Knowledge discovery and data mining. 2004. pp. 168–177.
33. Ko M.-C., Lin Z.-H. Cardbot: A chatbot for business card management. Proceedings of the 23rd International Conference on Intelligent User Interfaces Companion, ser. IUI '18 Companion. New York, NY, USA: Association for Computing Machinery, 2018. Available at: <https://doi.org/10.1145/3180308.3180313>. (accessed: 02.11.2022).
34. English tense system. Available at: <https://www.englishclub.com/grammar/verb-tenses-system.htm>. (accessed: 10.11.2022).
35. Narayanan R., Liu B., Choudhary A. Sentiment analysis of conditional sentences. Proceedings of the 2009 Conference on Empirical Methods in Natural Language Processing. Singapore: Association for Computational Linguistics. 2009. pp. 180–189.
36. Liu Q., Gao Z., Liu B., Zhang Y. Automated rule selection for aspect extraction in opinion mining,” in Proceedings of the 24th International Conference on Artificial Intelligence, ser. IJCAI'15. AAAI Press, 2015. p. 1291–1297.
37. Kasthuri E., Balaji S. A chatbot for changing lifestyle in education. Third International Conference on Intelligent Communication Technologies and Virtual Mobile Networks (ICICV). 2021. pp. 1317–1322.
38. Dharani M., Jyostna J., Sucharitha E., Likitha R., Manne S. Interactive transport enquiry with ai chatbot. 4th International Conference on Intelligent Computing and Control Systems (ICICCS). 2020. pp. 1271–1276.

Nguyen Viet Hung — Lecturer, East Asia University of Technology; Student, Faculty of information technology, Hanoi University of Science and Technology. Research interests: multimedia communications, network security, artificial intelligence, traffic engineering in next-generation networks, QoE/QoS guarantee for network services, green networking, applications. The number of publications — 11. hungnv@eaut.edu.vn; Ky Phu - Ky Anh, Ha Tinh, Viet Nam; office phone: +84(098)911-2079.

Nguyen Tan — Research assistant, East Asia University of Technology. Research interests: applications, data analysis. The number of publications — 1. tan25102000@gmail.com; Trung Dung - Tien Lu, Hung Yen, Viet Nam; office phone: +84(035)919-0216.

Nguyen Hong Quan — Research assistant, East Asia University of Technology. Research interests: applications, networks. The number of publications — 1. quan31nd@gmail.com; Xuan Thanh, Xuan Huong, Hanoi, Viet Nam; office phone: +84(082)220-6919.

Truong Thu Huong — Ph.D., Dr.Sci., Deputy head of the department, School of electrical and electronic engineering, Hanoi University of Science and Technology. Research interests: network security, artificial intelligence, traffic engineering in next-generation networks, QoE/QoS guarantee for network services, green networking, development of the internet of things ecosystems and applications. The number of publications — 84. huong.truongthu@hust.edu.vn; 1, Dai Co Viet St., Hanoi, Viet Nam; office phone: +84(243)869-2242.

Nguyen Huu Phat — Ph.D., Lecturer, Hanoi University of Science and Technology. Research interests: image and video processing, wireless networks, big data, intelligent transportation systems (ITS), the internet of things (IoT). The number of publications — 60. phat.nguyenhuu@hust.edu.vn; 1, Dai Co Viet St., Hanoi, Viet Nam; office phone: +84(243)869-2242.

Х.В. НГУЕН, Н. ТАН, Н.Х. КУАН, Ч.Т. ХЫОНГ, Н.Х. ПХАТ
**СОЗДАНИЕ СИСТЕМЫ ЧАТ-БОТОВ ДЛЯ АНАЛИЗА МНЕНИЙ
АНГЛОЯЗЫЧНЫХ КОММЕНТАРИЕВ**

Нгуен Х.В., Тан Н., Куан Н.Х., Хыонг Ч.Т., Пхат Н.Х. Создание системы чат-ботов для анализа мнений англоязычных комментариев.

Аннотация. Исследования чат-ботов значительно продвинулись за эти годы. Предприятия изучают, как улучшить производительность, принятие и внедрение этих инструментов, чтобы общаться с клиентами или внутренними командами через социальные сети. Кроме того, предприятия также хотят обращать внимание на качественные отзывы клиентов в социальных сетях о продуктах, доступных на рынке. Оттуда, пожалуйста, выберите новый метод для улучшения качества обслуживания своих продуктов, а затем отправьте его в издательские агентства для публикации на основе потребностей и оценки общества. Несмотря на то, что в последнее время было проведено множество исследований, не все из них затрагивают вопрос оценки мнений о системе чат-ботов. Основная цель исследования в этой статье — оценить человеческие комментарии на английском языке с помощью системы чат-ботов. Документы системы предварительно обрабатываются и сопоставляются мнения, чтобы предоставить заключения на основе комментариев на английском языке. Основанная на практических потребностях и социальных условиях, эта методология направлена на развитие контента чат-бота на основе взаимодействия с пользователем, что позволяет осуществлять циклический и контролируемый человеком процесс со следующими этапами оценки комментариев на английском языке. Сначала мы предварительно обрабатываем входные данные, собирая комментарии в социальных сетях, а затем наша система анализирует эти комментарии в соответствии с рейтингом просмотров по каждой затронутой теме. Наконец, данная система будет давать рейтинг и результат комментариев для каждого комментария, введенного в систему. Эксперименты показывают, что данный метод может повысить точность на 78,53% лучше, чем упомянутые методы.

Ключевые слова: чат-бот, оскорбительные комментарии, культура поведения, онлайн, онтология, анализ мнений, анализ настроений.

Литература

1. Scholz T., Conrad S. Opinion mining in newspaper articles by entropy-based word connections. Proceedings of the 2013 Conference on Empirical Methods in Natural Language Processing, Seattle, Washington, USA: Association for Computational Linguistics, 2013. pp. 1828–1839.
2. Feldman R. Techniques and applications for sentiment analysis. Communications of the ACM. 2013. vol. 56. pp. 82–89.
3. Mauldin M. Chatbot. Available at: <https://en.wikipedia.org/wiki/Chatbot>. (accessed: 10.11.2022).
4. Ganesan M., Deepika C., Hariavashini B., Krithikha A., Lokhratchana B. A survey on chatbots using artificial intelligence. International Conference on System, Computation, Automation and Networking (ICSCAN). 2020. pp. 1–5.
5. Singh S., Kaur M., Tanwar P., Sharma S. Design and development of conversational chatbot for covid-19 using nlp: an ai application. 6th International Conference on Computing Methodologies and Communication (ICCMC). 2022. pp. 1654–1658.

6. Liu F., Yang X., Wang Y. An interactive chatbot for university open day. *IEEE 13th International Conference on Software Engineering and Service Science (ICSESS)*. 2022. pp. 214–217.
7. Rawat B., Bist A., Rahardja U., Aini Q., Ayu Sanjaya Y. Recent deep learning based nlp techniques for chatbot development: An exhaustive survey. *10th International Conference on Cyber and IT Service Management (CITSM)*. 2022. pp. 1–4.
8. Garg R., Riya R., Thakur S., Tyagi N., Basha K., Vij D., Sodhi G. Nlp based chatbot for multiple restaurants. *10th International Conference on System Modeling Advancement in Research Trends (SMART)*. 2021. pp. 439–443.
9. Wahal A., Aggarwal M., Poongodi T. Iot based chatbots using nlp and svm algorithms. *3rd International Conference on Intelligent Engineering and Management (ICIEM)*. 2022. pp. 484–489.
10. Adam M., Wessel M., Benlian A. AI-based chatbots in customer service and their effects on user compliance. *Electronic Markets*. 2021. vol. 31. no. 2. pp. 427–445.
11. Ceuca L., Rednic A., Chifu E. Safer museum guide interaction during a pandemic and further using nlp in human interactive museum visits: Museum guide chatbot. *IEEE 17th International Conference on Intelligent Computer Communication and Processing (ICCP)*. 2021. pp. 313–318.
12. Smutny P., Schreiberova P. Chatbots for learning: A review of educational chatbots for the facebook messenger. *Computers Education*. 2020. vol. 151. p. 103862.
13. Purohit J., Bagwe A., Mehta R., Mangaonkar O., George E. Natural language processing based jaro-the interviewing chatbot. *3rd International Conference on Computing Methodologies and Communication (ICCMC)*. 2019. pp. 134–136.
14. Srivastava P., Singh N. Automatized medical chatbot (medibot). *International Conference on Power Electronics IoT Applications in Renewable Energy and its Control (PARC)*. 2020. pp. 351–354.
15. Polignano M., Narducci F., Iovine A., Musto C., De Gemmis M., Semeraro G., Healthassistantbot: A personal health assistant for the italian language. *IEEE Access*. 2020. vol. 8. pp. 107479–107497.
16. Kumari S., Naikwadi Z., Akole A., Darshankar P. Enhancing college chat bot assistant with the help of richer human computer interaction and speech recognition. *International Conference on Electronics and Sustainable Communication Systems (ICESC)*. 2020. pp. 427–433.
17. Ait-Mlouk A., Jiang L. Kbot: A knowledge graph based chatbot for natural language understanding over linked data. *IEEE Access*. 2020. vol. 8. pp. 149220–149230.
18. Makhkamova O., Lee K.-H., Do K., Kim D. Deep learning-based multi-chatbot broker for qa improvement of video tutoring assistant. *IEEE International Conference on Big Data and Smart Computing (BigComp)*. 2020. pp. 221–224.
19. Nguyen Huu P., Do Manh C., Nguyen Trong H. Proposing chatbot model for managing comments in vietnam,” in *Industrial Networks and Intelligent Systems*. Eds. N.-S. Vo, V.-P. Hoang, Q.-T. Vien. Cham: Springer International Publishing, 2021. pp. 287–297.
20. Sinha S., Basak S., Dey Y., Mondal A. An educational chatbot for answering queries. *Emerging technology in modelling and graphics*. Springer, 2020. pp. 55–60.
21. Selvi V., Saranya S., Chidida K., Abarna R. Chatbot and bullyfree chat. *IEEE International Conference on System, Computation, Automation and Networking (ICSCAN)*. 2019. pp. 1–5.
22. Miklosik A., Evans N., Qureshi A. The use of chatbots in digital business transformation: A systematic literature review. *IEEE Access*. 2021. vol. 9. pp. 106530–106539.
23. Daniel G., Cabot J., Deruelle L., Derras M. Xatkit: A multimodal low-code chatbot development framework,” *IEEE Access*, 2020. vol. 8. pp. 15332–15346.

24. Santos G., De Andrade G., Silva G., Duarte F., Costa J., De Sousa R. A conversation-driven approach for chatbot management. *IEEE Access*. 2022. vol. 10. pp. 8474–8486.
25. Medeiros L., Bosse T., Gerritsen C. Can a chatbot comfort humans? studying the impact of a supportive chatbot on users' self-perceived stress. *IEEE Transactions on Human-Machine Systems*. 2022. vol. 52. no. 3. pp. 343–353.
26. Abdellatif A., Badran K., Costa D., Shihab E. A comparison of natural language understanding platforms for chatbots in software engineering. *IEEE Transactions on Software Engineering*. 2022. vol. 48. no. 8. pp. 3087–3102.
27. Zhang L., Yang Y., Zhou J., Chen C., He L. Retrieval-polished response generation for chatbot. *IEEE Access*. 2020. vol. 8. pp. 123882–123890.
28. Ren R., Perez-Soler S., Castro J., Dieste O., Acuna S. Using the socio chatbot for uml modeling: A second family of experiments on usability in academic settings. *IEEE Access*. 2022. vol. 10. pp. 130542–130562.
29. The stanford natural language processing (nlp). Available at: <https://nlp.stanford.edu/software/tagger.shtml>. (accessed: 02.11.2022).
30. Nicolov N., Salvetti F., Ivanova S. Sentiment analysis: Does coreference matter. *AISB 2008 convention communication, interaction and social intelligence*. 2008. vol. 1. p. 37.
31. Kanayama H., Nasukawa T., Watanabe H. Deeper sentiment analysis using machine translation technology. *COLING 2004: Proceedings of the 20th International Conference on Computational Linguistics*. 2004. pp. 494–500.
32. Hu M., Liu B. Mining and summarizing customer reviews. *Proceedings of the tenth ACM SIGKDD international conference on Knowledge discovery and data mining*. 2004. pp. 168–177.
33. Ko M.-C., Lin Z.-H. Cardbot: A chatbot for business card management. *Proceedings of the 23rd International Conference on Intelligent User Interfaces Companion, ser. IUI '18 Companion*. New York, NY, USA: Association for Computing Machinery, 2018. Available at: <https://doi.org/10.1145/3180308.3180313>. (accessed: 02.11.2022).
34. English tense system. Available at: <https://www.englishclub.com/grammar/verb-tenses-system.htm>. (accessed: 10.11.2022).
35. Narayanan R., Liu B., Choudhary A. Sentiment analysis of conditional sentences. *Proceedings of the 2009 Conference on Empirical Methods in Natural Language Processing*. Singapore: Association for Computational Linguistics. 2009. pp. 180–189.
36. Liu Q., Gao Z., Liu B., Zhang Y. Automated rule selection for aspect extraction in opinion mining. *in Proceedings of the 24th International Conference on Artificial Intelligence, ser. IJCAI'15*. AAAI Press, 2015. p. 1291–1297.
37. Kasthuri E., Balaji S. A chatbot for changing lifestyle in education. *Third International Conference on Intelligent Communication Technologies and Virtual Mobile Networks (ICICV)*. 2021. pp. 1317–1322.
38. Dharani M., Jyostna J., Sucharitha E., Likitha R., Manne S. Interactive transport enquiry with ai chatbot. *4th International Conference on Intelligent Computing and Control Systems (ICICCS)*. 2020. pp. 1271–1276.

Нгуен Хунг Вьет — преподаватель, Восточноазиатский технологический университет; факультет информационных технологий, Ханойский университет науки и технологий. Область научных интересов: мультимедийные коммуникации, сетевая безопасность, искусственный интеллект, управление трафиком в сетях нового поколения, гарантия QoE/QoS для сетевых услуг, экологичные сети, приложения. Число научных публикаций — 10. hungnv@eaut.edu.vn; Ки Фу - Ки Ань, Хатинь, Вьетнам; р.т.: +84(098)911-2079.

Тан Нгуен — научный сотрудник, Восточноазиатский технологический университет. Область научных интересов: приложения, анализ данных. Число научных публикаций — 1. tan25102000@gmail.com; Чынг Зунг - Тиен Лу, Хынгйен, Вьетнам; р.т.: +84(035)919-0216.

Куан Нгуен Хонг — научный сотрудник, Восточноазиатский технологический университет. Область научных интересов: приложения, сети. Число научных публикаций — 1. quan31nd@gmail.com; Суан Тхань, Суан Хьонг, Ханой, Вьетнам; р.т.: +84(082)220-6919.

Хьонг Чьонг Тху — Ph.D., Dr.Sci., заместитель начальника отдела инженерных коммуникаций, школа электротехники и электронной инженерии, Ханойский университет науки и технологий. Область научных интересов: сетевая безопасность, искусственные интеллектуальные функции, управление трафиком в сетях следующего поколения, гарантия QoE/QoS для сети услуги, экологичные сети, развитие экосистем и приложений Интернета вещей. Число научных публикаций — 84. huong.truongthu@hust.edu.vn; улица Дай Ко Вьет, 1, Ханой, Вьетнам; р.т.: +84(243)869-2242.

Пхат Нгуен Хуу — Ph.D., преподаватель, Ханойский университет науки и технологий. Область научных интересов: обработка изображений и видео, беспроводные сети, большие данные, интеллектуальный транспорт, системы (ИТС), Интернет вещей (IoT). Число научных публикаций — 60. phat.nguyenhuu@hust.edu.vn; улица Дай Ко Вьет, 1, Ханой, Вьетнам; р.т.: +84(243)869-2242.

А.Н. ТРОФИМОВ, Ф.А. ТАУБИН
**АНАЛИЗ ЭФФЕКТИВНОСТИ КАСКАДНОГО КОДИРОВАНИЯ
ДЛЯ ПОВЫШЕНИЯ ВЫНОСЛИВОСТИ МНОГОУРОВНЕВОЙ
NAND ФЛЕШ-ПАМЯТИ**

Трофимов А.Н., Таубин Ф.А. Анализ эффективности каскадного кодирования для повышения выносливости многоуровневой NAND флеш-памяти.

Аннотация. Повышение плотности записи в современных чипах NAND флеш-памяти, достигаемое как за счет уменьшающегося физического размера ячейки, так и благодаря возрастающему количеству используемых состояний ячейки, сопровождается снижением надежности хранения данных – вероятности ошибки, выносливости (числа циклов перезаписи) и времени хранения. Стандартным решением, позволяющим повысить надежность хранения данных в многоуровневой флеш-памяти, является введение помехоустойчивого кодирования. Эффективность введения помехоустойчивого кодирования в существенной степени определяется адекватностью модели, формализующей основные процессы, связанные с записью и чтением данных. В работе приводится описание основных искажений, сопровождающих процесс записи/считывания в NAND флеш-памяти, и явный вид плотностей распределения результирующего шума. В качестве аппроксимации полученных плотностей распределения результирующего шума рассматривается модель на основе композиции гауссова распределения и распределения Лапласа, достаточно адекватно отражающая плотности распределения результирующего шума при большом числе циклов перезаписи. Для этой модели проводится анализ помехоустойчивости каскадных кодовых конструкций с внешним кодом Рида-Соломона и внутренним многоуровневым кодом, состоящим из двоичных компонентных кодов. Выполненный анализ позволяет получить обменные соотношения между вероятностью ошибки, плотностью записи и числом циклов перезаписи. Полученные обменные соотношения показывают, что предложенные конструкции позволяют за счет очень незначительного снижения плотности записи обеспечить увеличение граничного значения числа циклов перезаписи (определяемого производителем) в 2–2.5 раза при сохранении требуемого значения вероятности ошибки на бит.

Ключевые слова: многоуровневая NAND флеш-память, модель искажений в канале записи/считывания, NL распределение, каскадное кодирование, анализ помехоустойчивости, выносливость.

1. Введение. Энергонезависимая память (non-volatile memory, NVM) – класс памяти, который поддерживает сохраненные данные даже после отключения от источника питания. Современная энергонезависимая память обеспечивают более быстрый доступ к данным, пониженное энергопотребление и улучшенную отказоустойчивость по сравнению с обычными жесткими дисками. Как результат, устройства на основе энергонезависимой памяти стали основной заменой жестких дисков для различных новых приложений, связанных с хранением данных, включая аэрокосмические системы, персональную электронику, мобильные вычисления, автономные

транспортные средства, корпоративные хранилища, дата-центры, вычислительные системы с интенсивным использованием данных, медицинские и IoT устройства [1 – 9]. Существует множество типов энергонезависимой памяти, различающихся по своим потребительским свойствам, обусловленным способом построения запоминающих ячеек, и сферам применения. В настоящее время наиболее популярной технологией энергонезависимой памяти является NAND флеш-память, представляющая собой разновидность технологии электрически перепрограммируемой памяти и поддерживающая чтение, программирование и стирание данных в качестве основных операций. Ключевыми достоинствами NAND флеш-памяти являются компактность, дешевизна, механическая прочность, высокая емкость и скорость работы, низкое энергопотребление. Одной из основных целевых областей применения NAND флеш-памяти являются приложения, в которых данные постоянно записываются и стираются. Диапазон таких приложений весьма широк – от потребительской электроники до корпоративных центров обработки данных для облачных приложений и социальных сетей. Эффективное функционирование таких центров требует наличия быстрых и надежных механизмов хранения/обновления громадных объемов данных, и эта потребность удовлетворяется за счет использования флеш-памяти в виде твердотельных накопителей (SSD) [1, 6, 9]. Улучшения в технологии изготовления NAND флеш-памяти за последнее десятилетие привели к снижению удельной стоимости и увеличению плотности записи данных и, таким образом, сделали NAND флеш-память практически безальтернативным решением для ряда современных приложений.

Чип NAND флеш-памяти представляет собой микросборку из нескольких кристаллов. Каждый кристалл, в свою очередь, представляет собой массив ячеек, каждая из которых является полевым транзистором с плавающим затвором. Такая ячейка характеризуется одним из нескольких конечных состояний, определяемых величиной заряда в плавающем затворе. Флеш-память с числом состояний $q > 2$ называется многоуровневой (multi-level cell, MLC), в отличие от одноуровневой памяти (single-level cell, SLC), где возможен лишь один ненулевой уровень. Наиболее распространенными в настоящее время являются чипы многоуровневой NAND флеш-памяти с 4 и 8 уровнями хранимого электрического заряда в ячейке. На стадии разработки находятся чипы, позволяющие хранить в ячейке 4 бита [10, 11] и 5 бит [12].

Известно, что повышение плотности записи сопровождается как уменьшением максимально возможного числа циклов программирования и стирания (P/E cycling), так и снижением надежности хранения данных [13, 14]. В результате в типичных условиях эксплуатации срок службы чипа многоуровневой флеш-памяти вследствие повышенного износа может составить всего несколько месяцев [15]. Хорошо известно, что в результате повышения плотности записи, характеристики основных параметров многоуровневой NAND флеш-памяти – вероятность ошибки, выносливость (flash cell endurance) и долговечность хранения (data retention), как правило, оказываются неприемлемыми. Так, исходная вероятность ошибки в четырехуровневой флеш-памяти (raw error rate), даже при умеренных значениях числа циклов перезаписи и времени хранения, составляет порядка 10^{-4} ... 10^{-3} и более [16], тогда как требуемая вероятность ошибки лежит в диапазоне 10^{-12} ... 10^{-16} . Например, согласно рекомендациям Joint Electron Device Engineering Council (JEDEC), стандартное значение требуемой вероятности ошибки в твердотельных накопителях составляет 10^{-15} для клиентских приложений и 10^{-16} для корпоративных решений [17, 18]. Далее, твердотельные накопители на базе SLC флеш-памяти могли выдерживать 150 000 циклов программирования и стирания, тогда как современные твердотельные накопители с четырехуровневой флеш-памятью на основе техпроцесса 1x (т.е. 15–19 нм) могут выдержать только 3000 циклов программирования и стирания [19, 20]. При этом требуемая величина выносливости, т.е. допустимого числа циклов перезаписи, может достигать значения 10^5 , а требуемое время хранения данных составляет до 10 лет. В результате, производители и пользователи многоуровневой флеш-памяти (и в частности, твердотельных накопителей) сталкиваются с трудным выбором в соотношении цены, производительности, емкости и надежности, включающем в качестве важного компонента обеспечение увеличения выносливости при сохранении требуемой вероятности ошибки.

Важным фактором, влияющим на выносливость флеш-памяти, является то, что конструктивно чип NAND флеш-памяти состоит, как правило, из нескольких десятков расположенных друг над другом слоев двумерных массивов ячеек, называемых флеш-блоками. Флеш-блок есть своего рода неделимая единица при выполнении операции стирания (block granularity), т.е. в нулевое состояние переводятся одновременно все ячейки блока. Каждый флеш-блок делится на страницы, которые являются минимальными (по размеру) адресуемыми элементами при выполнении операций записи и/или

чтения (page granularity). В результате, при перезаписи некоторой страницы происходит ненужное (unnneeded) стирание остальных страниц флеш-блока (с последующим восстановлением стертых данных). Этот эффект – write amplification, который количественно определяется как отношение числа физически выполненных операций записи к числу логически выполненных записей и может достигать величин порядка 3...4, приводит, очевидно, к снижению производительности памяти и её выносливости. Многообещающим инструментом, позволяющим существенно снизить число «паразитных» циклов перезаписи, является использование перезаписываемых (re-write) кодов, и, в частности, WOM кодов [21, 22, 23]. WOM коды характеризуются тем, что при перезаписи только увеличивают уровень заряда в ячейке; в результате страница может быть перезаписана несколько раз (по меньшей мере дважды) перед тем, как потребуется стереть флеш-блок. Вместе с тем, известные WOM коды не позволяют обеспечить требуемую вероятность ошибки порядка $10^{-12}...10^{-16}$, а потенциально перспективное объединение перезаписываемых кодов и кодов, исправляющих ошибки, пока находится в начальной стадии исследования [24, 25].

Поэтому в настоящее время решением, позволяющим улучшить потребительские характеристики многоуровневой флеш-памяти, является введение помехоустойчивого кодирования [26 – 33]. Одна из основных проблем, возникающих при реализации помехоустойчивого кодирования в NAND флеш-памяти, связана с необходимостью учета наличия нескольких источников шума, характеристики которых изменяются при изменении как количества выполненных циклов программирования/стирания, так и времени хранения данных в памяти. Кроме того, параметры источников шума зависят от данных, так как увеличение величины заряда в ячейке сопровождается потенциально более интенсивным процессом утечки заряда. Таким образом, эффективность введения помехоустойчивого кодирования в существенной степени определяется адекватностью модели, формализующей основные процессы, связанные с записью и чтением данных. В разделе 2 приводится описание основных искажений, сопровождающих процесс записи/считывания в NAND флеш-памяти, и явный вид плотностей распределения результирующего шума. В качестве аппроксимации полученных плотностей распределения результирующего шума, рассматривается модель на основе композиции гауссова распределения и распределения Лапласа, достаточно адекватно отражающая

плотности распределения результирующего шума при большом числе циклов перезаписи. Применительно к этой модели искажений сигнала в NAND флеш-памяти, для ряда каскадных кодовых конструкций из [30, 31, 33], кратко представленных в разделе 3, в разделе 4 проводится анализ помехоустойчивости рассматриваемых кодовых конструкций, позволяющий получить обменные соотношения между вероятностью ошибки, плотностью записи и числом циклов перезаписи. Полученные обменные соотношения позволяют оценить эффективность использования рассматриваемых вариантов помехоустойчивого кодирования для увеличения выносливости NAND флеш-памяти при поддержании требуемой вероятности ошибки.

Некоторые предварительные результаты данной работы были частично представлены на XXV Международной научной конференции «Волновая электроника и инфокоммуникационные системы», Санкт-Петербург, 2022 [34].

2. Модель искажений в NAND флеш-памяти. К основным факторам, определяющим надежность флеш-памяти, обычно относят:

- 1) случайные изменения величины заряда, вводимого в плавающий затвор транзистора в процессе записи (program disturbance);
- 2) повторяющиеся в процессе эксплуатации циклы записи/стирания – P/E cycling;
- 3) взаимную интерференцию ячеек (cell-to-cell interference);
- 4) утечку заряда плавающего затвора с течением времени (retention problem).

Один из подходов, позволяющих формализовать влияние указанных факторов, состоит в следующем [35]. Случайное отклонение реального порогового уровня напряжения от целевого уровня (target value) рассматривается как сумма шумов (случайных величин), порождаемых процессами стирания (eraser noise), записи в ячейку (programming noise), перезаписи (random telegraph noise), старения памяти (retention noise) и взаимной интерференцией ячеек (cell-to-cell interference, CCI). Обозначим упорядоченную по возрастанию последовательность целевых уровней (состояний) напряжения в ячейке как $x_0, x_1, \dots, x_{q-1}$, где q – число уровней записи; наименьший по величине целевой уровень x_0 соответствует состоянию «стерто».

Плотность распределения шума ξ_e , возникающего при считывании состояния x_0 («стерто»), хорошо аппроксимируется, как показывают многочисленные эксперименты [35, 36], выражением:

$$f_e(\xi_e) = \frac{1}{\sqrt{2\pi\sigma_e^2}} \exp\left(-\frac{\xi_e^2}{2\sigma_e^2}\right), \quad (1)$$

где σ_e^2 – дисперсия шума в состоянии «стерто».

Плотность распределения шума ξ_p , возникающего в процессе записи, полагается равномерной на интервале $[0, \Delta]$, где Δ – величина шага приращения напряжения (the incremental step pulse programming, ISPP) при записи целевого уровня:

$$f_p(\xi_p) = \frac{1}{\Delta}, \quad 0 \leq \xi_p \leq \Delta. \quad (2)$$

Отметим, что при нахождении ячейки в состоянии «стерто» шум $\xi_p = 0$.

Плотность распределения шума ξ_t , возникающего вследствие перезаписи, хорошо аппроксимируется, согласно многочисленным измерениям [36], стандартным распределением Лапласа с нулевым средним:

$$f_t(\xi_t) = \frac{1}{2\theta} e^{-\frac{|\xi_t|}{\theta}}, \quad (3)$$

где значение параметра θ определяются числом циклов записи/стирания N в виде: $\theta = K_\theta N^\alpha$, где K_θ и α – константы, определяемые конкретной технологией изготовления чипа [36].

Плотность распределения шума ξ_r , возникающего вследствие старения памяти, может быть с разумной точностью аппроксимирована гауссовым распределением [35, 36], параметры которого – среднее значение и дисперсия, зависят от целевого уровня напряжения (состояния) ячейки. Условная плотность распределения шума ξ_r , возникающего вследствие старения памяти, при считывании уровня x_i , $i = 1, 2, \dots, q-1$, имеет вид:

$$f_r(\xi_r | x_i) = \frac{1}{\sqrt{2\pi\sigma_r^2(x_i)}} \exp\left(-\frac{(\xi_r - \mu_r(x_i))^2}{2\sigma_r^2(x_i)}\right), \quad (4)$$

где:

$$\mu_r(x_i) = -K_s(x_i - x_0)K_d N^{0.5} \ln\left(1 + \frac{T}{t_0}\right),$$

$$\sigma_r^2(x_i) = K_s(x_i - x_0)K_m N^{0.6} \ln\left(1 + \frac{T}{t_0}\right),$$

T – длительность хранения данных, K_s , K_d , K_m , t_0 – константы, определяемые конкретной технологией изготовления чипа [36]. Отметим, что среднее значение шума ξ_r : 1) отлично от нуля и отрицательно, что можно интерпретировать как «сжатие» диапазона целевых уровней, 2) различно для различных целевых уровней и 3) возрастает при увеличении целевого уровня. При нахождении ячейки в состоянии x_0 («стерто»), шум $\xi_r=0$.

Взаимная интерференция ячеек проявляется в том, величина заряда подвергающейся воздействию интерференции ячейки (victim cell) изменяется при выполнении операций записи или стирания в соседних ячейках (aggressor cells). В частности, при выполнении операции записи в соседних ячейках интерферирующая помеха увеличивает величину заряда в ячейке-жертве, тогда как при выполнении операции стирания (в соседних ячейках) величина заряда в ячейке-жертве уменьшается вследствие интерференции. Для ослабления влияния взаимной интерференции (CCI) используется компенсация интерференционной помехи в ячейке-жертве на основе удаления/добавления корректирующего заряда, величина которого есть линейная комбинация величин зарядов соседних ячеек с весами, определяемыми емкостными связями между ячейками [37]. Плотность распределения остаточного (после компенсации) интерференционного шума ξ_{cci} , как показывают результаты моделирования в [37], хорошо аппроксимируется «вплоть до хвостов» (down to the tails) гауссовым распределением с нулевым средним. Таким образом, плотность распределения интерференционного шума ξ_{cci} моделируется как:

$$f_{cci}(\xi_{cci}) = \frac{1}{\sqrt{2\pi\sigma_{cci}^2}} \exp\left(-\frac{\xi_{cci}^2}{2\sigma_{cci}^2}\right), \quad (5)$$

где σ_{cci}^2 – дисперсия остаточного (после компенсации) интерференционного шума.

Экспериментальные исследования [38, 39] показывают, что при эффективной компенсации взаимной интерференции в многоуровневой флеш-памяти доминирующим является распределение ошибок, близкое к равномерному на множестве блока ячеек, и при этом не наблюдается заметной тенденции к пакетированию. Последнее обстоятельство может служить аргументом для формального описания рассматриваемого канала записи/считывания блока ячеек флеш-памяти как канала без памяти. Поэтому, в рамках принятого допущения, описание модели сводится к отысканию плотности распределения сигнала, считываемого из ячейки флеш-памяти.

Указанные факторы, определяющие надежность флеш-памяти, характерны, главным образом, для 2D (планарной) NAND флеш-памяти. Архитектура 3D NAND флеш-памяти – трехмерное расположение ячеек чипа, и используемая в 3D NAND технология ловушки заряда (charge trap), т.е. хранение заряда в изолированной области ячейки из непроводящего материала типа нитрида кремния (вместо помещения заряда в плавающий затвор), позволяют существенно повысить плотность записи и снизить вероятность утечки заряда, но при этом порождают ряд дополнительных особенностей. К ним относятся, прежде всего: а) появление существенного отличия динамики изменения среднего значения и дисперсии шума ξ_r , возникающего вследствие старения памяти (early retention loss), и б) возникновение зависимости скорости утечки заряда из ячейки от величин зарядов, хранящихся в соседних ячейках (retention interference) [40, 41, 42]. Моделирование интерференционных явлений между «charge trap»-ячейками и оценки влияния интерференции на надежность памяти оказывается весьма сложной задачей. Поэтому формализация и апробирование модели шума старения в 3D NAND флеш-памяти пока, как представляется, ещё находятся в стадии анализа и накопления экспериментальных данных.

Приведенные выше источники шумов полагаются аддитивными и независимыми, поэтому условные плотности распределений $p_{y|x_0}(y|x_0)$, $p_{y|x_1}(y|x_1)$, ..., $p_{y|x_{q-1}}(y|x_{q-1})$ считываемого из ячейки сигнала y можно вычислить с использованием свертки распределений (1)–(5) отдельных слагаемых суммарного шума. Если ячейка находится в состоянии x_0 («стерто»), то считываемый из ячейки сигнал:

$$y = x_0 + \xi_e + \xi_t + \xi_{cci}, \quad (6)$$

т.е. суммарный шум равен $\xi_e + \xi_t + \xi_{cci}$, поэтому плотность распределения $p_{y|x_0}(y|x_0)$ считываемого из ячейки сигнала (6) вычисляется с использованием свертки распределений (1), (3) и (5). В результате получаем:

$$p_{y|x_0}(y|x_0) = \frac{1}{2\theta} e^{\frac{\delta_0^2}{2\theta^2}} \left(e^{-\frac{y-x_0}{\theta}} Q\left(\frac{\delta_0}{\theta} - \frac{y-x_0}{\delta_0}\right) + e^{\frac{y-x_0}{\theta}} Q\left(\frac{\delta_0}{\theta} + \frac{y-x_0}{\delta_0}\right) \right), \quad (7)$$

где $Q(z) = (1/\sqrt{2\pi}) \int_z^\infty \exp(-t^2/2) dt$, $\delta_0^2 = \sigma_e^2 + \sigma_{cci}^2$. Если ячейка находится в программируемом состоянии x_i , $i = 1, 2, \dots, q-1$, то считываемый из ячейки сигнал:

$$y = x_i + \xi_p + \xi_t + \xi_r + \xi_{cci}, \quad (8)$$

т.е. суммарный шум равен $\xi_p + \xi_t + \xi_r + \xi_{cci}$, поэтому плотность распределения $p_{y|x_i}(y|x_i)$ считываемого из ячейки сигнала (8) вычисляется с использованием свертки распределений (2) – (5). В результате получаем:

$$\begin{aligned} p_{y|x_i}(y|x_i) = & \frac{1}{2\Delta} \exp\left(\frac{\delta_i^2}{2\theta^2}\right) \times \\ & \left(\exp\left(\frac{y-x_i-\mu_r(x_i)}{\theta}\right) Q\left(\frac{\delta_i}{\theta} + \frac{y-x_i-\mu_r(x_i)}{\delta_i}\right) - \right. \\ & - \exp\left(-\frac{y-x_i-\mu_r(x_i)}{\theta}\right) Q\left(\frac{\delta_i}{\theta} - \frac{y-x_i-\mu_r(x_i)}{\delta_i}\right) + \\ & + \exp\left(-\frac{y-x_i-\Delta-\mu_r(x_i)}{\theta}\right) Q\left(\frac{\delta_i}{\theta} - \frac{y-x_i-\Delta-\mu_r(x_i)}{\delta_i}\right) - \\ & - \exp\left(\frac{y-x_i-\Delta-\mu_r(x_i)}{\theta}\right) Q\left(\frac{\delta_i}{\theta} + \frac{y-x_i-\Delta-\mu_r(x_i)}{\delta_i}\right) \Bigg) + \\ & + \frac{1}{\Delta} \left(Q\left(\frac{y-x_i-\Delta-\mu_r(x_i)}{\delta_i}\right) + Q\left(\frac{y-x_i-\mu_r(x_i)}{\delta_i}\right) \right), \end{aligned} \quad (9)$$

где $\delta_i^2 = \sigma_r^2(x_i) + \sigma_{cci}^2$.

Приведенные выражения (7) и (9), определяющие в замкнутой форме точное (в рамках принятой модели) аналитическое представление условных плотностей распределения $p_{y|x_i}(y|x_i)$, $i = 0, 1, \dots, q-1$, являются весьма громоздкими, что затрудняет их непосредственное использование как при анализе кодовых конструкций, так и при реализации процедур декодирования с использованием мягких решений. Указанное обстоятельство диктует необходимость достаточно простых аппроксимаций для условных плотностей распределения $p_{y|x_i}(y|x_i)$, $i = 0, 1, \dots, q-1$, хорошо согласующихся, при этом, с точными выражениями (7) и (9). Так, например, в работе [35] были рассмотрены возможности аппроксимации плотности распределения суммарного (результатирующего) шума гауссовым распределением, бета-распределением, гамма-распределением, логнормальным распределением и распределением Вейбулла. Среди этих распределений для рассмотренных в [35] сценариев работы флеш-памяти с умеренным числом циклов перезаписи гауссово распределение оказалось в большинстве случаев наиболее предпочтительным вариантом аппроксимации.

Гауссовская аппроксимация. Простейшей (и наиболее распространенной) аппроксимацией является аппроксимация выражений (7) и (9) гауссовым распределением [35, 38, 43, 44]. При этом средние значения и дисперсии аппроксимирующих гауссовых распределений полностью определяются первым и вторым моментами распределений (1)–(5) отдельных слагаемых суммарного шума. Так, для состояния x_0 имеем: среднее значение равно x_0 , дисперсия составляет $\sigma_e^2 + 2\theta^2 + \sigma_{cci}^2$. Для состояния x_i , $i = 1, 2, \dots, q-1$ (программируемое состояние) среднее значение равно $x_i + \Delta/2 - \mu_r(x_i)$, дисперсия составляет $2\theta^2 + \Delta^2/12 + \sigma_r^2(x_i) + \sigma_{cci}^2$. Таким образом, гауссова аппроксимация условных плотностей распределения сигнала, считываемого из ячейки, имеет следующий вид. Если ячейка находится в состоянии x_0 («стерто»), то:

$$p_{y|x_0}^{(g)}(y|x_0) = \frac{1}{\sqrt{2\pi(\sigma_e^2 + 2\theta^2 + \sigma_{cci}^2)}} \exp\left(-\frac{(y-x_0)^2}{2(\sigma_e^2 + 2\theta^2 + \sigma_{cci}^2)}\right). \quad (10)$$

Если ячейка находится в состоянии $x_i, i = 1, 2, \dots, q-1$, (одном из программируемых состояний), то:

$$p_{y|x_i}^{(g)}(y|x_i) = \frac{1}{\sqrt{2\pi(2\theta^2 + \Delta^2/12 + \sigma_r^2(x_i) + \sigma_{cci}^2)}} \times \exp\left(-\frac{(y - x_i - \Delta/2 + \mu_r(x_i))^2}{2(2\theta^2 + \Delta^2/12 + \sigma_r^2(x_i) + \sigma_{cci}^2)}\right). \quad (11)$$

Отметим, что среднее значение и дисперсия сигнала, считываемого из ячейки, находящейся в одном из программируемых состояний, зависят от этого состояния, поэтому рассматриваемая аппроксимация с условной плотностью распределения, задаваемой соотношениями (10) и (11), представляет собой гауссовскую модель с дисперсией аддитивного шума, зависящей от записанного значения (input-dependent additive Gaussian noise, ID-AGN).

Характеризовать в целом адекватность гауссовской аппроксимации, с учетом факторов, ограничивающие её применение, можно следующим образом. Среди перечисленных выше искажений, возникающих в процессе записи/считывания, искажения ξ_p , порождаемые при записи в ячейку (programming noise), и искажения ξ_i , порождаемые накопленными циклами перезаписи (random telegraph noise), имеют распределение, существенно отличающееся от гауссова. При малых значениях шага приращения напряжения Δ , что достаточно типично, влияние первого из этих двух видов искажений ξ_p , не является существенным. Влияние же искажений ξ_i , вызванных накопленным количеством циклов перезаписи N , довольно быстро усиливается с ростом N ; так, дисперсия случайной величины ξ_i возрастает пропорционально $N^{2\alpha}$, где показатель $\alpha = 0.5 \dots 1$. Таким образом, можно полагать, что гауссовская аппроксимация вполне приемлема для сценария, при котором влияние искажений ξ_i (оцениваемое, например, относительным весом дисперсии $2\theta^2$ в суммарной дисперсии $2\theta^2 + \Delta^2/12 + \sigma_r^2(x_i) + \sigma_{cci}^2, i = 1, 2, \dots, q-1$) сравнительно невелико. По мере увеличения накопленного количества циклов перезаписи N (и приближения к предельному значению, заявленному производителем), качество гауссовской аппроксимации плотностей (7) и (9) существенно снижается, что объясняется, главным

образом, появлением асимметрии и «утяжелением хвостов» в плотностях (7) и (9) [45]. Для учета указанных факторов – асимметрии и «утяжеления хвостов» (вследствие большего влияния распределения Лапласа в свертках соответствующих функций плотности вероятностей), в работе [46] было предложено использовать новую модель, представляющую собой композицию гауссова распределения и распределения Лапласа – NL распределение (normal-Laplace distribution).

Аппроксимация NL распределением. NL распределение представляет собой свертку гауссова распределения и распределения Лапласа. Это распределение «почти» гауссово в средней области, тогда как на краях области определения NL распределение «почти» совпадает с распределением Лапласа. Асимметрия может быть введена в NL распределение посредством использования в свертке перекошенного, или асимметричного, распределения Лапласа (skew-Laplace distribution).

Общее выражение для плотности NL распределения (с введенной асимметрией) имеет следующий вид. Пусть $\xi(m, \sigma)$ – гауссова случайная величина с плотностью распределения:

$$f_g(\xi; m, \sigma) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{(\xi - m)^2}{2\sigma^2}\right), \quad (12)$$

где m – среднее значение, σ^2 – дисперсия. Пусть $\zeta(\lambda, \mu)$ – случайная величина с асимметричной плотностью распределения Лапласа:

$$f_{sl}(\zeta; \lambda, \mu) = \begin{cases} \frac{\lambda\mu}{\lambda + \mu} \exp(-\lambda\zeta), & \zeta > 0, \\ \frac{\lambda\mu}{\lambda + \mu} \exp(\mu\zeta), & \zeta \leq 0, \end{cases} \quad (13)$$

где $\lambda > 0$, $\mu > 0$ – параметры перекошенного распределения Лапласа. Очевидно, что при $\lambda = \mu$ получаем обычное (симметричное) распределение Лапласа вида (3). Другие частные случаи $\mu = \infty$ и $\lambda = \infty$ приводят к экспоненциальным плотностям вида:

$$f_{sl}(\zeta; \lambda, \infty) = \begin{cases} \lambda \exp(-\lambda \zeta), & \zeta > 0, \\ 0, & \zeta \leq 0, \end{cases}$$

и:

$$f_{sl}(\zeta; \infty, \mu) = \begin{cases} 0, & \zeta > 0, \\ \mu \exp(\mu \zeta), & \zeta \leq 0. \end{cases}$$

Пусть $y = \xi(m, \sigma) + \zeta(\lambda, \mu)$; тогда случайная величина y имеет NL распределение с плотностью:

$$f_{NL}(y; m, \sigma, \lambda, \mu) = \frac{\lambda \mu}{\lambda + \mu} \left(e^{\frac{\lambda(-2y+2m+\lambda\sigma^2)}{2}} Q\left(-\frac{y-m-\lambda\sigma^2}{\sigma}\right) + e^{\frac{\mu(2y-2m+\mu\sigma^2)}{2}} Q\left(\frac{y-m+\mu\sigma^2}{\sigma}\right) \right). \quad (14)$$

В работе [46] предложена два варианта модели четырехуровневой ячейки флеш-памяти, названные «идеальный» и «смешанный». В идеальном варианте выходные значения ячейки памяти заданы следующим равенством:

$$y = \begin{cases} x_0 + \zeta(\lambda_0, \infty), & \text{если было записано значение } x_0, \\ x_1 + \xi(0, \sigma_1) + \zeta(\lambda_1, \mu), & \text{если было записано значение } x_1, \\ x_2 + \xi(0, \sigma_2) + \zeta(\lambda_2, \mu), & \text{если было записано значение } x_2, \\ x_3 + \zeta(\infty, \mu), & \text{если было записано значение } x_3. \end{cases}$$

С учетом того, что $x_1 + \xi(0, \sigma_1) = \xi(x_1, \sigma_1)$ и $x_2 + \xi(0, \sigma_2) = \xi(x_2, \sigma_2)$ можно записать, что:

$$y = \begin{cases} x_0 + \zeta(\lambda_0, \infty), & \text{если было записано значение } x_0, \\ \xi(x_1, \sigma_1) + \zeta(\lambda_1, \mu), & \text{если было записано значение } x_1, \\ \xi(x_2, \sigma_2) + \zeta(\lambda_2, \mu), & \text{если было записано значение } x_2, \\ x_3 + \zeta(\infty, \mu), & \text{если было записано значение } x_3, \end{cases} \quad (15)$$

где обозначения $\xi(m, \sigma)$ и $\zeta(\lambda, \mu)$ используются для указания случайных величин, распределенных в соответствии с плотностями вероятности, заданными равенствами (12) и (13) соответственно. Из равенств (15) следует, что условные функции плотности вероятности, определяющие «идеальную» модель, заданы как:

$$\begin{aligned} p_{y|x_0}^{(ideal)}(y|x_0) &= f_{sl}(y-x_0; \lambda_0, \infty), \\ p_{y|x_1}^{(ideal)}(y|x_1) &= f_{NL}(y; x_1, \sigma_1, \lambda_1, \mu), \\ p_{y|x_2}^{(ideal)}(y|x_2) &= f_{NL}(y; x_2, \sigma_2, \lambda_2, \mu), \\ p_{y|x_3}^{(ideal)}(y|x_3) &= f_{sl}(y-x_3; \infty, \mu). \end{aligned} \quad (16)$$

В работе [46] указано, что распределение, дающее более точное соответствие с результатами экспериментальных измерений, может быть получено как смесь распределений, задаваемых равенствами (16), вида:

$$\begin{bmatrix} p_{y|x_0}(y|x_0) \\ p_{y|x_1}(y|x_1) \\ p_{y|x_2}(y|x_2) \\ p_{y|x_3}(y|x_3) \end{bmatrix} = \begin{bmatrix} k_{00} & 0 & k_{02} & k_{03} \\ 0 & k_{11} & k_{12} & k_{13} \\ 0 & 0 & k_{22} & k_{23} \\ 0 & 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} p_{y|x_0}^{(ideal)}(y|x_0) \\ p_{y|x_1}^{(ideal)}(y|x_1) \\ p_{y|x_2}^{(ideal)}(y|x_2) \\ p_{y|x_3}^{(ideal)}(y|x_3) \end{bmatrix}, \quad (17)$$

где все величины $k_{00}, k_{01}, \dots, k_{33}$ положительны, и кроме того $k_{00} + k_{02} + k_{03} = 1, k_{11} + k_{12} + k_{13} = 1, k_{22} + k_{23} = 1$ и $k_{33} = 1$. Такое описание канала называется смешанной NL (mixed normal-Laplace) моделью. Конкретный вид функций плотности вероятности вида (17) определяется значениями пятнадцати параметров: четырех значений уровней $x_0, \dots, x_3$, шести значений параметров плотностей (16) $\lambda_0, \sigma_1, \lambda_1, \mu, \sigma_2, \lambda_2$ и пяти весовых коэффициентов $k_{02}, k_{03}, k_{12}, k_{13}, k_{23}$ в правой части равенства (17). В работе [46] указаны значения этих параметров в зависимости от отношения N/N_{nom} , где N – число циклов перезаписи, N_{nom} – число циклов перезаписи, гарантируемое производителем устройства памяти, или номинальное число циклов перезаписи. Эти зависимости указаны в таблице 1. Графики условных ф.п.в (17) в линейном и логарифмическом масштабе представлены на рисунках 1 и 2.

Таблица 1. Параметры смешанной NL модели ячейки флеш-памяти

Параметр модели	Функциональная зависимость от z , $z=N/N_{\text{ном}}$	c_2	c_1	c_0
k_{02} k_{03} k_{12} k_{13} k_{23}	$\exp(c_1 z + c_0)$		0.87 1.41 1.63 0.73 1.50	-11.89 -19.82 -19.22 -11.67 -17.69
x_0 x_1 x_2 x_3	$c_1 z^{c_2} + c_0$	1 0.17 0.26 0.39	1.52 12.96 11.69 6.83	-150.81 61.12 186.86 319.38
σ_1 σ_2	c_0			14.95 11.62
λ_0 λ_1 λ_2 μ	$c_1 z + c_0$		1.84 0.65 0.53 0.37	14.28 6.36 6.48 5.13

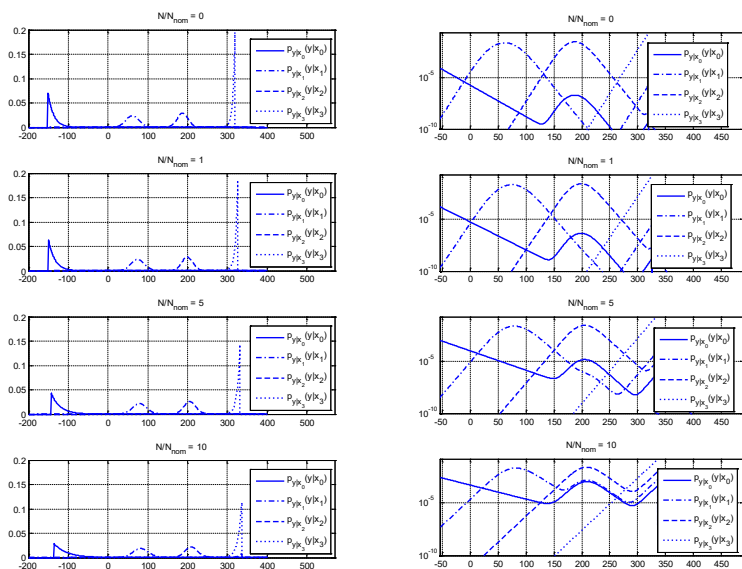


Рис. 1. Графики условных функций плотностей вероятности для смешанной NL модели в линейном и логарифмическом масштабах

3. Краткое описание рассматриваемой каскадной кодовой конструкции. Один из возможных эффективных подходов к организации введения помехоустойчивого кодирования связан с использованием каскадных конструкций [24]. Основной мотив целесообразности использования каскадной схемы кодирования для многоуровневой флеш-памяти состоит в том, чтобы с помощью подходящего сравнительно простого внутреннего кода обеспечить лучшее согласование мощного внешнего кода с расширенным каналом, включающим кодер и декодер внутреннего кода. Среди возможных вариантов внутреннего кода следует выделить многомерные сигнальные множества на основе целочисленных решеток, обладающие гибкой структурой с широким диапазоном варьирования параметров и допускающие сравнительно простую организацию мягкого декодирования, что позволяет существенно повысить эффективность внешнего кодирования. Каскадное кодирование (с многомерным сигнальным множеством в качестве внутреннего кода) рассматривалось в работах [47, 48]. В этих публикациях, рассматривавших в качестве возможных внешних кодов, соответственно, коды Рида-Соломона и LDPC коды, была продемонстрирована существенная эффективность такого рода каскадных конструкций. Весьма важным достоинством каскадной схемы в плане реализации является гибкость ее архитектуры, что позволяет, в частности, реализовать внутреннюю ступень кода, имеющую небольшую сложность, непосредственно на странице кристалла (чипа) флеш-памяти, содержащей защищаемые данные (on-chip implementation). Расширение класса приемлемых конструкций внутреннего кода с сопутствующим существенным расширением диапазона обменных соотношений «помехоустойчивость – плотность записи – сложность реализации», достигаемое путем использования произвольных mod-4 решеток и многоуровневых конструкций было представлено в работах [25, 26, 28].

Анализируемый далее класс каскадных многокомпонентных кодовых конструкций был предложен в работах [30, 31, 33]. Для удобства и связности изложения приведем далее краткое описание рассматриваемой класса каскадных многокомпонентных кодовых конструкций. В рассматриваемой схеме каскадного кодирования в качестве внешнего кода выбран расширенный (удлиненный на один символ) (N_1, K) код Рида-Соломона (РС) над полем $\mathbb{F}_{2^s}$, $2^s = N_1$. Внутренний код, обозначим его B_0 , представляет собой множество последовательностей, состоящих из n символов над алфавитом $A = \{0, 1, \dots, q-1\}$, где q – число уровней записи ячейки флеш-памяти.

Элементы алфавита A связаны с входными уровнями (уровнями записи) ячейки памяти посредством взаимно однозначного отображения I множества A на множество $\{x_0, x_1, \dots, x_{q-1}\}$ вида $I(i) = x_i$, $i = 0, 1, \dots, q-1$. Иными словами, алфавит A представляет собой множество индексов уровней записи q . Будем полагать далее, что $m = \log_2 q$ целое число, т.е. ячейка памяти может хранить m бит данных. Для исключения пакетирования ошибок на внешней ступени декодирования, кодовые символы внешнего кода (при необходимости) подвергаются блоковому перемежению (рисунок 2), а именно: h последовательных слов кода Рида-Соломона записываются в прямоугольную таблицу, содержащую h строк, после чего последовательно считываются по столбцам; каждый столбец этой таблицы отображается в один символ внутреннего кода.

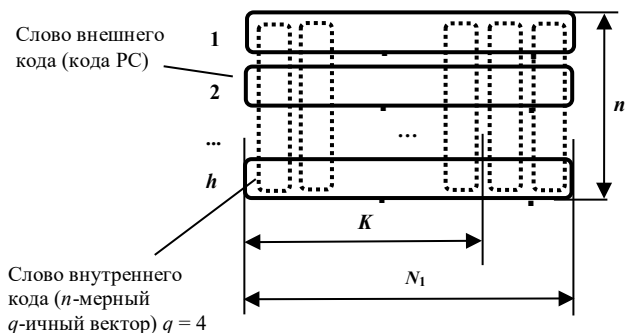


Рис. 2. Структура каскадного кода с перемежением между ступенями

Параметр h определяется объёмом внутреннего кода, обозначим его $|B_0|$, и размером алфавита кода Рида-Соломона, совпадающим с его длиной N_1 , в соответствии со следующим соотношением: $N_1^h = |B_0|$. При $h = 1$ один символ внешнего кода отображается в одно слово внутреннего кода, то есть имеет место обычная каскадная схема кодирования. Обозначим через R и R_0 скорости внешнего и внутреннего кода, соответственно; очевидно, что $R = K/N_1$, $R_0 = \log_2 |B_0|/n$. Общая скорость каскадного кода $R_{\text{общ}}$, определяющая значение плотности записи в ячейку флеш-памяти, есть произведение R и R_0 , т.е. $R_{\text{общ}} = RR_0 = K \log_2 |B_0| / (N_1 n)$ бит/ячейка.

Внутренний код B_0 определяется следующей формулой:

$$B_0 = \sum_{j=0}^{L-1} 2^j \psi \left(\left\{ \mathbf{a}_l^{(j)} \mathbf{G}_j \mid 1 \leq l \leq \exp_2(k_j) \right\} \right) + \sum_{j=L}^{m-1} 2^j \psi \left(\left\{ \mathbf{a}_l^{(j)} \mid 1 \leq l \leq \exp_2(n) \right\} \right),$$

где ψ – естественное вложение (ψ -вложение) $\mathbb{F}_2^n$ в n -мерную целочисленную решетку Z^n , $\{\mathbf{a}_l^{(j)} \mid 1 \leq l \leq \exp_2(k_j)\}$ – совокупность двоичных векторов размера k_j , $\{\mathbf{a}_l^{(j)} \mid 1 \leq l \leq \exp_2(n)\}$ – совокупность двоичных векторов размера n , $\mathbf{G}_j$ – порождающая матрица линейного двоичного кода C_j , имеющая размер $k_j \times n$, $j = 0, 1, \dots, L$. Скорость внутреннего многоуровневого кода $R_0 = (m - L) + \sum_{j=0}^{L-1} k_j / n$.

Компонентные коды $C_0, C_1, \dots, C_{L-1}$ внутреннего кода B_0 выбираются исходя из обменных соотношений между скоростью R_0 , минимальным евклидовым расстоянием δ_0 и сложностью декодирования. Отметим, что при выполнении ряда условий, указанных в [30, 31], внутренний код B_0 представляет собой конечное подмножество решетки Барнса-Уолла, что позволяет использовать для декодирования кода B_0 достаточно эффективную (хоть и не оптимальную) процедуру – bounded-distance decoding (BDD). Ограничимся далее рассмотрением двухуровневых ($L=2$) конструкций внутреннего кода для NAND флеш-памяти с четырьмя ранжированными по уровню состояниями.

Примеры каскадных конструкций с внутренним двухуровневым кодом B_0 представлены в таблицах 2 и 3, которые частично повторяют содержимое соответствующих таблиц, приведенных в [30, 31] и [33]. В таблице 2 перечислены варианты каскадного кодирования, в которых внутренний код построен на основе решеток Барнса-Уолла, а в таблице 3 – на основе циклически усеченного сверточного кода (tail-biting code, TB) C_0 и кода с проверкой на четность (single parity check code, SPC) C_1 . Номер варианта в общей нумерации показан в первом столбце таблиц, а выбор конкретного варианта(-ов) из нескольких, если он возможен, выделен полужирным шрифтом. Описание структуры кодов, перечисленных в таблицах 2 и 3, более подробно изложено в [30, 31] и [33] соответственно.

4. Анализ помехоустойчивости кодовой конструкции для смешанной NL модели. Приводимый далее анализ помехоустойчивости рассматриваемой кодовой конструкции позволяет получить обменные соотношения между вероятностью ошибки, плотностью записи и числом циклов перезаписи.

Таблица 2. Параметры каскадных конструкций с внутренним кодом на основе решеток Барнса-Уолла

№	Размерность решётки n	Исходная решетка Λ_0	Коды C_0 и C_1	Объём внутреннего кода $ B_0 $	Возможные значения h и длины кода РС N_1	Плотность записи, $R_{общ}$, бит/ячейка
1,2	8	E_8	$C_0 = (8,4,4)$ $C_1 = (8,8,1)$	2^{12}	$N_1 = 16, h = 3$ $N_1 = 64, h = 2$ $N_1 = 4096, h = 1$	1.5 R
3	8	RE_8	$C_0 = (8,1,8)$ $C_1 = (8,7,2)$	2^8	$N_1 = 256, h = 1$	1.0 R
4	16	Λ_{16}	$C_0 = (16,5,8)$ $C_1 = (16,15,2)$	2^{20}	$N_1 = 16, h = 5$ $N_1 = 32, h = 4$ $N_1 = 1024, h = 2$	1.25 R

Таблица 3. Параметры каскадных конструкций с внутренним двухуровневым кодом B_0 на основе циклически усечённого сверточного кода C_0 и кода с проверкой на четность C_1

№	Коды C_0 и C_1	Объём внутреннего кода $ B_0 $	Возможные значения h и длин кода РС N_1	Плотность записи, $R_{общ}$, бит/ячейка
5	$C_0 = (8,4,4)$ $C_1 = (8,7,2)$	2^{11}	$N_1 = 2048, h = 1$	1.375 R
6	$C_0 = (12,3,6)$ $C_1 = (12,11,2)$	2^{14}	$N_1 = 128, h = 2$	1.167 R
7	$C_0 = (18,9,6)$ $C_1 = (18,17,2)$	2^{26}	$N_1 = 8192, h = 2$	1.444 R

Полученные численным образом обменные соотношения позволяют оценить эффективность использования рассматриваемых вариантов помехоустойчивого кодирования для увеличения выносливости NAND флеш-памяти при поддержании требуемой вероятности ошибки. Оценка помехоустойчивости представленных каскадных кодовых конструкций включает в себя два этапа. Первый этап состоит в вычислении (оценке) вероятности ошибки декодирования по максимуму правдоподобия (МП) слова внутреннего кода. На втором этапе вычисляется (оценивается) вероятность ошибки декодирования слова внешнего кода (кода Рида-Соломона) с использованием результатов, полученных на первом этапе. Стандартная оценка вероятности ошибки на бит при декодировании недовичного внешнего кода в канале с независимыми ошибками может быть записана как [49]:

$$P_b \leq \frac{1}{2} \sum_{i=t+1}^{N_1} \frac{i+t}{N_1} C_{N_1}^i p_e^i (1-p_e)^{N_1-i}, \quad (18)$$

где p_e – вероятность ошибки декодирования символа внутреннего кода по МП, N_1 – длина внешнего кода, t – максимальное число ошибок, исправляемых кодом. Основная проблема при вычислении правой части оценки (18) состоит в вычислении аддитивной оценки вероятности ошибки декодирования внутреннего кода p_e . Она решается с использованием ранее разработанного подхода [30, 31, 50]. Итоговое выражение для верхней границы вероятности p_e при этом получается в виде:

$$p_e \leq \frac{1}{\pi M} \int_0^\infty \operatorname{Re} \frac{D(\alpha - j\beta) - M}{\beta + j\alpha} d\alpha, \quad (19)$$

где M – число слов внутреннего кода, и:

$$D(\omega) = \sum_{\mathbf{x}} \sum_{\mathbf{x}'} \prod_{l=1}^n c_z(\omega; x^{(l)}, x'^{(l)}). \quad (20)$$

Значение параметра β в правой части (19) выбирается в некоторых допустимых пределах; хорошим выбором для многих примеров может служить значение $\beta = 1/2$ [30, 31, 50]. В выражении (20):

$$c_{z(\omega; x, x')} = \overline{\exp(j\omega z(x, x'))}, \quad (21)$$

где $z(x, x') = \ln(p_{y|x'}(y|x') / p_{y|x}(y|x))$, черта сверху здесь обозначает усреднение по распределению, заданному условной плотностью вероятности $p_{y|x}(y|x)$, а суммирование по $\mathbf{x}$ и $\mathbf{x}'$ в правой части (20) выполняется по всем словам внутреннего кода. В качестве технических приемов для вычисления границы (19) используются: 1) решетчатое представление внутреннего кода, 2) рассмотрение произведения этой решетки на себя, 3) разметка решетки-произведения с использованием характеристических функций вида (21), 4) получение значений функции $D(\omega)$ (20), и 5) численное нахождение значения интеграла (19) при значении параметра $\beta = 1/2$. Подробное описание процесса вычисления значений функции $D(\omega)$ дано в работе [50].

Как следует из описания приведенного подхода, для вычисления оценки вероятности p_e требуется найти значения характеристических функций (21), которые можно выразить как:

$$c_{z(\omega; x, x')} = \overline{\exp(j\omega z(x, x'))} = \int_{-\infty}^{\infty} p_{y|x'}(y|x')^{j\omega} p_{y|x}(y|x)^{1-j\omega} dy, \quad (22)$$

для $\omega = \alpha - j\beta$, где $0 < \alpha < \infty$, $\beta = 1/2$ для различных условных ф.п.в $p_{y|x}(y|x)$, задающих модель канала, и всех входных значений канала x, x' . Ранее [30, 31, 50] этот подход применялся для случаев, когда выражения для характеристических функций (22) удавалось получить аналитически в замкнутом виде. К таким случаям, в частности, относится и ID-AGN модель канала флеш-памяти, которая рассматривалась в работах [30, 31]. Для рассматриваемой в настоящем исследовании смешанной NL модели получить выражение для характеристических функций (22) в замкнутом виде не удастся. Поэтому их вычисление выполнялось численным методом для ряда значений переменной α при соответствующем выборе пределов изменения этой величины. Набор значений характеристических функций (22), вычисленных для значений аргумента $\omega = \alpha - j\beta$, $\beta = 1/2$, далее использовался при вычислениях по формулам (20), (19) и (18).

В качестве примеров применения этой техники оценки характеристик корректирующего кодирования для канала, заданного смешанной NL моделью, были рассмотрены варианты кодовых схем, перечисленных в таблицах 2 и 3. Результаты вычисления границ (19) и (18) для семи вариантов кодовых схем, указанных в таблицах 2 и 3, в зависимости от значений нормированного числа циклов перезаписи $N/N_{\text{ном}}$ приведены на рисунках 3 – 6.

Приведенные на рисунках 3 – 6 графики в совокупности представляют собой искомые обменные соотношения между вероятностью ошибки P_b , плотностью записи $R_{\text{обц}}$ и числом циклов перезаписи N . Из рассмотрения данных, представленных на рисунках 3 – 6, следует:

1. Для каждой из рассмотренных конструкций, при числе циклов перезаписи N в диапазоне от $0.1N_{\text{ном}}$ до $N_{\text{ном}}$, незначительное увеличение плотности записи (это увеличение может быть даже несколько тысячных) сопровождается существенным ростом вероятности ошибки (на порядок и более).

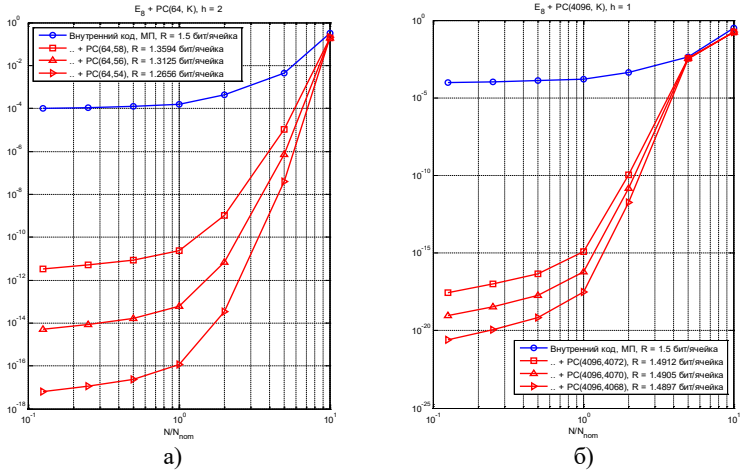


Рис. 3. Зависимость границы вероятности P_b от нормированного числа циклов перезаписи $N/N_{\text{ном}}$: а) кодовая схема № 1, б) кодовая схема № 2

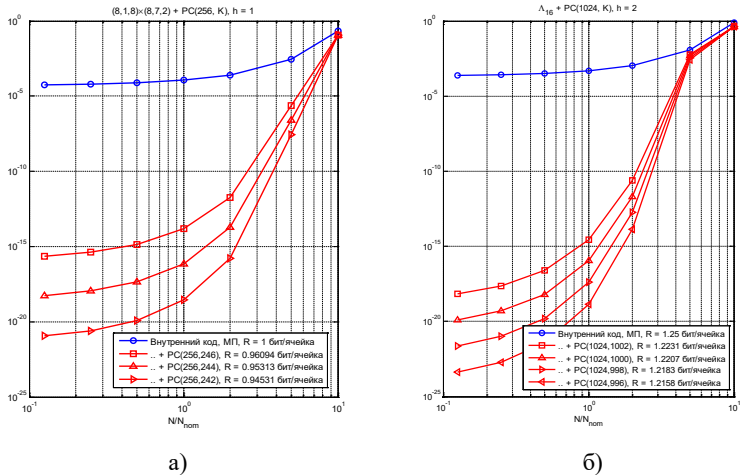


Рис. 4. Зависимость границы вероятности P_b от нормированного числа циклов перезаписи $N/N_{\text{ном}}$: а) кодовая схема № 3, б) кодовая схема № 4

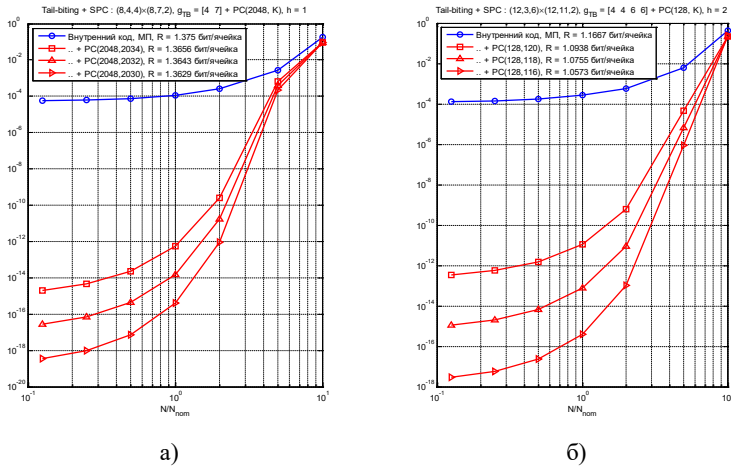


Рис. 5. Зависимость границы вероятности P_b от нормированного числа циклов перезаписи N/N_{nom} : а) кодовая схема № 5 (слева), б) кодовая схема № 6

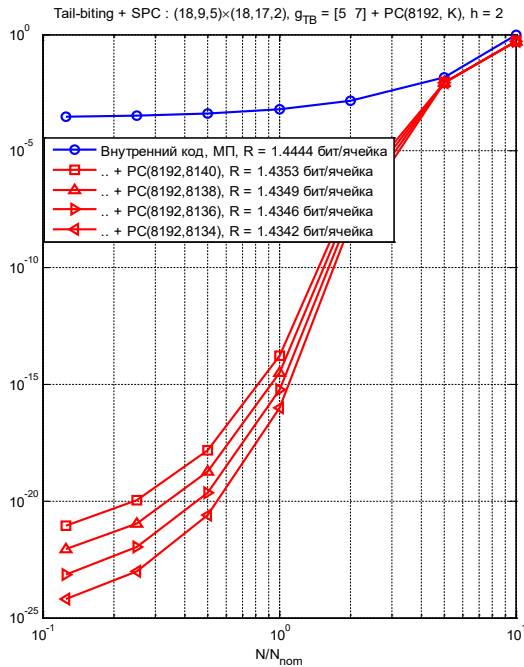


Рис. 6. Зависимость границы вероятности P_b от нормированного числа циклов перезаписи N/N_{nom} , кодовая схема № 7

2. Число циклов перезаписи $N = N_{\text{ном}}$ может действительно рассматриваться, в большинстве случаев, как своего рода граничное значение. Судя по приведенным графикам, увеличение числа циклов перезаписи N на порядок – от $0.1N_{\text{ном}}$ до $1.0N_{\text{ном}}$, приводит к росту вероятности ошибки на $1.5...2$ порядка, тогда как дальнейшее увеличение числа циклов перезаписи N на порядок – до $10N_{\text{ном}}$, сопровождается ростом вероятности ошибки на $12(!)$ и более порядков. Отметим, что при использовании гауссовской аппроксимации зависимость вероятности ошибки от числа циклов перезаписи N примерно линейна (в логарифмическом масштабе), поэтому выделить такого рода граничное значение в рамках гауссовой модели, видимо, не представляется возможным.

3. За счет очень незначительного снижения плотности записи $R_{\text{общ}}$ можно обеспечить увеличение допустимого числа циклов перезаписи в $2-2.5$ раза, относительно величины $N_{\text{ном}}$, при сохранении требуемого значения вероятности ошибки на бит. Ясно, что увеличение числа циклов перезаписи приводит к увеличению срока службы при неизменной интенсивности обменов с памятью, либо к увеличению интенсивности обмена данными при сохранении времени эксплуатации системы хранения. В обоих случаях получаем эффект увеличения выносливости флеш-памяти, что обеспечивает заметное улучшение эксплуатационных характеристик систем долговременного хранения данных. Одним из иллюстративных примеров такого обмена «плотность записи – число циклов перезаписи» может служить рисунок 5, из которого следует, что при снижении плотности записи с 1.2231 до 1.2158 бит/ячейка вероятность ошибки на бит остается в пределах $\approx 2 \cdot 10^{-15}$ при значении $N/N_{\text{ном}} = 2$.

4. Конструкции внутреннего кода вида TB/SPC могут обеспечить внушительный выигрыш (до нескольких порядков) по приемлемому числу циклов перезаписи, относительно кодов на основе решеток Барнса-Уолла, при равной плотности записи и надежности хранения данных. Отметим, что этот выигрыш имеет место как при гауссовской аппроксимации искажений сигнала записи/считывания (это было показано в работе [28]), так и при смешанной NL модели, причем в рамках смешанной NL модели выигрыш может достигать нескольких порядков.

5. Заключение. В настоящей работе представлены результаты применения разработанных авторами подходов к оценке эффективности корректирующего кодирования для каналов хранения данных многоуровневой NAND флеш-памяти. Эффективность корректирующего кодирования оценивается через обменное

соотношение, связывающее повышение надежности воспроизведения данных при их передаче или хранении за счет снижения скорости кодирования. Применительно к рассматриваемой модели канала флеш-памяти эти показатели выражаются в повышении числа циклов перезаписи (выносливости ячейки флеш-памяти) за счет снижения плотности записи при сохранении требуемого значения вероятности ошибки. Каскадный принцип построения корректирующего кода позволяет строить длинные коды с высокой корректирующей способностью. При этом достигается высокая корректирующая способность при умеренной сложности реализации, что порой имеет решающее значение при реализации систем кодирования для флеш-памяти.

Отличительной особенностью настоящей публикации является использование смешанной гауссово-лапласовой модели (NL модели) для описания искажений, возникающих при считывании из ячейки NAND флеш-памяти. Такая модель используется наряду с ID-AGN моделью и в некоторых случаях оказывается более предпочтительной, так как достаточно адекватно отражает искажения сигнала записи/считывания для широкого диапазона изменений основных параметров – числа циклов перезаписи и длительности хранения. В качестве схемы кодирования рассмотрена становящаяся все более популярной каскадная схема кодирования. В ней внутренняя ступень представляет собой многоуровневый код, построенный либо на основе решеток Барнса-Уолла, либо на основе циклически усеченного сверточного кода и кода с проверкой на четность, а в качестве внешней ступени используется код Рида-Соломона. Для этой конфигурации подразумевается мягкое декодирование внутреннего кода по МП и алгебраическое декодирование внешнего кода с исправлением сравнительно небольшого числа ошибок. Разработанный ранее подход [30, 31, 50] к анализу помехоустойчивости декодера внутреннего кода был применен в настоящей работе для анализа в условиях принятой смешанной NL модели. В ходе исследования оказалось, что особенности, обусловленные природой используемой модели, не позволяют получить в явном виде выражения для характеристических функций, используемых при вычислении границы вероятности ошибки декодирования внутреннего кода по МП. Это привело к определенным трудностям вычислительного характера, которые были преодолены, и в итоге были получены в численной форме обменные соотношения между вероятностью ошибки, плотностью записи и числом циклов

перезаписи. Полученные обменные соотношения позволили оценить эффективность использования представленных вариантов каскадного кодирования для увеличения выносливости NAND флеш-памяти при поддержании требуемой вероятности ошибки. Среди представленных в разделе 4 конечных результатов можно отметить два результата, важных в прикладном плане. Во-первых, в рамках смешанной NL модели можно достаточно определенно выделить такую важную характеристику, как номинальное число циклов перезаписи, представляющего собой своего рода граничное значение, при превышении которого помехоустойчивость резко ухудшается. Во-вторых, предложенные конструкции позволяют за счет очень незначительного снижения плотности записи обеспечить увеличение этого граничного значения числа циклов перезаписи в 2–2.5 раза при сохранении требуемого значения вероятности ошибки на бит.

Литература

1. Advances in Non-Volatile Memory and Storage Technology. Second Edition / Eds.: Magyari-Köpe B., Nishi Y. // Amsterdam.: Woodhead Publishing. 2019. 662 p.
2. Gao B. Emerging Non-Volatile Memories for Computation-in-Memory // 25th Asia and South Pacific Design Automation Conference (ASP-DAC). 2020. pp. 381–384. DOI: 10.1109/ASP-DAC47756.2020.9045394.
3. Ishimaru K. Future of Non-Volatile Memory — From Storage to Computing // IEEE International Electron Devices Meeting (IEDM). 2019. pp. 1.3.1–1.3.6. DOI: 10.1109/IEDM19573.2019.8993609.
4. Ishimaru K. Non-Volatile Memory Technology for Data Age // 14th IEEE International Conference on Solid-State and Integrated Circuit Technology (ICSICT). 2018. pp. 1–4. DOI: 10.1109/ICSICT.2018.8564815.
5. Gerardin S., Paccagnella A. Present and future non-volatile memories for space // IEEE Transactions on Nuclear Science. 2010. vol. 57. no. 6. pp. 3016–3039. DOI: 10.1109/TNS.2010.2084101.
6. Nonvolatile Memory Technologies with Emphasis on Flash: A Comprehensive Guide to Understanding and Using Flash Memory Devices / Eds.: Brewer J., Jill M. // Wiley–IEEE Press. 2008. 792 p.
7. Kang J., Huang P., Han R., Xiang Y., Cui X., Liu X. Flash-based Computing in-Memory Scheme for IOT // Proceedings of the 2019 IEEE 13th International Conference on ASIC (ASICON). 2019. pp. 1–4. DOI: 10.1109/ASICON47005.2019.8983502.
8. Bennett S., Sullivan J. NAND Flash Memory and Its Place in IoT // Proceedings of the 2021 32nd Irish Signals and Systems Conference (ISSC). 2021. pp. 1–6. DOI: 10.1109/ISSC52156.2021.9467859.
9. Aritome S. NAND Flash Memory Technologies // Hoboken.: Wiley. 2016. 432 p.
10. Ohshima S.J. Empowering Next-Generation Applications through FLASH Innovation // Proceedings of the 2020 IEEE Symposium on VLSI Technology 2020. pp. 1–4. DOI: 10.1109/VLSITechnology18217.2020.9265031.
11. Janukowicz J. How New QLC SSDs Will Change the Storage Landscape. IDC White Paper. 2018. Available at: <https://www.micron.com/-/media/client/global/documents/products/white->

- paper/how_new_qlc_ssds_will_change_the_storage_landscape.pdf?la=en (accessed: 13.10.2022).
12. Goda A. Recent Progress on 3D NAND Flash Technologies // *Electronics*. 2021. vol. 10. no. 24. pp. 3156. DOI: 10.3390/electronics10243156.
 13. Luo Y., Ghose S., Cai Y., Haratsch E., Mutlu O. Enabling Accurate and Practical Online Flash Channel Modeling for Modern MLC NAND Flash Memory // *IEEE Journal on Selected Areas in Communications*. 2016. vol. 34. no. 9. pp. 2294–2311. DOI: 10.1109/JSAC.2016.2603608.
 14. Liu W. et al., Modeling of Threshold Voltage Distribution in 3D NAND Flash Memory // 2021 Design, Automation & Test in Europe Conference & Exhibition (DATE). 2021. pp. 1729–1732. DOI: 10.23919/DATE51398.2021.9473974.
 15. Grupp L., Davis J., Swanson S. The bleak future of NAND flash memory // *Proceedings of the 10th USENIX Conference on File and Storage Technologies (FAST'12)*. 2012. pp. 2.
 16. Mielke N. et al. Bit error rate in NAND flash memories // *Proceedings of IEEE International Reliability Physics Symposium*. 2008. pp. 9–19.
 17. Liu J., Hsu C., Wang I., Hou T. Categorization of multilevel-cell storage-class memory: an RRAM example // *IEEE Transactions on Electron Devices*. 2015. vol. 62. no. 8. pp. 2510–2516. DOI: 10.1109/TED.2015.2444663.
 18. Solid-State Drive (SSD) Requirements and Endurance Test Method (JESD218) // JEDEC Solid State Technology Association. 2010.
 19. Yoon J., Tressler G. Advanced Flash Technology Status, Scaling Trends & Implications to Enterprise SSD Technology Enablement // *Flash Memory Summit*. 2012.
 20. Maislos A. A New Era in Embedded Flash Memory // *Flash Memory Summit*. 2011.
 21. Fan B., Qin M., Siegel P. Enhancing the Expected Lifetime of NAND Flash by Short q-Ary WOM Codes // *IEEE Communications Letters*. 2018. vol. 22. no. 7. pp. 1302–1305. DOI: 10.1109/LCOMM.2017.2776200.
 22. Chee Y., Kiah H., Vardy A., Yaakobi E. Explicit and Efficient WOM Codes of Finite Length. // *IEEE Transactions on Information Theory*. 2020. vol. 66. no. 5. pp. 2669–2682. DOI: 10.1109/TIT.2019.2946483.
 23. Yaakobi E., Yucovich A., Maor G., Yadgar G. When do WOM codes improve the erasure factor in flash memories? *IEEE International Symposium on Information Theory (ISIT)*. 2015. pp. 2091–2095. DOI: 10.1109/ISIT.2015.7282824.
 24. Jiang A., Li Y., Gad E., Langberg M., Bruck J. Joint rewriting and error correction in write-once memories // *IEEE International Symposium on Information Theory (ISIT)*. 2013. pp. 1067–1071. DOI: 10.1109/ISIT.2013.6620390.
 25. Solomon A., Cassuto Y. Error-Correcting WOM Codes: Concatenation and Joint Design // *IEEE Transactions on Information Theory*. 2019. vol. 65. no. 9. pp. 5529–5546. DOI: 10.1109/TIT.2019.2917519.
 26. Micheloni R., Marelli A., Ravasio R. Error Correction Codes for Non-Volatile Memories // *Springer Science & Business Media*. 2008. 338 p.
 27. Li S., Zhang T. Improving multi-level NAND flash memory storage reliability using concatenated BCH- TCM coding // *IEEE Transactions on Very Large Scale Integration (VLSI) Systems*. 2010. vol. 18. no. 10. pp. 1412–1420. DOI: 10.1109/TVLSI.2009.2024154.
 28. Dong G., Xie N., Zhang T. On the Use of Soft-Decision Error-Correction Codes in NAND Flash Memory // *IEEE Transactions on Circuits and Systems I: Regular Papers*. 2011. vol. 58. no. 2. pp. 429–439. DOI: 10.1109/TCSI.2010.2071990.
 29. Dolecek L., Cassuto Y. Channel coding for nonvolatile memory technologies: Theoretical advances and practical considerations // *Proceedings of the IEEE*. 2017. vol. 105. no. 9. pp. 1705–1724. DOI: 10.1109/JPROC.2017.2694613.

30. Таубин Ф.А., Трофимов А.Н. Каскадное кодирование на основе многомерных решеток и кодов Рида-Соломона для многоуровневой флеш-памяти // Труды СПИИРАН. 2018. Вып. 2(57). С. 75–103. DOI: 10.15622/sp.57.4.
31. Таубин Ф.А., Трофимов А.Н. Каскадное кодирование для многоуровневой флеш-памяти с исправлением ошибок малой кратности во внешней ступени // Труды СПИИРАН. 2019. Вып. 18(5). С. 1149–1181. DOI: 10.15622/sp.2019.18.5.1149-1181.
32. IEEE Std 1890-2018 // IEEE Standard for Error Correction Coding of Flash Memory Using Low-Density Parity Check Codes. 2019. pp. 1–51.
33. Таубин Ф.А., Трофимов А.Н. Каскадное кодирование с внутренним двухуровневым tail-biting/parity check кодом для многоуровневой flash памяти // XXIII международная научная конференция Волновая электроника и инфокоммуникационные системы: Сб. научн. тр. конференции. 2020. С. 354–361.
34. Трофимов А.Н., Таубин Ф.А. Анализ каскадного кодирования для многоуровневой флеш-памяти с использованием смешанной Normal-Laplace модели // XXV международная научная конференция Волновая электроника и инфокоммуникационные системы: Сб. научн. тр. конференции. 2022. С. 109–113.
35. Cai Y., Ghose S., Haratsch E., Luo Y., Mutlu O. Error characterization, mitigation, and recovery in Flash Memory-Based solid-state drives // Proceedings of IEEE. 2017. vol. 105. no. 9. pp. 1666–1704. DOI: 10.1109/JPROC.2017.2713127.
36. Dong G., Pan Y., Xie N., Varanasi C., Zhang T. Estimating information-theoretical NAND flash memory storage capacity and its implication to memory system design space exploration // IEEE Transactions on Very Large Scale Integration (VLSI) Systems. 2012. vol. 20. no. 9. pp. 1705–1714. DOI: 10.1109/TVLSI.2011.2160747.
37. Park S., Moon J. Characterization of Inter-Cell Interference in 3D NAND Flash Memory // IEEE Transactions on Circuits and Systems I: Regular Papers. 2021. vol. 68. no. 3. pp. 1183–1192. DOI: 10.1109/TCSI.2020.3047484.
38. Moon J., No J., Lee S., Kim S., Choi S., Song Y. Statistical Characterization of Noise and Interference in NAND Flash Memory // IEEE Transactions on Circuits and Systems I: Regular Papers. 2013. vol. 60. no. 8. pp. 2153–2164. DOI: 10.1109/TCSI.2013.2239116.
39. Wang X., Dong G., Pan L., Zhou R. Error Correction Codes and Signal Processing in Flash Memory. / Ed.: Igor Stievano // IntechOpen. 2011. pp. 57–82. Available at: www.intechopen.com/books/flash-memories/error-correction-codes-and-signal-processing-in-flash-memory (accessed 13.10.2022).
40. Wang K., Du G., Lun Z., Liu X. Investigation of Retention Noise for 3-D TLC NAND Flash Memory // IEEE Journal of the Electron Devices Society. 2019. vol. 7. pp. 150–157. DOI: 10.1109/JEDS.2018.2886359.
41. Luo Y. et al. Improving 3D NAND Flash Memory Lifetime by Tolerating Early Retention Loss and Process Variation // Proceedings of the ACM on Measurement and Analysis of Computing Systems. 2018. vol. 2. no. 3. pp 1–48. DOI: 10.1145/3224432.
42. Liu W. et al. Characterization Summary of Performance, Reliability, and Threshold Voltage Distribution of 3D Charge-Trap NAND Flash Memory // ACM Transactions on Storage. 2022. vol. 18. no. 2. pp 1–25. DOI: 10.1145/3491230.
43. Cai Y., Haratsch E., Mutlu O., Mai K. Threshold voltage distribution in MLC NAND flash memory: Characterization, analysis, and modeling // Proceedings of Design, Automation and Test in Europe Conference. 2013. pp. 1285–1290. DOI: 10.7873/DATE.2013.266.
44. Li Q., Jiang A., Haratsch E. Noise modeling and capacity analysis for NAND flash memories // Proceedings of IEEE International Symposium on Information Theory. 2014. pp. 2262–2266. DOI: 10.1109/ISIT.2014.6875236.

45. Ashrafi R., Arslan S., Pusane A. On the distribution of the threshold voltage in multi-level cell flash memories // *Physical Communication*. 2019. vol. 36. no. 1–2. pp. 1–21. DOI: 10.1016/j.phycom.2019.100747.
46. Parnell T., Papandreou N., Mittelholzer T., Pozidis H. Modelling of the Threshold Voltage Distributions of Sub-20nm NAND Flash Memory // *IEEE Global Communications Conference*. 2014. pp. 2351–2356. DOI: 10.1109/GLOCOM.2014.7037159.
47. Xu Q., Gong P., Chen T.M. Concatenated LDPC-TCM coding for reliable storage in multi-level flash memories // *Proceedings of the 9th International Symposium on Communication System, Networks & Digital Signal Processing (CSNDSP' 2014)*. 2014. pp. 166–170.
48. Kurkoski B.M. Coded modulation using lattices and Reed-Solomon codes, with applications to flash memories // *IEEE Transactions on Selected Areas in Communications*. 2014. vol. 32. no. 5. pp. 900–908. DOI: 10.1109/JSAC.2014.140510.
49. Кларк Дж., Кейн Дж. Кодирование с исправлением ошибок в системах цифровой связи / Под ред. Б.С. Цыбакова // М.: Радио и связь. 1987. 392 с.
50. Трофимов А.Н., Таубин Ф.А. Вычисление аддитивной границы вероятности ошибки декодирования с использованием характеристических функций // *Информационно-управляющие системы*. 2021. № 4. С. 71–85. DOI:10.31799/1684-8853-2021-4-71-85.

Трофимов Андрей Николаевич — канд. техн. наук, доцент, институт радиотехники и инфокоммуникационных технологий, кафедра инфокоммуникационных технологий и систем связи, Санкт-Петербургский государственный университет аэрокосмического приборостроения (СПбГУАП). Область научных интересов: теория цифровой связи, теория информации, методы помехоустойчивого кодирования. Число научных публикаций — 66. andrei.trofimov@vu.spb.ru; улица Большая Морская, 67, 190000, Санкт-Петербург, Россия; р.т.: +7(812)494-7052.

Таубин Феликс Александрович — д-р техн. наук, профессор, институт аэрокосмических приборов и систем, кафедра аэрокосмических компьютерных и программных систем, Санкт-Петербургский государственный университет аэрокосмического приборостроения (СПбГУАП). Область научных интересов: цифровые системы связи, методы помехоустойчивого кодирования, широкополосные системы, беспроводная связь. Число научных публикаций — 102. ftaubin@yahoo.com; улица Большая Морская, 67, 190000, Санкт-Петербург, Россия; р.т.: +7(812)494-7051.

Поддержка исследований. Работа выполнена при финансовой поддержке Министерства науки и высшего образования Российской Федерации, соглашение № FSRF-2020-0004.

A. TROFIMOV, F. TAUBIN

PERFORMANCE ANALYSIS OF CONCATENATED CODING TO INCREASE THE ENDURANCE OF MULTILEVEL NAND FLASH MEMORY*Trofimov A., Taubin F. Performance Analysis of Concatenated Coding to Increase the Endurance of Multilevel NAND Flash Memory.*

Abstract. The increasing storage density of modern NAND flash memory chips, achieved both due to scaling down the cell size, and due to the increasing number of used cell states, leads to a decrease in data storage reliability, namely, error probability, endurance (number of P/E cycling) and retention time. Error correction codes are often used to improve the reliability of data storage in multilevel flash memory. The effectiveness of using error correction codes is largely determined by the model accuracy that exhibits the basic processes associated with writing and reading data. The paper describes the main sources of disturbances for a flash cell that affect the threshold voltage of the cell in NAND flash memory, and represents an explicit form of the threshold voltage distribution. As an approximation of the obtained threshold voltage distribution, a Normal-Laplace mixture model was shown to be a good fit in multilevel flash memories for a large number of rewriting cycles. For this model, a performance analysis of the concatenated coding scheme with an outer Reed-Solomon code and an inner multilevel code consisting of binary component codes is carried out. The performed analysis makes it possible to obtain tradeoffs between the error probability, storage density, and the number of P/E cycling. The resulting tradeoffs show that the considered concatenated coding schemes allow, due to a very slight decrease in the storage density, to increase the number of P/E cycling up to 2–2.5 times than their nominal endurance specification while maintaining the required value of the bit error probability.

Keywords: multilevel NAND flash memory, threshold voltage distribution, Normal-Laplace mixture model, concatenated coding, performance analysis, endurance.

References

1. Advances in Non-Volatile Memory and Storage Technology. Second Edition. Eds.: Magyari-Köpe B., Nishi Y. Amsterdam.: Woodhead Publishing. 2019. 662 p.
2. Gao B. Emerging Non-Volatile Memories for Computation-in-Memory. 25th Asia and South Pacific Design Automation Conference (ASP-DAC). 2020. pp. 381–384. DOI: 10.1109/ASP-DAC47756.2020.9045394.
3. Ishimaru K. Future of Non-Volatile Memory – From Storage to Computing. IEEE International Electron Devices Meeting (IEDM). 2019. pp. 1.3.1–1.3.6. DOI: 10.1109/IEDM19573.2019.8993609.
4. Ishimaru K. Non-Volatile Memory Technology for Data Age. 14th IEEE International Conference on Solid-State and Integrated Circuit Technology (ICSICT). 2018. pp. 1–4. DOI: 10.1109/ICSICT.2018.8564815.
5. Gerardin S., Paccagnella A. Present and future non-volatile memories for space. IEEE Transactions on Nuclear Science. 2010. vol. 57. no. 6. pp. 3016–3039. DOI: 10.1109/TNS.2010.2084101.
6. Nonvolatile Memory Technologies with Emphasis on Flash: A Comprehensive Guide to Understanding and Using Flash Memory Devices. Eds.: Brewer J., Jill M. Wiley–IEEE Press. 2008. 792 p.
7. Kang J., Huang P., Han R., Xiang Y., Cui X., Liu X. Flash-based Computing in-Memory Scheme for IOT. Proceedings of the 2019 IEEE 13th International

- Conference on ASIC (ASICON). 2019. pp. 1–4. DOI: 10.1109/ASICON47005.2019.8983502.
8. Bennett S., Sullivan J. NAND Flash Memory and Its Place in IoT. Proceedings of the 2021 32nd Irish Signals and Systems Conference (ISSC). 2021. pp. 1–6. DOI: 10.1109/ISSC52156.2021.9467859.
9. Aritome S. NAND Flash Memory Technologies. Hoboken.: Wiley. 2016. 432 p.
10. Ohshima S.J. Empowering Next-Generation Applications through FLASH Innovation. Proceedings of the 2020 IEEE Symposium on VLSI Technology 2020. pp. 1–4. DOI: 10.1109/VLSITechnology18217.2020.9265031.
11. Janukowicz J. How New QLC SSDs Will Change the Storage Landscape. IDC White Paper. 2018. Available at: https://www.micron.com/-/media/client/global/documents/products/white-paper/how_new_qlc_ssds_will_change_the_storage_landscape.pdf?la=en (accessed 13.10.2022).
12. Goda A. Recent Progress on 3D NAND Flash Technologies. Electronics. 2021. vol. 10. no. 24. pp. 3156. DOI: 10.3390/electronics10243156.
13. Luo Y., Ghose S., Cai Y., Haratsch E., Mutlu O. Enabling Accurate and Practical Online Flash Channel Modeling for Modern MLC NAND Flash Memory. IEEE Journal on Selected Areas in Communications. 2016. vol. 34. no. 9. pp. 2294–2311. DOI: 10.1109/JSAC.2016.2603608.
14. Liu W. et al. Modeling of Threshold Voltage Distribution in 3D NAND Flash Memory. Design, Automation & Test in Europe Conference & Exhibition (DATE). 2021. pp. 1729–1732. DOI: 10.23919/DATE51398.2021.9473974.
15. Grupp L., Davis J., Swanson S. The bleak future of NAND flash memory. Proceedings of the 10th USENIX Conference on File and Storage Technologies (FAST'12). 2012. pp. 2.
16. Mielke N. et al. Bit error rate in NAND flash memories. Proceedings of IEEE International Reliability Physics Symposium. 2008. pp. 9–19.
17. Liu J., Hsu C., Wang L., Hou T. Categorization of multilevel-cell storage-class memory: an RRAM example. IEEE Transactions on Electron Devices. 2015. vol. 62. no. 8. pp. 2510–2516. DOI: 10.1109/TED.2015.2444663.
18. Solid-State Drive (SSD) Requirements and Endurance Test Method (JESD218). JEDEC Solid State Technology Association. 2010.
19. Yoon J., Tressler G. Advanced Flash Technology Status, Scaling Trends & Implications to Enterprise SSD Technology Enablement. Flash Memory Summit. 2012.
20. Maislos A. A New Era in Embedded Flash Memory. Flash Memory Summit. 2011.
21. Fan B., Qin M., Siegel P. Enhancing the Expected Lifetime of NAND Flash by Short q-Ary WOM Codes. IEEE Communications Letters. 2018. vol. 22. no. 7. pp. 1302–1305. DOI: 10.1109/LCOMM.2017.2776200.
22. Chee Y., Kiah H., Vardy A., Yaakobi E. Explicit and Efficient WOM Codes of Finite Length. IEEE Transactions on Information Theory. 2020. vol. 66. no. 5. pp. 2669–2682. DOI: 10.1109/TIT.2019.2946483.
23. Yaakobi E., Yucovich A., Maor G., Yadgar G. When do WOM codes improve the erasure factor in flash memories? IEEE International Symposium on Information Theory (ISIT). 2015. pp. 2091–2095. DOI: 10.1109/ISIT.2015.7282824.
24. Jiang A., Li Y., Gad E., Langberg M., Bruck J. Joint rewriting and error correction in write-once memories. IEEE International Symposium on Information Theory (ISIT). 2013. pp. 1067–1071. DOI: 10.1109/ISIT.2013.6620390.
25. Solomon A., Cassuto Y. Error-Correcting WOM Codes: Concatenation and Joint Design. IEEE Transactions on Information Theory. 2019. vol. 65. no. 9. pp. 5529–5546. DOI: 10.1109/TIT.2019.2917519.
26. Micheloni R., Marelli A., Ravasio R. Error Correction Codes for Non-Volatile Memories. Springer Science & Business Media. 2008. 338 p.

27. Li S., Zhang T. Improving multi-level NAND flash memory storage reliability using concatenated BCH- TCM coding. *IEEE Transactions on Very Large Scale Integration (VLSI) Systems*. 2010. vol. 18. no. 10. pp. 1412–1420. DOI: 10.1109/TVLSI.2009.2024154.
28. Dong G., Xie N., Zhang T. On the Use of Soft-Decision Error-Correction Codes in NAND Flash Memory. *IEEE Transactions on Circuits and Systems I: Regular Papers*. 2011. vol. 58. no. 2. pp. 429–439. DOI: 10.1109/TCSI.2010.2071990.
29. Dolecek L., Cassuto Y. Channel coding for nonvolatile memory technologies: Theoretical advances and practical considerations. *Proceedings of the IEEE*. 2017. vol. 105. no. 9. pp. 1705–1724. DOI: 10.1109/JPROC.2017.2694613.
30. Taubin F.A., Trofimov A.N. [Concatenated Reed–Solomon/lattice coding for multi-level flash memory]. *Trudy SPIIRAN – SPIIRAS Proceedings*. 2018. vol. 2(57). pp. 75–103. DOI: 10.15622/sp.57.4. (In Russ.).
31. Taubin F.A., Trofimov A.N. [Concatenated coding for multi-level flash memory with low error correction capabilities in outer stage]. *Trudy SPIIRAN – SPIIRAS Proceedings*. 2019. vol. 18. no. 5. pp. 1149–1181. DOI: 10.15622/sp.2019.18.5.1149-1181. (In Russ.).
32. IEEE Std 1890-2018. *IEEE Standard for Error Correction Coding of Flash Memory Using Low-Density Parity Check Codes*. 2019. pp. 1–51.
33. Taubin F.A., Trofimov A.N. [Concatenated coding with inner two-level TB/SPC code for multi-level flash memory]. *Volnovaja jelektronika i infokommunikacionnye sistemy: Sb. nauchn. tr. XXIII mezhdunarodnoj konf. [XXIII International Conference Wave electronics and infocommunication systems: Collected papers]*. 2020. pp. 354–361. (In Russ.).
34. Trofimov A.N., Taubin F.A. [Analysis of concatenated coding scheme for the multilevel flash memory using normal-Laplace mixture]. *Volnovaja jelektronika i infokommunikacionnye sistemy: Sb. nauchn. tr. XXV mezhdunarodnoj konf. [XXV International Conference Wave electronics and infocommunication systems: Collected papers]*. 2022. pp. 109–113. (In Russ.).
35. Cai Y., Ghose S., Haratsch E., Luo Y., Mutlu O. Error characterization, mitigation, and recovery in Flash Memory-Based solid-state drives. *Proceedings of IEEE*. 2017. vol. 105. no. 9. pp. 1666–1704. DOI: 10.1109/JPROC.2017.2713127.
36. Dong G., Pan Y., Xie N., Varanasi C., Zhang T. Estimating information-theoretical NAND flash memory storage capacity and its implication to memory system design space exploration. *IEEE Transactions on Very Large Scale Integration (VLSI) Systems*. 2012. vol. 20. no. 9. pp. 1705–1714. DOI: 10.1109/TVLSI.2011.2160747.
37. Park S., Moon J. Characterization of Inter-Cell Interference in 3D NAND Flash Memory. *IEEE Transactions on Circuits and Systems I: Regular Papers*. 2021. vol. 68. no. 3. pp. 1183–1192. DOI: 10.1109/TCSI.2020.3047484.
38. Moon J., No J., Lee S., Kim S., Choi S., Song Y. Statistical Characterization of Noise and Interference in NAND Flash Memory. *IEEE Transactions on Circuits and Systems I: Regular Papers*. 2013. vol. 60. no. 8. pp. 2153–2164. DOI: 10.1109/TCSI.2013.2239116.
39. Wang X., Dong G., Pan L., Zhou R. Error Correction Codes and Signal Processing in Flash Memory. Ed.: Igor Stievano. *IntechOpen*. 2011. pp. 57–82. Available at www.intechopen.com/books/flash-memories/error-correction-codes-and-signal-processing-in-flash-memory (accessed 13.10.2022).
40. Wang K., Du G., Lun Z., Liu X. Investigation of Retention Noise for 3-D TLC NAND Flash Memory. *IEEE Journal of the Electron Devices Society*. 2019. vol. 7. pp. 150–157. DOI: 10.1109/JEDS.2018.2886359.
41. Luo Y. et al. Improving 3D NAND Flash Memory Lifetime by Tolerating Early Retention Loss and Process Variation. *Proceedings of the ACM on Measurement and Analysis of Computing Systems*. 2018. vol. 2. no. 3. pp. 1–48. DOI: 10.1145/3224432.

42. Liu W. et al. Characterization Summary of Performance, Reliability, and Threshold Voltage Distribution of 3D Charge-Trap NAND Flash Memory. *ACM Transactions on Storage*. 2022. vol. 18. no. 2. pp 1–25. DOI: 10.1145/3491230.
43. Cai Y., Haratsch E., Mutlu O., Mai K. Threshold voltage distribution in MLC NAND flash memory: Characterization, analysis, and modeling. *Proceedings of Design, Automatin and Test in Europe Conference*. 2013. pp. 1285–1290. DOI: 10.7873/DATE.2013.266.
44. Li Q., Jiang A., Haratsch E. Noise modeling and capacity analysis for NAND flash memories. *Proceedings of IEEE International Symposium on Information Theory*. 2014. pp. 2262–2266. DOI: 10.1109/ISIT.2014.6875236.
45. Ashrafi R., Arslan S., Pusane A. On the distribution of the threshold voltage in multi-level cell flash memories. *Physical Communication*. 2019. vol. 36. no. 1–2. pp. 1–21. DOI: 10.1016/j.phycom.2019.100747.
46. Parnell T., Papandreou N., Mittelholz T., Pozidis H. Modelling of the Threshold Voltage Distributions of Sub-20nm NAND Flash Memory. *IEEE Global Communications Conference*. 2014. pp. 2351–2356. DOI: 10.1109/GLOCOM.2014.7037159.
47. Xu Q., Gong P., Chen T. Concatenated LDPC-TCM coding for reliable storage in multi-level flash memories. *Proceedings of the 9th International Symposium on Communication System, Networks & Digital Signal Processing (CSNDSP' 2014)*. 2014. pp. 166–170.
48. Kurkoski B. Coded modulation using lattices and Reed-Solomon codes, with applications to flash memories. *IEEE Transactions on Selected Areas in Communications*. 2014. vol. 32. no. 5. pp. 900–908. DOI: 10.1109/JSAC.2014.140510.
49. Clark G., Cain J. *Error-Correction Coding for Digital Communications*. Plenum Press, 1982. 432 p. (Russ. ed.: Klark Dzh., Kejn Dzh. Kodirovanie s isprav-leniem oshibok v sistemah cifrovoy svyazi. Moscow, Radio i svyaz' Publ. 1987. 392 p.)
50. Trofimov A.N., Taubin F.A. [Evaluation of the union bound for the decoding error probability using characteristic functions]. *Informatsionno-upravliaiushchie sistemy – Information and Control Systems*. 2021. no. 4. pp. 71–85. DOI: 10.31799/1684-8853-2021-4-71-85. (In Russ.).

Trofimov Andrey — Ph.D., Associate professor, Department of infocommunication technologies and communication systems, institute of radio engineering and infocommunication technologies, Saint Petersburg State University of Aerospace Instrumentation (SUAI). Research interests: communication theory, error-correcting coding, information theory. The number of publications — 66. andrei.trofimov@vu.spb.ru; 67, Bolshaya Morskaya St., 190000, St. Petersburg, Russia; office phone: +7(812)494-7052.

Taubin Feliks — Ph.D., Dr.Sci., Professor, Institute of aerospace instruments and systems, department of aerospace computer and software systems, Saint Petersburg State University of Aerospace Instrumentation (SUAI). Research interests: communication theory, error-correcting coding, spread spectrum systems, wireless communication. The number of publications — 102. ftaubin@yahoo.com; 67, Bolshaya Morskaya St., 190000, St. Petersburg, Russia; office phone: +7(812)494-7051.

Acknowledgements. This research is supported by the Ministry of Science and Higher Education of the Russian Federation, agreement no FSRF-2020-0004.

Д.В. ЕФАНОВ, Т.С. ПОГОДИНА
**ИССЛЕДОВАНИЕ СВОЙСТВ САМОДВОЙСТВЕННЫХ
КОМБИНАЦИОННЫХ УСТРОЙСТВ С КОНТРОЛЕМ
ВЫЧИСЛЕНИЙ НА ОСНОВЕ КОДОВ ХЭММИНГА**

Ефанов Д.В., Погодина Т.С. **Исследование свойств самодвойственных комбинационных устройств с контролем вычислений на основе кодов Хэмминга.**

Аннотация. Рассматривается новый подход к синтезу самопроверяемых устройств, основанный на контроле вычислений контролируемыми объектами с помощью кодов Хэмминга, проверочные символы (контрольные биты) которых описываются самодвойственными функциями. При этом структура работает в импульсном режиме, что фактически основано на внесении временной избыточности при построении самопроверяемого устройства. Это, к сожалению, приводит к некоторому снижению быстродействия, однако существенно повышает характеристики контролепригодности, что особенно актуально для устройств и систем критического применения, входные данные для которых изменяются не столь часто. Дается краткий обзор методов построения схем встроенного контроля на основе свойства самодвойственности вычисляемых функций. Приведены основные структуры организации схем встроенного контроля. Отмечены предполагаемые пути развития теории синтеза схем встроенного контроля на основе проверки принадлежности вычисляемых функций классу самодвойственных булевых функций. Установлены все возможные значения числа информационных символов для кодов Хэмминга, которые будут обладать свойством самодвойственности функций, описывающих контрольные биты. Кодеры таких кодов Хэмминга будут являться самодвойственными устройствами. Так как функции, описывающие контрольные биты кодов Хэмминга, являются линейными, то для того, чтобы они были самодвойственными необходимо, чтобы в каждой из них использовалось нечетное количество аргументов. Доказано, что число разрядов кодовых слов кодов Хэмминга с самодвойственными контрольными функциями равно $n=3+4l$, $l \in N_0$. Приводятся результаты моделирования самодвойственных устройств со схемами встроенного контроля по двум диагностическим признакам в среде Multisim. Предложен способ модификации структуры контроля вычислений по двум диагностическим признакам, позволяющий использовать любой линейный блочный код (не обязательно код Хэмминга). Он основан на дооснащении кодера устройством преобразования функций в самодвойственные. Фактически это устройство для формирования модифицированного кода. Доказано, что для получения модифицированного кода Хэмминга с самодвойственными контрольными функциями для случаев $n \neq 3+4l$, $l \in N_0$, достаточно сложить по модулю $M=2$ несамодвойственную контрольную функцию с функцией старшего информационного бита.

Ключевые слова: самопроверяемое комбинационное устройство, схема встроенного контроля, контроль вычислений на выходах комбинационных устройств, линейный блочный код, контроль вычислений по двум диагностическим признакам, контроль самодвойственности, контроль вычислений по кодам Хэмминга.

1. Введение. Для реализации высоконадежных и безопасных устройств автоматики и вычислительной техники требуется своевременно обнаруживать возникающие в процессе их эксплуатации

неисправности (устойчивые отказы и сбои), а также парировать их проявления в виде ошибок на линиях схем и неверных сигналов на выходах блоков и узлов. Для этого проектируемые устройства реализуют с контролепригодными, самопроверяемыми и отказоустойчивыми структурами, что требует специальных подходов к их разработке [1 – 5].

Для обнаружения неисправностей в процессе функционирования исходные устройства (назовем их объектами диагностирования) снабжаются дополнительными средствами технического диагностирования в виде схем встроенного контроля (СВК) [6, 7]. СВК косвенно контролируют в процессе эксплуатации возникновение неисправностей по результатам вычислений значений на рабочих выходах объектов диагностирования, либо же в специально выбранных контрольных точках. При фиксации факта возникновения неисправности отказавшее устройство (блок, подсистема и пр.) отключается от последующих каскадов, сигналы блокируются, а неверно вычисленные данные в последующем не используются. Устройства, снабжаемые СВК, фактически входят в состав более сложных устройств, позволяющих отключать объекты, осуществлять реконфигурацию архитектуры, а также запускать процесс восстановления после перезагрузки, либо ремонта объекта с выявленной неисправностью. Такое обустройство отказоустойчивой системы требует внесения существенной избыточности. При этом повсеместно используются методы информационного, временного и структурного резервирования на различных уровнях реализации отказоустойчивых систем: от микроуровня и резервирования самих элементарных составляющих до макроуровня защиты целых блоков и узлов [8].

При синтезе СВК применяются разнообразные подходы: от дублирования для сопоставления сигналов на одноименных выходах различных копий устройств до применения методов помехозащищенного и помехоустойчивого кодирования [9]. В процессе синтеза СВК ориентируются на заранее установленную модель неисправностей, которая с некоторой вероятностью покрывает реальное множество дефектов [10]. Например, модель одиночной константной неисправности (*stuck-at fault*) покрывает от 80 до 95 % реальных физических дефектов при использовании КМОП-технологии для реализации устройств [11].

В настоящей статье читателю предлагаются новые результаты в исследовании методов синтеза СВК с контролем вычислений сразу

же по двум диагностическим признакам – контролю принадлежности формируемых значений заранее выбранным линейным блоковым кодам, а также принадлежности каждой реализуемой функции установленному особому классу булевых функций. В качестве кода выбраны всем известные коды Хэмминга [12], а в качестве особого класса булевых функций рассматриваются самодвойственные функции [13]. В статье освещены особенности синтеза СВК с применением кодов Хэмминга, проверочные символы (контрольные биты) которых описываются самодвойственными булевыми функциями.

2. Основные результаты в теории синтеза схем встроенного контроля по признаку самодвойственности функций. Рассматриваемая в настоящей работе задача следует из более чем полувекового опыта ученых всего мира в части разработки методов синтеза СВК с применением временного и пространственного кодирования. Широко распространенными методами в практике синтеза СВК являются дублирование, методы, подразумевающие использование различных блоковых кодов (кодовые методы), методы, основанные на контроле принадлежности вычисляемых функций особым классам булевых функций [9]. Остановимся на кратком обзоре достижений ученых-диагностов в области синтеза СВК с контролем самодвойственности вычисляемых функций.

В работе [14] обсуждаются вопросы обнаружения ошибок в устройствах автоматики и вычислительной техники с помощью временной избыточности и применения свойств самодвойственных функций. Авторами данной статьи установлены свойства структуры, которые позволяют обнаруживать любые одиночные неисправности. Устройства автоматики и вычислительной техники, выходы которых описываются самодвойственными функциями, называются *самодвойственными устройствами* [15]. Некоторые типовые устройства являются самодвойственными, например, полный сумматор или сумматор по модулю $M=2$ [16]. Они исследуются и в современном периоде развития микроэлектроники. Например, в [17] моделируются простейшие устройства сложения двоичных чисел. В [18] показано, что логические элементы, реализующие самодвойственные функции, могут быть эффективно реализованы с применением реконфигурируемых нанотехнологий.

Известно [19], что любую булеву функцию можно преобразовать в самодвойственную с использованием всего одной избыточной переменной. Поэтому путем модернизации структуры

(ресинтеза) можно любое устройство преобразовать в самодвойственное.

Методы преобразования булевых функций в самодвойственные описаны, например, в [20, 21]. Один из методов состоит в замене в структуре исходного устройства несамодвойственных элементов на самодвойственные аналоги. Другой метод базируется на использовании альтернативного сигнала a и известного преобразования К.Э. Шеннона, заключающегося в том, что любую булеву функцию можно разложить по любой переменной по формуле [13]:

$$\begin{aligned} f(x_1, \dots, x_j, \dots, x_t) = \\ = x_j f(x_1, \dots, 1, \dots, x_t) \vee \overline{x_j} f(x_1, \dots, 0, \dots, x_t), j = \overline{1, t}. \end{aligned} \quad (1)$$

Самодвойственная функция f_{SD} получается по формуле:

$$f_{SD} = \overline{a} f \vee a g, \quad (2)$$

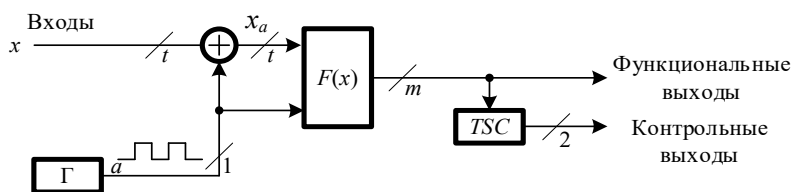
где f – исходная функция, а g – двойственная к ней функция.

В [20] приводятся примеры самодвойственных преобразований, а также даются результаты экспериментов по преобразованиям тестовых комбинационных схем [22, 23] в самодвойственные схемы. При этом отмечается, что в среднем избыточность самодвойственной схемы возрастает для представленной выборки схем на 171 %, что ниже, чем при использовании дублирования. Однако различные схемы обладают различным индексом самодвойственности I_{SD} , показывающим, насколько исходное устройство близко к самодвойственному, поэтому для устройств с $I_{SD} \geq 0,5$ в среднем усложнение гораздо ниже – 146 % относительно показателя сложности технической реализации исходной схемы.

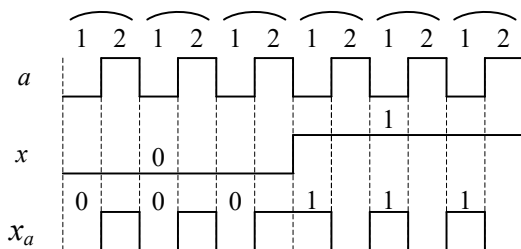
В [19] рассматривается метод синтеза самопроверяемых устройств, который подразумевает использование инвертирования данных и самодвойственного представления функций, реализуемых устройством. Автором данной статьи распространяется идея самодвойственного контроля комбинационных устройств на устройства с памятью, при этом установлены условия, при которых конечные автоматы, описывающие устройства с памятью, будут самодвойственными; приводятся простейшие модификации триггеров

и условия реализации СВК для использования метода инвертирования данных.

В дальнейшем теория синтеза СВК на основе свойств самодвойственных функций развивается научной школой под руководством проф. В.В. Сапожникова и Вл.В. Сапожникова. В уже отмеченной выше работе [20] приведена структура организации СВК для комбинационной схемы, основанная на использовании импульсного режима работы и контроля принадлежности функций к классу самодвойственных. Она изображена на рисунке 1. С помощью генератора Γ формируется импульсная последовательность a , которая позволяет сигналы с каждого входа $x_1, x_2, \dots, x_t$ преобразовать в импульсную последовательность и реализовать самодвойственное устройство по формуле (2). Контроль самодвойственности каждой выходной функции $f_1, f_2, \dots, f_m$ осуществляется с помощью тестера *TSC (totally self-checking checker)*, который, в данном случае, представляет собой каскад тестеров самодвойственного сигнала (тестеров самодвойственности) *SSC (self-checking self-dual checker)* и самопроверяемого компаратора. Схема *SSC* приведена на рисунке 2 [24]. Он имеет два входа: f и a , на который подаются функциональный и альтернативный сигналы соответственно. Для формирования на внутренних линиях *SSC* двухфазного сигнала в точках v_1 и v_2 из сигнала f используется элемент временной задержки τ . Величина задержки определяется как длина одного такта импульсной последовательности a . На выходах z^0 и z^1 формируется контрольный сигнал. Если на входы подан самодвойственный сигнал, то на выходах будет присутствовать парафазный сигнал $\langle 01 \rangle$ или $\langle 10 \rangle$, при нарушении самодвойственности сигнала на входе f и при собственных неисправностях *SSC* выдаст непарафазный сигнал. Компаратор реализуется в виде схемы сжатия парафазных сигналов на основе стандартных модулей сжатия парафазных сигналов *TRC (two-rail checker)*, структуры которых даны в [25]. В качестве альтернативного устройства для синтеза компаратора может использоваться схема из [26].



а)



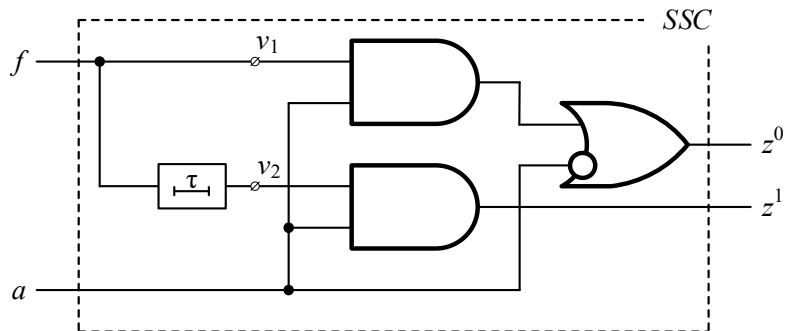
б)

Рис. 1. Контроль самодвойственных схем: а) структура; б) представление сигналов

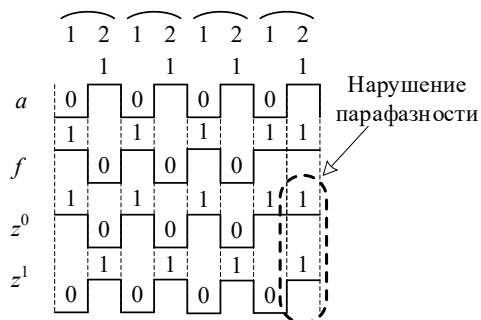
В [27, 28] развивается подход к синтезу СВК с контролем вычислений по признаку самодвойственности. Предложена структура организации СВК, основанная на принципе логической коррекции сигналов (названном авторами впоследствии *логическим дополнением* [29] и развиваемом, в том числе, в работах современников [30, 31]), приведенная на рисунке 3. В данной структуре подразумевается преобразование сигналов с выходов устройства $F(x)$ в самодвойственные сигналы с помощью элементов сложения по модулю $M=2$ по формуле:

$$h_i = f_i \oplus \delta_i, \quad i = \overline{1, m}, \quad (3)$$

где δ_i – функции логической коррекции, вычисляемые блоком $\Delta(x)$.



а)



б)

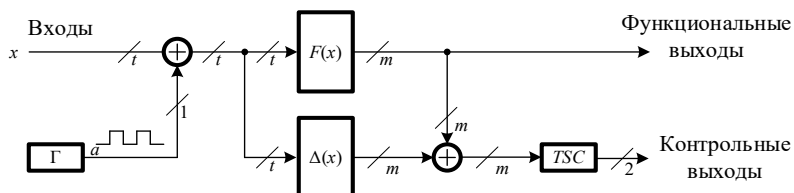
 Рис. 2. Тестер самодвойственности: а) структура; б) фиксация несамодвойственности сигнала f на временной диаграмме работы


Рис. 3. Структура организации СВК на основе логической коррекции сигналов

В этой же работе приводятся методы расчета функций логической коррекции, и описывается структура организации контроля вычислений по методу *самодвойственного паритета*. Данный метод

подразумевает, что предварительно сигналы со всех выходов объекта диагностирования $F(x)$ будут сжаты с помощью свертки по модулю $M=2$ в один сигнал, который впоследствии будет преобразован в самодвойственный с помощью единственной функции коррекции по формуле (3). В эксперименте с тестовыми комбинационными схемами показано, что такой подход к организации СВК позволяет получать выигрыш даже по сравнению с использованием традиционной структуры контроля по паритету, описанной в [32].

В [33, 34] предлагаются различные модификации структур, основанных на использовании принципа логической коррекции сигналов: структура *самодвойственного дублирования* и структура с выделением групп выходов объекта диагностирования со свертками по модулю $M=2$ с последующим самодвойственным преобразованием. В [15] исследуются особенности синтеза самопроверяемых комбинационных устройств с контролем вычислений по признаку самодвойственности формируемых функций. Авторами данной работы предложен метод преобразования структур комбинационных устройств в защищенные от неисправностей, основанный на реализации схем с монотонным проявлением неисправностей на выходах. В [35] исследованы вопросы реализации СВК для устройств с памятью.

Результаты многолетних исследований ученых в части организации СВК с контролем самодвойственности вычисляемых функций обобщены в трех монографиях [21, 36, 37].

Дальнейшее же развитие методов синтеза СВК с контролем самодвойственности вычисляемых функций, по нашему мнению, связано с применением нескольких диагностических признаков – с комбинированием кодовых методов синтеза СВК и метода контроля самодвойственности сигналов.

На рисунке 4 изображена структура СВК, в которой применено устройство $G(f)$ для преобразования m сигналов в k . Устройство $G(f)$ является кодером заранее выбранного блокового кода [9]. Следует отметить, что, фактически, приведенная на рисунке 4 структура получена путем развития структуры организации СВК по методу «самодвойственного паритета». Подобная структура может использоваться без элементов логической коррекции, если кодер $G(f)$ будет описываться самодвойственными функциями, а может быть модифицирована в структуру с дополнительным контролем вычислений по признаку принадлежности кодовых векторов, формируемых в СВК, заранее выбранным блоковым кодам.

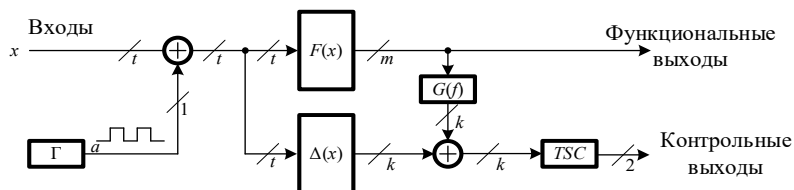


Рис. 4. Структура организации СВК на основе логической коррекции сигналов и применения кодовых методов для контроля вычислений

В качестве кодов, которые могут применяться для преобразований в структуре рисунка 4, могут выступать любые блочные коды. Так как различные блочные коды обладают различными характеристиками обнаружения ошибок, то существуют и определенные особенности их применения [9]. Например, использование линейных блочных кодов позволяет применять помимо свойства самодвойственности еще и свойство линейности реализуемых функций. Среди таких кодов могут эффективно применяться классические коды Хэмминга, алгебраические коды и их разнообразные модификации [12, 38 – 41]. Может быть использовано свойство реализации монотонных функций, которым обладают классические коды с суммированием (коды Бергера), равновесные коды и некоторые их модификации [9, 42, 43]. Таким образом, синтез СВК может быть осуществлен с контролем вычислений сразу же по нескольким диагностическим признакам. В ряде статей, например, в упомянутой выше [24], приводится способ организации СВК, который подразумевает контроль вычислений по равновесным кодам « r из $2r$ » (r – вес кодового слова) с самодвойственной реализацией каждой функции, описывающей биты равновесных кодов. Отмечается, что вместо равновесных кодов могут быть применены и другие блочные коды (при этом, однако, не любой код подходит для этих целей). Использование при контроле вычислений двух диагностических признаков и импульсного режима работы позволяет повышать число тестовых комбинаций для проверки неисправностей, что особенно актуально для диагностирования устройств с редко меняющимися входными данными [44 – 46].

Настоящая работа относится к ветви исследований, охватывающей вопросы изучения особенностей применения кодов Хэмминга при синтезе СВК с контролем самодвойственности сигналов. Применению кодов Хэмминга при построении СВК посвящено достаточно большое количество работ, направленных как

на практические приложения при синтезе самопроверяемых устройств [47, 48], так и на исследования характеристик самих кодов, проявляющихся при их построении [49 – 51]. Однако работы, посвященные исследованию применения при синтезе СВК кодов Хэмминга, контрольные биты которых описываются самодвойственными функциями, отсутствуют. Представленная статья восполняет данный пробел.

3. Постановка задачи. Целью представленного в настоящей работе исследования является изучение особенностей реализации СВК с контролем вычислений по признаку самодвойственности булевых функций, описывающих проверочные символы кодов Хэмминга.

Для достижения цели решаются следующие задачи:

1. Разработка структуры СВК по кодам Хэмминга, контрольные биты которых описываются самодвойственными функциями.

2. Исследование особенностей кодов Хэмминга и поиск тех длин кодовых слов, при которых все контрольные биты будут описываться самодвойственными булевыми функциями.

3. Моделирование простейших цифровых устройств с контролем вычислений по двум обозначенным выше диагностическим признакам в целях подтверждения эффективности в части повышения характеристик контролепригодности.

4. Разработка структуры СВК по произвольным кодам Хэмминга с преобразователем сигналов от кодера в самодвойственные сигналы.

4. Коды Хэмминга. Код Хэмминга строится следующим образом. Формируется проверочная матрица вида:

$$H_n = \begin{pmatrix} 0 & 0 & \dots & 1 & 1 \\ 0 & 0 & \dots & 1 & 1 \\ \dots & \dots & \dots & \dots & \dots \\ 0 & 1 & \dots & 1 & 1 \\ 1 & 0 & \dots & 0 & 1 \end{pmatrix}. \quad (4)$$

В матрице H_n перечисляются подряд слева направо все возможные двоичные числа от числа $[00\dots 01]_2$ до числа $[11\dots 11]_2$. Столбцы, соответствующие двоичным числам с десятичными аналогами $n = 2^i$, $i \in \mathbb{N}_0$, отводятся под проверочные символы

(им соответствуют столбцы с одним символом «1»). Остальные столбцы соответствуют информационным символам. Значение функции g_j , $j = \overline{1, k}$, описывающей j -ый проверочный символ, получается, как сумма по модулю $M=2$ сигналов тех информационных бит f_i , на пересечении столбцов которых с j -ой строкой стоит единица. Число проверочных символов определяется как ближайшее целое, удовлетворяющее неравенству: $m+1 \leq 2^k - k$, где $m=n-k$ – число информационных символов.

Часто проверочная матрица (4) для удобства представляется в форме $k \times (k+m)$, где часть, соответствующая проверочным символам, отделена от части, соответствующей информационным символам:

$$H_n = \left(\begin{array}{ccccc|ccccc} 0 & 0 & \dots & 0 & 1 & 0 & 0 & \dots & 1 & 1 \\ 0 & 0 & \dots & 1 & 0 & 0 & 0 & \dots & 1 & 1 \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ 0 & 1 & \dots & 0 & 0 & 1 & 1 & \dots & 1 & 1 \\ 1 & 0 & \dots & 0 & 0 & 1 & 1 & \dots & 0 & 1 \end{array} \right). \quad (5)$$

Примеры получения кодовых слов, принадлежащих кодам Хэмминга, здесь приводить не будем, они есть в большом числе источников, включая работы [48 – 51].

5. Структура самопроверяемого устройства с контролем вычислений по кодам Хэмминга с самодвойственными функциями. Рассмотрим структуру, приведенную на рисунке 5. В отличие от структуры рисунка 4 в представленной структуре не используется логической коррекции сигналов в СВК. Здесь подразумевается следующее: кодер $G(f)$ является самодвойственным устройством, преобразующим m информационных сигналов в k контрольных. Их можно напрямую контролировать с помощью каскада тестеров самодвойственности $kSSC1$, преобразующих k сигналов в один контрольный сигнал, а также снабдить кодер компаратором $kTRC1$, а СВК дооборудовать блоком $G(x)$ – вычисления контрольных функций по входным воздействиям, выходы которого также подключить к компаратору. Это позволит контролировать вычисления по заранее выбранному коду, проверочные символы которого описываются самодвойственными булевыми функциями. Устройства $G(f)$ и $kTRC1$ образуют тестер выбранного кода

(устройство TSC). Если рассматривать только эту схемную часть, исключая каскад тестеров самодвойственности и входной каскад получения самодвойственных сигналов, то будет представлена классическая структура организации СВК [6, 7, 9]. Добавление каскада тестеров самодвойственности и входного каскада получения самодвойственных функций дает возможность дополнительного контроля еще одного диагностического признака. Выходы устройств TSC и $kSSC1$ подключаются к входам одного модуля сжатия парафазных сигналов TRC , выходы которого уже являются контрольными выходами СВК.

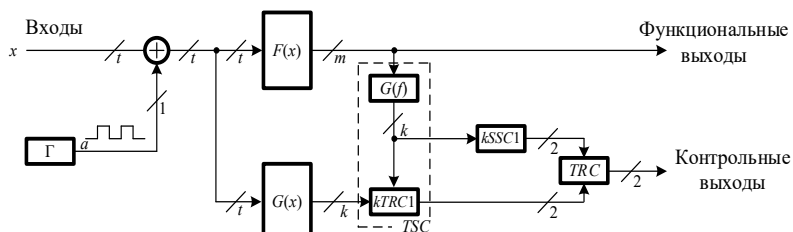


Рис. 5. Структура самопроверяемого устройства с контролем вычислений по блоковым разделимым кодам с самодвойственными контрольными функциями

Остается определить те коды (здесь ограничимся рассмотрением только кодов Хэмминга), которые будут обладать свойством самодвойственности функций, описывающих проверочные символы.

Коды Хэмминга строятся с использованием контрольных сумм по модулю $M=2$ некоторого подмножества информационных бит. Сами контрольные функции являются линейными. Для линейной функции известно следующее свойство [52].

Теорема 1. *Любая линейная булева функция является самодвойственной только при нечетном числе аргументов, от которых она зависит существенно.*

Таким образом, если каждый проверочный символ кода Хэмминга описывается функцией с нечетным числом аргументов, то его кодер будет самодвойственным устройством. Это условие выполняется не для любого числа информационных бит.

6. Определение числа информационных бит кодов Хэмминга, для которых функции, описывающие проверочные символы, будут самодвойственными. Определим те значения m ,

для которых строятся коды Хэмминга с самодвойственными функциями, описывающими проверочные символы.

Вначале определим число суммируемых слагаемых при выполнении j -ой контрольной проверки (при вычислении j -го проверочного символа), $j = 1, \lceil \log_2(n+1) \rceil$, где n – десятичный эквивалент двоичного числа, записанного на вертикали матрицы. При этом в каждой сумме E_j каждой j -ой строки будем учитывать и значения, характеризующие сами контрольные функции (напомним читателю, что столбцы матрицы, для которых $n = 2^i$, $i \in \mathbb{N}$, отводятся под проверочные символы, а нулевой столбец не используется). Далее этот факт учтем.

Рассмотрим матрицу, представленную таблицей 1. Она может служить в качестве проверочной матрицы для кода Хэмминга (15, 11, 3) при $n=15$ и для укороченных кодов Хэмминга с длиной кодового слова $n < 15$ (при удалении некоторого количества столбцов). В ней по возрастающей от $n=1$ до $n=15$ записаны четырехбитные двоичные кодовые векторы от $[0001]_2$ до $[1111]_2$. Определим, как заполнена матрица.

Таблица 1. Пример проверочной матрицы для кода Хэмминга

g_i	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
g_4	0	0	0	0	0	0	0	1	1	1	1	1	1	1	1
g_3	0	0	0	1	1	1	1	0	0	0	0	1	1	1	1
g_2	0	1	1	0	0	1	1	0	0	1	1	0	0	1	1
g_1	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1

Заполнение матрицы производится с периодичностью, определяемой величиной:

$$\tau_j = 2^j. \quad (6)$$

При этом чередуются 2^{j-1} нулей и столько же единиц в каждом периоде. К примеру, для $j=3$ длина периода составляет $\tau_3 = 2^3 = 8$, чередуется $2^{3-1} = 4$ нуля и столько же единиц. Исключение составляет

только первый период для каждой строки, так как отсутствует число $[0000]_2$. Будем далее это учитывать.

Введем следующее обозначение: E_j – сумма единиц по j строке проверочной матрицы. Формализуем получение чисел E_j для задаваемых значений n и j .

Алгоритм. Последовательность получения чисел E_j :

1. Задается число n .
2. Определяется число полных периодов $k_{\tau_j}^F$ для заданного числа n . Оно равно:

$$k_{\tau_j}^F = \left\lfloor \frac{n+1}{2^j} \right\rfloor. \quad (7)$$

К примеру, для $n=13$ имеется $j=1, 2, 3, 4$, число полных периодов для которых равно: $k_{\tau_1}^F = \left\lfloor \frac{13+1}{2^1} \right\rfloor = 7$, $k_{\tau_2}^F = \left\lfloor \frac{13+1}{2^2} \right\rfloor = 3$,

$$k_{\tau_3}^F = \left\lfloor \frac{13+1}{2^3} \right\rfloor = 1, \quad k_{\tau_4}^F = \left\lfloor \frac{13+1}{2^4} \right\rfloor = 0.$$

3. Определяется число единиц для всех полных периодов τ_j .

Сумма единиц в каждом полном периоде равна 2^{j-1} . Соответственно, для всех полных периодов она составляет:

$$2^{j-1} k_{\tau_j}^F = 2^{j-1} \left\lfloor \frac{n+1}{2^j} \right\rfloor. \quad (8)$$

К примеру, для $n=13$ и $j=2$ получаем в полных периодах $2^{2-1} \cdot 3 = 6$ единиц. К слову, число нулей определяется аналогично.

4. Определяется число единиц в неполном периоде. Неполный период включает в себя:

$$(n+1) - \tau_j k_{\tau_j}^F = (n+1) - 2^j \left\lfloor \frac{n+1}{2^j} \right\rfloor \text{ нулей и единиц.} \quad (9)$$

К примеру, для $n=13$ и $j=2$ получаем: $14 - \tau_2 k_{\tau_2}^F = 14 - 2^2 \cdot 3 = 14 - 12 = 2$.

В неполном периоде имеется следующее количество нулей:

$$\begin{cases} 2^{j-1} & \text{при } (n+1) - \tau_j k_{\tau_j}^F > 2^{j-1}; \\ (n+1) - \tau_j k_{\tau_j}^F & \text{при } (n+1) - \tau_j k_{\tau_j}^F \leq 2^{j-1}. \end{cases} \quad (10)$$

Остальное – единицы.

К примеру, для $n=13$ и $j=2$ имеем: $(n+1) - \tau_j k_{\tau_j}^F = 2$ и $2^{j-1} = 2^{2-1} = 2$, откуда следует, что число нулей равно 2, а число единиц – 0; для $n=13$ и $j=3$ имеем: $(n+1) - \tau_j k_{\tau_j}^F = 14 - 8 \cdot 1 = 6$ и $2^{j-1} = 2^{3-1} = 4$, откуда следует, что число нулей равно 4, а число единиц – 2.

Отбросим число нулей из неполного периода, получив в нем сумму единиц:

$$\begin{cases} (n+1) - 2^j \left\lfloor \frac{n+1}{2^j} \right\rfloor - 2^{j-1} & \text{при } (n+1) - \tau_j k_{\tau_j}^F > 2^{j-1}; \\ 0 & \text{при } (n+1) - \tau_j k_{\tau_j}^F \leq 2^{j-1}. \end{cases} \quad (11)$$

5. Определяется сумма единиц в j строке.

Итак, сумма единиц по j строке будет равна:

$$E_j = \begin{cases} 2^{j-1} \left\lfloor \frac{n+1}{2^j} \right\rfloor + (n+1) - 2^j \left\lfloor \frac{n+1}{2^j} \right\rfloor - 2^{j-1} & \text{при } (n+1) - 2^j \left\lfloor \frac{n+1}{2^j} \right\rfloor > 2^{j-1}; \\ 2^{j-1} \left\lfloor \frac{n+1}{2^j} \right\rfloor & \text{при } (n+1) - 2^j \left\lfloor \frac{n+1}{2^j} \right\rfloor \leq 2^{j-1}. \end{cases} \quad (12)$$

Упростим выражение (12):

$$E_j = \begin{cases} (n+1) - 2^{j-1} \left(\left\lfloor \frac{n+1}{2^j} \right\rfloor + 1 \right) & \text{при } (n+1) - 2^j \left\lfloor \frac{n+1}{2^j} \right\rfloor > 2^{j-1}; \\ 2^{j-1} \left\lfloor \frac{n+1}{2^j} \right\rfloor & \text{при } (n+1) - 2^j \left\lfloor \frac{n+1}{2^j} \right\rfloor \leq 2^{j-1}. \end{cases} \quad (13)$$

К примеру, для $n=13$ и $j=3$ имеем:
 $(n+1) - 2^j \left\lfloor \frac{n+1}{2^j} \right\rfloor = (13+1) - 2^3 \cdot \left\lfloor \frac{13+1}{2^3} \right\rfloor = 14 - 8 = 6, \quad 2^{j-1} = 2^{3-1} = 4,$

откуда следует, что для подсчетов нужно использовать верхнее выражение в формуле (13):

$$E_j = (n+1) - 2^{j-1} \left(\left\lfloor \frac{n+1}{2^j} \right\rfloor + 1 \right) = (13+1) - 2^{3-1} \left(\left\lfloor \frac{13+1}{2^3} \right\rfloor + 1 \right) = 14 - 4 \cdot 2 = 6.$$

В таблице 2 приводятся значения чисел E_j для $n=0...31$ (для демонстрации свойств сумм E_j рассмотрено также и число $n=0$, имеющее смысл лишь с точки зрения наглядности представления закономерности). Выделены столбцы, соответствующие числам n , для которых каждая контрольная сумма будет иметь нечетное количество слагаемых.

Матрица, в которой позиции $n = 2^i$, $i \in \mathbb{N}_0$, отведены также под информационные символы, описывает известный [51] модифицированный код Хэмминга. Такая матрица содержит в каждом столбце слева направо возрастающие двоичные числа. Каждый столбец соответствует информационному символу. Таким образом, при построении модифицированного кода Хэмминга значение j -го проверочного символа, $j = 1, \overline{\log_2(m+1)}$ (здесь именно m – размерность кода), получается, как сумма по модулю $M=2$ тех информационных символов, которым соответствуют столбцы в проверочной матрице, в которых на пересечении j -ой строки записана единица. Отметим, что длина такого кода будет равна $n = m + \overline{\log_2(m+1)}$, а минимальное расстояние $d_{\min}=2$.

Таблица 2. Значения чисел E_j для $n=0...31$

g_i	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
g_5	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0
g_4	0	0	0	0	0	0	0	0	1	1	1	1	1	1	1	1
g_3	0	0	0	0	1	1	1	1	0	0	0	0	1	1	1	1
g_2	0	0	1	1	0	0	1	1	0	0	1	1	0	0	1	1
g_1	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1
E_j																
g_i	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
g_5	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0
g_4	0	0	0	0	0	0	0	0	1	2	3	4	5	6	7	8
g_3	0	0	0	0	1	2	3	4	4	4	4	4	5	6	7	8
g_2	0	0	1	2	2	2	3	4	4	4	5	6	6	6	7	8
g_1	0	1	1	2	2	3	3	4	4	5	5	6	6	7	7	8

g_i	16	17	18	19	20	21	22	23	24	25	26	27	28	29	30	31
g_5	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
g_4	0	0	0	0	0	0	0	0	1	1	1	1	1	1	1	1
g_3	0	0	0	0	1	1	1	1	0	0	0	0	1	1	1	1
g_2	0	0	1	1	0	0	1	1	0	0	1	1	0	0	1	1
g_1	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1
E_j																
g_i	16	17	18	19	20	21	22	23	24	25	26	27	28	29	30	31
g_5	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16
g_4	8	8	8	8	8	8	8	8	9	10	11	12	13	14	15	16
g_3	8	8	8	8	9	10	11	12	12	12	12	12	13	14	15	16
g_2	8	8	9	10	10	10	11	12	12	12	13	14	14	14	15	16
g_1	8	9	9	10	10	11	11	12	12	13	13	14	14	15	15	16

Теорема 2. Для модифицированного кода Хэмминга при числе информационных символов $t = n = 2^k - 2$, $k \in \{2, 3, \dots\}$, все числа E_j являются нечетными и равными $E_j = 2^{k-1} - 1$.

Доказательство. Рассмотрим элементарный случай $k=3$ (случаи с меньшими k не показательны, так как модифицированный код Хэмминга имеет смысл строить для $k \geq 3$). Для данного случая существует $2^k = 2^3 = 8$ различных $n = \overline{0,7}$. Всего для каждой контрольной функции $j = \overline{1,3}$ имеется $2^{k-1} = 2^{3-1} = 4$ нуля и столько же единиц. Максимальная сумма E_j для каждого j достигается для случая $n=7$ и равна $E_1 = E_2 = E_3 = 4$. Это число четное. Если рассмотреть случай $n=6$, то для него следует вычесть по одной единице для каждого E_j , ведь десятичному числу $[7]_{10}$ соответствует двоичное $[111]_2$, в котором все разряды равны единице. Таким образом, для $n=6$ $E_1 = E_2 = E_3 = 3$. Это число нечетное. Если рассмотреть любое меньшее число n , то потребуется отнять от нечетного числа $E_1 = E_2 = E_3 = 3$ для каждого j различное число единиц. Для $n=5$ для двух из трех E_j (E_2 и E_3 , поскольку для $n=6$ двоичный эквивалент равен $[110]_2$) потребуется отнять по единице. Поэтому $E_1 = 3$ и $E_2 = E_3 = 2$. Далее, при формировании сумм числа $n=4$ нужно от числа $[101]_2$ перейти к числу $[100]_2$, что потребует отнять от $E_1 = 3$ и $E_3 = 2$ по единице: $E_1 = E_2 = 2$ и $E_3 = 1$. Числа $n < 4$ рассматривать не следует, так как для этих чисел $k=2$. Таким образом, из всех возможных вариантов для $k=3$ только вариант $m = n = 2^k - 2 = 2^3 - 2 = 6$ дает все нечетные E_j . При этом, они в точности будут равны $E_j = 2^{k-1} - 1 = 2^{3-1} - 1 = 3$.

Рассмотрим произвольное значение k . Для него существует k различных сумм E_j , $j = \overline{1,k}$. Для каждой контрольной функции максимальная сумма $E_j = 2^{k-1}$. Это четное число. Число n для такой суммы будет в двоичном виде выглядеть так: $[11...11]_2$. И только для числа $n-1$ одновременно от каждого $E_j = 2^{k-1}$ будет отнята единица. Новые $E_j = 2^{k-1} - 1$, и они все будут нечетными. Число $n-1$ в двоичном виде будет равно: $[11...10]_2$. Следующее за ним число по убыванию $n-2$ в двоичном виде будет равно: $[11...01]_2$. Для его получения потребуется от всех $E_j = 2^{k-1} - 1$, кроме $E_1 = 2^{k-1} - 1$ отнять по единице. Новые суммы станут равными $E_j = 2^{k-1} - 2$, $j \in \{2, 3, \dots, k\}$

и $E_1 = 2^{k-1} - 1$. Все числа, кроме E_1 , четные. Далее при получении числа $n-3$ изменится четность у всех $E_j = 2^{k-1} - 2, j \in \{1, 2, \dots, k\} \setminus \{2\}$. Все числа $E_j = 2^{k-1} - 3, j \in \{3, 4, \dots, k\}$ будут нечетными, а $E_1 = E_2 = 2^{k-1} - 2$ – четными. Далее процесс продолжится, так как для числа $n-4$ четность функций старших разрядов должна поменяться. В итоге, предельный случай для числа $\frac{n}{2} + 1$ будет соответствовать двоичному числу $[10\dots 00]_2$, для которого $E_j = 2^{k-2}, j = \overline{1, k-1}$ и $E_k = 1$. Условие теоремы выполняется. Выполняется оно и для случая $k+1$. Теорема доказана.

Теорема 2 относится к модифицированным кодам Хэмминга. При построении классических кодов Хэмминга столбцы $n = 2^i, i \in \mathbb{N}_0$, отводятся под проверочные символы (таблица 3). Определим число суммируемых слагаемых при выполнении j -ой контрольной проверки, $j = \overline{1, k}$, для классического кода Хэмминга.

Таблица 3. Пример проверочной матрицы кода Хэмминга с выделением столбцов, соответствующих контрольным символам

g_i	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
g_4	0	0	0	0	0	0	0	1	1	1	1	1	1	1	1
g_3	0	0	0	1	1	1	1	0	0	0	0	1	1	1	1
g_2	0	1	1	0	0	1	1	0	0	1	1	0	0	1	1
g_1	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1

Число E_j (формула 13) определяет сумму единиц без учета отбрасываемых единиц для чисел $n = 2^i, i \in \mathbb{N}_0$.

Новая проверка появляется только тогда, когда формируется столбец для числа j . Для каждого числа j требуется удалить одну «лишнюю» единицу. Получаем:

$$E_j^* = E_j - 1. \quad (14)$$

В таблице 4 приводятся значения чисел E_j для $n=0\dots 31$ для классического кода Хэмминга. Аналогично таблице 2, в таблице 4 выделены те столбцы, для которых каждая полученная контрольная

сумма будет иметь нечетное количество слагаемых. В таблице 4, в отличие от таблицы 2, выделены уже другие столбцы, соответствующие числам n . Анализ таблицы 4 позволяет установить закономерность в выделении столбцов.

Таблица 4. Значения чисел E_j для $n=0...31$ для классических кодов Хэмминга

g_i	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
g_5	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0
g_4	0	0	0	0	0	0	0	0	1	1	1	1	1	1	1	1
g_3	0	0	0	0	1	1	1	1	0	0	0	0	1	1	1	1
g_2	0	0	1	1	0	0	1	1	0	0	1	1	0	0	1	1
g_1	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1
E_j																
g_i	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
g_5	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0
g_4	0	0	0	0	0	0	0	0	0	1	2	3	4	5	6	7
g_3	0	0	0	0	0	1	2	3	3	3	3	3	4	5	6	7
g_2	0	0	0	1	1	1	2	3	3	3	4	5	5	5	6	7
g_1	0	0	0	1	1	2	2	3	3	4	4	5	5	6	6	7
E_j																
g_i	16	17	18	19	20	21	22	23	24	25	26	27	28	29	30	31
g_5	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
g_4	0	0	0	0	0	0	0	0	1	1	1	1	1	1	1	1
g_3	0	0	0	0	1	1	1	1	0	0	0	0	1	1	1	1
g_2	0	0	1	1	0	0	1	1	0	0	1	1	0	0	1	1
g_1	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1
E_j																
g_i	16	17	18	19	20	21	22	23	24	25	26	27	28	29	30	31
g_5	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
g_4	7	7	7	7	7	7	7	7	8	9	10	11	12	13	14	15
g_3	7	7	7	7	8	9	10	11	11	11	11	11	12	13	14	15
g_2	7	7	8	9	9	9	10	11	11	11	12	13	13	13	14	15
g_1	7	8	8	9	9	10	10	11	11	12	12	13	13	14	14	15

Теорема 3. Для классического кода Хэмминга при числе разрядов в кодовых словах $n=3+4l$, $l \in \mathbb{N}_0$, все числа E_j являются нечетными.

Доказательство. Рассмотрим первое из чисел $n=3$, для которого должно выполняться условие теоремы 3. При формировании каждого из чисел E_j , $j=1, 2$, требуется сложить ровно $E_j = 2^{\log_2(n+1)-1}$ единиц: от двоичных чисел $[1]_2$, $[10]_2$ и собственно числа $[11]_2$. Однако

при формировании матрицы было оговорено, что числа n , равные степени числа 2, не должны учитываться, и из числа E_j каждый раз при добавлении новой контрольной функции с ростом n должна вычитаться единица. Числа $[1]_2$ и $[10]_2$ есть степени числа 2. Вычтем из каждого E_j , $j = 1, 2$, по единице и получим $E_1 = E_2 = 1$.

Далее появляется группа из четырех чисел, для которой вводится еще один контрольный разряд: $[100]_2$, $[101]_2$, $[110]_2$, $[111]_2$. Первое число исключается как степень числа 2. Для формирования второго числа $[101]_2$ нужно добавить к $E_1 = 1$ единицу и то же к $E_3 = 1$. Число E_1 станет равным $E_1 = 2$. Это четное число. Далее следует число $[110]_2$. Нужно добавить к $E_2 = 1$ и к $E_3 = 1$ по единице. Таким образом, $E_2 = E_3 = 2$. Все три суммы четные. Следующее число формируется при добавлении к каждой сумме по единице. Все суммы становятся нечетными. Это случай $n=3+4=7$.

Далее следует группа чисел $[1000]_2$, $[1001]_2$, $[1010]_2$, $[1011]_2$. К нечетным числам $E_1 = E_2 = 3$ при формировании последнего из представленных двоичных чисел добавляется ровно по 2 единицы, что не меняет их четности. К числу E_3 не прибавляется ничего. Число $E_4 = 1$. Десятичное число, соответствующее двоичному $[1011]_2$, равно $n=3+4+4=7+4=11$.

Далее рассматривается следующая «четверка» чисел. Всякий раз при формировании четвертого из них в каждую сумму либо не добавляется ничего, либо добавляется четное число 2, а с ростом n – либо 2, либо 4 (таблица 5).

Таблица 5. Значения чисел E_j для $n = 3 + 4l$, $l \in \{0, 1, \dots, 15\}$

3	7	11	15	19	23	27	31	35	39	43	47	51	55	59	63
0	0	0	0	0	0	0	0	3	7	11	15	19	23	27	31
0	0	0	0	3	7	11	15	15	15	15	15	19	23	27	31
0	0	3	7	7	7	11	15	15	15	19	23	23	23	27	31
0	3	3	7	7	11	11	15	15	19	19	23	23	27	27	31
1	3	5	7	9	11	13	15	17	19	21	23	25	27	29	31
1	3	5	7	9	11	13	15	17	19	21	23	25	27	29	31

В итоге, как видно, при рассмотрении каждого четвертого числа после 3, четность сумм E_j не меняется. Теорема доказана.

Определим числа m для тех n , которые характеризуют нечетные суммы E_j для всех j , что важно с практической точки зрения, так как дают возможность выбора конкретного кода Хэмминга.

Задача заключается в определении числа $m = n - k$ для каждого случая $n = 3 + 4l$, $l \in \mathbb{N}_0$.

Для каждого числа k существует предельное значение $m = m^*$:

$$m^* = (2^k - 1) - k, \quad k \in \mathbb{N}. \quad (15)$$

Все значения m для данного k (то есть числа $m > m^{**} = (2^{k-1} - 1) - (k - 1) - 1$) могут быть получены по формуле:

$$m = m^* - 4l, \quad l \in \left\{ 0, 1, \dots, \left\lfloor \frac{m^* - m^{**}}{4} \right\rfloor \right\}. \quad (16)$$

К примеру, получим по формуле (16) все возможные значения m для кодов с $k=5$:

$$m^* = (2^5 - 1) - 5 = 26,$$

$$m^{**} = (2^{5-1} - 1) - (5 - 1) = 11,$$

$$m = 26 - 4l, \quad l = \overline{0, 3}: \quad m = 26 - 4 \cdot 0 = 26, \quad m = 26 - 4 \cdot 1 = 22, \\ m = 26 - 4 \cdot 2 = 18, \quad m = 26 - 4 \cdot 3 = 14.$$

Таким образом, пользуясь формулой (16), можно получить для каждого значения k все возможные значения m , для которых E_j для всех j .

Можно сделать следующие выводы. Во-первых, не любой код Хэмминга может быть использован для контроля вычислений на выходах самодвойственных устройств. Для кодов Хэмминга со значениями длин $n \neq 3 + 4l$, $l \in \mathbb{N}_0$ кодеры не будут являться самодвойственными устройствами. Во-вторых, организовать контроль вычислений на выходах самодвойственного устройства по представленному методу можно только путем выделения групп

выходов с мощностью $n = 3 + 4l$, $l \in \mathbb{N}_0$. При построении СВК выходы отдельных СВК для образованных групп выходов исходного устройства подключаются к входам самопроверяемого компаратора, выходы которого уже являются контрольными выходами устройства.

7. Моделирование самодвойственных устройств. Для демонстрации особенностей тестирования самодвойственных устройств с применением структуры, приведенной на рисунке 5, проведем моделирование в среде Multisim, широко используемой для отладки и симуляции цифровых схем [53, 54]. Для этого рассмотрим элементарное комбинационное устройство, реализованное в базисе стандартных функциональных элементов (рисунок 6).

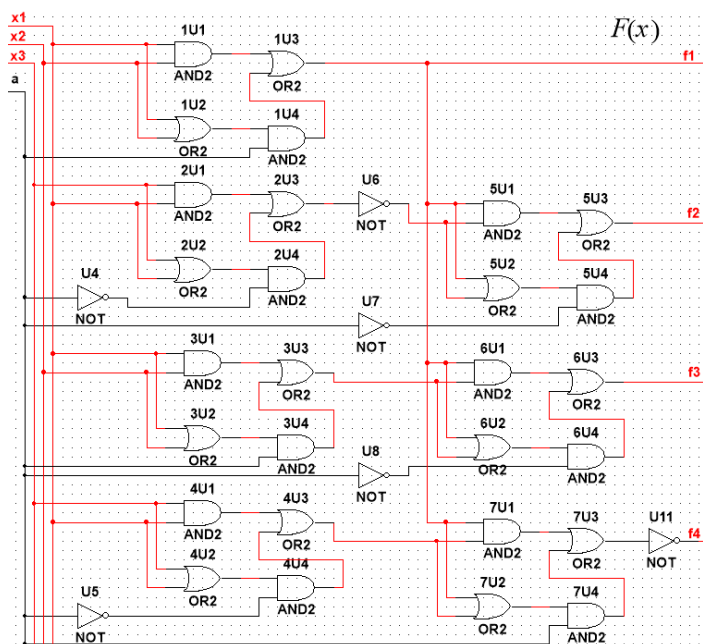


Рис. 6. Схема комбинационного устройства, реализованная в Multisim

Выбранное комбинационное устройство имеет три входа x_1 , x_2 , x_3 и один дополнительный вход, на который подается альтернативный сигнал a для реализации импульсного режима работы, а также снабжено четырьмя выходами f_1, f_2, f_3, f_4 .

Для реализации структуры рисунка 5 был синтезирован блок контрольной логики $G(x)$. Сами шаги процедуры синтеза здесь не приведены, так как использованы известные методики [11]. Остальные элементы структуры рисунка 5 являются типовыми: кодер кода Хэмминга $G(f)$, тестеры самодвойственности SSC и модули сжатия парафазных сигналов TRC . Синтезированное устройство со схемой встроенного контроля по двум диагностическим признакам приведено на рисунке 7.

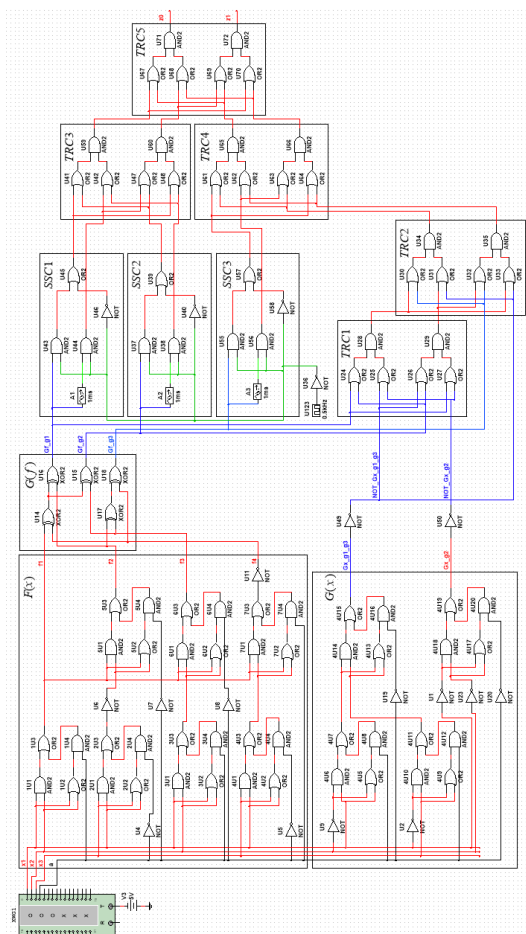


Рис. 7. Самопроверяемое устройство для рассматриваемого примера

На рисунке 8 показана работа устройства и выведены сигналы с тех точек, в которых наблюдаются значения рабочих выходов устройств $F(x)$ и $G(x)$ и контрольных сигналов в СВК. Генератор кодовых слов (XWG1) настраивается таким образом, чтобы на входы устройства в каждой паре тактов поступали ортогональные по всем переменным входные комбинации (рисунок 8а). Использована подача пар входных комбинаций по возрастающей в их последовательности от (0000, 1111) к (0111, 1000). Всего 8 пар входных комбинаций. На рисунках 8б и 8в демонстрируются временные диаграммы работы схем при подаче всех пар входных комбинаций. Читатель может обратить внимание на то, что на каждой паре входных комбинаций значения на выходах парафазны (рисунок 8б), как парафазны и значения на парах тестовых выходов проверяющих элементов (рисунок 8в). Это характеризует работу исправного устройства.

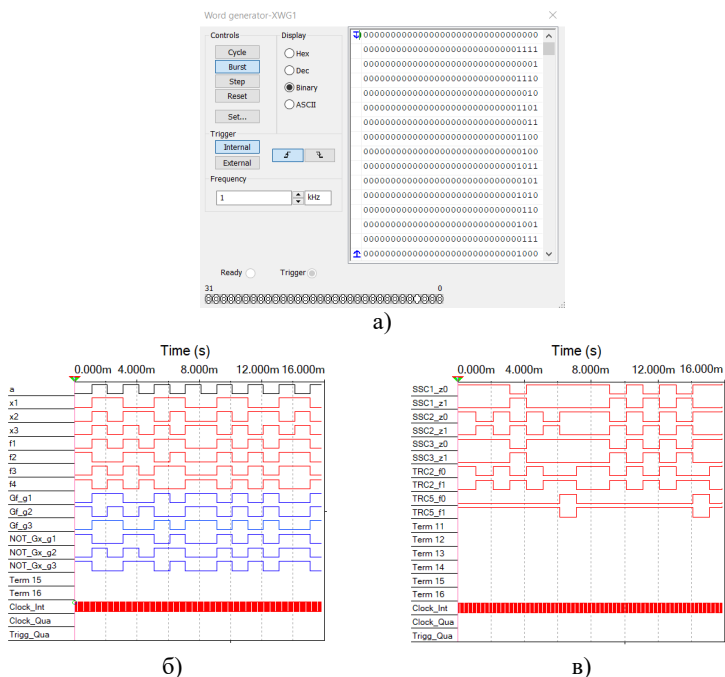


Рис. 8. Моделирование работы схемы при исправности объекта диагностирования: а) настройка генератора кодовых слов; б) временная диаграмма работы устройств $F(x)$, $G(f)$, $G(x)$; в) временная диаграмма работы устройств SSC и TRC

На рисунке 9 отдельно приводятся схема кодера $G(f)$, а также временные диаграммы его работы в отсутствие и при наличии неисправности «константа 0» на верхнем входе его элемента U3. Из сравнения рисунков 9б и 9в видно, что при неисправности в схеме на каждой паре входных комбинаций парафазность сохраняется, однако не на всех комбинациях генерируются корректные контрольные векторы (пары 2, 4, 6 и 8 (обе комбинации в паре тестовые); 3 – первая комбинация в паре тестовая).

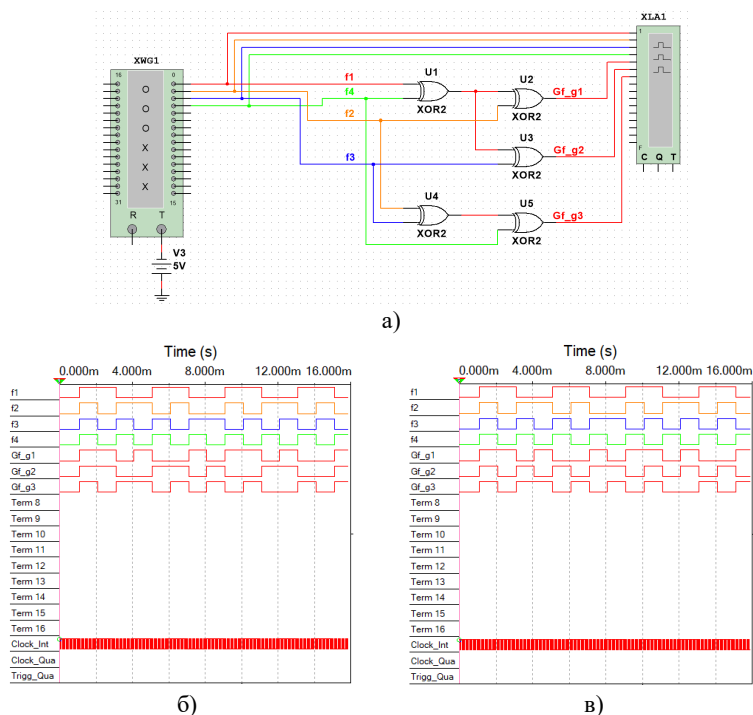


Рис. 9. Моделирование работы кодера кода Хэмминга: а) схема кодера с подключенным анализатором XLA1; б) временная диаграмма работы кодера в исправном состоянии объекта диагностирования; в) временная диаграмма работы кодера при неисправности «константа 0» на верхнем входе его элемента U3

Схема тестера самодвойственности и результаты моделирования его работы представлены на рисунке 10. На вход f

тестера подается самодвойственный сигнал, который в самом тестере уже преобразуется в двухфазный с использованием линии задержки.

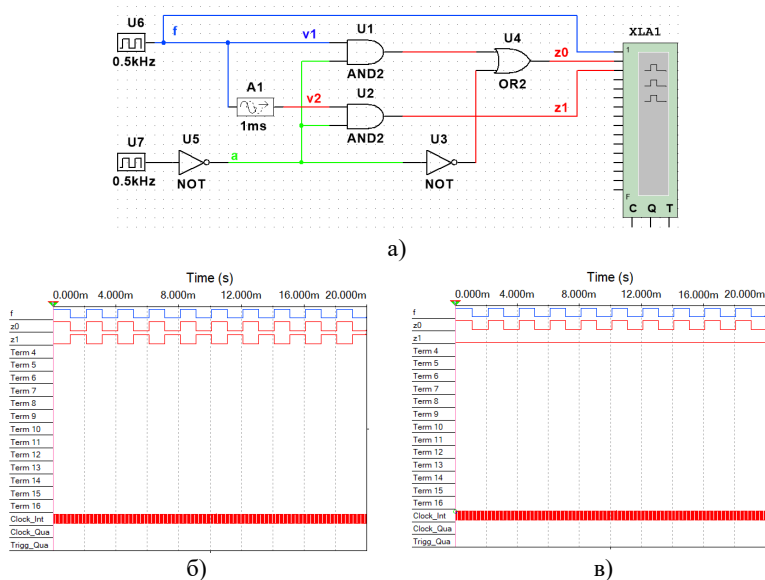


Рис. 10. Моделирование работы тестера самодвойственности: а) схема SSC с подключенным анализатором XLA1; б) временная диаграмма работы SSC в штатном режиме; в) временная диаграмма работы при неисправности «константа 0» на нижнем входе его элемента U2

Задержка на элементе A1 выбирается равной одному такту импульсной последовательности a , которая подается на соответствующий вход тестера. Частота последовательности равна

$$0,5 \text{ кГц. Поэтому задержка равна } \tau = \frac{1/\nu}{2} = \frac{1/0,5 \text{ кГц}}{2} = 1 \text{ мс.}$$

На рисунке 10б показана работа тестера самодвойственности в отсутствии неисправностей. На рисунке 10в демонстрируется работа при возникновении неисправности «константа 0» на нижнем входе элемента U2. Из анализа диаграмм следует, что вторая комбинация в каждой паре входных комбинаций является тестовой.

Сигналы с одноименных выходов блоков $G(f)$ и $G(x)$ поступают на входы двух элементов TRC – TRC1 и TRC2 на схеме рисунка 7. Так как на одноименных выходах обоих блоков реализуются

одинаковые функции, сигналы от блока контрольной логики инвертируются на элементах U49 и U50. Выходы *TRC2* фактически представляют собой контрольные выходы тестера кода Хэмминга. Каскад тестеров *TRC3 – TRC5* служит для сжатия четырех парафазных сигналов – с выходов *SSC1 – SSC3* и *TRC2*. Выходы устройства *TRC5* представляют собой контрольные выходы z^0 и z^1 всей СВК по двум диагностическим признакам.

С целью демонстрации эффективности использования контроля вычислений по двум диагностическим признакам по сравнению с контролем по одному из них, смоделируем работу всего устройства с СВК при возникновении неисправностей в объекте диагностирования $F(x)$. Особенностью схемы устройства $F(x)$ является то, что только элементы 1U1...1U4 связаны путями со всеми четырьмя ее выходами f_1, f_2, f_3 и f_4 ; остальные элементы схемы путями связаны только с каким-либо одним из ее выходов. Поэтому рассмотрим для примера особенности фиксации ошибок, вызванных неисправностями «константа 0» и константа 1» только элемента 1U3 (рисунок 6). Временные диаграммы работы самопроверяемого устройства представлены на рисунке 11. Анализируются они на каждой паре входных комбинаций, показанной вертикальными пунктирными отсечками.

Результаты моделирования сведены в таблицу 6, где указаны те входные комбинации в каждой паре, которые являются тестовыми для заданных неисправностей. Приведены данные для выходов устройств *TRC2* и *TRC5*. Использование контроля по двум диагностическим признакам позволило увеличить число тестовых комбинаций, что улучшает свойство контролепригодности самого устройства.

Таблица 6. Результаты моделирования неисправностей в самопроверяемом устройстве

Пара входных комбинаций	Номер такта	Номер тестовой комбинации в паре			
		Константа 0, 1U3		Константа 1, 1U3	
		<i>TRC2</i>	<i>TRC5</i>	<i>TRC2</i>	<i>TRC5</i>
(0000, 1111)	1	2	2	1	1, 2
(0001, 1110)	2	2	2	1	1, 2
(0010, 1101)	3	2	2	1	1, 2
(0011, 1100)	4	1	1, 2	2	2
(0100, 1011)	5	2	2	1	1, 2
(0101, 1010)	6	2	2	1	1, 2
(0110, 1001)	7	2	2	1	1, 2
(0111, 1000)	8	2	1, 2	2	2

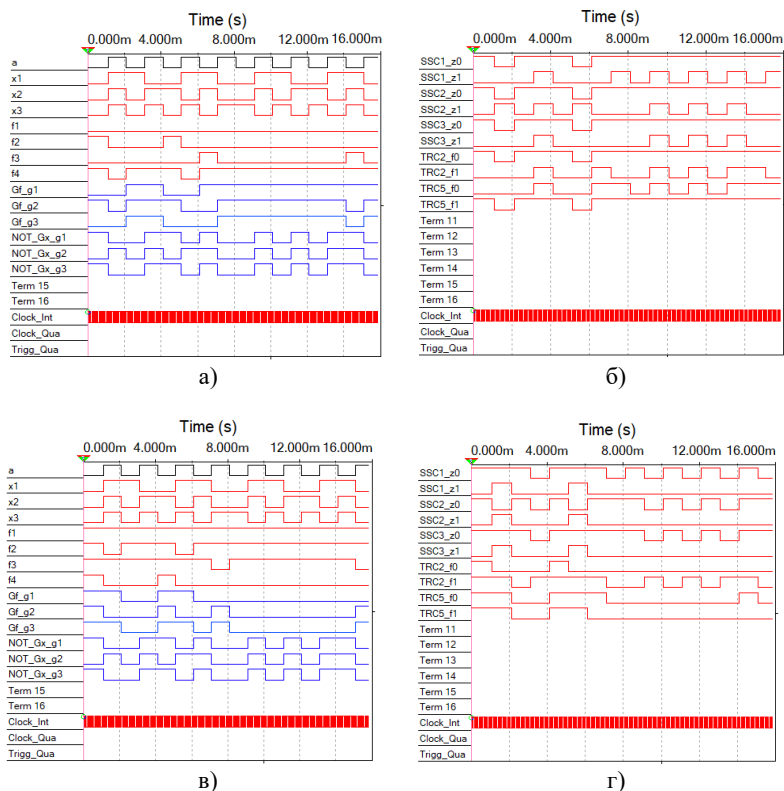


Рис. 11. Моделирование работы самопроверяемого устройства при возникновении неисправностей: а) временная диаграмма работы устройств $F(x)$, $G(f)$, $G(x)$ при неисправности «константа 0» на выходе элемента 1U3; б) временная диаграмма работы устройств SSC и TRC при неисправности «константа 0» на выходе элемента 1U3; в) временная диаграмма работы устройств $F(x)$, $G(f)$, $G(x)$ при неисправности «константа 1» на выходе элемента 1U3; г) временная диаграмма работы устройств SSC и TRC при неисправности «константа 1» на выходе элемента 1U3

8. Использование δ -преобразователей для получения самодвойственных функций. Ранее было доказано, что для кодов Хэмминга со значениями длин $n \neq 3 + 4l$, $l \in \mathbb{N}_0$ кодеры не будут являться самодвойственными устройствами. Возможна организация СВК таким образом, чтобы и для несамодвойственных кодеров $G(f)$ организовать контроль по двум диагностическим признакам. Для этого

каждый классический код Хэмминга модифицируем в некоторый δ -код, проверочные символы которого будут всегда описываться самодвойственными функциями.

Внесем следующие изменения в структуру рисунка 5. Во-первых, заменим блок $G(f)$ на блок $\Delta(f)$, образованный последовательным соединением исходного несамодвойственного блока $G(f)$ и преобразователя несамодвойственного сигнала в самодвойственный $\Delta(g)$. Устройство $\Delta(f)$ фактически будет кодером модифицированного δ -кода. Во-вторых, заменим блок $G(x)$ на блок $\Delta(x)$, который будет формировать проверочные символы δ -кода по значениям входов для устройства $F(x)$. Остальное оставим без изменений (рисунок 12).

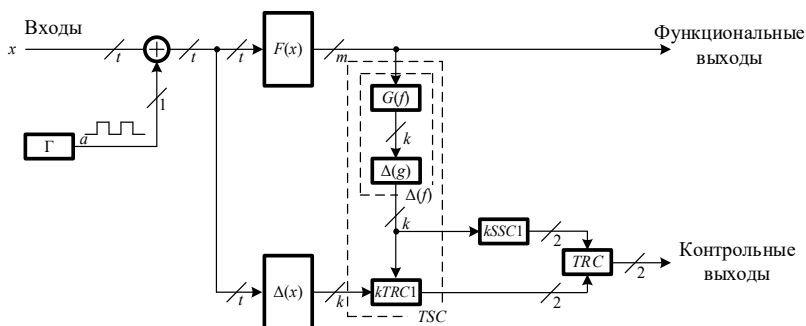


Рис. 12. Структура самопроверяемого устройства с контролем вычислений по модифицированным кодам в коды с самодвойственными контрольными функциями

Процесс синтеза блока $\Delta(x)$ не отличается от процесса синтеза блока $G(x)$. Интерес представляет процесс синтеза преобразователя $\Delta(g)$.

Теорема 4. Для того, чтобы функция, описывающая выход кодера и имеющая четное число аргументов, стала самодвойственной необходимо и достаточно сложить ее по модулю $M=2$ с функцией, описывающей старший разряд в векторе аргументов.

Доказательство. Для рассматриваемой функции q четно. Функция, описывающая старший разряд в векторе аргументов – функция f_1 , – на половине наборов равна 0 и на половине равна 1. Для тех наборов, на которых она равна 0, имеем: $\delta = g \oplus f_1 = g$. Для тех наборов, на которых она равна 1, справедливо: $\delta = g \oplus f_1 = \bar{g}$.

Таким образом, при инвертировании всех аргументов инвертируется значение функции. Функция будет самодвойственной, **что и требовалось доказать.**

Из теорем 1 и 4 следуют правила построения кодера линейного блочного кода от m переменных, снабженного k самодвойственными выходами:

$$g_j = f_{i_1} \oplus f_{i_2} \oplus \dots \oplus f_{i_q}, i_1, i_2, \dots, i_q \in \{1, 2, \dots, m\}, j = \overline{1, k}. \quad (17)$$

$$\delta_j = \begin{cases} g_j, & \text{если } q - \text{нечетно,} \\ g_j \oplus f_1, & \text{если } q - \text{четно.} \end{cases} \quad (18)$$

Выражения (17) и (18) фактически описывают проверочные символы модифицированного δ -кода.

Формула (18) применима к любому линейному блочному коду. Приведем здесь пример ее использования для формирования δ -кода из классического кода Хэмминга со значением $m=6$ (матрицы 4 и 5).

Из правил построения кода Хэмминга получаем:

$$\begin{cases} g_1 = f_1 \oplus f_2 \oplus f_4 \oplus f_5; \\ g_2 = f_1 \oplus f_3 \oplus f_4 \oplus f_6; \\ g_3 = f_2 \oplus f_3 \oplus f_4; \\ g_4 = f_5 \oplus f_6. \end{cases} \quad (19)$$

В формуле (19) $E_1 = 4, E_2 = 4, E_3 = 3, E_4 = 2$.

Необходимо преобразование g_1, g_2 и g_4 :

$$\begin{cases} \delta_1 = g_1 \oplus f_1 = f_2 \oplus f_4 \oplus f_5; \\ \delta_2 = g_2 \oplus f_1 = f_2 \oplus f_4 \oplus f_6; \\ \delta_3 = g_3; \\ \delta_4 = g_4 \oplus f_1. \end{cases} \quad (20)$$

Контрольные разряды δ -кода будут описываться системой функций (20).

Устройство $\Delta(g)$ является довольно простым. И содержит не более k элементов XOR . В таблице 7 приводятся значения числа элементов в преобразователе $\Delta(g)$ для кодов Хэмминга со значениями $k=3, 4$ и 5 . Отметим, что максимальное количество элементов XOR , равное $p=k$, используется только в одном случае числа n для каждого k : $n = 2^k - 2$. Другими словами, максимальное число элементов XOR преобразователь $\Delta(g)$ будет иметь при $n=6, 14, 30$ и т. д. Для всех остальных значений n число $p \leq k-1$.

Таблица 7. Число элементов XOR в преобразователе $\Delta(g)$

n	m	k	p					
			0	1	2	3	4	5
4	—	—						
5	2	3		×				
6	3	3				×		
7	4	3	×					
8	—	—						
9	5	4		×				
10	6	4				×		
11	7	4	×					
12	8	4			×			
13	9	4		×				
14	10	4					×	
15	11	4	×					
16	—	—						
17	12	5		×				
18	13	5			×			
19	14	5	×					
20	15	5			×			
21	16	5		×				
22	17	5					×	
23	18	5	×					
24	19	5			×			
25	20	5		×				
26	21	5					×	
27	22	5	×					
28	23	5				×		
29	24	5		×				
30	25	5						×
31	26	5	×					

Примечания. 1. Знаком «×» обозначено число элементов XOR для каждого n , необходимое для получения δ -кода. 2. Для чисел $n = 2^i, i \in \{2, 3, \dots\}$, коды Хэмминга не строятся.

В заключение отметим, что работу структуры, приведенной на рисунке 12, можно также смоделировать в Multisim и показать эффективность ее функционирования. Кроме того, как показано выше, в ее основе может использоваться любой линейный блочный код, к которому всегда можно применить преобразование (18). Это расширяет число способов организации СВК по двум диагностическим признакам.

9. Заключение. Коды Хэмминга и их модификации могут эффективно использоваться при СВК по двум диагностическим признакам – с контролем принадлежности формируемых кодовых слов выбранному коду и с контролем самодвойственности каждой вычисляемой функции. При этом использование кодов Хэмминга позволяет обнаруживать одно- и двукратные ошибки в кодовом векторе, а дополнительный контроль самодвойственности повышает число рабочих воздействий с тестовыми свойствами. Таким образом, использование предложенной новой структуры самопроверяемого устройства с контролем вычислений по кодам Хэмминга с самодвойственными функциями позволяет повышать характеристики контролепригодности – увеличивать число рабочих входных комбинаций, которые одновременно являются и тестовыми для рассматриваемого класса неисправностей. С точки зрения контролепригодности это приводит к улучшению свойства наблюдаемости.

По сравнению с традиционным использованием кодовых методов для контроля вычислений описанный в статье метод обладает следующим основным недостатком: за счет использования временной избыточности и импульсного режима работы снижается быстродействие схем. Помимо этого требуется использование тестеров самодвойственности сигналов. Однако это незначительно влияет на увеличение структурной избыточности СВК. Достоинством является повышение показателей контролепригодности и увеличение числа рабочих входных комбинаций, которые одновременно являются и тестовыми, что позволяет при редкой смене входных комбинаций осуществлять проверку работоспособности самопроверяемого устройства. Это является весомым преимуществом даже в том случае, когда применение дублирования оказывается более простым по характеристикам структурной избыточности, поскольку повышает показатели контролепригодности и позволяет фиксировать ошибки в вычислениях на большем числе рабочих комбинаций.

Представленный в статье подход к синтезу СВК по двум диагностическим признакам может быть распространен на использование других избыточных кодов (например, часто применяемых для организации контроля вычислений устройствами кодов паритета, Рида-Маллера, Сяо, алгебраических кодов и др. [55 – 62]), а также иных разделимых блочных кодов с учетом их характеристик.

Литература

1. Ubar R., Raik J., Vierhaus H. Design and Test Technology for Dependable Systems-on-Chip (Premier Reference Source). Information Science Reference. 2011. 578 p.
2. Дрозд А.В., Харченко В.С., Антошук С.Г., Дрозд Ю.В., Дрозд М.А., Сулима Ю.Ю. Рабочее диагностирование безопасных информационно-управляющих систем. Под ред. Дрозда А.В. и Харченко В.С. Харьков: Национальный аэрокосмический университет им. Н.Е. Жуковского «ХАИ», 2012. 614 с.
3. Dubrova E. Fault-Tolerant Design. Springer Science+Business Media New York. 2013. 185 p. DOI: 10.1007/978-1-4614-2113-9.
4. Hahanov V. Cyber Physical Computing for IoT-driven Services. New York, Springer International Publishing AG, 2018. 279 p.
5. Ярмолик В.Н. Контроль и диагностика вычислительных систем. Минск: «Бестпринт», 2019. 387 с.
6. Согомонян Е.С., Слабаков Е.В. Самопроверяемые устройства и отказоустойчивые системы. М.: Радио и связь. 1989. 208 с.
7. Mitra S., McCluskey E. Which Concurrent Error Detection Scheme to Choose? // Proceedings of International Test Conference. 2000. pp. 985-994. DOI: 10.1109/TEST.2000.894311.
8. Гаврилов М.А., Остиану В.М., Потехин А.И. Надежность дискретных систем // Итоги науки и техники. Серия «Теория вероятностей. Математическая статистика. Теоретическая кибернетика». 1969, 1970. С. 7-104.
9. Сапожников В.В., Сапожников Вл.В., Ефанов Д.В. Теория синтеза самопроверяемых цифровых систем на основе кодов с суммированием. Санкт-Петербург: «Лань», 2021. 580 с.
10. Багхдади А.А.А., Хаханов В.И., Литвинова Е.И. Методы анализа и диагностирования цифровых устройств (аналитический обзор) // Автоматизированные системы управления и приборы автоматики. 2014. № 166. С. 59-74.
11. Сапожников В.В., Сапожников Вл.В., Ефанов Д.В. Основы теории надежности и технической диагностики. Санкт-Петербург: Издательство «Лань», 2019. 588 с.
12. Hamming R. Error Detecting and Correcting Codes // Bell System Technical Journal. 1950. vol. 29 (2). pp. 147-160. DOI: 10.1002/j.1538-7305.1950.tb00463.x.
13. Яблонский С.В. Введение в дискретную математику. Под ред. В.А. Садовничева. М.: «Высшая школа», 2003. 384 с.
14. Reynolds D. Meize G. Fault Detection Capabilities of Alternating Logic // IEEE Transactions on Computers. 1978. vol. C-27(12). pp. 1093-1098. DOI: 10.1109/TC.1978.1675011.

15. Гессель М., Морозов А.А., Сапожников В.В., Сапожников Вл.В. Построение самопроверяемых комбинационных схем на основе свойств самодвойственных функций // Автоматика и телемеханика. 2000. № 2. С. 151-163.
16. Takeda K., Tohma J. Logic Design of Fault-Tolerant Arithmetic Units Based on the Data Complementation Strategy // 10th International Symposium on Fault-Tolerant Computing (FTCS'10). 1980. p. 348.
17. Biernat J. Self-Dual Modules in Design of Dependable Digital Devices // International Conference on Dependability of Computer Systems. 2006. DOI: 10.1109/DEPCOS-RELCOMEX.2006.50.
18. Rai S., Raitza M., Sahoo S., Kumar A. DiSCERN: Distilling Standard-Cells for Emerging Reconfigurable Nanotechnologies // Design, Automation & Test in Europe Conference & Exhibition (DATE). 2020. DOI: 10.23919/DATE48585.2020.9116216.
19. Аксенова Г.П. Восстановление в дублированных устройствах методом инвертирования данных // Автоматика и телемеханика. 1987. № 10. С. 144-153.
20. Гессель М., Мошанин В.И., Сапожников В.В., Сапожников Вл.В. Обнаружение неисправностей в самопроверяемых комбинационных схемах с использованием свойств самодвойственных функций // Автоматика и телемеханика. 1997. № 12. С. 193-200.
21. Сапожников В.В., Сапожников Вл.В., Гессель М. Самодвойственные дискретные устройства. СПб: Энергоатомиздат, 2001. 331 с.
22. Sentovich E., Singh K., Moon C., Savoj H., Brayton R., Sangiovanni-Vincentelli A. Sequential Circuit Design Using Synthesis and Optimization // Proceedings IEEE International Conference on Computer Design: VLSI in Computers & Processors. 1992. pp. 328-333. DOI: 10.1109/ICCD.1992.276282.
23. Sentovich E., Singh K., Lavagno L., Moon C., Murgai R., Saldanha A., Savoj H., Stephan P., Brayton R., Sangiovanni-Vincentelli A. SIS: A System for Sequential Circuit Synthesis // Electronics Research Laboratory, Department of Electrical Engineering and Computer Science. 1992. 45 p.
24. Efanov D., Sapozhnikov V., Sapozhnikov Vl., Osadchy G., Pivovarov D. Self-Dual Complement Method up to Constant-Weight Codes for Arrangement of Combinational Logical Circuits Concurrent Error-Detection Systems // Proceedings of 17th IEEE East-West Design & Test Symposium (EWDTS'2019). 2019. pp. 136-143. DOI: 10.1109/EWDTS.2019.8884398.
25. Carter W., Duke K., Schneider P. Self-Checking Error Checker for Two-Rail Coded Data // United States Patent Office. 1968. no. 747533. 10 p.
26. Сапожников В.В., Сапожников Вл.В. Самопроверяемый компаратор с дополнительным импульсным входом // Автоматика и телемеханика. 1997. № 6. С. 200-208.
27. Sapozhnikov Vl.V., Dmitriev A., Goessel M., Saposhnikov V.V. Self-Dual Parity Checking a New Method for on Line Testing // Proceedings of 14th IEEE VLSI Test Symposium. 1996. pp. 162-168. DOI: 10.1109/VTEST.1996.510852.
28. Гессель М., Дмитриев А.В., Сапожников В.В., Сапожников Вл.В. Самотестируемая структура для функционального обнаружения отказов в комбинационных схемах // Автоматика и телемеханика. 1999. № 11. С. 162-174.
29. Гессель М., Морозов А.В., Сапожников В.В., Сапожников Вл.В. Логическое дополнение новый метод контроля комбинационных схем // Автоматика и телемеханика. 2003. № 1. С. 167-176.

30. Пивоваров Д.В. Построение систем функционального контроля многовыходных комбинационных схем методом логического дополнения по равновесным кодам // Автоматика на транспорте. 2018. Т. 4. № 1. С. 131-149.
31. Пашуков А.В. Применение взвешенных кодов с суммированием при синтезе схем встроенного контроля по методу логического дополнения // Автоматика на транспорте. 2022. Т. 8. № 1. С. 101-114. DOI: 10.20295/2412-9186-2022-8-01-101-114.
32. Аксёнова Г.П. Метод синтеза схем встроенного контроля для автоматов с памятью // Автоматика и телемеханика. 1973. № 2. С. 109-116.
33. Гессель М., Дмитриев А.В., Сапожников В.В., Сапожников Вл.В. Обнаружение неисправностей в комбинационных схемах с помощью самодвойственного контроля // Автоматика и телемеханика. 2000. № 7. С. 140-149.
34. Saposhnikov V.I., Moshanin V., Saposhnikov V., Goessel M. Experimental Results for Self-Dual Multi-Output Combinational Circuits // Journal of Electronic Testing: Theory and Applications. 1999. vol. 14(3). pp. 295-300. DOI: 10.1023/A:1008370405607.
35. Гессель М., Дмитриев А.В., Сапожников В.В., Сапожников Вл.В. Исследование свойств самодвойственных самопроверяемых многотактных схем // Автоматика и телемеханика. 2001. № 4. С. 148-159.
36. Сапожников В.В., Сапожников Вл.В., Валиев Р.Ш. Синтез самодвойственных дискретных систем. СПб: Элмор, 2006. 220 с.
37. Göessel M., Ocheretny V., Sogomonyan E., Marienfeld D. New Methods of Concurrent Checking. Frontiers in Electronic Testing, Springer. 2008. vol. 42. 184 p.
38. Тельпухов Д.В., Жукова Т.Д., Деменева А.И., Гуров С.И. Схема функционального контроля для комбинационных схем на основе R-кода // Проблемы разработки перспективных микро- и нанoeлектронных систем (МЭС). 2018. № 4. С. 98-104. DOI: 10.31114/2078-7707-2018-4-98-104.
39. Стемповский А.Л., Тельпухов Д.В., Жукова Т.Д., Деменева А.И., Надоленко В.В., Гуров С.И. Синтез схемы функционального контроля на основе спектрального R-кода с разбиением выходов на группы // Микроэлектроника. 2019. Т. 48. № 4. С. 284-294. DOI: 10.1134/S0544126919040094.
40. Стемповский А.Л., Тельпухов Д.В., Гуров С.И., Жукова Т.Д., Щелоков А.Н., Новиков А.Д. Синтез СФК на основе LDPC кода с использованием мажоритарного декодирования // Известия ЮФУ. Технические науки. 2019. № 4(206). С. 195-206. DOI: 10.23683/2311-3103-2019-4-195-206.
41. Абдуллаев Р.Б. Синтез полностью самопроверяемых схем встроенного контроля на основе полиномиальных кодов для комбинационных логических устройств // Автоматика на транспорте. 2021. Т. 7. № 3. С. 452-476. DOI: 10.20295/2412-9186-2021-7-3-452-476.
42. Berger J.M. A Note on Error Detection Codes for Asymmetric Channels // Information and Control. 1961. vol. 4(1). pp. 68-73. DOI: 10.1016/S0019-9958(61)80037-5.
43. Piestrak S.J. Design of Self-Testing Checkers for Unidirectional Error Detecting Codes. Wrocław: Oficyna Wydawnicza Politechniki Wrocławskiej, 1995. 111 p.
44. Drozd O., Antoniuk V., Nikul V., Drozd M. Hidden Faults in FPGA-Built Digital Components of Safety-Related Systems // Proceedings of the 14th International Conference "TCSET" 2018. 2018. pp. 805-809. DOI: 10.1109/TCSET.2018.8336320.
45. Drozd O., Rucinski A., Zascholkina K., Martynyuk O., Drozd J. Resilient Development of Models and Methods in Computing Space // Proceedings of 19th IEEE East-West

- Design & Test Symposium (EWDTS'2021). 2021. pp. 70-75. DOI: 10.1109/EWDTS52692.2021.9581002.
46. Сапожников Вл.В. Синтез систем управления движением поездов на железнодорожных станциях с исключением опасных отказов. М.: Наука, 2021. 229 с.
47. Tshagharyan G., Harutyunyan G., Shoukourian S., Zorian Y. Experimental Study on Hamming and Hsiao Codes in the Context of Embedded Applications // Proceedings of 15th IEEE East-West Design & Test Symposium (EWDTS'2017). 2017. pp. 25-28. DOI: 10.1109/EWDTS.2017.8110065.
48. Тельпухов Д.В., Жукова Т.Д., Щелоков А.Н., Крестинина П.Д. Применение кода Хэмминга в задаче повышения сбоеустойчивости комбинационных схем // Известия ЮФУ. Технические науки. 2021. № 4(221). С. 220-231. DOI: 10.18522/2311-3103-2021-4-220-231.
49. Ефанов Д.В. Пределные свойства кода Хэмминга в схемах функционального диагностирования // Информатика и системы управления. 2011. № 3. С. 70-79.
50. Сапожников В.В., Сапожников Вл.В., Ефанов Д.В. Особенности применения кодов Хэмминга при организации самопроверяемых схем встроенного контроля // Известия высших учебных заведений. Приборостроение. 2018. Т. 61. № 1. С. 47-59. DOI: 10.17586/0021-3454-2018-61-1-47-59.
51. Сапожников В.В., Сапожников Вл.В., Ефанов Д.В. Коды Хэмминга в системах функционального контроля логических устройств. СПб.: Наука, 2018. 151 с.
52. Ефанов Д.В., Погодина Т.С. Самодвойственный контроль комбинационных схем с применением кодов Хэмминга // Проблемы разработки перспективных микро- и наноэлектронных систем (МЭС). 2022. № 3. С. 113-122. DOI: 10.31114/2078-7707-2022-3-113-122.
53. Zhang C., Liu Y., Jiang T., Mao W., Wang J. Multisim-Based Digital Clock Design // 2020 IEEE 9th Joint International Information Technology and Artificial Intelligence Conference (ITAIC). 2020. DOI: 10.1109/ITAIC49862.2020.9338902.
54. Chen Y., Zhang M., Hao J. The Circuit Design of Voltage-controlled Color Changing Lamp Based on Multisim // IEEE International Conference on Power, Intelligent Computing and Systems (ICPICS). 2020. DOI: 10.1109/ICPICS50287.2020.9202148.
55. Richter M., Goessel M. Concurrent Checking with Split-Parity Codes // 15th IEEE International On-Line Testing Symposium. 2009. pp. 159-163, DOI: 10.1109/IOLTS.2009.5196001.
56. Sogomonyan E., Weidling S., Goessel M. A New Method for Correcting Time and Soft Errors in Combinational Circuits // IEEE 16th International Symposium on Design and Diagnostics of Electronic Circuits & Systems (DDECS). 2013. pp. 283-286. DOI: 10.1109/DDECS.2013.6549835.
57. Gopi S., Kopparty S., Oliveira R., Ron-Zewi N., Saraf S. Locally Testable and Locally Correctable Codes Approaching the Gilbert-Varshamov Bound // IEEE Transactions on Information Theory. 2018. vol. 64(8). pp. 5813-5831. DOI: 10.1109/TIT.2018.2809788.
58. Harsha P., Srinivasan S. Robust Multiplication-Based Tests for Reed-Muller Codes // IEEE Transactions on Information Theory. 2019. vol. 65(1). pp. 184-197. DOI: 10.1109/TIT.2018.2863713.
59. Mandry H., Herkle A., Kürzinger L., Muelich S., Becker J., Fischer R., Ortmanns M. Modular PUF Coding Chain with High-Speed Reed-Muller Decoder // IEEE International Symposium on Circuits and Systems (ISCAS). 2019. DOI: 10.1109/ISCAS.2019.8702484.

60. Sim M., Zhuang Y. Design of Two Interleaved Error Detection and Corrections Using Hsiao Code and CRC // IECON 2020 The 46th Annual Conference of the IEEE Industrial Electronics Society. 2020. DOI: 10.1109/IECON43393.2020.9254837.
61. Абдуллаев Р.Б. Вероятностные характеристики полиномиальных кодов в системах технического диагностирования // Автоматика на транспорте. 2020. Т. 6. № 1. С. 64-88. DOI: 10.20295/2412-9186-2020-6-1-64-88.
62. Mishra N., Naresh N., Acharya A. Parallel Field Test Architecture for Boot-ROMs in Safety-Critical SoCs // 2021 IEEE International Test Conference India (ITC India). 2021. DOI: 10.1109/ITCIndia52672.2021.9532633.

Ефанов Дмитрий Викторович — д-р техн. наук, доцент, профессор кафедры, кафедра автоматики, телемеханики и связи на железнодорожном транспорте, Российский университет транспорта (МИИТ); профессор, высшая школа транспорта института машиностроения, материалов и транспорта, Санкт-Петербургский политехнический университет Петра Великого. Область научных интересов: дискретная математика, надежность, безопасность и техническая диагностика дискретных систем, методы непрерывного мониторинга систем автоматического управления и сложных инженерных конструкций и сооружений. Число научных публикаций — 500. TrES-4b@yandex.ru; улица Образцова, 9/9, 127994, Москва, Россия; р.т.: +7(911)709-2164.

Погодина Татьяна Сергеевна — студент, кафедра автоматики, телемеханики и связи на железнодорожном транспорте, Российский университет транспорта (МИИТ). Область научных интересов: дискретная математика, надежность и техническая диагностика дискретных систем. Число научных публикаций — 4. pogodina-ts@mail.ru; улица Образцова, 9/9, 127994, Москва, Россия; р.т.: +7(977)404-7953.

D. EFANOV, T. POGODINA
**PROPERTIES INVESTIGATION OF SELF-DUAL
COMBINATIONAL DEVICES WITH CALCULATION CONTROL
BASED ON HAMMING CODES**

Efanov D., Pogodina T. Properties Investigation of Self-Dual Combinational Devices with Calculation Control Based on Hamming Codes.

Abstract. A new approach to the synthesis of self-checking devices is considered, based on the control of calculations in testing objects using Hamming codes, the check bits of which are described by self-dual functions. In this case, the structure operates in a pulsed mode, which is actually based on the introduction of temporal redundancy when building a self-checking device. This, unfortunately, leads to some decrease in performance, however, it significantly improves the characteristics of controllability, which is especially important for devices and systems of critical use, the input data for which does not change so often. A brief review of methods for constructing built-in control circuits based on the self-duality property of calculated functions is given. The basic structures of the organization of built-in control circuits are given. The proposed ways of developing the theory of synthesis of built-in control circuits are based on checking whether or not the calculated functions belong to a class of self-dual Boolean functions. All possible values of the number of data bits for Hamming codes have been established. They will have the property of the self-duality of functions describing control bits. En-coders of such Hamming codes will be self-dual devices. Since the functions of the check bits of Hamming codes are linear, in order for them to be self-dual, it is necessary that an odd number of arguments be used in each of them. It is proved that the number of bits of code words of Hamming codes with self-dual check functions is equal to $n=3+4l$, $l \in \mathbb{N}_0$. The results of the simulations self-dual devices with built-in control circuits along two diagnostic parameters in the Multisim environment are presented. A method is proposed for modification of the structure of calculation control along two diagnostic parameters, which allows to use any linear block code (not necessarily Hamming code). It is based on retrofitting the encoder with a device for converting functions into self-dual ones. In fact, this is a code modification device. It is proved that to obtain a modified Hamming code with self-dual control functions for $n=3+4l$, $l \in \mathbb{N}_0$; cases, it is enough to add modulo $M=2$ the non-self-dual control function with the function of the high data bit.

Keywords: self-checking combinational device, integrated control circuit, checking of calculations at the outputs of combinational devices, linear block code, checking of calculations by two diagnostic parameters, control of self-duality, checking of calculations by Hamming codes.

References

1. Ubar R., Raik J., Vierhaus H. Design and Test Technology for Dependable Systems-on-Chip (Premier Reference Source). Information Science Reference. 2011. 578 p.
2. Drozd A.V., Kharchenko V.S., Antoshchuk S.G., Drozd Yu.V., Drozd M.A., Sulima Yu.Yu. Rabochee diagnostirovanie bezopasnykh informacionno-upravljajushhih sistem [Working Diagnostics of Safe Information and Control Systems]. Eds: A.V. Drozd, V.S. Kharchenko. Khar'kov: N.E. Zhukovsky National Aerospace University «KAI», 2012. 614 p. (In Russ.).

3. Dubrova E. Fault-Tolerant Design. Springer Science+Business Media New York. 2013. 185 p. DOI: 10.1007/978-1-4614-2113-9.
4. Hahanov V. Cyber Physical Computing for IoT-driven Services. New York, Springer International Publishing AG, 2018. 279 p.
5. Yarmolik V.N. Kontrol' i diagnostika vychislitel'nyh sistem [Control and Diagnostics of Computer Systems]. Minsk: «Bestprint», 2019. 387 p. (In Russ.).
6. Sogomonyan E.S., Slabakov E.V. Samoproverjaemye ustrojstva i otkazoustojchivye sistemy [The Self-Checked Devices and Failure-Safe Systems]. M.: Radio and Communications, 1989. 208 p. (In Russ.).
7. Mitra S., McCluskey E. Which Concurrent Error Detection Scheme to Choose? // Proceedings of International Test Conference. 2000. pp. 985-994. DOI: 10.1109/TEST.2000.894311.
8. Gavrilov M A., Ostianu V.M., Potekhin A.I. [Reliability of discrete systems]. Itogi Nauki. Seriya "Teoriya Veroyatnostei. Matematicheskaya Statistika. Teoreticheskaya Kibernetika Probability Theory. Mathematical Statistics. Theoretical Cybernetics. 1969, 1970. pp. 7-104. (In Russ.).
9. Sapozhnikov V.V., Sapozhnikov V.I., Efanov D.V. Teoriya sinteza samoproverjaemyh cifrovyyh sistem na osnove kodov s summirovaniem [Theory of Synthesis of Self-Checking Digital Systems Based on Sum Codes]. St. Petersburg: Lan Publishing House, 2021. 580 p. (In Russ.).
10. Baghdadi A.A.A., Hahanov V.I., Litvinova E.I. [Digital System Analysis and Diagnosis Methods (Analytical Review)]. Avtomatizirovannyye sistemy upravleniya i pribory avtomatiki Management Information System and Devices. 2014. no. 166. pp. 59-74. (In Russ.).
11. Sapozhnikov V.V., Sapozhnikov V.I., Efanov D.V. Osnovy teorii nadezhnosti i tehniceskoy diagnostiki [Fundamentals of the Theory of Reliability and Technical Diagnostics]. St. Petersburg: Lan Publishing House, 2019. 588 p. (In Russ.).
12. Hamming R. Error Detecting and Correcting Codes // Bell System Technical Journal. 1950. vol. 29 (2). pp. 147-160. DOI: 10.1002/j.1538-7305.1950.tb00463.x.
13. Yablonsky S.V. Vvedenie v diskretnuyu matematiku [Introduction to Discrete Mathematics]. Proc. manual for universities, Ed.: V.A. Sadovnichev. Moscow, Higher School, 2003. 384 p. (In Russ.).
14. Reynolds D. Meize G. Fault Detection Capabilities of Alternating Logic. IEEE Transactions on Computers. 1978. vol. C-27(12). pp. 1093-1098. DOI: 10.1109/TC.1978.1675011.
15. Göessel M., Morozov A.A., Sapozhnikov V.V., Sapozhnikov V.I. [Self-Testing Combinational Circuits: Their Design Through the Use of the Properties of Self-Dual Functions]. Avtomatika i telemekhanika Automation and Remote Control. 2000. no. 2. pp. 151-163. (In Russ.).
16. Takeda K., Tohma J. Logic Design of Fault-Tolerant Arithmetic Units Based on the Data Complementation Strategy. 10th International Symposium on Fault-Tolerant Computing (FTCS'10). 1980. p. 348.
17. Biernat J. Self-Dual Modules in Design of Dependable Digital Devices // International Conference on Dependability of Computer Systems. 2006. DOI: 10.1109/DEPCOS-RELCOMEX.2006.50.
18. Rai S., Raitza M., Sahoo S., Kumar A. DiSCERN: Distilling Standard-Cells for Emerging Reconfigurable Nanotechnologies. Design, Automation & Test in Europe Conference & Exhibition (DATE). 2020. DOI: 10.23919/DATE48585.2020.9116216.

19. Aksenova G.P. [Restoration in Duplicated Units by the Method of Data Inversion]. *Avtomatika i telemekhanika Automation and Remote Control*. 1987. no. 10. pp. 144-153. (In Russ.).
20. Göessel M., Moshanin V.I., Sapozhnikov V.V., Sapozhnikov V.I.V. [Fault Detection in Self-Test Combination Circuits Using the Properties of Self-Dual Functions]. *Avtomatika i telemekhanika Automation and Remote Control*. 1997. no. 12. pp. 193-200. (In Russ.).
21. Sapozhnikov V.V., Sapozhnikov V.I.V., Göessel M. *Samodvoystvennyye diskretnyye ustroystva [Self-dual Digital Devices]*. St. Petersburg: Energoatomizdat, 2001. 331 p. (In Russ.).
22. Sentovich E., Singh K., Moon C., Savoj H., Brayton R., Sangiovanni-Vincentelli A. Sequential Circuit Design Using Synthesis and Optimization. *Proceedings IEEE International Conference on Computer Design: VLSI in Computers & Processors*. 1992. pp. 328-333. DOI: 10.1109/ICCD.1992.276282.
23. Sentovich E., Singh K., Lavagno L., Moon C., Murgai R., Saldanha A., Savoj H., Stephan P., Brayton R., Sangiovanni-Vincentelli A. SIS: A System for Sequential Circuit Synthesis. *Electronics Research Laboratory, Department of Electrical Engineering and Computer Science*. 1992. 45 p.
24. Efanov D., Sapozhnikov V., Sapozhnikov V.I., Osadchy G., Pivovarov D. Self-Dual Complement Method up to Constant-Weight Codes for Arrangement of Combinational Logical Circuits Concurrent Error-Detection Systems. *Proceedings of 17th IEEE East-West Design & Test Symposium (EWDTS'2019)*. 2019. pp. 136-143. DOI: 10.1109/EWDTS.2019.8884398.
25. Carter W.C., Duke K.A., Schneider P.R. Self-Checking Error Checker for Two-Rail Coded Data. *United States Patent Office*, filed July 25, 1968, ser. No. 747533, patented Jan. 26, 1971, N. Y., 10 p.
26. Saposhnikov V.V., Saposhnikov V.I.V. [A Self-Checking Comparator with Additional Pulse Input]. *Avtomatika i telemekhanika Automation and Remote Control*. 1997. no. 6. pp. 200-208. (In Russ.).
27. Saposhnikov V.I.V., Dmitriev A., Goessel M., Saposhnikov V.V. Self-Dual Parity Checking a New Method for on Line Testing. *Proceedings of 14th IEEE VLSI Test Symposium*. 1996. pp. 162-168. DOI: 10.1109/VTEST.1996.510852.
28. Göessel M., Dmitriev A.V., Sapozhnikov V.V., Sapozhnikov V.I.V. [A Functional Fault-Detection Self-Test for Combinational Circuits]. *Avtomatika i telemekhanika – Automation and Remote Control*. 1999. no. 11. pp. 162-174. (In Russ.).
29. Gessel M., Morozov A.V., Sapozhnikov V.V., Sapozhnikov V.I.V. [Logic Complement, a New Method of Checking the Combinational Circuits]. *Avtomatika i telemekhanika Automation and Remote Control*. 2003. no. 1. pp. 167-176. (In Russ.).
30. Pivovarov D.V. [Formation of Concurrent Error Detection Systems in Multiple-Output Combinational Circuits Using the Boolean Complement Method Based on Constant-Weight Codes]. *Avtomatika na transporte – Transport Automation Research*. 2018. vol. 4. no. 1. pp. 131-149. (In Russ.).
31. Pashukov A.V. [Application of Weight-Based Sum Codes at the Synthesis of Circuits for Built-In Control by Boolean Complement Method]. *Avtomatika na transporte – Transport Automation Research*. 2022. vol. 8. no. 1. pp. 101-114. DOI: 10.20295/2412-9186-2022-8-01-101-114. (In Russ.).
32. Aksjonova G.P. [Method of Synthesizing Built-In Monitoring Arrangements for Automata with Memory]. *Avtomatika i telemekhanika – Automation and Remote Control*. 1973. no. 2. pp. 109-116. (In Russ.).

33. Göessel M., Dmitriev A.V., Sapozhnikov V.V., Sapozhnikov V.I.V. [Detection of Faults in Combinational Circuits by a Self-Dual Test]. *Avtomatika i telemekhanika Automation and Remote Control*. 2000. no. 7. pp. 140-149. (In Russ.).
34. Saposhnikov V.I.V., Moshanin V., Saposhnikov V.V., Goessel M. Experimental Results for Self-Dual Multi-Output Combinational Circuits. *Journal of Electronic Testing: Theory and Applications*. 1999. vol. 14(3). pp. 295-300. DOI: 10.1023/A:1008370405607.
35. Göessel M., Dmitriev A.V., Sapozhnikov V.V., Sapozhnikov V.I.V. [Malfunctioning Detection in Combination Circuits Via Self-Dual Duplication]. *Avtomatika i telemekhanika Automation and Remote Control*. 2001. no. 4. pp. 148-159. (In Russ.).
36. Sapozhnikov V.V., Sapozhnikov V.I.V., Valiev R.Sh. Sintez samodvoystvennykh diskretnykh sistem [Synthesis of Self-Dual Digital Systems]. St. Petersburg: Elmor, 2006. 220 p. (In Russ.).
37. Göessel M., Ocheretny V., Sogomonyan E., Marienfeld D. New Methods of Concurrent Checking. *Frontiers in Electronic Testing*, Springer. 2008. vol. 42. 184 p.
38. Telpukhov D.V., Zhukova T.D., Demeneva A.I., Gurov S.I. [Circuit of Functional Control for Combinational Circuits Based on R-Code]. *Problemy razrabotki perspektivnykh mikro- i nanojelektronnykh sistem (MJeS) Problems of Advanced Micro- and Nanoelectronic Systems Development (MES)*. 2018. no. 4. pp. 98-104. DOI: 10.31114/2078-7707-2018-4-98-104. (In Russ.).
39. Stempkovskii A.L., Tel'pukhov D.V., Zhukova T.D., Demeneva A.I., Nadolenko V.V., Gurov S.I. [Synthesis of a Concurrent Error Detection Circuit Based on the Spectral R-Code with the Partitioning of Outputs into Groups]. *Mikroelektronika – Russian Microelectronics*. 2019. vol. 48. no. 4. pp. 240-249. DOI: 10.1134/S0544126919040094. (In Russ.).
40. Stempkovskiy A.L., Telpukhov D.V., Gurov S.I., Zhukova T.D., Schelokov A.N., Novikov A.D. [Synthesis Method of Fault-Tolerant Combination Circuits with CED Based on LDPC Code]. *Izvestiya JuFU. Tehnicheskie nauki Izvestiya SFedU. Engineering Sciences*. 2019. no. 4 (206). pp. 195-206. DOI: 10.23683/2311-3103-2019-4-195-206. (In Russ.).
41. Abdullaev R.B. [Synthesis of Fully Self-Checked Schemes Built-In Control Based on Polynomial Codes for Combination Logic Devices]. *Avtomatika na transporte – Transport Automation Research*. 2021. vol. 7. no. 3. pp. 452-476. DOI: 10.20295/2412-9186-2021-7-3-452-476. (In Russ.).
42. Berger J.M. A Note on Error Detection Codes for Asymmetric Channels. *Information and Control*. 1961. vol. 4(1). pp. 68-73. DOI: 10.1016/S0019-9958(61)80037-5.
43. Piestrak S.J. Design of Self-Testing Checkers for Unidirectional Error Detecting Codes. Wrocław: Oficyna Wydawnicza Politechniki Wrocławskiej, 1995. 111 p.
44. Drozd O., Antoniuk V., Nikul V., Drozd M. Hidden Faults in FPGA-Built Digital Components of Safety-Related Systems. *Proceedings of the 14th International Conference “TCSET’2018*. 2018. pp. 805-809. DOI: 10.1109/TCSET.2018.8336320.
45. Drozd O., Rucinski A., Zascholkina K., Martynyuk O., Drozd J. Resilient Development of Models and Methods in Computing Space // *Proceedings of 19th IEEE East-West Design & Test Symposium (EWDTS’2021)*. 2021. pp. 70-75. DOI: 10.1109/EWDTS52692.2021.9581002.
46. Sapozhnikov V.I.V. Sintez sistem upravleniya dvizheniem poezdov na zheleznodorozhnykh stantsiyah s iskljucheniem opasnykh otkazov [Synthesis of train traffic control system at railway stations with the exception of dangerous failures]. M.: Nauka, 2021. 229 p. (In Russ.).

47. Tshagharyan G., Harutyunyan G., Shoukourian S., Zorian Y. Experimental Study on Hamming and Hsiao Codes in the Context of Embedded Applications. Proceedings of 15th IEEE East-West Design & Test Symposium (EWDTS'2017). 2017. pp. 25-28. DOI: 10.1109/EWDTS.2017.8110065.
48. Telpukhov D.V., Zhukova T.D., Schelokov A.N., Kretinina P.D. [Application of the Hamming Code in the Problem of Increasing Fault Tolerance of Logic Circuits]. Izvestija JuFU. Tehnicheskie nauki Izvestiya SFedU. Engineering Sciences. 2021. no. 4(221). pp. 220-231. DOI: 10.18522/2311-3103-2021-4-220-231. (In Russ.).
49. Efanov D.V. [The Hamming Code's Limit Properties in Functional Control Scheme]. Informatika i sistemy upravleniya Information Science and Control Systems. 2011. no. 3. pp. 70-79. (In Russ.).
50. Sapozhnikov V.V., Sapozhnikov V.I., Efanov D.V. [Features of Hamming Codes Application in Self-Checking Test Circuit Organization]. Izvestiya vysshih uchebnykh zavedenij. Priborostroenie Journal of Instrument Engineering. 2018. vol. 61. no. 1. pp. 47-59. DOI: 10.17586/0021-3454-2018-61-1-47-59. (In Russ.).
51. Sapozhnikov V.V., Sapozhnikov V.I., Efanov D.V. Kody Hjemminga v sistemah funkcional'nogo kontrolja logicheskikh ustrojstv [Hamming Codes in Concurrent Error Detection Systems of Logic Devices]. SPb.: Nauka, 2018. 151 p. (In Russ.).
52. Efanov D.V., Pogodina T.S. [Self-Dual Control of Combinational Circuits with Using Hamming Codes]. Problemy razrabotki perspektivnykh mikro- i nanoelektronnykh sistem (MJeS) Problems of Advanced Micro- and Nanoelectronic Systems Development (MES). 2022. no. 3. pp. 113-122. DOI: 10.31114/2078-7707-2022-3-113-122. (In Russ.).
53. Zhang C., Liu Y., Jiang T., Mao W., Wang J. Multisim-Based Digital Clock Design. 2020 IEEE 9th Joint International Information Technology and Artificial Intelligence Conference (ITAIC). 2020. DOI: 10.1109/ITAIC49862.2020.9338902.
54. Chen Y., Zhang M., Hao J. The Circuit Design of Voltage-controlled Color Changing Lamp Based on Multisim. IEEE International Conference on Power, Intelligent Computing and Systems (ICPICS). 2020. DOI: 10.1109/ICPICS50287.2020.9202148.
55. Richter M., Goessel M. Concurrent Checking with Split-Parity Codes. 15th IEEE International On-Line Testing Symposium. 2009. pp. 159-163, DOI: 10.1109/IOLTS.2009.5196001.
56. Sogomonyan E., Weidling S., Goessel M. A New Method for Correcting Time and Soft Errors in Combinational Circuits. IEEE 16th International Symposium on Design and Diagnostics of Electronic Circuits & Systems (DDECS). 2013. pp. 283-286. DOI: 10.1109/DDECS.2013.6549835.
57. Gopi S., Kopparty S., Oliveira R., Ron-Zewi N., Saraf S. Locally Testable and Locally Correctable Codes Approaching the Gilbert-Varshamov Bound. IEEE Transactions on Information Theory. 2018. vol. 64(8). pp. 5813-5831. DOI: 10.1109/TIT.2018.2809788.
58. Harsha P., Srinivasan S. Robust Multiplication-Based Tests for Reed-Muller Codes. IEEE Transactions on Information Theory. 2019. vol. 65(1). pp. 184-197. DOI: 10.1109/TIT.2018.2863713.
59. Mandry H., Herkle A., Kürzinger L., Muelich S., Becker J., Fischer R., Ortmanns M. Modular PUF Coding Chain with High-Speed Reed-Muller Decoder. IEEE International Symposium on Circuits and Systems (ISCAS). 2019. DOI: 10.1109/ISCAS.2019.8702484.
60. Sim M., Zhuang Y. Design of Two Interleaved Error Detection and Corrections Using Hsiao Code and CRC. IECON 2020 The 46th Annual Conference of the IEEE Industrial Electronics Society. 2020. DOI: 10.1109/IECON43393.2020.9254837.

61. Abdullaev R.B. [Probabilistic Features of Polynomial Codes in Technical Diagnosis Systems]. *Avtomatika na transporte – Transport Automation Research*. 2020. vol. 6 no. 1. pp. 64-88. DOI: 10.20295/2412-9186-2020-6-1-64-88. (In Russ.).
62. Mishra N., Naresh N., Acharya A. Parallel Field Test Architecture for Boot-ROMs in Safety-Critical SoCs. 2021 IEEE International Test Conference India (ITC India). 2021. DOI: 10.1109/ITCIndia52672.2021.9532633.

Efanov Dmitry — Ph.D., Dr.Sci., Associate Professor, Professor of the department, Department of automation, remote control and communication on railway transport, Russian University of Transport; Professor, Transport higher school of mechanical engineering, material and transport institute, Peter the Great Saint Petersburg Polytechnic University. Research interests: discrete mathematics, reliability, safety and technical diagnostics of discrete devices, methods of health monitoring of automatic control systems and complex engineering structures and facilities. The number of publications — 500. TrES-4b@yandex.ru; 9/9, Obraztsova St., 127994, Moscow, Russia; office phone: +7(911)709-2164.

Pogodina Tatiana — Student, Department of automation, remote control and communication on railway transport, Russian University of Transport. Research interests: discrete mathematics, reliability and technical diagnostics of discrete devices. The number of publications — 4. pogodina-ts@mail.ru; 9/9, Obraztsova St., 127994, Moscow, Russia; office phone: +7(977)404-7953.

YU. ZELENKOV

OPTIMIZATION OF THE REGRESSION ENSEMBLE SIZE

Zelenkov Yu. Optimization of the Regression Ensemble Size.

Abstract. Ensemble learning algorithms such as bagging often generate unnecessarily large models, which consume extra computational resources and may degrade the generalization ability. Pruning can potentially reduce ensemble size as well as improve performance; however, researchers have previously focused more on pruning classifiers rather than regressors. This is because, in general, ensemble pruning is based on two metrics: diversity and accuracy. Many diversity metrics are known for problems dealing with a finite set of classes defined by discrete labels. Therefore, most of the work on ensemble pruning is focused on such problems: classification, clustering, and feature selection. For the regression problem, it is much more difficult to introduce a diversity metric. In fact, the only such metric known to date is a correlation matrix based on regressor predictions. This study seeks to address this gap. First, we introduce the mathematical condition that allows checking whether the regression ensemble includes redundant estimators, i.e., estimators, whose removal improves the ensemble performance. Developing this approach, we propose a new ambiguity-based pruning (AP) algorithm that bases on error-ambiguity decomposition formulated for a regression problem. To check the quality of AP, we compare it with the two methods that directly minimize the error by sequentially including and excluding regressors, as well as with the state-of-art Ordered Aggregation algorithm. Experimental studies confirm that the proposed approach allows reducing the size of the regression ensemble with simultaneous improvement in its performance and surpasses all compared methods.

Keywords: ensemble pruning, regression, ensemble learning, error-ambiguity decomposition, diversity of regressors.

1. Introduction. Ensemble learning is a method that combines several models, which are obtained by applying a learning process to a given problem. The main idea of this approach is that models view the problem from different points. Therefore, their combination improves robustness and accuracy either in classification or regression. However, the existing ensemble learning algorithms often generate unnecessarily large ensembles, which consume extra computational resources and may degrade the generalization ability [1]. There are theoretical and empirical publications that have shown that small ensembles can be better than large ensembles [1, 2].

The ensemble learning process can be described as the overproduce-and-choose approach [3]. The overproduction phase is aimed to produce a large set $\mathcal{F}_0 = [f_i, i = 1 \dots M_0]$ of candidate base models f_i . The choice phase is intended to select the subset of models $\mathcal{F} \subseteq \mathcal{F}_0$ that can be combined to achieve optimal performance.

In general, there are two ways to realize the choice phase. The first is a sequential selection when the algorithm starts from an empty set and

sequentially adds models according to some metric. Often the selection is combined with model generation. The second is pruning, in that case, the ensemble includes all candidate models, and the goal is to choose their optimal subset according to some metric.

Both the selection and pruning have the potential advantage of reducing ensemble size, and improving performance [4]. However, the selection and pruning of classifiers, rather than regressors, has previously received more attention from researchers [5, 6]. Some of these methods have been adapted to the regression task [7], but there is a lack of theoretical and empirical works dedicated exclusively to the regression problem.

There are theories considering the specifics of regression, in particular, these are the error-ambiguity decomposition [8, 9], which can be applied to develop a pruning algorithm. Here we present an ambiguity-based pruning algorithm that sequentially removes regressors with the worst generalization ability. We compare the performance of this algorithm with a state-of-the-art Ordered Aggregation [10] method also as with two algorithms based on direct optimization of the quality metric.

The rest of the paper organizes as follows. After the literature review, we introduce the mathematical condition that allows checking whether the regression ensemble includes redundant estimators, i.e., estimators, whose removal improves the ensemble performance. Next, on the basis of this approach, we propose the Ambiguity-based Pruning (AP) algorithm. In the last part of the paper, we present the results of experiments on real datasets that confirm that the proposed approach outperforms known methods in terms of accuracy and model complexity.

2. Literature Review. We consider the typical regression problem, and for a clear presentation, we establish the notation that will be used below. Take X to be the vector space of all possible inputs, and $Y \in \mathbb{R}$ to be the vector space of all possible outputs and there exists some unknown probability distribution over the product space $X \times Y$. The training set $D_{train} = \{(\mathbf{x}_1, y_1), (\mathbf{x}_2, y_2), \dots, (\mathbf{x}_N, y_N)\}$ is made up of N samples from this probability distribution. Every $\mathbf{x}_i$ is an input vector from the training data, and y_i is the corresponding output. The goal is to induce on the basis of the training set a function $f: X \rightarrow Y$ that approximates an unknown true function such as $f(\mathbf{x}) \sim y$. The quality of the approximation is given by the generalization error, which usually is a mean squared error:

$$MSE(f) = \mathbb{E}[(f(\mathbf{x}) - y)^2].$$

Because it is not possible to determine this true error of a model f , the error is estimated on a different set of data D_{test} , containing K samples:

$$MSE(f) \approx \frac{1}{K} \sum_{i=1}^K [(f(\mathbf{x}_i) - y_i)]^2.$$

We will consider a regression ensemble, i.e., the combination of a few models since this should improve robustness and accuracy. In regression problems, ensemble integration most often is performed using a linear combination of the base models [3, 6].

$$f_E(\mathbf{x}) = \sum_{i=1}^M [w_i(\mathbf{x}) * f_i(\mathbf{x})],$$

where $w_i(\mathbf{x})$'s are the weighting functions, and M is a number of models $f_i(\mathbf{x}) \in \mathcal{F}$ in the ensemble (models f_i are often also referred to as estimators, predictors, regressors or learners). It follows from this that the problem of the models' selection is closely related to choosing the optimal weights. From now on, we will use notation f instead of $f(\mathbf{x})$ and w instead of $w(\mathbf{x})$ for simplicity.

Study [8] proposed the ambiguity decomposition of ensemble error that separates the weighted average error of the individual regressors and variability among their estimations at an arbitrary single data point:

$$(f_E - y)^2 = \sum_{i=1}^M w_i (f_i - y)^2 - \sum_{i=1}^M w_i (f_i - f_E)^2. \quad (1)$$

The first term $w_i (f_i - y)^2$ is the weighted error of the i -th ensemble member. The second, $w_i (f_i - f_E)^2$, is the ambiguity term, measuring the amount of variability among the ensemble member answers for this pattern. This equation explains why the quadratic error of the ensemble is less than or equal to the average quadratic error of the component estimators. Note that this decomposition is valid only for convex ensembles [9], i.e., when $f_E = \sum_{i=1}^M w_i f_i$ and $\sum_{i=1}^M w_i = 1, w_i \geq 0$.

An essential assumption of ensemble learning is that the base models should be sensitive to variations in the training set, so Decision Trees (DT) and Neural Networks (NN) usually are used.

The most popular ensemble learning algorithms for regression are Bagging, Random Forest, Negative Correlation and Gradient Boosting. The *Bagging* algorithm [11] employs bootstrap sampling to generate many training sets from the original training set and then trains a model for each of those training sets. The component predictions are combined via simple averaging for regression tasks. Bagging can be used both with DT and NN.

The *Random Forest* [12] algorithm is similar to Bagging in that they both resample the data. However, Random Forest is based exclusively on DT, when it performs splitting, a random sample of the features is also selected. In *Negative Correlation* learning [13], all the individual estimators in the ensemble are trained simultaneously and interactively through the correlation penalty terms in their error functions. This approach is used exclusively with NN since in that case there is a possibility to include a penalty in the formula for weights tuning in the backpropagation method. The *Gradient Boosting* algorithm [14] on each iteration computes pseudo-residuals and trains a new model using them as a target. Thus, each new estimator attempts to correct the error of its predecessors. The weight of each member is found in the process of a linear search.

The stochastic nature of Bagging and Random Forest leads to ensembles that can be significantly improved by pruning. Many authors used this fact in their research [7]. The family of Boosting methods (including AdaBoost) produces more balanced ensembles in general. However, some researchers report on successful applications of pruning especially in case of the classification problem solved with the AdaBoost algorithm [1, 2].

Different authors proposed different classification schemes of pruning algorithms. In study [6] the authors classify them as partitioning-based and as search-based. Partitioning-based methods divide the pool of models into subgroups. Then, for each subgroup, one or more models are selected using a given selection criterion. Search-based algorithms, in turn, are divided into (1) exponential that search the complete space of models, (2) randomized that use stochastic methods, such as evolutionary algorithms, and (3) sequential that search for a subset of the original pool by iteratively adding or removing models.

In study [1] the authors split pruning algorithms into two categories, (1) selection-based that do not weight each model by a weighting coefficient and either select or reject the learner, and (2) weight-based algorithms that improve the generalization performance of the ensemble by tuning the weight on each ensemble member.

In paper [6] the authors reviewed regression ensemble pruning approaches published before 2008, here we will consider some recent publications on the basis of the approach to the classification proposed in [1]. First, we review some selection-based algorithms.

Study [5] reviewed a family of pruning methods based on modifying the order of estimators in a Bagging ensemble. This order in the original Bagging algorithm is unspecified, and the error of the ensemble generally exhibits a monotonic decrease as a function of the number of estimators.

According to pruning strategies based on ordered aggregation, from the subensemble $\mathcal{F}_{L-1}$ of size $L-1$, the subensemble $\mathcal{F}_L$ of size L is constructed by incorporating a single estimator selected from the set $\mathcal{F}_0 \setminus \mathcal{F}_{L-1}$, which contains the estimators from the original ensemble not included in $\mathcal{F}_{L-1}$. This estimator is identified using a rule that attempts to optimize the performance of the augmented ensemble $\mathcal{F}_L$. The ordered ensemble that includes $L < M$ estimators generally exhibits the error that is below the error of the complete bagging ensemble.

Assuming that the generalization error of the regression ensemble can be expressed as:

$$E = \frac{1}{M^2} \sum_{i=1}^M \sum_{j=1}^M C_{ij}, C_{ij} = \frac{1}{N} \sum_{n=1}^N [(f_i(\mathbf{x}_n) - y_n)(f_j(\mathbf{x}_n) - y_n)], \quad (2)$$

where the correlation matrix C is estimated over a training dataset. In paper [10] the authors proposed Ordered Aggregation (OA) algorithm. The algorithm starts with an empty ensemble and then selects at each iteration the regressor that, when incorporated, reduces the training error (2) of the new ensemble the most.

As for the disadvantages of this method, we can state the following. First, this approach based on the assumption that minimizing training error leads to the minimization of generalization error but in fact, this usually leads to overfitting. Second, time complexity grows exponentially. Third, the number of ensemble members is an external factor; there is no internal stopping criterion.

Later the same authors [7] proposed to use Semidefinite Programming (SDP) introduced in [15] for the classification task. In that case, it is necessary to find a sub-ensemble for which the sum of the elements in the corresponding sub-matrix of C is as low as possible. Note, that it is also NP-hard computational problem.

Authors reported that the minimum of test error obtained either with OA and SDP-pruning is generally below the asymptotic error of the complete bagging ensemble, and pruned ensembles obtained by retaining only 20% of the original bagging ensemble have the best overall performance. The main conclusion of [7] is that the key to improvement in generalization performance is the selection of subsets of regressors whose bias is low and whose correlations are small or negative.

Study [4] extended the OA approach using dynamic ensemble selection technique. Their algorithm consists of two steps. First, the base regressors are trained on bootstrap samples of the training dataset, and the regressor order is found for every instance in the training set. In the second stage, the regressor order that is associated with the training instance closest

to the test instance is retrieved. To find the closest training instance the k-Nearest Neighbors method is used. Empirical testing on several data sets showed that in most cases this approach outperforms original OA-pruning.

There are also examples of the use of evolutionary methods for selecting the optimal set of estimators. For example, a genetic algorithm that searches in the space of candidate base models $\mathcal{F}_0$. In that case, binary fixed-length strings $\{0,1\}^{M_0}$ where $M_0 = |\mathcal{F}_0|$ represent ensembles that form the evolving population. The pruned ensemble includes only estimators that have a value of 1 in the corresponding position of the coding string. In study [16] the authors used such an approach for pruning classification ensembles obtained by the AdaBoost algorithm.

In paper [17] the authors generalized this approach as a multi-objective optimization problem; they proposed simultaneously to minimize two variables – the generalization error of the ensemble and its size.

Some authors claim that weight-based pruning is a more general approach than selection-based [1]. According to [2], the optimal weights of the regression ensemble can be obtained as:

$$w_i = \frac{\sum_{j=1}^M (C^{-1})_{ij}}{\sum_{k=1}^M \sum_{j=1}^M (C^{-1})_{kj}}.$$

However, in real-world applications, some estimators can be quite similar, which makes the correlation matrix C (2) ill-conditioned [2]. The second problem of this formulation is that the optimal combination of weights is computed from the training set, which can lead to overfitting [1].

In paper [1] the authors presented the ensemble pruning algorithm by expectation propagation that approximates the posterior estimation of the weight vector. It produces a «sparse» combination of weights, most of which are zeros. For experiments with the regression, authors used Bagging and Random Forest algorithms with 100 Decision Trees, they reported that the size of the pruned ensemble was reduced, on average, approximately ten times.

In study [18] the authors explored two other weight-based pruning techniques: one based on a cocktail ensemble (CE) algorithm [19] and the second on stacking generalization [20].

CE that was designed for generating the ensemble of ensembles is the following. Since the combination of multiple regressors is an NP-hard problem, the authors of [19] proposed to use the pair-wise combinations of estimators. Through a linear combination of models f_1 and f_2 , a new ensemble is formed:

$$f_E(\mathbf{x}) = w_1 f_1(\mathbf{x}) + (1 - w_1) f_2(\mathbf{x}) \text{ w.r.t. } w_1 \in [0,1].$$

Following the error-ambiguity decomposition [8], in study [19] the authors proved that given E_1 and E_2 as generalization errors of f_1 and f_2 respectively, the optimal weight of f_1 is $w_1 = (E_2 - E_1)/2\Delta + 0.5$, where $\Delta = \mathbb{E}[(f_1(\mathbf{x}) - f_2(\mathbf{x}))^2]$ is the squared output difference of the two ensembles. E_1, E_2 and Δ can be estimated from training data. CE algorithm starts from the selection base model with the lowest error; then each subsequently selected regressor is the one which reduces the combined estimated error the most.

Stacking is a popular ensemble learning strategy, where the weights of the base models are the regression coefficients of the meta-level regressor [20]. The authors of [18] used the linear regression as a meta-level; the final ensemble consists only of models with greater than zero weight. According to the experimental results, in some cases, the stacked-based approach provides results with less error than standard Bagging and CE and generates the shortest ensembles [18].

Other authors use various randomized algorithms to search for the optimal combination of weights in the ensemble. Study [2] utilized the genetic algorithm with a floating coding scheme to represent weights. Authors reported that their approach outperforms on regression task original Bagging and AdaBoost.R2 algorithms with NN as the base model, an average number of networks after pruning was less than 4.

All approaches discussed above are based on the set of models obtained by ensemble methods that produce estimators using a single learning algorithm. An alternative approach is the usage of heterogeneous models. In this case, a large number of models is generated using various learning algorithms, and the goal is to select those, which produce better generalization. This task also can be formulated as a pruning problem, there is an ensemble that includes all models, and it is necessary to remove non-effective ones. Most of the works in this scientific flow are dedicated to the selection problem [21].

To conclude this review we can note that, in general, ensemble pruning is based on two metrics: diversity and accuracy [22, 23]. Formally the problem is to find a subset of candidate estimators $\{f_k\} \subset \mathcal{F}_0$ with both high accuracy (i.e., lower loss $L(\{f_k\})$) and maximum diversity $D(\{f_k\})$:

$$f_M = \operatorname{argmin}_k L(\{f_k\}) \text{ w.r.t. } D(\{f_k\}) \rightarrow \max.$$

Many diversity metrics are known for problems dealing with a finite set of classes defined by discrete labels [22, 24]. Therefore, much work on

ensemble pruning is focused on tasks of classification [25, 26], clustering [27, 28], and selection of an optimal set of features [29, 30].

For the regression problem, it is much more difficult to introduce a diversity metric. In fact, the only such metric known to date is the correlation matrix (2) proposed in [10]. This approach is used in recent papers: [31], but in general the number of studies on regression ensemble pruning published after 2010 is extremely small.

We propose a new approach to regressor selection based on error-ambiguity decomposition, which was introduced in [8], studied in detail in [9], and generalized to all supervised learning problems in [32]. Many authors have pointed out that this decomposition potentially opens ways of optimizing ensembles [33, 34]. Other applications of the error-ambiguity decomposition include, for example, adaptive sampling [35, 36].

3. Ensemble pruning on the basis of Error-Ambiguity decomposition. Study [37] proposed the Reduce-Error pruning algorithm. The first learner incorporated into the ensemble is the one with the lowest error, as estimated on the selection set. The remaining estimators are then sequentially incorporated in the ensemble, one at a time, in such a way that the error of the partial subensemble, estimated on the selection set, is as low as possible. This method was proposed for the classification problem, but we can adapt it to the regression task. So, the regressor f_M that should be incorporated into the subensemble f_E^{M-1} to construct f_E^M is selected by the rule:

$$f_M = \operatorname{argmin}_k \sum_{(x,y) \in D} \text{MSE}(f_E^{M-1}(x) \cup f_k(x), y), k \in \mathcal{F}_0 \setminus f_E^{M-1}. \quad (3)$$

Instead of the selection, we can also use a sequential reduction process, i.e., starting from an ensemble that includes all regressors, at each step, one of them whose removal reduces error maximally is deleted. This regressor is selected by the rule:

$$f_M = \operatorname{argmin}_k \sum_{(x,y) \in D} \text{MSE}(f_E^M(x) \setminus f_k(x), y), k \in f_E^M. \quad (4)$$

In what follows, we will refer to the algorithm presented by Equation (3) as a Direct Reduce Error (DR), and to the algorithm Equation (4) as a Reverse Reduce Error (RR).

Consider the conditions under which a regressor selected according to Equation (4) must meet to create the maximum effect. Averaging the error – ambiguity decomposition (1) for all N samples in the dataset, we obtain the following formula:

$$\begin{aligned}\mathbb{E}[(f_E^M - y)^2] &= \sum_{i=1}^M w_i^M \mathbb{E}[(f_i - y)^2] - \sum_{i=1}^M w_i^M \mathbb{E}[(f_i - f_E^M)^2] \\ &= \sum_{i=1}^M w_i^M (\mathbb{E}[(f_i - y)^2] - \mathbb{E}[(f_i - f_E^M)^2]),\end{aligned}\quad (5)$$

where w_i^M is the optimal weight of the i -th estimator in the ensemble of size M .

The first term $E_i = \mathbb{E}[(f_i - y)^2]$ is the mean squared error of regressors in the ensemble, and the second one $A_i^M = \mathbb{E}[(f_i - f_E^M)^2]$ is their ambiguity (note that its value depends on all ensemble members). Thus, to ensure $MSE(f_E^M) = 0$, it is enough to fulfill the conditions:

$$\forall i: E_i = A_i^M, i = 1, \dots, M. \quad (6)$$

This can be explained in terms of diversity. Formula (6) means that the i -th estimator, which satisfies this condition, differs from other members of the ensemble (i.e., its outputs are correlated with outputs of other members) in such a way that its error is compensated by this diversity.

Using Equation (5), we can find the condition when the removal of one regressor will not lead to a growth of the total ensemble error, i.e.:

$$\mathbb{E}[(f_E^M - y)^2] - \mathbb{E}[(f_E^{M-1} - y)^2] \geq 0.$$

Let the regressor that is to be removed have index M . Since the ensemble can be presented as:

$$f_E^M = \sum_{i=1}^M w_i^M f_i = \sum_{i=1}^{M-1} w_i^M f_i + w_M^M f_M,$$

after some algebra, we get the condition that the removed regressor must match in order not to reduce the overall performance:

$$w_M^M (E_M - A_M^M) - \left[\sum_{i=1}^{M-1} w_i^{M-1} (E_i - A_i^{M-1}) - \sum_{i=1}^{M-1} w_i^M (E_i - A_i^M) \right] \geq 0. \quad (7)$$

Accordingly, this regressor must have a maximal positive value defined by Equation (7). An essential consequence of this formula is that it allows us to evaluate the potential of the current ensemble for pruning. If

there are no regressors in the ensemble for which the value of Equation (7) is positive, then the pruning of such an ensemble is impossible.

Formula (7) is derived from Equation (4), but Equation (4) presents the algorithm that is more effective computationally since it requires just the computation of the MSE of each regressor. It can be done once before the pruning cycle since the errors of individual estimators do not depend on ensemble composition. The procedure presented by Equation (7) requires the computation of ambiguity value, which depends on the current ensemble composition and therefore changes on each step.

The term in square brackets in Equation (7) represents the difference between the errors of the pruned subensemble and the original ensemble, from which the prediction of deleted regressor is excluded. This value determines the threshold that the value of $w_M^M(E_M - A_M^M)$ of deleted regressor must exceed. But, since we determine these values based on the training set, we are not sure that this takes into account all the information about the actual distribution of data. Thus, the assumption is reasonable that we can omit the term in square brackets and present the rule for choosing a regressor to be deleted as:

$$\max_i [w_i^M (E_i - A_i^M) \geq 0], i = 1, \dots, M. \quad (8)$$

On the one hand, this simplification reduces computational complexity since we exclude the A_i^{M-1} term. On the other hand, it could facilitate the selection of subensemble with greater generalization since it introduces stochasticity in the selection process. In most cases, the regressor that is selected for removal from the ensemble by rule Equation (8) will coincide with the one chosen by Equation (4). Discrepancies can only be observed when the ensemble is already well optimized, and differences defined by Equation (7) tend to zero.

Since for the ensembles based on averaging like Bagging and Random Forest $w_i^M = 1/M$, this rule can be simplified as:

$$\max_i [(E_i - A_i^M) \geq 0], i = 1, \dots, M.$$

We can show that this is equivalent to choosing a regressor with a maximum distance to the line defined by Equation (4) on the half-plane $E_i \geq A_i^M$. The distance from the point $M(x_M, y_M)$ to the line determined by the equation $Ax + By + C = 0$ is defined as $d = |Ax_M + By_M + C| / \sqrt{A^2 + B^2}$. Substituting into this formula the

coordinate values of the i th regressor (E_i , A_i^M) and the coefficients of the equation of the straight line $E - A^M = 0$, we obtain $d_i = |E_i - A_i^M|/\sqrt{2}$.

Figure 1 presents the example of the averaging ensemble of 4 regressors trained on the airfoil dataset (its parameters will be given below). The left chart shows the values of E_i and A_i^M of these regressors; a solid line presents Equation (6). The right graph shows the corresponding values of the two terms of Equation (7). As it follows from the data presented on the right graph best candidate for removing is estimator number 2 according to Equation (7) and estimator number 0 according to Equation (8). Corresponding points are marked by red and green circles on the left chart, respectively. The exclusion of estimator 2 will lead to better performance on the training set, but, as said above, it can also lead to a loss of the generalization since the pruned model will be overfitted to training data. The ambiguity of the remaining ensemble members changes when any estimator is removed; therefore, the use of equations (7) and (8) can lead to different final structures. In the next sections, we will present empirical data, which confirms that rule Equation (8) allows us to generate more effective ensembles.

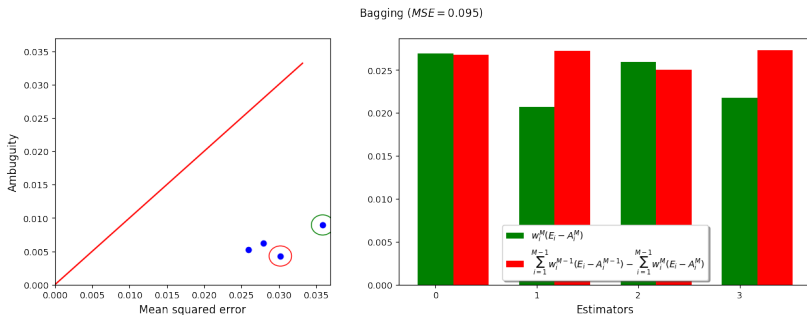


Fig. 1. Bagging ensemble of 4 regressors

The next issue that should be discussed is the data for the creation of an initial pool of regressors and their subsequent selection. Some authors use separated training and selection datasets. Researchers often use this approach to avoid overfitting, but it also can lead to degradation of the generalizing ability of the model since it reduces data available for the model generation. On the other hand, paper [10] reported that the size of the subensembles minimizing the error on the training data tends to be smaller than the optimal subensembles when the error is estimated on an independent selection set. Besides, the selection of the optimal subensemble is an NP-hard problem and not generally feasible in practice.

Our experiments in the preparation of this article confirm that the separation of data into training and selection sets does not improve overall performance. On the contrary, this often leads to a decrease in productivity since candidate models obtained on reduced training datasets do not have sufficient predictive ability.

We compared three different approaches. First, we divided all the data available for training D_{train} into independent sets: the first one for generating models and second for selecting D_{sel} . Table 1 shows the results for cases when the size of the D_{sel} includes 50% and 33% of the data available. In the third experiment, we performed the generation and selection on the same set D_{train} , which included all available data.

Table 1. The pruning model performance for various combinations of training and selection data

Dataset splitting	Bagging	AP	RR	OA	DR
$ D_{sel} = 0.5 D_{train} $	0.273	0.235	0.243	0.259	0.243
$ D_{sel} = 0.33 D_{train} $	0.250	0.228	0.226	0.234	0.219
Training and selection on D_{train}	0.223	0.189	0.189	0.202	0.197

Table 1 shows the best average MSE obtained with 10-fold cross-validation on the airfoil dataset. We used a single-layer artificial neural network (NN) as a base estimator. At each iteration, we first optimized the number of neurons in the NN, and then we generated the initial pool of 100 estimates using the Bagging algorithm and, finally, applied reduction and selection algorithms. We tested two pruning algorithms Reverse Reduce Error (RR) (4) and Ambiguity-based Pruning (AP) (8) and two selection methods Direct Reduced Error (DR) (3) and Ordered Aggregation (OA) [10].

As we can see, the critical factor that defines the pruned model performance is the performance of the initial ensemble that improves when the amount of data available for training grows. Therefore, here and below, we propose to use the single dataset for the generating of candidate models and the selection of optimal subsets.

Summing up the above, the proposed algorithm, which we call Ambiguity-Based Pruning (AP), is presented in Table 2. Since instead of using the exact expression (7) to select the regressor for reduction, we use an approximate value (8), we must stop the algorithm when the current MSE value begins to rise to control convergence and prevent the loss of generalization (lines 9-13).

Table 2. Ambiguity Based Pruning Algorithm

INPUT	Training dataset $D_{train} = \{(x_i, y_i)\}_{i=1}^N$, pool f_E of regressors fitted on D_{train} .
1	Initialize $e_E = MSE(f_E)$
2	Repeat
3	$ensPruned$ is <i>False</i>
4	For each f_i in f_E do
5	$e_i = w_i^M (E_i - A_i^M)$
6	End for
7	$j = \text{argmax}(e)$
8	$f_{new} = f_E \setminus f_j$
9	If $MSE(f_{new}) \leq e_E$
10	$e_E = MSE(f_{new})$
11	$f_E = f_{new}$
12	$ensPruned$ is <i>True</i>
13	End if
14	Until $ensPruned$ is <i>True</i>
RETURN	f_E

4. Empirical Analysis and Evaluation. In this section, we assess the generalization performance of Ambiguity-Based Pruning (AP) and compare this technique with other algorithms. As a state-of-art method that sets the baseline for the performance of pruning algorithms, we use Ordered Aggregation (OA). In study [7] the authors conducted an experimental comparison of OA with conventional ensemble methods such as Bagging, AdaBoost.R2, and Negative Correlation, and showed that it surpasses them in performance. Also, according to [7], OA outperforms other known pruning methods such as hybrid ensembles [38], GASEN [2], regularized stacked generalization [39, 40], and performs slightly better than SDP-pruning [7], although the difference between average ranks of OA and SDP is not statistically significant.

In addition to OA, we also included in the set of methods for comparing two algorithms discussed in Section 3, Direct Reduce Error (DR), and Reverse Reduce Error (RR).

It should be noted that the authors of [7] defined the number of estimators in the OA-pruned ensemble as an external parameter. In our experiments, we used a slightly modified OA algorithm: the selection of regressors continues while the error in the training data does not increase (lines 9-13 in Table 2). Such an approach helps to identify the optimal size of the final subensemble. A similar condition was applied to other algorithms tested (RR and DR).

The experiments are carried out on 20 regression problems from the UCI-Repository and other sources. They include real-world problems from different fields of application: industry, biology, urban management, etc. All datasets were scaled to center around zero and have unit variance. Table 3 displays the number of instances (N), the number of attributes (K) and the source of the different datasets considered.

Table 3. Datasets used in the experiments

Dataset	Description	N / K	Source
abalone	Predicting the age of abalone	4177/8	KEEL
airfoil	Aerodynamic and acoustic tests of airfoil	1503/5	UCI
bank8FM	Simulation of how customers choose bank	8192/8	OML
bike	Count of rental bikes	17379/12	UCI
california	Median house value of the block groups	20640/8	KEEL
CASP	Properties of protein tertiary structure	45730/9	UCI
CCPP	Combined Cycle Power Plant	9568/4	UCI
compactiv	Computer systems activity measures	8192/21	KEEL
concrete	Predicting the concrete compressive	1030/8	KEEL
egrid	Electrical grid stability simulated data	10000/12	UCI
elevators	Action taken on the elevators of the aircraft	16599/18	KEEL
facebook	Facebook comment volume	40949/53	UCI
forestFires	Burned area of forest fires in Portugal	517/12	KEEL
house	Median house price in the regions of USA	22784/16	KEEL
kin8nm	Forward kinematics of an 8-link robot arm	8192/8	OML
laser	Far-Infrared-Laser in a chaotic state	993/4	KEEL
stock	Daily stock prices for aerospace companies	950/9	KEEL
supercond	Superconductors and their relevant features	21263/81	UCI
treasury	Economic data of USA	1049/15	KEEL
wankara	The weather information of Ankara	1609/9	KEEL
Data sources: – KEEL – KEEL dataset: https://sci2s.ugr.es/keel/datasets.php – OML – OpenML: https://www.openml.org/search?type=data – UCI – UC Irvine Machine Learning Repository: http://archive.ics.uci.edu/ml/			

The experimental protocol is the following. We used 10-fold cross-validation to estimate the mean squared error. The experiments involve building a bagging ensemble of $M = 100$ predictors. Following the work [7],

we use the feed-forward neural network (NN) with a single hidden layer of sigmoidal neurons and a linear unit in the output layer as a base learner. The networks are trained during 1000 epochs using the quasi-Newton optimization method BFGS. Before the generation of the bagging ensemble, the optimal number of hidden units in NN was explored using Bayesian optimization and a separate 5-fold cross-validation procedure. Once the best architecture of NN is found, a bagging ensemble is generated using neural networks with these hyper-parameters. Next, the bagging ensemble is pruned using various pruned algorithms. Finally, we computed the mean value of MSE on all testing folds for all compared algorithms, also as bagging ensemble, to use the last as a benchmark for performance comparison. All computations were performed using the Python programming language and scikit-learn and scikit-optimize libraries.

Table 4 shows the average mean squared error and its standard deviation estimated by 10-fold cross-validation for each dataset and each prediction method.

Table 4. Average mean squared error normalized by the corresponding scaling factor

Dataset	Scaling factor	Bagging	AP	DR	RR	OA
abalone	10^{-1}	4.309 (0.272)	4.318 (0.272)	4.315 (0.272)	4.324 (0.273)	4.319 (0.272)
airfoil	10^{-1}	2.225 (0.112)	1.890 (0.099)	1.973(0.104)	1.887 (0.094)	2.023 (0.109)
bank8FM	10^{-2}	3.470 (0.002)	3.443 (0.002)	3.440 (0.002)	3.442 (0.002)	3.467 (0.002)
bike	10^{-1}	1.448 (0.091)	1.154 (0.065)	1.163 (0.066)	1.155 (0.065)	1.209 (0.071)
cadata	10^{-1}	2.775 (0.096)	2.685 (0.088)	2.684 (0.089)	2.683 (0.088)	2.756 (0.095)
CASP	10^{-1}	5.609 (0.012)	5.522 (0.013)	5.521 (0.013)	5.521 (0.012)	5.601 (0.012)
CCPP	10^{-2}	5.889 (0.005)	5.833 (0.005)	5.833 (0.005)	5.833 (0.005)	5.889 (0.005)
compactiv	10^{-2}	1.738 (0.001)	1.674 (0.001)	1.672 (0.001)	1.674 (0.001)	1.707 (0.001)
concrete	10^{-2}	8.629 (0.020)	8.206 (0.019)	8.139 (0.019)	8.140 (0.019)	8.333 (0.021)
egrid	10^{-2}	4.773 (0.003)	4.523 (0.003)	4.528 (0.003)	4.529 (0.003)	4.623 (0.003)
elevators	10^{-2}	8.829 (0.014)	8.575 (0.013)	8.575 (0.013)	8.575 (0.013)	8.754 (0.014)
facebook	10^{-1}	4.599 (0.188)	4.621 (0.181)	4.672 (0.177)	4.659 (0.180)	4.498 (0.181)
forestFires	10^0	1.218 (1.664)	1.281 (1.585)	2.722 (2.429)	1.408 (1.553)	1.533 (1.614)
house	10^{-1}	4.207 (0.057)	4.138 (0.055)	4.140 (0.056)	4.140 (0.056)	4.193 (0.057)
kin8nm	10^{-2}	9.233 (0.008)	8.648 (0.008)	8.652 (0.008)	8.646 (0.008)	8.816 (0.008)
laser	10^{-2}	1.362 (0.015)	1.360 (0.015)	1.411 (0.016)	1.430 (0.017)	1.369 (0.015)
stock	10^{-1}	1.486 (0.112)	1.498 (0.126)	1.541 (0.128)	1.514 (0.126)	1.535 (0.121)
supercond	10^{-1}	1.759 (0.128)	1.737 (0.123)	1.741 (0.123)	1.735 (0.123)	1.754 (0.127)
treasury	10^{-3}	4.027 (0.002)	3.818 (0.002)	3.822 (0.002)	3.804 (0.002)	4.003 (0.002)
wankara	10^{-3}	6.099 (0.001)	6.022 (0.001)	6.043 (0.001)	6.081 (0.001)	6.152 (0.001)

The figures displayed are scaled by a factor shown in the first column of the table. To check whether the observed improvements in error are statistically significant, a paired Wilcoxon test was performed on cross-validation data for each dataset. Error values that are significantly better than bagging at an α -value of 0.05 are highlighted in boldface. Error values that are significantly better than OA at an α -value of 0.05 are underlined. As it follows from Table 4, there were no algorithms that perform analysis statistically worse than bagging and OA on the datasets selected for the experiment.

Following the framework proposed by J. Demšar [41], we also conducted the Friedman test to compare the overall performance of different methods in the collection. The results obtained ($F_F = 59.14$, the corresponding p -value is $2E-11$, and the critical value of χ^2 distribution is 11.07) confirm that the hypothesis of the equivalent performance of all algorithms should be rejected at $\alpha = 0.05$. If the null hypothesis is rejected, we can proceed with a post-hoc Nemenyi test. Figure 2 displays the results of this test for $\alpha = 0.05$. The differences in performance between algorithms whose average ranks are further than a critical distance (CD) are statistically significant. The obtained value of the CD is 1.686. In Figure 2, algorithms whose differences in performance are not statistically significant are connected with a solid line.

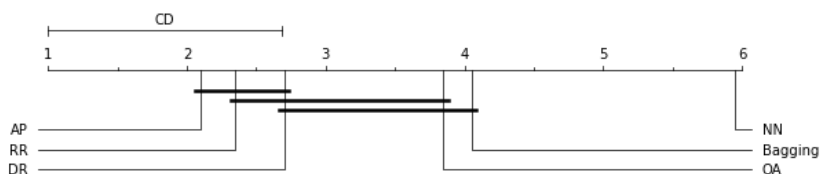


Fig. 2. Comparison of algorithms' performance against each other with the Nemenyi test. Groups of methods that are not significantly different (at $\alpha = 0.05$) are connected with a solid black line

As follows from the data presented in Figure 2, four groups of models are distinguished, and models inside one group have statistically comparable results. The first group includes algorithms of pruning AP, RR, and DR, while AP has a slight advantage within this group. The second group consists of the algorithms RR, DR, and OA, the OA; this group showed the worst results. The third group comprises the bagging algorithm, as well as methods for pruning the DR and OA. The fourth group is represented only by a single-layer feed-forward neural network (NN).

The following conclusions can be drawn from the presented results. First, pruning methods can improve the performance of the initial set of regressors generated using the bagging algorithm. Data in Table 4 confirms this conclusion generally, but we should note that there are a few exceptions, namely results for datasets: abalone, forestFires, and stock. Second, methods based on the sequential exclusion of estimators from the initial ensemble (AP and RR) are superior to sequential inclusion algorithms (DR and OA). One possible explanation is that in the case of the exclusion, the algorithm has more information about the interaction of all estimators in the ensemble. In the case of the AP algorithm, this information is presented explicitly in the ambiguity term (1). When the algorithm implements an inclusion scenario, only information about the pair interaction of estimators is available explicitly given by correlation matrix Equation (2) in the OA algorithm.

Another critical point associated with the quality of pruning algorithms is the size of the ensemble obtained after the algorithm application. Table 5 presents these data.

Table 5. Average number of regressors in the pruned sub ensemble

Data	NN	AP	DR	RR	OA
abalone	10.5 (3.4)	12.6 (2.1)	11.2 (2.5)	11.9 (1.9)	98.2 (5.7)
airfoil	10.5 (3.5)	9.8 (3.0)	7.6 (1.8)	8.8 (2.3)	23.9 (5.2)
bank8FM	17.1 (1.8)	9.8 (2.0)	9.7 (1.9)	9.6 (2.1)	94.8 (4.0)
bike	14.9 (3.1)	11.0 (6.6)	7.2 (1.9)	7.7 (1.6)	22.3 (3.5)
cadata	13.2 (5.3)	8.8 (2.0)	6.5 (1.8)	8.0 (2.4)	65.2 (10.4)
CASP	19.1 (0.9)	10.8 (2.1)	10.4 (1.9)	10.3 (1.9)	89.7 (4.9)
CCPP	13.5 (2.0)	5.1 (1.5)	4.2 (1.5)	5.1 (1.5)	99.5 (1.0)
compactiv	11.2 (4.9)	9.9 (2.9)	8.5 (2.6)	8.5 (2.1)	53.4 (5.2)
concrete	16.1 (1.7)	14.7 (3.2)	13.5 (2.5)	13.9 (2.8)	27.2 (5.8)
egrid	18.9 (1.0)	14.4 (2.8)	11.1 (1.3)	12.4 (1.6)	34.2 (5.8)
elevators	7.9 (2.6)	11.8 (3.3)	6.7 (2.6)	9.7 (1.7)	70.5 (8.2)
facebook	11.9 (2.4)	8.7 (2.1)	8.3 (1.6)	8.1 (1.4)	31.9 (4.8)
forestFires	9.1 (3.1)	43.1 (10.1)	13.1 (9.8)	23.7 (3.8)	7.9 (3.3)
house	16.2 (2.3)	14.6 (3.0)	12.3 (1.9)	13.4 (2.9)	60.2 (7.4)
kin8nm	18.0 (0.9)	12.4 (2.4)	9.3 (1.3)	11.2 (1.3)	39.3 (4.2)
laser	11.1 (3.9)	13.8 (5.5)	7.8 (2.9)	8.1 (1.5)	24.0 (5.8)
stock	13.1 (5.0)	12.0 (3.4)	9.7 (3.4)	10.9 (3.1)	28.4 (3.8)
supercond	9.2 (4.4)	9.5 (2.0)	8.9 (1.7)	9.2 (1.8)	55.4 (8.6)
treasury	11.8 (4.4)	7.3 (2.2)	5.5 (1.6)	6.7 (1.6)	45.9 (7.5)
wankara	9.4 (3.7)	13.6 (3.3)	9.9 (2.2)	11.0 (1.8)	66.9 (10.5)

The first column of Table 5 (NN) contains the mean and the standard deviation of the number of neurons obtained during the optimization of NN at each iteration in the process of 10-fold cross-validation. Other columns show the mean and standard deviation of the number of base estimators (NN) obtained from the 10-fold cross-validation for each algorithm. Better value in each line is highlighted with bold. As we can see, the best value is achieved by methods that directly minimize the prediction error (DR and RR), AP creates somewhat larger ensembles, while significantly ahead of OA.

We also conducted the Friedman and Nemenyi tests for these data. Results confirm that the sizes of ensembles obtained with pruning algorithms are statistically different at $\alpha=0.05$. Figure 3 presents the results of the Nemenyi test.

The ideal pruning algorithm should reduce both the error and size of the ensemble relative to the original bagging ensemble. This condition is satisfied in most cases; however, Table 4 shows that pruning algorithms increase the error for some datasets. However, this degradation is not statistically significant (Table 4), and the pruned model can be ten or more times smaller than the original bagging ensemble, which can be crucial for devices with limited memory (e.g., smart sensors). Distributions of AP, DR, and RR results are very similar (with a slight advantage of AP in terms of error); OA in our experiments demonstrates the worst performance in terms of error and model size.

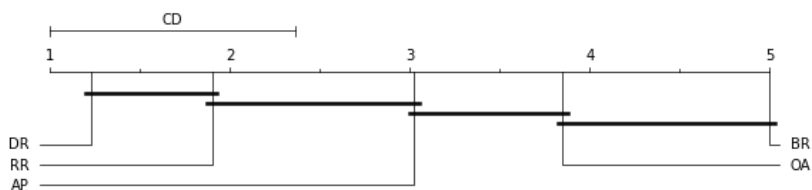


Fig. 3. Comparison of the size of ensembles with the Nemenyi test. $CD = 1.364$

5. Conclusions. In this work, we investigated regression ensembles pruning methods since previous research did not pay enough attention to this area.

The first of the main contributions of our work is the formal condition (7), which allows estimating the potential to reduce the size of the convex regression ensemble. If there are no regressors in the ensemble for which the value of Equation (7) is positive, then the reduction of such an ensemble is impossible.

Next, we proposed an Ambiguity-based pruning algorithm based on well-known error-ambiguity decomposition and compared its performance with other pruning techniques such as minimization of error using direct and reverse approaches. The results of experiments with real datasets show that Ambiguity-based pruning outperforms in most cases other algorithms also the state-of-art Ordered Aggregation algorithm.

References

1. Chen H., Tiño P., Yao X. Predictive ensemble pruning by expectation propagation. *IEEE Transactions on Knowledge & Data Engineering*. 2009. vol. 21. no. 7. pp. 999–1013.
2. Zhou Z., Wu J., Tang W. Ensembling neural networks: many could be better than all. *Artificial Intelligence*. 2002. vol. 137. no. 1–2. pp. 239–263.
3. Sagi O., Rokach L. Ensemble learning: A survey. *WIREs Data Mining and Knowledge Discovery*. 2018. vol. 8. no. 4. e1249.
4. Dias K., Windeatt T. Dynamic ensemble selection and instantaneous pruning for regression. *Proc. of the ESANN*. Bruges, 2014. pp. 643–648.
5. Martínez-Muñoz G., Hernández-Lobato D., Suárez A. An analysis of ensemble pruning techniques based on ordered aggregation. *IEEE Transactions on Pattern Analysis and Machine Intelligence*. 2009. vol. 31. no. 2. pp. 245–259.
6. Mendes-Moreira J., Soares C., Jorge A.M., de Sousa J.F. Ensemble approaches for regression: A survey. *ACM Computing Surveys*. 2012. vol. 45. no. 1. Article 10.
7. Hernández-Lobato D., Martínez-Muñoz G., Suárez A. Empirical Analysis and Evaluation of Approximate Techniques for Pruning Regression Bagging Ensembles. *Neurocomputing*. 2011. vol. 74. no. 12–13. pp. 2250–2264.
8. Krogh A., Vedelsby J. Neural network ensembles, cross validation, and active learning. *Advances in neural information processing systems*. 1995. pp. 231–238.
9. Brown G., Wyatt J.L., Tino P. Managing diversity in regression ensembles. *Journal of Machine Learning Research*. 2005. vol. 6. pp. 1621–1650.
10. Hernández-Lobato D., Martínez-Muñoz G., Suárez A. Pruning in ordered regression bagging ensembles. *Proceedings of the International Joint Conference on Neural Networks*, Vancouver, 2006. pp. 1266–1273.
11. Breiman L. Bagging predictors. *Machine Learning*. 1996. vol. 24. no. 2. pp. 123–140.
12. Breiman L. Random forests. *Machine learning*. 2001. vol. 45. no. 1. pp. 5–32.
13. Liu Y., Yao X. Ensemble learning via negative correlation. *Neural networks*. 1999. vol. 12. no. 10. pp. 1399–404.
14. Friedman J.H. Greedy function approximation: A gradient boosting machine. *Annals of statistics*. 2001. vol. 29. no. 5. pp. 1189–1232.
15. Zhang Y., Burer S., Street W.N. Ensemble pruning via semidefinite programming. *Journal of Machine Learning Research*. 2006. vol. 7. pp. 1315–1338.
16. Hernández-Lobato D., Hernández-Lobato J.M., Ruiz-Torribiano R., Valle Á. Pruning adaptive boosting ensembles by means of a genetic algorithm. *International Conference on Intelligent Data Engineering and Automated Learning*. Springer, 2006. pp. 322–329.
17. Qian C., Yu Y., Zhou Z. Pareto Ensemble Pruning. *Proceedings of the 29th AAAI Conference on Artificial Intelligence*. Austin, 2015. pp. 2935–2941.
18. Sun Q., Pfahringer B. Bagging ensemble selection for regression. *Australasian Joint Conference on Artificial Intelligence*. Sydney, 2012. pp. 695–706.
19. Yu Y., Zhou Z.H., Ting K.M. Cocktail ensemble for regression. *Proceedings of ICDM'07*, 2007. pp. 721–726.
20. Wolpert D.H. Stacked generalization. *Neural Networks*. 1992. vol. 5. pp. 241–259.
21. Caruana R., Niculescu-Mozil A., Crew G., Skikes A. Ensemble selection from libraries of models. *Proceedings of the ICML'04. Banf*, 2004. pp. 18–25.

22. Bian Y., Wang Y., Yao Y., Chen H. Ensemble pruning based on objection maximization with a general distributed framework. *IEEE Transactions on Neural Networks and Learning Systems*. 2020. vol. 31. no. 9. pp. 3766–3774.
23. Mao S., Chen J., Jiao L., Gou S., Wang R. Maximizing diversity by transformed ensemble learning. *Applied Soft Computing*. 2019. vol. 82. p. 105580.
24. Zhou Z. *Machine learning*. Springer, 2021. 472 p.
25. Guo H., Liu H., Li R., Wu C., Guo Y., Xu M. Margin & diversity based ordering ensemble pruning. *Neurocomputing*. 2018. vol. 275. pp. 237–246.
26. Lustosa Filho J.A.S., Canuto A.M., Santiago R.H.N. Investigating the impact of selection criteria in dynamic ensemble selection methods. *Expert Systems with Applications*. 2018. vol. 106. pp. 141–153.
27. Fan Y., Tao L., Zhou Q., Han X. Cluster ensemble selection with constraints. *Neurocomputing*. 2017. vol. 235. pp. 59–70.
28. Golalipour K., Akbari E., Hamidi S.S., Lee M., Enayatifar R. From clustering to clustering ensemble selection: A review. *Engineering Applications of Artificial Intelligence*. 2021. vol. 104. p. 104388.
29. Zhang C., Wu Y., Zhu M. Pruning variable selection ensembles. *Statistical Analysis and Data Mining: The ASA Data Science Journal*. 2019. vol. 12. no. 3. pp. 168–184.
30. Baron G. Greedy selection of attributes to be discretized. (Ed.: Hassanien A.) *Machine Learning Paradigms: Theory and Application. Studies in Computational Intelligence*. Springer, Cham, 2019. vol. 801. pp. 45–67.
31. Khairalla M.A.E. Metaheuristic ensemble pruning via greedy-based optimization selection. *International Journal of Applied Metaheuristic Computing*. 2022. vol. 13. no. 1. pp. 1–22.
32. Jiang Z., Liu H., Fu B., Wu Z. Generalized ambiguity decompositions for classification with applications in active learning and unsupervised ensemble pruning. *Proceedings of the Thirty-First AAAI Conference on Artificial Intelligence (AAAI-17)*, 2017. pp. 2073–2079.
33. Dong X., Yu Z., Cao W., Shi Y., Ma Q. A survey on ensemble learning. *Frontiers of Computer Science*. 2020. vol. 14. no. 2. pp. 241–258.
34. Shahhosseini M., Hu G., Pham H. Optimizing ensemble weights and hyperparameters of machine learning models for regression problems. *Machine Learning with Applications*. 2022. vol. 7. p. 100251.
35. Fuhg J., Fau A., Nackenhorst U. State-of-the-Art and Comparative Review of Adaptive Sampling Methods for Kriging. *Archives of Computational Methods in Engineering*. 2021. vol. 28. pp. 2689–2747.
36. Liu H., Ong Y.-S., Cai J. A survey of adaptive sampling for global metamodeling in support of simulation-based complex engineering design. *Structural and Multidisciplinary Optimization*. 2018. vol. 57. no. 1. pp. 393–416.
37. Margineantu D.D., Dietterich T.G. Pruning adaptive boosting. *Proc. of 14th International Conference on Machine Learning. ICML, 1997*. pp. 211–218.
38. Hsu K.W. A theoretical analysis of why hybrid ensembles work. *Computational Intelligence and Neuroscience*. 2017. vol. 2017. p. 1930702.
39. Yao Y., Pirš G., Vehtari A., Gelman A. Bayesian hierarchical stacking: Some models are (somewhere) useful. *Bayesian Analysis*. 2022. vol. 17. no. 4. pp. 1043–1071.
40. Nuzhny A.S. Bayes regularization in the selection of weight coefficients in the predictor ensembles. *Proc. ISP RAS*, 2019. vol. 31. no. 4. pp. 113–120. (in Russ.).
41. Demšar J. Statistical comparisons of classifiers over multiple data sets. *Journal of Machine Learning Research*. 2006. vol. 7. pp. 1–30.

Zelenkov Yuri — Ph.D., Dr.Sci., Professor, Graduate school of business, HSE University. Research interests: modeling of soft systems, machine learning, knowledge management. The number of publications — 80. yzelenkov@hse.ru; 28/11, Shabolovka St., 119049, Moscow, Russia; office phone: +7(495)771-3232.

Ю.А. ЗЕЛЕНКОВ
ОПТИМИЗАЦИЯ РАЗМЕРА АНСАМБЛЯ РЕГРЕССОРОВ

Зеленков Ю.А. Оптимизация размера ансамбля регрессоров.

Аннотация. Алгоритмы обучения ансамблей, такие как bagging, часто генерируют неоправданно большие композиции, которые, помимо потребления вычислительных ресурсов, могут ухудшить обобщающую способность. Обрезка (pruning) потенциально может уменьшить размер ансамбля и повысить точность; однако большинство исследований сегодня сосредоточены на использовании этого подхода при решении задачи классификации, а не регрессии. Это связано с тем, что в общем случае обрезка ансамблей основывается на двух метриках: разнообразии и точности. Многие метрики разнообразия разработаны для задач, связанных с конечным набором классов, определяемых дискретными метками. Поэтому большинство работ по обрезке ансамблей сосредоточено на таких проблемах: классификация, кластеризация и выбор оптимального подмножества признаков. Для проблемы регрессии гораздо сложнее ввести метрику разнообразия. Фактически, единственной известной на сегодняшний день такой метрикой является корреляционная матрица, построенная на предсказаниях регрессоров. Данное исследование направлено на устранение этого пробела. Предложено условие, позволяющее проверить, включает ли регрессионный ансамбль избыточные модели, т. е. модели, удаление которых улучшает производительность. На базе этого условия предложен новый алгоритм обрезки, который основан на декомпозиции ошибки ансамбля регрессоров на сумму индивидуальных ошибок регрессоров и их рассогласованность. Предложенный метод сравнивается с двумя подходами, которые напрямую минимизируют ошибку путем последовательного включения и исключения регрессоров, а также с алгоритмом упорядоченного агрегирования (Ordered Aggregation). Эксперименты подтверждают, что предложенный метод позволяет уменьшить размер ансамбля регрессоров с одновременным улучшением его производительности и превосходит все сравниваемые методы.

Ключевые слова: обрезка ансамбля, ансамбль регрессоров, обучение ансамбля, декомпозиция ошибка-разнообразие, разнообразие регрессоров.

Литература

1. Chen H., Tiño P., Yao X. Predictive ensemble pruning by expectation propagation. *IEEE Transactions on Knowledge & Data Engineering*. 2009. vol. 21. no. 7. pp. 999–1013.
2. Zhou Z., Wu J., Tang W. Ensembling neural networks: many could be better than all. *Artificial Intelligence*. 2002. vol. 137. no. 1–2. pp. 239–263.
3. Sagi O., Rokach L. Ensemble learning: A survey. *WIREs Data Mining and Knowledge Discovery*. 2018. vol. 8. no. 4. e1249.
4. Dias K., Windeatt T. Dynamic ensemble selection and instantaneous pruning for regression. *Proc. of the ESANN*. Bruges, 2014. pp. 643–648.
5. Martínez-Muñoz G., Hernández-Lobato D., Suárez A. An analysis of ensemble pruning techniques based on ordered aggregation. *IEEE Transactions on Pattern Analysis and Machine Intelligence*. 2009. vol. 31. no. 2. pp. 245–259.
6. Mendes-Moreira J., Soares C., Jorge A.M., de Sousa J.F. Ensemble approaches for regression: A survey. *ACM Computing Surveys*. 2012. vol. 45, no. 1. Article 10.

7. Hernández-Lobato D., Martínez-Muñoz G., Suárez A. Empirical Analysis and Evaluation of Approximate Techniques for Pruning Regression Bagging Ensembles. *Neurocomputing*. 2011. vol. 74. no. 12–13. pp. 2250–2264.
8. Krogh A., Vedelsby J. Neural network ensembles, cross validation, and active learning. *Advances in neural information processing systems*. 1995. pp. 231–238.
9. Brown G., Wyatt J.L., Tino P. Managing diversity in regression ensembles. *Journal of Machine Learning Research*. 2005. vol. 6. pp. 1621–1650.
10. Hernández-Lobato D., Martínez-Muñoz G., Suárez A. Pruning in ordered regression bagging ensembles. *Proceedings of the International Joint Conference on Neural Networks, Vancouver, 2006*. pp. 1266–1273.
11. Breiman L. Bagging predictors. *Machine Learning*. 1996. vol. 24. no. 2. pp. 123–140.
12. Breiman L. Random forests. *Machine learning*. 2001. vol. 45. no. 1. pp. 5–32.
13. Liu Y., Yao X. Ensemble learning via negative correlation. *Neural networks*. 1999. vol. 12. no. 10. pp. 1399–404.
14. Friedman J.H. Greedy function approximation: A gradient boosting machine. *Annals of statistics*. 2001. vol. 29. no. 5. pp. 1189–1232.
15. Zhang Y., Burer S., Street W.N. Ensemble pruning via semidefinite programming. *Journal of Machine Learning Research*. 2006. vol. 7. pp. 1315–1338.
16. Hernández-Lobato D., Hernández-Lobato J.M., Ruiz-Torribiano R., Valle Á. Pruning adaptive boosting ensembles by means of a genetic algorithm. *International Conference on Intelligent Data Engineering and Automated Learning*. Springer, 2006. pp. 322–329.
17. Qian C., Yu Y., Zhou Z. Pareto Ensemble Pruning. *Proceedings of the 29th AAAI Conference on Artificial Intelligence*. Austin, 2015. pp. 2935–2941.
18. Sun Q., Pfahringer B. Bagging ensemble selection for regression. *Australasian Joint Conference on Artificial Intelligence*. Sydney, 2012. pp. 695–706.
19. Yu Y., Zhou Z.H., Ting K.M. Cocktail ensemble for regression. *Proceedings of ICDM'07, 2007*. pp. 721–726.
20. Wolpert D.H. Stacked generalization. *Neural Networks*. 1992. vol. 5. pp. 241–259.
21. Caruana R., Niculescu-Mozil A., Crew G., Ksikes A. Ensemble selection from libraries of models. *Proceedings of the ICML'04, Banf, 2004*. pp. 18–25.
22. Bian Y., Wang Y., Yao Y., Chen H. Ensemble pruning based on objection maximization with a general distributed framework. *IEEE Transactions on Neural Networks and Learning Systems*. 2020. vol. 31. no. 9. pp. 3766–3774.
23. Mao S., Chen J., Jiao L., Gou S., Wang R. Maximizing diversity by transformed ensemble learning. *Applied Soft Computing*. 2019. vol. 82. p. 105580.
24. Zhou Z. *Machine learning*. Springer, 2021. 472 p.
25. Guo H., Liu H., Li R., Wu C., Guo Y., Xu M. Margin & diversity based ordering ensemble pruning. *Neurocomputing*. 2018. vol. 275. pp. 237–246.
26. Lustosa Filho J.A.S., Canuto A.M., Santiago R.H.N. Investigating the impact of selection criteria in dynamic ensemble selection methods. *Expert Systems with Applications*. 2018. vol. 106. pp. 141–153.
27. Fan Y., Tao L., Zhou Q., Han X. Cluster ensemble selection with constraints. *Neurocomputing*. 2017. vol. 235. pp. 59–70.
28. Golalipour K., Akbari E., Hamidi S.S., Lee M., Enayatifar R. From clustering to clustering ensemble selection: A review. *Engineering Applications of Artificial Intelligence*. 2021. vol. 104. p. 104388.
29. Zhang C., Wu Y., Zhu M. Pruning variable selection ensembles. *Statistical Analysis and Data Mining: The ASA Data Science Journal*. 2019. vol. 12. no. 3. pp. 168–184.
30. Baron G. Greedy selection of attributes to be discretized. (Ed.: Hassanien A.) *Machine Learning Paradigms: Theory and Application*. Studies in Computational Intelligence. Springer, Cham, 2019. vol. 801. pp. 45–67.

31. Khairalla M.A.E. Metaheuristic ensemble pruning via greedy-based optimization selection. *International Journal of Applied Metaheuristic Computing*. 2022. vol. 13. no. 1. pp. 1–22.
32. Jiang Z., Liu H., Fu B., Wu Z. Generalized ambiguity decompositions for classification with applications in active learning and unsupervised ensemble pruning. *Proceedings of the Thirty-First AAAI Conference on Artificial Intelligence (AAAI-17)*, 2017. pp. 2073–2079.
33. Dong X., Yu Z., Cao W., Shi Y., Ma Q. A survey on ensemble learning. *Frontiers of Computer Science*. 2020. vol. 14. no. 2. pp. 241–258.
34. Shahhosseini M., Hu G., Pham H. Optimizing ensemble weights and hyperparameters of machine learning models for regression problems. *Machine Learning with Applications*. 2022. vol. 7. p. 100251.
35. Fuhg J., Fau A., Nackenhorst U. State-of-the-Art and Comparative Review of Adaptive Sampling Methods for Kriging. *Archives of Computational Methods in Engineering*. 2021. vol. 28. pp. 2689–2747.
36. Liu H., Ong Y.-S., Cai J. A survey of adaptive sampling for global metamodeling in support of simulation-based complex engineering design. *Structural and Multidisciplinary Optimization*. 2018. vol. 57. no. 1. pp. 393–416.
37. Margineantu D.D., Dietterich T.G. Pruning adaptive boosting. *Proc. of 14th International Conference on Machine Learning. ICML, 1997*. pp. 211–218.
38. Hsu K.W. A theoretical analysis of why hybrid ensembles work. *Computational Intelligence and Neuroscience*. 2017. vol. 2017. p. 1930702.
39. Yao Y., Pirš G., Vehtari A., Gelman A. Bayesian hierarchical stacking: Some models are (somewhere) useful. *Bayesian Analysis*. 2022. vol. 17. no. 4. pp. 1043–1071.
40. Нужный А.С. Регуляризация Байеса при подборе весовых коэффициентов в ансамблях предикторов. *Труды ИСП РАН*. Т. 31. № 4. 2019.
41. Demšar J. Statistical comparisons of classifiers over multiple data sets. *Journal of Machine Learning Research*. 2006. vol. 7. pp. 1–30.

Зеленков Юрий Александрович — д-р техн. наук, профессор, высшая школа бизнеса, НИУ «Высшая школа экономики». Область научных интересов: моделирование социально-экономических систем, машинное обучение, управление знаниями. Число научных публикаций — 80. yzelenkov@hse.ru; улица Шаболовка, 28/11, 119049, Москва, Россия; р.т.: +7(495)771-3232.

Е.К. СЕМЕНЕНКО, А.Г. БЕЛОЛИПЕЦКАЯ, Р.Н. ЮРЬЕВ, А.П. АЛОДЖАНЦ,
И.А. БЕССМЕРТНЫЙ, И.А. СУРОВ
**ВЫЯВЛЕНИЕ ЭКОНОМИЧЕСКИХ СГОВОРОВ МЕТРИКАМИ
КВАНТОВОЙ ЗАПУТАННОСТИ**

Семенов Е.К., Белоліпецкая А.Г., Юрьев Р.Н., Алоджанц А.П., Бессмертный И.А., Суров И.А. Выявление экономических сговоров метриками квантовой запутанности.

Аннотация. Эффективность экономики обусловлена оперативностью пресечения незаконного поведения хозяйствующих субъектов. В условиях ускорения деловой активности важной частью данного условия становится выявление рыночных сговоров на основе статистики электронных следов. В статье представлено решение этой задачи на основе кванто-теоретического подхода к моделированию принятия решений. А именно, когнитивные состояния субъектов представляются комплекснозначными векторами в пространстве, образованном базисными поведенческими альтернативами, тогда как вероятности принятия решений определяются проекциями этих состояний на соответствующие направления. Согласованность многостороннего поведения при этом соответствует запутанности порождающего когнитивного состояния, степень которой измеряется стандартными кванто-теоретическими метриками. Высокое значение метрики свидетельствует о вероятном наличии сговора между рассматриваемыми субъектами. Полученный таким образом метод выявления поведенческой координации апробирован на открытых данных об участии юридических лиц в государственных закупках за период с 2015 по 2020 годы, доступных на федеральном портале <https://zakupki.gov.ru>. Для использованной выборки построены квантовые модели примерно 80 тысяч уникальных пар и 10 миллионов уникальных троек ИНН. Достоверность выявления сговоров определялась сравнением подозреваемых с открытыми данными Федеральной антимонопольной службы <https://br.fas.gov.ru>. Согласно полученным функциям ошибок, половина известных парных сговоров выявляется с достоверностью более 50%, что сравнимо с методами выявления на основе классической корреляции и классической взаимной информации. В трёхстороннем случае, напротив, квантовая модель оказывается практически безальтернативной в силу ограниченности классических метрик двусторонней корреляцией. Половина таких сговоров выявляется с достоверностью 40%. Полученные результаты свидетельствуют об эффективности кванто-вероятностного подхода к моделированию многостороннего экономического поведения. Разработанные метрики могут быть использованы в качестве информативных признаков для аналитических систем и алгоритмов машинного обучения подобной направленности.

Ключевые слова: сговор, картель, принятие решений, квантовая вероятность, квантовая запутанность, поведенческое моделирование, рекомендательные системы.

1. Введение. Машинный анализ данных с целью обеспечения правопорядка является актуальной областью прикладных исследований. Он используется для выявления прошлых и прогнозирования будущих актов противоправного поведения на основе анализа разнородных поведенческих данных, включая информацию в социальных сетях, финансовые транзакции и др. [1 – 4]. Одна из таких задач – выявление незаконных сговоров между субъектами экономической деятельности, в том числе между участниками государственных закупок товаров и услуг, опубликованных в единой информационной системе <https://zakupki.gov.ru>; таким образом закупаются, например, лекарства и оборудование, услуги Интернет, телефонии и ЖКХ, строительства и ремонта и т.д. В целях повышения прибыли недобросовестные подрядчики договариваются о согласованном поведении на торгах, вследствие чего товар или услуга приобретается из бюджетных средств по завышенной цене.

Такое поведение хозяйствующих субъектов – называемое *сговором* или *картелем* – незаконно и преследуется по статьям административного и уголовного кодексов РФ [5]. Для сокрытия своих действий правонарушители привлекают к торгам подставных участников, которые играют по заранее распределённым ролям для создания видимости конкуренции. Выявление сложных схем согласованного поведения представляет большую практическую проблему, которая обостряется с ускорением экономической активности [6]. Растущая скорость транзакций в электронной среде превышает возможности ручной обработки, которые, однако, расширяются с помощью машинных средств.

Первая модель автоматического выявления сговоров, также известная как эконометрический скрининг, была предложена в 2003 г. [7]. Эта модель позволяет рассчитать естественное, конкурентно-рыночное распределение заявок, систематические отклонения от которого являются признаками недобросовестного поведения. Недостатком модели является необходимость знания экономических параметров участников торгов, многие из которых представляют собой коммерческую тайну. В результате большая часть необходимых данных должна дополнительно определяться отраслевыми экспертами, что приводит к потере точности. Аналогичный подход [8] использует анализ аномальных дисперсий в распределении заявок, предполагая, что они следуют равномерному распределению.

В последние годы разработан ряд методов выявления сговоров на основе машинного обучения. В первом исследовании такого типа [9] создан классификатор заявок по наличию либо отсутствию сговора на

строительстве швейцарских железных дорог. С четырьмя достоверно установленными картелями, на малой выборке из 584 торгов правильно классифицированы 84% из 3799 заявок. Широкое исследование методов машинного обучения проведено для больших массивов данных о торгах из пяти стран [10]. Усреднённая по странам сбалансированная точность 11 рассмотренных алгоритмов находится в диапазоне от 48% до 86%.

Недостатком большинства известных подходов является их ограниченность двусторонним поведением, т.е. рассмотрением только сговоров между двумя субъектами поведения. Помимо наибольшей простоты, основанием для этой ограниченности служит неявное предположение о том, что многосторонние связи сводятся к комбинации двусторонних. Соответственно, классические монографии по статистике, психометрике, экономике и социологии [11 – 14] рассматривают двусторонний случай как единственно заслуживающий внимания.

В реальной практике социальных взаимодействий, однако, для такого допущения оснований нет [15 – 18]; напротив, установлено, что в качестве единицы межсубъектных отношения более оправдано рассмотрение не пар, а троек индивидов [19, 20]. В парадигме машинного обучения эта задача может быть решена аналогично двустороннему случаю, когда признаками классифицируемого объекта являются параметры многостороннего поведения [21]. Нахождение наиболее информативных параметров такого поведения, однако, есть теоретическая задача, требующая отдельного решения.

Недостаточность классической статистики в этом отношении указывает на актуальность рассмотрения альтернативных подходов к поведенческому моделированию. В настоящей работе в качестве такового исследуется развитый в последние десятилетия подход на основе квантовой теории. Среди других направлений социофизики [22 – 24], этот подход выделяется удобством моделирования связанного поведения, играющего ключевую роль в динамике социальных систем [25]. Это свойство позволяет, в частности, описывать коллективное принятие решений [26, 30], неклассические равновесия в играх [27–29, 31, 32], иррациональную динамику финансовых рынков [33 – 40], коллективные социальные возбуждения [41 – 43], и др. [44, 46 – 48]. При этом, согласованность принятия индивидуальных решений соответствует свойству *запутанности* многочастичных квантовых систем, не имеющему аналога в классических поведенческих моделях.

В настоящей статье квантовая теория вероятности использована для выявления экономических сговоров. Развивая намеченный ранее подход [6], в разделе 2 представлена соответствующая модель

многостороннего поведения, на основе которой построен ряд квантово-теоретических метрик поведенческой согласованности. В разделе 3 описана методика эксперимента по апробации модели на открытых данных об участии юридических лиц в государственных закупках. В разделе 4 представлены результаты эксперимента, включая достоверность выявления сговоров и сравнение квантово-теоретических метрик корреляции с существующими классическими аналогами. В разделе 5 обсуждаются перспективные направления развития представленного подхода.

2. Теория

2.1. Принципы поведенческого моделирования на основе квантовой теории. В соответствии с принципами квантового моделирования [49], когнитивное состояние субъекта данной поведенческой неопределённости представляется нормированным вектором $|\psi\rangle$ в комплекснозначном (Гильбертовом) пространстве $\mathcal{H}$. Базис пространства образован ортонормированными векторами:

$$|1\rangle = \begin{bmatrix} 1 \\ 0 \\ \vdots \\ 0 \end{bmatrix}, \quad |2\rangle = \begin{bmatrix} 0 \\ 1 \\ \vdots \\ 0 \end{bmatrix}, \quad \dots \quad |N\rangle = \begin{bmatrix} 0 \\ 0 \\ \vdots \\ 1 \end{bmatrix}, \quad (1)$$

соответствующими N альтернативным взаимоисключающим вариантам поведения. Квантово-когнитивное состояние субъекта тогда описывается вектором:

$$|\psi\rangle = \sum_{i=k}^N \gamma_k |k\rangle = \begin{bmatrix} \gamma_1 \\ \gamma_2 \\ \vdots \\ \gamma_N \end{bmatrix}, \quad (2)$$

где комплекснозначные компоненты разложения γ_k называются также *амплитудами* (вероятности). Эти амплитуды можно представить в виде скалярного произведения $\gamma_k = \langle k|\psi\rangle$, где вектор-строка $\langle k|$ есть комплексно-сопряжённая транспозиция вектор-столбца $|k\rangle$. Квадрат модуля этой величины определяет вероятность осуществления соответствующей поведенческой альтернативы:

$$|\langle k|\psi\rangle|^2 = |\gamma_k|^2 = p_k, \quad k = 1 \dots N. \quad (3)$$

Соответственно, нормировка:

$$\langle \psi | \psi \rangle = [\gamma_1^* \quad \gamma_2^* \quad \dots \quad \gamma_N^*] \begin{bmatrix} \gamma_1 \\ \gamma_2 \\ \vdots \\ \gamma_N \end{bmatrix} = \sum_{k=1}^N p_k = 1, \quad (4)$$

выражает полноту базисной системы поведенческих альтернатив.

Классическая и квантовая неопределённость. Состояние (2) коренным образом отличается от нормированного распределения вероятности $[p_1, p_2, \dots, p_N]$, описывающего неопределённость наблюдаемых величин в классической теории вероятности Колмогорова. Распределения этого типа описывают незнание величины, объективно находящейся в одном из N состояний. В случае $N = 2$ примером такой ситуации является незнание человеком того, какой стороной вверх лежит монета под салфеткой. Эта *субъективная* неопределённость снимается измерением, не меняющим объективное состояние величины.

Вектор состояния (2), напротив, описывает состояние неустранимой неопределённости, когда величина ещё не находится ни в одном из базисных состояний. В отличие от неопределённости, обусловленной ограниченной информированностью субъекта как описано выше, данная неопределённость носит *объективный* характер. Информацию о её значении можно получить только путём эксперимента, приводящего величину в одно из N базисных состояний [50, 51].

В поведенческих процессах таким экспериментом является каждый акт принятия решения – свободный выбор человека, не обусловленный объективными факторами. Алгоритмически не предопределённый, спонтанный характер таких решений относит поведенческие неопределённости в рассматриваемой задаче к квантовому типу [52, 53]. Как и в физике, соответствующие статистические закономерности моделируются с помощью линейной алгебры вектор-состояний в комплекснозначных (Гильбертовых) пространствах как описано выше.

2.2. Модель принятия многосторонних двухвариантных решений. В рассматриваемой задаче субъектом поведения является группа из K юридических лиц, принимающих решение относительно участия «1» либо неучастия «0» государственном тендере на поставку товаров, оказание услуг или выполнение работ по заранее объявленным в документации условиям (далее торг). В результате имеет место $N = 2^K$ взаимоисключающих поведенческих альтернатив, одна из которых

осуществляется в процессе принятия решения. Эти альтернативы кодируются последовательностями нулей и единиц. Базисное состояние $|10\rangle$, например, соответствует случаю, когда из $K = 2$ лиц A и B участие в торге приняло только лицо A . Согласно первой части выражения (2) общий вид когнитивного состояния такой группы описывается вектором:

$$|\psi_{AB}\rangle = \gamma_{00}|00\rangle + \gamma_{01}|01\rangle + \gamma_{10}|10\rangle + \gamma_{11}|11\rangle, \quad (5)$$

нормировка (4) которого требует:

$$\sum_{a,b} p_{ab} = \sum_{a,b} |\gamma_{ab}|^2 = 1.$$

Аналогично, когнитивное состояние группы из трёх лиц A , B , C описывается вектором:

$$|\psi_{ABC}\rangle = \sum_{a,b,c} \gamma_{abc}|abc\rangle, \quad \sum_{a,b,c} |\gamma_{abc}|^2 = 1. \quad (6)$$

Как отмечено во введении, эти два случая $K = 2$ и $K = 3$ представляют наибольший практический интерес.

Комплекснозначные амплитуды γ удобно представить в виде модулей r и фазовых множителей $e^{i\phi}$ как:

$$\gamma = r e^{i\phi}, \quad r \in [0, 1], \quad \phi \in [0, 2\pi). \quad (7)$$

Согласно (3), модули однозначно определяются вероятностями принимаемых решений, так что для случая $K = 2$:

$$r_{ab} = \sqrt{p_{ab}}, \quad p_{ab} = \frac{n_{ab}}{n}, \quad (8)$$

где n_{ab} есть число решений «ab», а $n = \sum n_{ab}$ - полное число решений в имеющихся данных. Фазы ϕ_{ab} , напротив, являются свободными параметрами, составляющими особенность квантовой модели.

2.3. Метрики двусторонней согласованности. Согласованность частей системы при разрешении неопределённости квантового типа известна в физике как феномен квантовой запутанности [54]. В рассматриваемой задаче запутанность выражается в согласованном принятии решений юридическими лицами, являющимися частями распределенной поведенческой системы.

Запутанность двухчастичного (двустороннего) состояния (5) выражается в невозможности представить его в виде произведения:

$$\begin{aligned} |\psi_A\rangle \otimes |\psi_B\rangle &= \\ &= (\alpha_0|0_A\rangle + \alpha_1|1_A\rangle) \otimes (\beta_0|0_B\rangle + \beta_1|1_B\rangle) = \\ &= \alpha_0\beta_0|00\rangle + \alpha_0\beta_1|01\rangle + \alpha_1\beta_0|10\rangle + \alpha_1\beta_1|11\rangle, \end{aligned} \quad (9)$$

где $|\psi_A\rangle$ и $|\psi_B\rangle$ есть индивидуальные когнитивные состояния лиц A и B . Соответственно, мерой запутанности является величина:

$$Q_{AB} = 2 |\gamma_{11}\gamma_{00} - \gamma_{01}\gamma_{10}|, \quad 0 \leq Q \leq 1, \quad (10)$$

называемая *конкурентс*¹ [55]. Конкурентс обращается в ноль для факторизуемых состояний (9), соответствующих несогласованному поведению субъектов A и B , т.е. независимости их решений. Наибольшее значение, равное единице, достигается для максимально запутанных состояний:

$$\frac{|00\rangle + e^{i\phi}|11\rangle}{\sqrt{2}}, \quad \frac{|01\rangle + e^{i\phi}|10\rangle}{\sqrt{2}},$$

описывающих максимальную согласованность (корреляцию и антикорреляцию) двусторонних решений, моделируемую в рамках квантовой теории.

Соотношения (7) и (8) позволяют выразить величину (10) через наблюдаемые вероятности p_{ab} :

$$\begin{aligned} Q_{AB} &= 2\sqrt{p_{00}p_{11} + p_{01}p_{10} - 2\sqrt{p_{00}p_{01}p_{10}p_{11}} \cos \Delta}, \\ \Delta &= \phi_{00} + \phi_{11} - \phi_{10} - \phi_{01}, \end{aligned} \quad (11)$$

где параметр Δ определяет влияние фазовых параметров исходного состояния (5) на значение Q_{AB} . Данная форма метрики конкурентс наиболее удобна для практического использования.

¹ В квантовой информатике используются и другие меры двухчастичной запутанности, в частности запутанность формирования [58], (логарифмическая) отрицательность [59] и энтропия (взаимная информация) фон-Неймана [60]. Для рассматриваемой задачи отличие этих величин от конкурентс [61, 62] незначительно.

Классическими аналогами конкуренса (11) является модуль коэффициента сопряженности [11]:

$$C_{AB} = \left| \frac{p_{00}p_{11} - p_{01}p_{10}}{\sqrt{(p_{00} + p_{01})(p_{00} + p_{10})(p_{11} + p_{10})(p_{11} + p_{01})}} \right|, \quad (12)$$

$$0 \leq C \leq 1$$

и взаимная информация [63]:

$$I_c = \sum_{a,b} p_{ab} \log \frac{p_{ab}}{p_a p_b} > 0, \quad (13)$$

где $p_i = \sum_j p_{ij}$ есть маргинальные распределения двухчастичной вероятности p_{ab} . Как и в случае конкуренса, наименьшие и наибольшие значения этих метрик означают отсутствие корреляции и наиболее сильную корреляцию, двустороннего поведения, соответственно.

2.4. Метрики трёхсторонней согласованности. Для измерения согласованности трёхстороннего поведения была использована стандартная метрика трехчастичной запутанности [64]:

$$\tau = 4 |d_1 - 2d_2 + 4d_3|, \quad (14)$$

где:

$$\begin{aligned} d_1 &= \gamma_{000}^2 \gamma_{111}^2 + \gamma_{001}^2 \gamma_{110}^2 + \gamma_{010}^2 \gamma_{101}^2 + \gamma_{100}^2 \gamma_{011}^2, \\ d_2 &= \gamma_{000} \gamma_{111} \gamma_{011} \gamma_{100} + \gamma_{000} \gamma_{111} \gamma_{101} \gamma_{010} + \\ &\quad \gamma_{000} \gamma_{111} \gamma_{110} \gamma_{001} + \gamma_{011} \gamma_{100} \gamma_{101} \gamma_{010} + \\ &\quad \gamma_{011} \gamma_{100} \gamma_{110} \gamma_{001} + \gamma_{101} \gamma_{010} \gamma_{110} \gamma_{001} + \\ d_3 &= \gamma_{000} \gamma_{110} \gamma_{101} \gamma_{011} + \gamma_{111} \gamma_{001} \gamma_{010} \gamma_{100}. \end{aligned} \quad (15)$$

Кроме того, был использован трёхчастичный аналог выражения для конкуренса (10):

$$\begin{aligned} Q_3 &= |\gamma_{111} \gamma_{000} - \gamma_{001} \gamma_{110}| + |\gamma_{111} \gamma_{000} - \gamma_{010} \gamma_{101}| + \\ &\quad |\gamma_{111} \gamma_{000} - \gamma_{100} \gamma_{011}| + |\gamma_{001} \gamma_{110} - \gamma_{010} \gamma_{101}| + \\ &\quad |\gamma_{001} \gamma_{110} - \gamma_{100} \gamma_{011}| + |\gamma_{010} \gamma_{101} - \gamma_{100} \gamma_{011}|. \end{aligned} \quad (16)$$

Аналогично двухчастичному случаю (9)-(10), обращение в ноль этой величины гарантирует факторизацию трёхчастичного состояния (6) тремя

индивидуальными состояниями $|\psi_A\rangle, |\psi_B\rangle, |\psi_C\rangle^2$. Как и для метрики (14), это значение соответствует независимому принятию решений, т.е. нулевой согласованности.

В отличие от двухчастичного случая (11), метрики (14) и (16) зависят от большого числа фазовых параметров ϕ_{abc} (7), варьирование которых не представлялось возможным в силу ограниченности доступного вычислительного времени. По этой причине все такие параметры брались равными нулю, что соответствует действительным значениям амплитуд γ_{abc} в состоянии (6). Аналогично двухчастичному случаю (8), эти величины вычислялись на основе экспериментальных данных как:

$$\gamma_{abc} = \sqrt{p_{abc}}, \quad p_{abc} = \frac{n_{abc}}{n}, \quad (17)$$

где например n_{100} есть полное число торгов, в которых среди трех рассматриваемых субъектов A, B, C участие принимал только первый.

Не-квантовой метрикой трёхсторонней согласованности является величина:

$$\theta_{123} = \log \frac{p_{111}p_{100}p_{010}p_{001}}{p_{110}p_{101}p_{011}p_{000}}, \quad (18)$$

использованная для анализа корреляции нейронных всплесков [66]. Отсутствию корреляции соответствует нулевая величина этой метрики, принимающей как положительные, так и отрицательные значения. Как и квантовая метрика (14), эта величина отражает подлинно-трехстороннюю корреляцию, нечувствительную к двухсторонним корреляциям в парах, входящих в тройку (там же).

3. Эксперимент. Представленная модель использована для выявления согласованного экономического поведения юридических лиц на основе информации, опубликованной в сети Интернет. Состав и методы получения данных описаны в Разделе 3.1. Раздел 3.2 описывает процедуру их численной обработки. Результаты эксперимента для двусторонних и трёхсторонних метрик представлены в следующем разделе 4.

3.1. Данные. Используемые в работе данные делятся на два типа. Первый тип - это данные об участии юридических лиц в государственных закупках, на основе которых строится модель поведения, описанная в Разделе 2. Второй тип - это данные о фактах согласованного поведения,

²Это условие достаточно, но не необходимо. Альтернативой использованному является определение трёхстороннего конкуренса, данное в [65].

установленных Федеральной Антимонопольной Службой (ФАС). Эти данные использованы в качестве экспертной оценки, позволившей оценить точность разработанной модели согласованного поведения. Данные обоих типов приводились в соответствие на основе индивидуального налогового номера (ИНН), однозначно характеризующим каждое юридическое лицо.

Данные системы ФАС. Источником данных этого типа являлась база судебных решений и правовых актов ФАС, расположенная в открытом доступе по адресу <https://br.fas.gov.ru>. Были использованы документы, опубликованные Управлением по борьбе с картелями в текстовом формате, допускающем автоматическую обработку. За период с 2015 по 2020 г. таким образом извлечено 734 протокола, содержащих информацию о поведении юридических лиц общим числом 1457 ИНН. Для рабочей выборки ИНН, описанной в следующем параграфе, число пар находящихся в сговоре ИНН составило 2038.

Данные системы госзакупок. Данные этого типа получены из единой информационной системы <https://zakupki.gov.ru>, предоставляющей открытый доступ к базам данных об участии юридических лиц в государственных закупках через интерфейс FTP³. Полученные данные имеют формат таблицы, в которой строки соответствуют отдельным ИНН, а столбцы - отдельным торгам. А каждой ячейке находится число 1 или 0, обозначающее участие либо неучастие данного ИНН в данном торге. За период с 2015 по 2020 год по региону Москва данная таблица содержит 151228 ИНН и 952856 торгов.

В силу большого объема данных, а также недостоверности оценки вероятностей входящих в выражения (11) и (12) при числе участий меньше 10, использование полной таблицы указанного размера нецелесообразно. Рабочая выборка ИНН была сформирована из подозреваемой части и фона. Подозреваемая часть, состоящая из 210 ИНН, получена по следующим критериям:

- количество участий данного ИНН в торгах более 20;
- по информации ФАС, данный ИНН находится в сговоре как минимум с одним другим ИНН из выборки.

³Использовался адрес <ftp://ftp.zakupki.gov.ru> с логином/паролем free/free и IP-адресом сервера 95.167.245.94. В корневом каталоге в папке «fcs_regions» содержатся данные по отдельным регионам России. Каталог «protocols» содержит заархивированные папки, каждая из которых содержит файлы XML содержащие протоколы торгов. Файлы XML сгруппированы по годам по 1000 файлов в одной папке. Необходимая информация содержится в полях с тегами и «purchaseNumber» и «INN». На момент обращения (декабрь 2021 г.) были доступны данные с марта 2014 г. по октябрь 2021 г.

Фоновая часть из 190 ИНН была выбрана случайным образом из оставшихся первичных данных. Полученная в результате таблица из 400 ИНН и 86831 торгов составила рабочую выборку объёмом 66,3 Мб. Распределение ИНН по числу участия показано на рисунке 1.

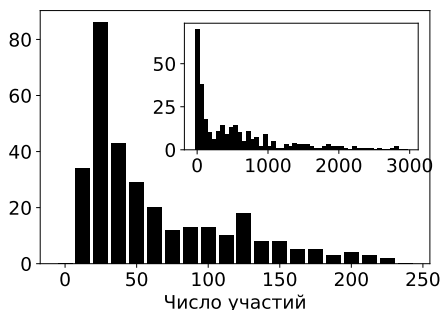


Рис. 1. Распределение 400 ИНН по числу участий в 86831 торгах рассматриваемой выборки. Вставка: то же самое после преобразования таблицы с учётом окружений (раздел 3.2)

3.2. Методика. Двусторонний случай. Полученная выборка из 400 уникальных ИНН позволила проанализировать $400 \times 399/2 = 79800$ уникальных пар ИНН. Каждая такая пара соответствует двум строкам таблицы данных, отражающих участие рассматриваемых ИНН в 86831 торгах. Каждый торг при этом имеет четыре взаимоисключающих исхода, представляемых базисными состояниями $|00\rangle$, $|01\rangle$, $|10\rangle$, $|11\rangle$, охватывающими все варианты совместного участия. Распределение торгов по данным исходам определяется числами n_{ab} , в сумме дающими полное их число $n = 86831$.

В соответствии с формулой (8), на основе величин n_{ab} определяется когнитивное состояние каждой пары субъектов (5). Это состояние далее используется для вычисления метрики запутанности конкуренс (11) для различных значений фазового параметра Δ , в том числе 0° и 180° . Для сравнения также вычислялись значения классического коррелятора (12) и классической взаимной информации (13). Данный набор значений для каждой пары ИНН является результатом теоретического моделирования. Для оценки точности выявления, эти результаты сравнивались с данными ФАС, представляемыми в виде строки длиной 79800, состоящей из нулей и 2038 единиц, описывающих отсутствие либо наличие сговора для данной пары ИНН (рисунок 2).

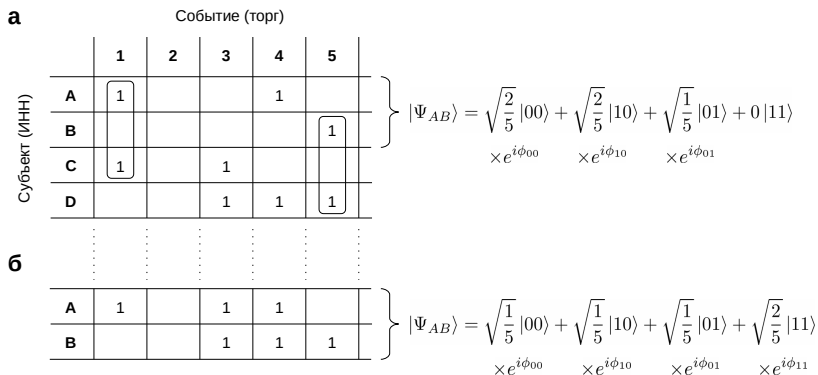


Рис. 2. Построение двухчастичного когнитивного состояния (5) субъектов А и В на основе данных об участии в пяти торгах. а: Базовый метод, применённый к исходной таблице данных. Факты участия отмечены единицами, пустые клетки обозначают неучастие субъекта в данном торге. В соответствии с (7) и (8), амплитуды компонент состояния есть квадратные корни статистических вероятностей соответствующих исходов, дополненные фазовыми множителями $e^{i\phi_{ij}}$. б: Модификация метода с учётом окружений. В каждую строку (А и В) добавляются единицы для торгов, в которых участвовал кто-либо из окружений данных субъектов (С и D, соответственно). Когнитивное состояние пары АВ далее вычисляется на основе модифицированной таблицы данных согласно базовому алгоритму

Для исследования влияния фазового параметра Δ на точность моделирования, в дополнение к вышеуказанным метрикам также вычислялось значение Q_{fit} метрики конкуренс с оптимизированной величиной фазового параметра. А именно, для пар находящихся и не находящихся в сговоре использовались значения $\Delta = 180^\circ$ и 0° , соответственно максимизирующие и минимизирующие значение метрики (11) аналогично работе [56].

Трёхсторонний случай. Та же самая исходная выборка из 400 ИНН позволила проанализировать $400 * 399 * 398 / 6 = 10586800$ уникальных троек ИНН. По данным ФАС, в сговоре среди них подозревалось 22820 троек.

Аналогично двухчастичному случаю, распределение торгов по восьми возможным исходам определяется счётчиками n_{abc} , дающими в сумме полное число торгов 86831. В соответствии с (17), амплитуды γ_{abc} трёхсторонних состояний (6) определялись как квадратные корни

статистических вероятностей. Метрики (16), (14) и (17) далее вычислялись на основе полученных действительныхзначных амплитуд.

Модификация методики с учётом окружений. Несмотря на отсечку в 20 и более торгов для каждого ИНН, матрица исходных данных является разреженной. Всего в них содержится 115369 событий участия, так что в матрице $400 * 86831$ единицами занято порядка 0,003 от полного числа ячеек. Соответственно, в двухчастичных и трехчастичных состояниях (5), (6) доминируют компоненты $|00\rangle$ и $|000\rangle$, тогда как члены с двумя, и тем более с тремя единицами практически не встречаются. Обращение в ноль соответствующих вероятностей и амплитуд переводит рассматриваемые метрики согласованности в вырожденный режим, малоэффективный для решаемой задачи.

Для решения этой проблемы вышеописанная методика модифицирована следующим образом. Считается, что субъект участвовал во всех торгах, в которых участвовали все субъекты, с которыми он соучаствовал в каком-либо торге в рамках рассматриваемой выборки. Множество таких соучастников есть *окружение*, участвовавшее в много большем числе торгов чем исходный субъект - центр данного окружения (вставка на рисунке 1). Соответственно, каждая строка исходной матрицы данных пересекается со строками окружения как показано на рисунке 2, после чего применяется стандартный метод расчёта амплитуд.

4. Результаты

4.1. Результаты двустороннего моделирования. Значение метрики конкуренс с фиксированным значением фазового параметра $\Delta = 0^\circ$ для исследованных пар ИНН, отсортированных по этой величине, показано слева на рисунке 3(а) красной линией. Вертикальными линиями отмечены пары, состоящие в сговоре по данным ФАС. Серой кривой показана плотность этих линий, рассчитанная как бегущее среднее с шириной окна 500. Корреляция Пирсона между этой функцией и значением метрики составляет 0,93. Аналогичный результат для фазы $\Delta = 180^\circ$ показан синей линией на панели (б). Для удобства сравнения, значение метрики в обоих случаях умножено на постоянный коэффициент ≈ 8 , не влияющий на величину корреляций.

Рисунки 3(г) соответствуют взаимной информации (13) и классическому коррелятору (12). Как и предыдущие метрики, эти величины коррелируют с плотностью сговоров, однако коэффициенты корреляции (0,33 и 0,73 соответственно) уступают метрике конкуренс.

Работа тех же метрик с учётом окружений представлена на рисунке 3 (с окружениями).

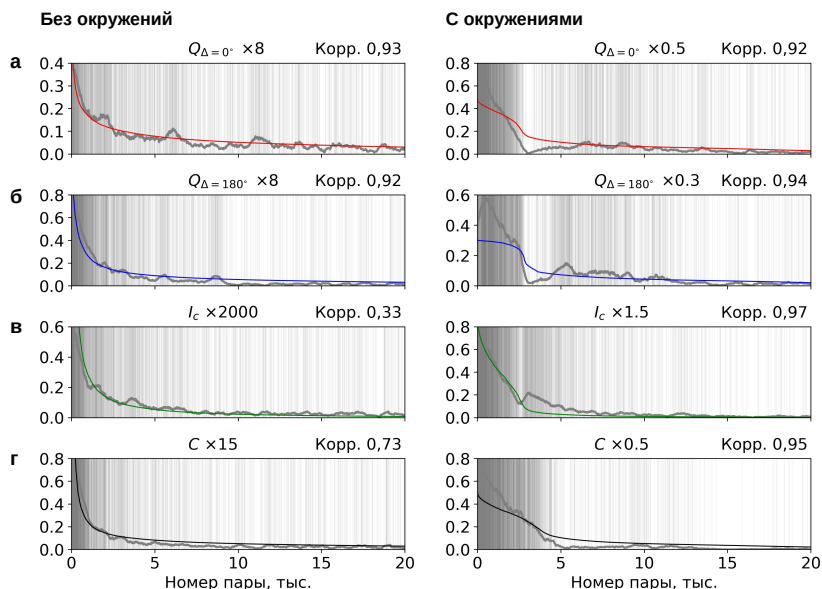


Рис. 3. Значения метрик согласованности для первых 20000 пар ИНН, отсортированных по значению каждой метрики в порядке убывания. Вертикальными линиями отмечены 2038 пар, находящиеся в сговоре по данным ФАС. Локальная плотность линий показана серой кривой. Наименования метрик, а также их корреляции с локальными плотностями сговоров в каждом случае указаны над графиками. Верхние два ряда: метрика конкуренс (11) при фиксированных значениях фазы $\Delta = 0^\circ$ (а) и $\Delta = 180^\circ$ (б). Нижние ряды: взаимная информация (13) (в) и классический коррелятор (12) (г). Слева: расчёт без учёта окружений. Справа: расчёт с учётом окружений (рисунок 2)

Как видно по расположению вертикальных линий, группировка сговоров на высоких значениях метрик становится более плотной. Пропорциональность между метриками и плотностями сговоров, характерная для случая “без окружений”, нарушается. Корреляция классических метрик (в, г) при этом возрастает до уровня 0,95 - 0,97, превосходя квантово-теоретические метрики (а, б).

Качество сортировки пар ИНН рассматриваемыми метриками в той же цветовой кодировке показано на рисунке 4. Линии показывают положение 2038 сговоров в массиве пар, отсортированном по каждой из метрик. Для каждой метрики, точка с горизонтальной координатой 1, например, показывает положение сговора с наибольшим значением метрики. На соответствующей панели рисунка 3, этот сговор показан

первой слева серой линией, тогда как ее положение на оси X есть вертикальная координата точки на рисунке 4. Точечная прямая на вставке соответствует идеальной метрике, которая расположила бы все 2038 сговоров на первых 2038 позициях в сортировке.

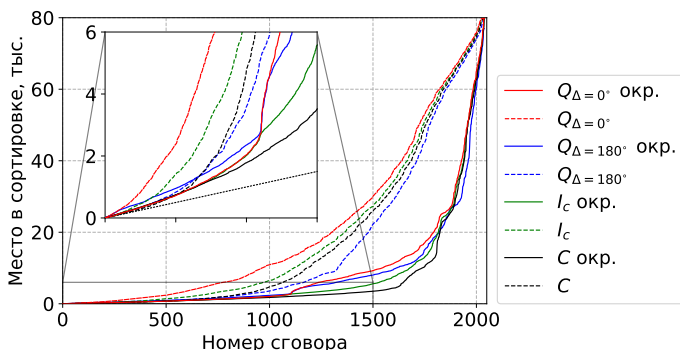


Рис. 4. Расположение 2038 сговоров в массиве пар ИНН, отсортированном по различным метрикам согласованности. Цветовая кодировка метрик совпадает с рисунком 3, случаи с учётом и без учёта окружений показаны сплошными и пунктирными линиями. Точечная прямая на вставке соответствует идеальной сортировке, располагающей сговоры на первых 2038 местах

Для каждой метрики результаты с учётом и без учёта окружений показаны сплошными и пунктирными линиями. За исключением небольшого диапазона в начале горизонтальной оси, сплошные линии во всех случаях располагаются ниже пунктирных линий того же цвета. Таким образом, учёт окружений делает сортировку более качественной, группируя сговоры в начале отсортированного массива пар более плотно.

Учёт окружений также меняет взаимное расположение линий на рисунке 4. Без учёта окружений, наиболее и наименее плотные группировки достигаются метриками конкуренс с фазами $\Delta = 180^\circ$ и $\Delta = 0^\circ$ соответственно, тогда как классические метрики занимают промежуточное положение. Как отмечено выше, учёт окружений существенно улучшает работу последних, так что сплошные зеленая и чёрная линии на рисунке 4 лежат ниже синей и красной в диапазоне $\sim 1100 - 1800$. Тысяча сговоров при этом группируется на первых ~ 2000 местах сортировки всеми метриками.

4.2. Результаты трёхстороннего моделирования. Метрики трёхсторонней запутанности вычислялись на основе трёхчастичных состояний как описано в разделе 3.2. Значения трёхчастичной

модификации метрики конкуренс (16), квантовой метрики (14) и классической метрики (18) показаны на рисунке 5а, 5б и 5в для троек ИНН, отсортированных по значению соответствующей метрики. Для метрик Q_3 и τ показаны первые 10 тысяч из полного числа 10586800 троек.

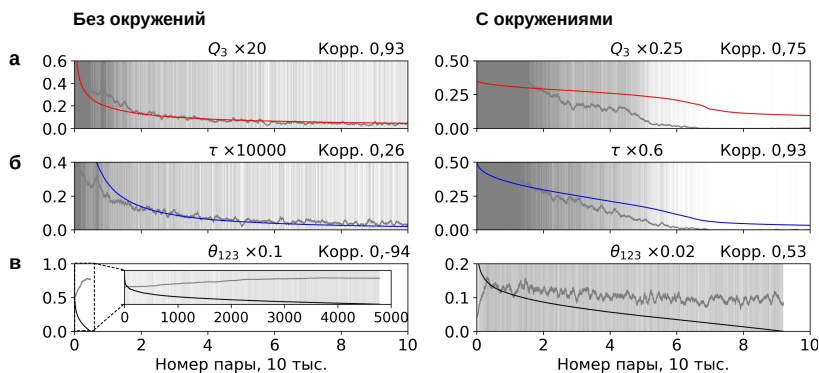


Рис. 5. Значения метрик согласованности для первых 10^5 из 10586800 троек ИНН, отсортированных по значению каждой метрики в порядке убывания: а) трехчастичная модификация метрики конкуренс (16); б) квантовая метрика (14); в) классическая метрика (18). Как и на рисунке 3, вертикальными линиями отмечены тройки, находящиеся в сговоре по данным ФАС, их локальная плотность показана серыми линиями. Слева: расчёты без учёта окружений. Справа: расчёты с учётом окружений

Как и в двустороннем случае, учёт окружений увеличивает плотность группировки сговоров (серые линии) на высоких значениях метрик Q_3 и τ . При этом, корреляция плотности сговоров со значением метрики Q_3 падает с 0,93 до 0,75, тогда как для метрики τ та же величина возрастает с 0,26 до 0,53. Как и на рисунке 3, коэффициенты 20, 0,25, 10000 и 0,6, не меняющие значение корреляции, подобраны для удобства визуального сравнения цветных и серых линий.

Без учёта окружений классическая метрика θ_{123} определена лишь для 4781 троек, для которых знаменатель (18) отличен от нуля. Как показано на рисунке 5в слева, метрика коррелирует с плотностью сговоров отрицательно с коэффициентом Пирсона -0,94. Учёт окружений увеличивает число определённых троек до 92014, причём корреляция меняется на положительную с коэффициентом 0,53.

Качество сортировки трёхсторонних сговоров рассматриваемыми метриками в той же цветовой кодировке показано на рисунке 6.

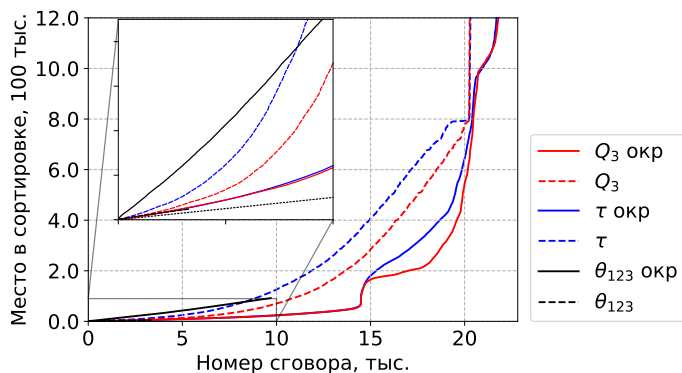


Рис. 6. Аналог рисунка 4 для трёхстороннего случая: расположение 22820 тройных сговоров в массиве троек ИНН, отсортированных по значениям трёхчастичных метрик согласованности. Точечная линия на вставке соответствует идеальной сортировке, располагающей 10^4 сговоров на 10^4 первых мест в сортировке

Как и на рисунке 4, сплошные и пунктирные линии соответствуют алгоритмам с учётом и без учёта окружений, соответственно. Метрики Q_3 (красный) и τ (синий) определены для всех 22820 тройных сговоров, тогда как метрика θ_{123} (чёрный) определена для 3303 сговоров без учёта окружений и для 9695 сговоров с учётом окружений.

Для первых двух метрик использование окружений приводит к улучшению качества сортировки на всём диапазоне значений. Наилучший результат показан трёхсторонним аналогом метрики конкуренс (16), превосходящей квантовую метрику (14) для выявления от 15 до 20 тыс. (70 – 90%) сговоров.

4.3. Достоверность выявления сговоров. Разработанные методы измерения согласованности поведения можно использовать для выявления возможных сговоров на основе открытых данных. Для этого каждая пара либо тройка ИНН относится к классу “Сговор есть” либо “Сговора нет” на основе значения соответствующей метрики. В этом отношении графики на рисунках 4 и 6 эквивалентны функциям ошибок (точность - полнота, ROC), используемым для оценки качества бинарной категоризации [67].

Для практического использования интерес представляет вероятность наличия сговора для пар / троек ИНН, отсеянных из общей выборки по критерию превышения метрикой некоторого порогового уровня. Эта вероятность вычисляется на основе рисунков 4 и 6 как отношение номера сговора к месту в сортировке, тогда как различным

пороговым значениям соответствуют различные точки на кривых. Полученная таким образом зависимость достоверности выявления двусторонних сговоров от числа подозреваемых для соответствующих метрик показана на рисунке 7.

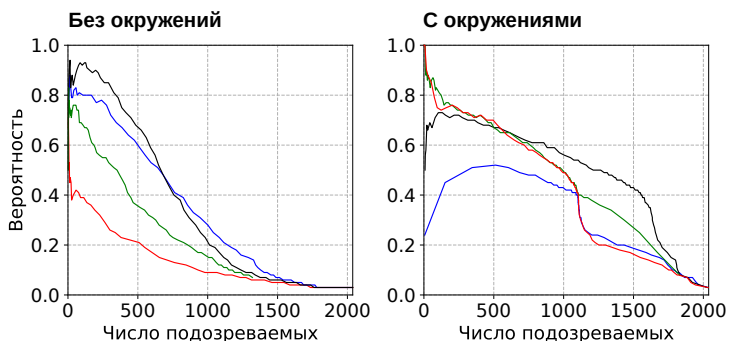


Рис. 7. Достоверность выявления двусторонних сговоров при отборе определённого числа подозреваемых пар с наибольшим значением метрик. Цветовая кодировка как на рисунках 3 и 4

Без учёта окружений и для небольшого числа подозреваемых наибольшую точность показывает метод на основе классического коррелятора (чёрная линия на левой панели). 350 сговоровившихся пар, т.е. 17% от полного числа двусторонних сговоров, например, можно выявить с достоверностью около 80%. В случае, когда требуется рекомендация по одной трети от полного числа сговоров и более, большую точность показывает метрика конкуренс с фазой $\Delta = 180^\circ$ (синяя линия). Методы на основе конкуренс с фазой $\Delta = 0^\circ$ и классической взаимной информации (красная и зеленая линии) менее достоверны.

При учёте окружений предпочтительность метрик меняется. Наибольшую точность выявления до четверти сговоров показывают методы на основе конкуренс с фазой $\Delta = 0^\circ$ и классической взаимной информации (красная и зеленая линии), причём для малого числа подозреваемых достоверность стремится к единице. Классический коррелятор, напротив, становится выгоден для большого числа подозреваемых, позволяя выявить две трети сговоров с достоверностью около 50%.

Аналогичный результат для тройных сговоров показан на рисунке 8. В соответствии с рисунком 6 квантово-теоретические метрики Q_3 и τ показывают сходную достоверность. Помимо отмеченного выше небольшого преимущества первой в диапазоне 15-20 тыс. сговоров, на

данном графике виден ещё один подобный интервал до 1.5 тыс. (6%) сговоров.

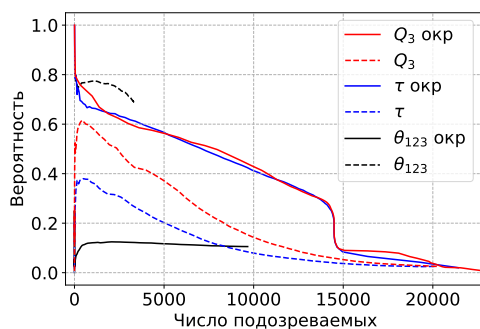


Рис. 8. Трёхсторонний аналог рисунка 7: достоверность выявления трёхсторонних сговоров при отборе определённого числа подозреваемых пар с наибольшим значением метрик. Цветовая кодировка как на рисунках 5 и 6

В отличие от квантовых метрик, классическая величина θ_{123} показывает лучшую достоверность без учёта окружений. В частности, до 2000 сговоров (10% от полного числа) могут быть выявлены с достоверностью более 75%. Более 3303 (14%) сговоров, однако, с помощью данной метрики выявить невозможно в силу обращения в ноль знаменателя в (18). В данном диапазоне оптимальными являются рассмотренные выше квантово-теоретические метрики.

Сговор с заказчиком. Помимо рассмотренного сговора между исполнителями, распространённым правонарушением является сговор между исполнителем и заказчиком. При этом заказчик обычно усложняет условия торга с тем, чтобы они подходили только для предполагаемого подрядчика. В силу открытости информации об участниках, выявление сговоров этого типа достаточно просто. Если в торге участвует, например, единственный подрядчик, то по законодательству РФ для заключения государственного контракта необходимо получить согласие ФАС, предусматривающее индивидуальное рассмотрение каждого случая в “ручном” режиме.

Именно поэтому участники сговора часто выбирают такой способ сокрытия своей договорённости, как привлечение других участников торгов. В целях создания видимости конкурентной ситуации, эти участники играют по заранее распределённым ролям, которые, как и состав участников, могут меняться [6]. Разработанный метод направлен именно на этот тип сговоров, наиболее трудный для обнаружения. Сговор

между подрядчиками при этом часто указывает и на сговор с заказчиком, что также требует отдельного рассмотрения в каждом случае.

5. Обсуждение. Представленные результаты подтверждают эффективность методов анализа данных для выявления скрытых особенностей коллективного поведения. Информация, извлекаемая из открытых данных с помощью сравнительно простых алгоритмов, может быть использована в целях информационного противоборства, конкурентной разведки и обеспечения общественной безопасности [68 – 73]. Разработанный алгоритм расширяет семейство таких методов на новую предметную область.

Сравнение классических и квантово-теоретических метрик согласованности поведения не выявило безусловной предпочтительности тех или других. Для двустороннего случая, в частности, оптимальный выбор зависит от доли подозреваемых и желаемой достоверности оценки (рисунок 7). Выявление трёхсторонних связей, напротив, в большинстве случаев оказалось значительно более точным при использовании квантовых метрик (рисунок 8). Данный результат указывает на перспективность дальнейшей разработки квантовых моделей коллективного поведения. Для моделирования запутанных состояний многочастичных ансамблей разработаны открытые библиотеки [74, 75], предоставляющие широкий выбор квантово-теоретических методов анализа данных.

В развитие представленного подхода интерес представляет, во-первых, возможность психологической интерпретации фазовых параметров поведенческих состояний (5), (6), в настоящей работе использовавшихся формально. Данные параметры могут быть использованы для кодирования эмоционально-смысловых состояний субъектов принятия решений [76], играющих важнейшую роль в коллективном поведении [77]. Как показано в задачах оценки смысловой близости языковых понятий [56] и семантического анализа численных данных [78], использование фазовых параметров позволяет существенно повысить точность результата. Данное направление таким образом сопрягается с задачами моделирования коллективных психологических процессов [42, 79, 80]. Семантическая интерпретация квантовой фазы, в частности, позволяет использовать для этой цели существующие модели и методы [81, 82]. Кроме того, интерес представляет введение в модель дополнительных параметров, описывающих субъектность акторов рассматриваемого поведения. Соответствующие степени свободы также имеются в квантовых моделях принятия решений.

Литература

1. Ferguson A. Policing predictive policing // *Washington University Law Review*. 2017. vol. 94. no. 5. p. 1109.
2. Yang F. *Oxford Research Encyclopedia of Criminology and Criminal Justice* // Oxford University Press. 2019. vol. 44. no. 1. pp. 57–61.
3. McDaniel J., Pease K. *Predictive Policing and Artificial Intelligence* // Routledge, Taylor Francis Group. 2021. 330 p.
4. Berk R. Artificial Intelligence, Predictive Policing, and Risk Assessment for Law Enforcement // *Annual Review of Criminology*. 2021. vol. 4. no. 1. pp. 209–237.
5. Официальный сайт Федеральная Антимонопольная служба. URL: fas.gov.ru (дата обращения: 02.09.2022).
6. Юрьев Р.Н., Алоджанц А.П. Проблема сговора участников торгов и пути ее решения в рамках парадигмы цифровой экономики с применением квантовой теории вероятностей // *Современная наука: актуальные проблемы теории и практики. Серия: Естественные и технические науки*. 2021. №. 10. С. 139–149.
7. Bajari P., Ye L. Deciding Between Competition and Collusion // *Review of Economics and Statistics*. 2003. vol. 85. no. 4. pp. 971–989.
8. Ballesteros-Perez P., Skitmore M., Das R., del Campo-Hitschfeld M. Quick Abnormal-Bid-Detection Method for Construction Contract Auctions // *Journal of Construction, Engineering and Management*. 2015. vol. 141. no. 7. p. 04015010.
9. Huber M., Imhof D. Machine learning with screens for detecting bid-rigging cartels // *International Journal of Industrial Organization*. 2019. vol. 65. pp. 277–301.
10. Garcia Rodriguez M., Rodriguez-Montequin V., Ballesteros-Perez P., Love P., Signor R. Collusion detection in public procurement auctions with machine learning algorithms // *Automation in Construction*. 2022. vol. 133. p. 104047.
11. Наследов А.Д. Математические методы психологического исследования. Анализ и интерпретация данных // *Речь*. 2004. 392 с.
12. Харченко М.А. Корреляционный анализ // *ВГУ*. 2008. 31с.
13. Марченко В.М., Можей Н.П., Шинкевич Е.А. Эконометрика и экономико-математические методы и модели // Минск: БГТУ. 2011. 157 с.
14. Кремер Н.Ш. Теория вероятностей и математическая статистика // М: Юнити. 2012.
15. Killworth P., Russell H. Informant Accuracy in Social Network Data III: A Comparison of Triadic Structure in Behavioral and Cognitive Data // *Social Networks*. 1979. vol. 2. pp. 19–46.
16. Morgenstern O., Schwodiauer G. Competition and collusion in bilateral markets // *Zeitschrift fur Nationalokonomie*. 1976. vol. 36. no. 3–4. pp. 217–245.
17. Thomas C., Wilson B. A Comparison of Auctions and Multilateral Negotiations // *The RAND Journal of Economics*. 2002. vol. 33. no. 1. p. 140.
18. Uddin S., Hossain L. Dyad and Triad Census Analysis of Crisis Communication Network // *Social Networking*. 2013. vol. 2. no. 1. pp. 32–41.
19. Holland P., Leinhardt S. The Statistical Analysis of Local Structure in Social Networks // *Sociological Methodology*. 1974. p. 45.
20. Martean L. The Triangle and the Eye inside the Circle: Dyadic and Triadic Dynamics in the Group // *Group Analysis*. 2014. vol. 47. no. 1. pp. 42–61.
21. Razmi P., Oloomi Buygi M., Esmalifalak M. A Machine Learning Approach for Collusion Detection in Electricity Markets Based on Nash Equilibrium Theory // *Group Analysis*. 2021. vol. 9. no. 1. pp. 170–180.

22. Ball P. The physical modelling of society: a historical perspective // *Physica A: Statistical Mechanics and its Applications*. 2002. vol. 314. no. 1–4. pp. 1–14.
23. Jorion P. Accounting for human activity through physics // *Cybernetics and Systems*. 2004. vol. 35. no. 2–3. pp. 275–284.
24. Galam S. Sociophysics. A Physicist's Modeling of Psycho-political Phenomena // Boston, MA: Springer US. 2012. p. 439.
25. Maldonado C.E. Quantum Theory and the Social Sciences // *Momento*. 2019. no. 59E. pp. 34–47.
26. Meghdadi A., Akbarzadeh-T., Javidan K. A Quantum-Like Model for Predicting Human Decisions in the Entangled Social Systems // *IEEE Transactions on Cybernetics*. 2022. pp. 1–11.
27. Meyer D. Quantum Strategies // *Physical Review Letters*. 1999. vol. 82. no. 5. pp. 1052–1055.
28. Eisert J., Wilkens M., Lewenstein M. Quantum Games and Quantum Strategies // *Physical Review Letters*. 1999. vol. 83. no. 15. pp. 3077–3080.
29. Marinatto L., Weber T. A quantum approach to static games of complete information // *Physics Letters A*. 2000. vol. 272. pp. 291–303.
30. Yukalov V., Yukalova E., Sornette D. Role of collective information in networks of quantum operating agents // *Physica A*. 2022. vol. 598. p. 127365.
31. Pothos E., Perry G., Corr P., Matthew M., Busemeyer J. Understanding cooperation in the Prisoner's Dilemma game // *Personality and Individual Differences*. 2011. vol. 51. no. 3. pp. 210–215.
32. Pelosse Y. The Intrinsic Quantum Nature of Nash Equilibrium Mixtures // *Journal of Philosophical Logic*. 2017. vol. 45. no. 1. pp. 25–64.
33. Baatique B. Quantum finance. Path Integrals and Hamiltonians for Options and Interest Rates // Cambridge. 1998.
34. Khrennikov A. Quantum-psychological model of the stock market // *Problems and Perspectives in Management*. 2003. pp. 136–148.
35. Bagarello F. Stock markets and quantum dynamics: A second quantized description // *Physica A: Statistical Mechanics and its Applications*. 2007. vol. 386. no. 1. pp. 283–302.
36. Choustova O. Quantum probability and financial market // *Information Sciences*. 2009. vol. 179. no. 5. pp. 478–484.
37. Goncalves C.P. Quantum financial economics - risk and returns // *Journal of Systems Science and Complexity*. 2013. vol. 26. no. 2. pp. 187–200.
38. Tahmasebi F., Meskinimood S., Namaki A., Vashghani Farahani S., Jalalzadeh S., Jafari G.R. Financial market images: A practical approach owing to the secret quantum potential // *EPL (Europhysics Letters)*. 2015. vol. 109. no. 3. p. 30001.
39. Orrell D. A quantum model of supply and demand // *Physica A: Statistical Mechanics and its Applications*. 2020. vol. 539. p. 122928.
40. Athalye V., Haven E. Socio-Economic Sciences: Beyond Quantum Math-like Formalisms // *Quantum Reports*. 2021. vol. 3. no. 4. pp. 656–663.
41. Khrennikov A. Social laser model: from color revolutions to Brexit and election of Donald Trump // *Kybernetes*. 2018. vol. 47. no. 2. pp. 273–288.
42. Tsarev D., Trofimova A., Alodjants A., Khrennikov A. Phase transitions, collective emotions and decision-making problem in heterogeneous social systems // *Scientific Reports*. 2019. vol. 9. no. 1. p. 18039.
43. Alodjants A., Bazhenov A., Khrennikov A., Bukhanovsky A. Mean-field theory of social laser // *Scientific Reports*. 2022. vol. 12. no. 1. p. 8566.

44. Словохотов Ю.Л. Физика и социофизика. Ч. 2. Сети социальных взаимодействий. Эконофизика // Проблемы управления. 2012. № 2. С. 2–31.
 45. Haven E., Khrennikov A. Quantum Social Science // NY: Cambridge University Press. 2013. 297 p.
 46. Orrell D. A Quantum Theory of Money and Value // Economic Thought. 2017. vol. 5. no. 2. pp. 19–28.
 47. Khrennikov A., Haven E. Quantum-like Modeling: from Economics to Social Laser // Asian Journal of Economics and Banking. 2020. vol. 4. no. 1. pp. 87–99.
 48. Orrell D. The value of value: A quantum approach to economics, security and international relations // Security Dialogue. 2020. vol. 51. no. 5. pp. 482–498.
 49. Суров И.А., Алоджанц А.П. Модели принятия решений в квантовой когнитивистике // СПб.: Университет ИТМО. 2018. 63 с.
 50. Peres A. Unperformed experiments have no results // American Journal of Physics. 1978. vol. 46. no. 7. pp. 745–747.
 51. Bell J. Against “measurement” // Physics World. 1990. vol. 3. pp. 32–41.
 52. Ballentine L. Propensity, Probability, and Quantum Theory // Foundations of Physics. 2017. vol. 46. no. 8. pp. 973–1005.
 53. Surov I. Quantum Cognitive Triad: Semantic Geometry of Context Representation // Foundations of Science. 2020. vol. 26. no. 4. pp. 947–975.
 54. Horodecki R., Horodecki P., Horodecki M., Horodecki K. Quantum entanglement // Reviews of Modern Physics. 2009. vol. 81. no. 2. pp. 865–942.
 55. Hill S., Wootters W. Entanglement of a Pair of Quantum Bits // Physical Review Letters. 1997. vol. 78. no. 26. pp. 5022–5025.
 56. Surov I., Semenenko E., Platonov A., Bessmertny I., Galofaro F., Toffano Z., Khrennikov A., Alodjants A. Quantum semantics of text perception // Scientific Reports. 2021. vol. 11. no. 1. p. 4193.
 57. Caves C., Fuchs C., Rungta P. Entanglement of Formation of an Arbitrary State of Two Qubits // Foundations of Physics Letters. 2001. vol. 14. no. 3. pp. 199–212.
 58. Wootters W. Entanglement of formation of an arbitrary state of two qubits // Physical Review Letters. 1998. vol. 80. no. 10. pp. 2245–2248.
 59. Vidal G., Werner R. Computable measure of entanglement // Physical Review A. 2002. vol. 65. no. 3. p. 032314.
 60. Vedral V. The role of relative entropy in quantum information theory // Reviews of Modern Physics. 2002. vol. 74. no. 1. pp. 197–234.
 61. Eisert J., Plenio M. A comparison of entanglement measures // Journal of Modern Optics. 1999. vol. 46. no. 1. pp. 145–154.
 62. Miranowicz A., Grudka A. A comparative study of relative entropy of entanglement, concurrence and negativity // Journal of Optics B: Quantum and Semiclassical Optics. 2004. vol. 6. no. 12. pp. 542–548.
 63. Верещагин Н.К., Шепин Е.В. Информация, кодирование и предсказание // М.: ФМОП МЦНМО. 2012. 236 с.
 64. Coffman V., Kundu J., Wootters W. Distributed entanglement // Physical Review A. 2000. vol. 61. no. 5.
 65. Gao X., Fei S., Wu K. Lower bounds of concurrence for tripartite quantum systems // Physical Review A. 2006. vol. 74. no. 5. pp. 1–9.
 66. Nakahara H., Amari S. Information-Geometric Measure for Neural Spikes // Neural Computation. 2002. vol. 14. no. 10. pp. 2269–2316.
 67. Fawcett T. An introduction to ROC analysis // Pattern Recognition Letters. 2006. vol. 27. no. 8. pp. 861–874.
- 438 Информатика и автоматизация. 2023. Том 22 № 2. ISSN 2713-3192 (печ.)
ISSN 2713-3206 (онлайн) www.ia.spcras.ru

68. Губанов Д.А., Новиков Д.А., Чхатишвили А. Социальные сети: модели информационного влияния, управления и противоборства // М: Физматлит. 2010. 228 с.
69. Vitali S., Glattfelder J., Battiston S. The Network of Global Corporate Control // PLoS ONE. 2011. vol. 6. no. 10.
70. Седаков Д., Филонов П. Разведка сетью: как система Avalanche помогает спецслужбам и бизнесу // Forbes. 2015.
71. Дорофеев А.В., Марков А.С. Структурированный мониторинг открытых персональных данных в сети интернет // Мониторинг правоприменения. 2017. № 18. С. 30–39.
72. Пилькевич С.В., Мажников П.В. Современные исследования в области мониторинга и анализа данных социальных сетей // Защита информации. Инсайд. 2018. № 70. С. 41–53.
73. Масалович А.И. Верона (англ. Verona) – компьютерная программа интеллектуального мониторинга сети Интернет и экспресс-анализа открытых данных № RU 2021660918 // 2021.
74. Johansson J., Nation P., Nori F. QuTiP: An open-source Python framework for the dynamics of open quantum systems // Computer Physics Communications. 2012. vol. 183. no. 8. pp. 1760–1772.
75. Aleksandrowicz G. et al. Qiskit: An Open-source Framework for Quantum Computing // Zenodo. 2019. DOI: 10.5281/zenodo.2562111.
76. Surov I.A. Quantum core affect. Color-emotion structure of semantic atom // Frontiers in Psychology. 2022. vol. 13.
77. Лебон Г. Психология народов и масс // Академический проект. 2021. 272 с.
78. Kozhisseri S., Surov I. Quantum-probabilistic SVD: complex-valued factorization of matrix data // Scientific and Technical Journal of Information Technologies, Mechanics and Optics. 2022. vol. 22. no. 3. pp. 567–573.
79. Гнидко К.О., Ломако А.Г. Моделирование Индивидуального и группового поведения субъектов массовой коммуникации в р-адических системах координат для индикации уровня контаминации сознания // Вопросы Кибербезопасности. 2017. № 15. С. 54–68.
80. Иванов О.С., Пилькевич С.В., Гнидко К.О., Лохвицкий В.А., Дудкин А.С., Сабиров Т.Р. Обоснование терминологического базиса исследований форм проявления контаминации психики человека // Вестник Российского нового университета. Серия: Сложные системы: модели, анализ и управление. 2019. С. 69–76.
81. Яньшин П.В. Цветосоциометрия. Исследование эмоционального состояния группы // Сборник научных трудов ученых Московского городского педагогического университета и Бакинского славянского университета. ред. Мыльников М.А. 2010. С. 278–288.
82. Петренко В.Ф. Основы психосемантики // М.: Эксмо. 2010. 480 с.

Семененко Евгений Константинович — аспирант, Университет ИТМО. Область научных интересов: информационный поиск, анализ текстов, когнитивно-поведенческое моделирование, квантовая семантика. Число научных публикаций — 6. semenenko.e.k@yandex.ru; Кронверкский проспект, 49А, 197101, Санкт-Петербург, Россия; р.т.: +7(931)247-7376.

Белопицкая Анна Геннадьевна — аспирант, Университет ИТМО. Область научных интересов: квантовые графы, теория рассеяния, спектральная теория линейных операторов. Число научных публикаций — 5. annabell1502@mail.ru; Кронверкский проспект, 49А, 197101, Санкт-Петербург, Россия; р.т.: +7(952)399-5142.

Юрьев Родион Николаевич — аспирант, Университет ИТМО. Область научных интересов: искусственный интеллект, обработка неструктурированных данных, экономика. Число научных публикаций — 8. juryev7@gmail.com; Кронверкский проспект, 49А, 197101, Санкт-Петербург, Россия; р.т.: +7(921)908-1432.

Алоджанц Александр Павлович — д-р физ.-мат. наук, профессор, Университет ИТМО. Область научных интересов: квантовая информация, квантовые алгоритмы, задачи принятия решения, анализ текстов и семантика. Число научных публикаций — 125. alexander_ap@list.ru; Кронверкский проспект, 49А, 197101, Санкт-Петербург, Россия; р.т.: +7(921)885-4734.

Бессмертный Игорь Александрович — д-р техн. наук, профессор, Университет ИТМО. Область научных интересов: искусственный интеллект, производственные системы, информационный поиск, компьютерная лингвистика. Число научных публикаций — 80. bessmertny@itmo.ru; Кронверкский проспект, 49А, 197101, Санкт-Петербург, Россия; р.т.: +7(812)233-2476.

Суров Илья Алексеевич — канд. физ.-мат. наук, доцент, Университет ИТМО. Область научных интересов: когнитивно-поведенческое моделирование, квантовая семантика. Число научных публикаций — 25. ilya.a.surov@itmo.ru; Кронверкский проспект, 49А, 197101, Санкт-Петербург, Россия; р.т.: +7(812)232-1467.

Поддержка исследований. Исследование выполнено при финансовой поддержке в рамках гос. задания № 2019-1339 Министерства науки и высшего образования РФ, а также гранта Фонда содействия инновациям № 569ГУЦЭС8-D3/62160.

E. SEMENENKO , A. BELOLIPETSKAYA , R. YURIEV , A. ALODJANTS ,
I. BESSMERTNY , I. SUROV
**DISCOVERY OF ECONOMIC COLLUSION BY METRICS OF
QUANTUM ENTANGLEMENT**

Semenenko E., Belolipetskaya A., Yuriev R., Alodjants A., Bessmertny I., Surov I. **Discovery of Economic Collusion by Metrics of Quantum Entanglement.**

Abstract. An effective economy requires prompt prevention of misconduct of legal entities. With the ever-increasing transaction rate, an important part of this work is finding market collusions based on statistics of electronic traces. We report a solution to this problem based on a quantum-theoretical approach to behavioral modeling. In particular, cognitive states of economic subjects are represented by complex-valued vectors in space formed by the basis of decision alternatives, while decision probabilities are defined by projections of these states to the corresponding directions. Coordination of multilateral behavior then corresponds to entanglement of the joint cognitive state, measured by standard metrics of quantum theory. A high score of these metrics indicates the likelihood of collusion between the considered subjects. The resulting method for collusion discovery was tested with open data on the participation of legal entities in public procurement between 2015 and 2020 available at the federal portal <https://zakupki.gov.ru>. Quantum models are built for about 80 thousand unique pairs and 10 million unique triples of agents in the obtained dataset. The reliability of collusion discovery was defined by comparison with open data of Federal antimonopoly service available at <https://br.fas.gov.ru>. The achieved performance allows the discovery of about one-half of known pairwise collusions with a reliability of more than 50%, which is comparable with detection based on classical correlation and mutual information. For three-sided behavior, in contrast, the quantum model is practically the only available option since classical measures are typically limited to the bilateral case. Half of such collusions are detected with a reliability of 40%. The obtained results indicate the efficiency of the quantum-probabilistic approach to modeling economic behavior. The developed metrics can be used as informative features in analytic systems and algorithms of machine learning for this field.

Keywords: collusion, cartel, decision making, quantum cognition, quantum entanglement, behavioral modeling, recommendation systems.

References

1. Ferguson A.G. Policing predictive policing. Washington University Law Review. 2017. vol. 94. no. 5. p. 1109.
2. Yang F. Oxford Research Encyclopedia of Criminology and Criminal Justice Oxford University Press. 2019. vol. 44. no. 1. pp. 57–61.
3. McDaniel J., Pease K. Predictive Policing and Artificial Intelligence Routledge, Taylor Francis Group. 2021. 330 p.
4. Berk R. Artificial Intelligence, Predictive Policing, and Risk Assessment for Law Enforcement Annual Review of Criminology. 2021. vol. 4. no. 1. pp. 209–237.
5. Oficial'nyj sayt Federal' naja Antimonopol' naja sluzhba. Available at: fas.gov.ru (accessed: 09.2022). (In Russ.).
6. Yuriev R.N., Alodjants A.P. [The problem of collusion of bidders and ways to solve it in the framework of digital economics paradogms and by using quantum probability theory]. Sovremennaya nauka: aktualnye problemy teorii i praktiki – Modern Science: Actual

- Problems of Theory and Practice. Series: Natural and technical sciences. 2021. no. 10. pp. 139–149. (In Russ.).
7. Bajari P., Ye L. Deciding Between Competition and Collusion. *Review of Economics and Statistics*. 2003. vol. 85. no. 4. pp. 971–989.
 8. Ballesteros-Perez P., Skitmore M., Das R., del Campo-Hitschfeld M. Quick Abnormal-Bid-Detection Method for Construction Contract Auctions. *Journal of Construction Engineering and Management*. 2015. vol. 141. no. 7. p. 04015010.
 9. Huber M., Imhof D. Machine learning with screens for detecting bid-rigging cartels. *International Journal of Industrial Organization*. 2019. vol. 65. pp. 277–301.
 10. Garcia Rodriguez M., Rodriguez-Montequin V., Ballesteros-Perez P., Love P., Signor R. Collusion detection in public procurement auctions with machine learning algorithms. *Automation in Construction*. 2022. vol. 133. p. 104047.
 11. Nasledov A.D. *Matematicheskie metody psihologicheskogo issledovanija. Analiz i interpretacija dannyh* [Mathematical methods of psychological research. Data analysis and interpretation]. Rech', 2004. p. 392. (In Russ.).
 12. Harchenko M.A. *Korreljacionnyj analiz* [Correlation analysis]. VGU, 2008. 31p. (In Russ.).
 13. Marchenko V.M., Mozhej N.P., Shinkevich E.A. *Jekometrika i jekonomiko-matematicheskie metody i modeli* [Econometrics and economic-mathematical methods and models]. Minsk: BGTU, 2011. p. 157. (In Russ.).
 14. Kremer N.Sh. *Teorija verojatnostej i matematicheskaja statistika* [Theory of Probability and Mathematical Statistics]. M: Yuniti. 2012. (In Russ.).
 15. Killworth P., Russell H. Informant Accuracy in Social Network Data III: A Comparison of Triadic Structure in Behavioral and Cognitive Data. *Social Networks*. 1979. vol. 2. pp. 19–46.
 16. Morgenstern O., Schwodiauer G. Competition and collusion in bilateral markets. *Zeitschrift fur Nationalokonomie*. 1976. vol. 36. no. 3–4. pp. 217–245.
 17. Thomas C., Wilson B. A Comparison of Auctions and Multilateral Negotiations. *The RAND Journal of Economics*. 2002. vol. 33. no. 1. p. 140.
 18. Uddin S., Hossain L. Dyad and Triad Census Analysis of Crisis Communication Network. *Social Networking*. 2013. vol. 2. no. 1. pp. 32–41.
 19. Holland P., Leinhardt S. The Statistical Analysis of Local Structure in Social Networks. *Sociological Methodology*. 1974. p. 45.
 20. Martean L. The Triangle and the Eye inside the Circle: Dyadic and Triadic Dynamics in the Group. *Group Analysis*. 2014. vol. 47. no. 1. pp. 42–61.
 21. Razmi P., Oloomi Buygi M., Esmalifalak M. A Machine Learning Approach for Collusion Detection in Electricity Markets Based on Nash Equilibrium Theory. *Group Analysis*. vol. 9. no. 1. pp. 170–180.
 22. Ball P. The physical modelling of society: a historical perspective. *Physica A: Statistical Mechanics and its Applications*. 2002. vol. 314. no. 1–4. pp. 1–14.
 23. Jorion P. Accounting for human activity through physics. *Cybernetics and Systems*. 2004. vol. 35. no. 2–3. pp. 275–284.
 24. Galam S. *Sociophysics. A Physicist's Modeling of Psycho-political Phenomena* Boston, MA: Springer US. 2012. p. 439.
 25. Maldonado C. *Quantum Theory and the Social Sciences Memento*. 2019. no. 59E. pp. 34–47.
 26. Meghdadi A., Akbarzadeh-T., Javidan K. A Quantum-Like Model for Predicting Human Decisions in the Entangled Social Systems *IEEE Transactions on Cybernetics*. pp. 1–11.

27. Meyer D.A. Quantum Strategies. *Physical Review Letters*. 1999. vol. 82. no. 5. pp. 1052–1055.
28. Eisert J., Wilkens M., Lewenstein M. Quantum Games and Quantum Strategies. *Physical Review Letters*. 1999. vol. 83. no. 15. pp. 3077–3080.
29. Marinatto L., Weber T. A quantum approach to static games of complete information. *Physics Letters A*. 2000. vol. 272. pp. 291–303.
30. Yukalov V., Yukalova E., Sornette D. Role of collective information in networks of quantum operating agents. *Physica A*. 2022. vol. 598. p. 127365.
31. Pothos E., Perry G., Corr P., Matthew M., Busmeyer J. Understanding cooperation in the Prisoner's Dilemma game. *Personality and Individual Differences*. vol. 51. no. 3. pp. 210–215.
32. Pelosse Y. The Intrinsic Quantum Nature of Nash Equilibrium Mixtures. *Journal of Philosophical Logic*. 2016. vol. 45. no. 1. pp. 25–64.
33. Baatique B.E. Quantum finance. *Path Integrals and Hamiltonians for Options and Interest Rates*. Cambridge. 1998.
34. Khrennikov A. Quantum-psychological model of the stock market Problems and Perspectives in Management. 2003. pp. 136–148.
35. Bagarello F. Stock markets and quantum dynamics: A second quantized description. *Physica A: Statistical Mechanics and its Applications*. 2007. vol. 386. no. 1. pp. 283–302.
36. Choustova O. Quantum probability and financial market *Information Sciences*. 2009. vol. no. 5. pp. 478–484.
37. Goncalves C.P. Quantum financial economics – risk and returns *Journal of Systems Science and Complexity*. 2013. vol. 26. no. 2. pp. 187–200.
38. Tahmasebi F., Meskinimood S., Namaki A., Vasheghani Farahani S., Jalalzadeh S., Jafari G. Financial market images: A practical approach owing to the secret quantum potential. *EPL (Europhysics Letters)*. 2015. vol. 109. no. 3. p. 30001.
39. Orrell D. A quantum model of supply and demand *Physica A: Statistical Mechanics and its Applications*. 2020. vol. 539. p. 122928.
40. Athalye V., Haven E. Socio-Economic Sciences: Beyond Quantum Math-like Formalisms *Quantum Reports*. 2021. vol. 3. no. 4. pp. 656–663.
41. Khrennikov A. Social laser model: from color revolutions to Brexit and election of Donald Trump *Kybernetes*. 2018. vol. 47. no. 2. pp. 273–288.
42. Tsarev D., Trofimova A., Alodjants A., Khrennikov A. Phase transitions, collective emotions and decision-making problem in heterogeneous social systems *Scientific Reports*. 2019. vol. 9. no. 1. p. 18039.
43. Alodjants A., Bazhenov A., Khrennikov A., Bukhanovsky A. Mean-field theory of social laser *Scientific Reports*. 2022. vol. 12. no. 1. p. 8566.
44. Slovohtov Ju.L. [Physics and sociophysics. Part 2. Networks of social interactions. *Econophysics*]. *Problemy upravleniya – Management issues*. 2012. no. 2. pp. 2–31. (In Russ.).
45. Haven E., Khrennikov A. *Quantum Social Science*. NY: Cambridge University Press.
46. Orrell D. A Quantum Theory of Money and Value *Economic Thought*. 2016. vol. 5. no. 2. pp. 19–28.
47. Khrennikov A., Haven E. Quantum-like Modeling: from Economics to Social Laser. *Asian Journal of Economics and Banking*. 2020. vol. 4. no. 1. pp. 87–99.
48. Orrell D. The value of value: A quantum approach to economics, security and international relations. *Security Dialogue*. 2020. vol. 51. no. 5. pp. 482–498.
49. Surov I.A., Alodjants A.P. Modeli prinjatija reshenij v kvantovoj kognitivistike [Decision models in quantum cognitive science]. SPb.: Universitet ITMO, 2018. 63 p. (In Russ.).

50. Peres A. Unperformed experiments have no results *American Journal of Physics*. 1978. vol. 46. no. 7. pp. 745–747.
51. Bell J.S. Against “measurement”. *Physics World*. 1990. vol. 3. pp. 32–41.
52. Ballentine L. Propensity, Probability, and Quantum Theory. *Foundations of Physics*. vol. 46. no. 8. pp. 973–1005.
53. Surov I. Quantum Cognitive Triad: Semantic Geometry of Context Representation. *Foundations of Science*. 2020. vol. 26. no. 4. pp. 947–975.
54. Horodecki R., Horodecki P., Horodecki M., Horodecki K. Quantum entanglement. *Reviews of Modern Physics*. 2009. vol. 81. no. 2. pp. 865–942.
55. Hill S., Wootters W. Entanglement of a Pair of Quantum Bits. *Physical Review Letters*. vol. 78. no. 26. pp. 5022–5025.
56. Surov I., Semenenko E., Platonov A., Bessmertny I., Galofaro F., Toffano Z., Khrennikov A., Alodjants A. Quantum semantics of text perception. *Scientific Reports*. vol. 11. no. 1. p. 4193.
57. Caves C., Fuchs C., Rungta P. Entanglement of Formation of an Arbitrary State of Two Rebits. *Foundations of Physics Letters*. 2001. vol. 14. no. 3. pp. 199–212.
58. Wootters W. Entanglement of formation of an arbitrary state of two qubits. *Physical Review Letters*. 1998. vol. 80. no. 10. pp. 2245–2248.
59. Vidal G., Werner R. Computable measure of entanglement. *Physical Review A*. 2002. vol. 65. no. 3. p. 032314.
60. Vedral V. The role of relative entropy in quantum information theory. *Reviews of Modern Physics*. 2002. vol. 74. no. 1. pp. 197–234.
61. Eisert J., Plenio M. A comparison of entanglement measures. *Journal of Modern Optics*. vol. 46. no. 1. pp. 145–154.
62. Miranowicz A., Grudka A. A comparative study of relative entropy of entanglement, concurrence and negativity. *Journal of Optics B: Quantum and Semiclassical Optics*. vol. 6. no. 12. pp. 542–548.
63. Vereshhagin N.K., Shhepin E.V. Информација, кодиrowanie i predskazanie [Information, coding and prediction]. M.: FMOP MCNMO, 2012. p. 236. (In Russ.).
64. Coffman V., Kundu J., Wootters W. Distributed entanglement. *Physical Review A*. 2000. vol. 61. no. 5.
65. Gao X. Fei S. Wu K. Lower bounds of concurrence for tripartite quantum systems. *Physical Review A*. 2006. vol. 74. no. 5.
66. Nakahara H., Amari S. Information-Geometric Measure for Neural Spikes. *Neural Computation*. 2002. vol. 14. no. 10. pp. 2269–2316.
67. Fawcett T. An introduction to ROC analysis. *Pattern Recognition Letters*. 2006. vol. 27. no. 8. pp. 861–874.
68. Gubanov D.A., Novikov D.A., Chhartishvili A. Social’nye seti: modeli informacionnogo vlijanija, upravljenija i protivoborstva [Social networks: models of information influence, control and confrontation]. M.: Fizmatlit. 2010. p. 228. (In Russ.).
69. Vitali S., Glattfelder J.B., Battiston S. The Network of Global Corporate Control. *PLoS ONE*. 2011. vol. 6. no. 10.
70. Sedakov D., Filonov P. Razvedka set’ju: kak sistema Avalanche pomogaet specslozhabam i biznesu [Network intelligence: how the Avalanche system helps intelligence agencies and businesses]. *Forbes*. 2015. (In Russ.).
71. Dorofeev A.V., Markov A.S. [Structured monitoring of open personal data on the Internet]. Monitoring pravoprimenenija – Law enforcement monitoring. 2017. no. 18. pp. 30–39. (In Russ.).

72. Pil'kevich S.V., Mazhnikov P.V. [Modern research in the field of monitoring and data analysis of social networks]. *Zashhita informacii. Insajd – Protection of information*. Inside. 2018. no. 70. pp. 41–53. (In Russ.).
73. Masalovich A.I. [Verona is a computer program for intelligent monitoring of the Internet and express analysis of open data No. RU 2021660918]. 2021. (In Russ.).
74. Johansson J., Nation P., Nori F. QuTiP: An open-source Python framework for the dynamics of open quantum systems. *Computer Physics Communications*. 2012. vol. 183. no. 8. pp. 1760–1772.
75. Aleksandrowicz G. et al. Qiskit: An Open-source Framework for Quantum Computing. *Zenodo*. 2019. DOI: 10.5281/zenodo.2562111.
76. Surov I.A. Quantum core affect. Color-emotion structure of semantic atom. *Frontiers in Psychology*. 2022. vol. 13.
77. Lebon G. [Psychology of peoples and masses]. *Akademicheskij proekt – Academic project*. 2021. p. 272. (In Russ.).
78. Kozhisseri S., Surov I. Quantum-probabilistic SVD: complex-valued factorization of matrix data. *Scientific and Technical Journal of Information Technologies, Mechanics and Optics*. 2022. vol. 22. no. 3. pp. 567–573.
79. Gnidko K.O., Lomako A.G. [Modeling Individual and Group Behaviors subjects of mass communication in p-adic coordinate systems for indicating the level of contamination of consciousness]. *Voprosy Kiberbezopasnosti – Issues of Cybersecurity*. 2017. no. 15. pp. 54–68. (In Russ.).
80. Ivanov O.S., Pil'kevich S.V., Gnidko K.O., Lohvickij V.A., Dudkin A.S., Sabirov T.R. [Substantiation of the terminological basis for research on the forms of manifestation of contamination of the human psyche]. *Vestnik Rossijskogo novogo universiteta. Serija: Slozhnye sistemy: modeli, analiz i upravlenie – Bulletin of the Russian New University. Series: Complex systems: models, analysis and control*. 2019. pp. 69–76. (In Russ.).
81. Yanshin P.V. [Study of the emotional state of the group]. *Sbornik nauchnyh trudov uchenyh Moskovskogo gorodskogo pedagogicheskogo universiteta i Bakinskogo slavyanskogo universiteta – Collection of scientific works of scientists of the Moscow City Pedagogical University and Baku Slavic University*. Ed.: Mylnikov M. 2010. pp. 278–288. (In Russ.).
82. Petrenko V.F. *Osnovy psihoosemantiki [Fundamentals of psychosemantics]*. M.: Eksmo. 2010. (In Russ.).

Semenenko Evgeny — Postgraduate student, ITMO University. Research interests: information retrieval, text analysis, cognitive behavioral modeling, quantum semantics. The number of publications — 6. semenenko.e.k@yandex.ru; 49A, Kronverksky Av., 197101, St. Petersburg, Russia; office phone: +7(931)247-7376.

Belolipetskaya Anna — Postgraduate student, ITMO University. Research interests: quantum graphs, scattering theory, spectral theory of linear operators. The number of publications — 5. annabell1502@mail.ru; 49A, Kronverksky Av., 197101, St. Petersburg, Russia; office phone: +7(952)399-5142.

Yuriev Rodion — Postgraduate student, ITMO University. Research interests: artificial intelligence, unstructured data processing, economics. The number of publications — 8. juryev7@gmail.com; 49A, Kronverksky Av., 197101, St. Petersburg, Russia; office phone: +7(921)908-1432.

Alodjants Alexander — Ph.D., Dr.Sci., Professor, ITMO University. Research interests: quantum information, quantum algorithms, decision-making tasks, text analysis and semantics. The number of publications — 125. alexander_ap@list.ru; 49A, Kronverksky Av., 197101, St. Petersburg, Russia; office phone: +7(921)885-4734.

Bessmertny Igor — Ph.D., Dr.Sci., Professor, ITMO University. Research interests: artificial intelligence, production systems, information retrieval, computational linguistics. The number of

publications — 80. bessmertny@itmo.ru; 49A, Kronverksky Av., 197101, St. Petersburg, Russia; office phone: +7(812)233-2476.

Surov Ilya — Ph.D., Associate professor, ITMO University. Research interests: cognitive-behavioral modeling, quantum semantics. The number of publications — 25. ilya.a.surov@itmo.ru; 49A, Kronverksky Av., 197101, St. Petersburg, Russia; office phone: +7(812)232-1467.

Acknowledgements. The work was supported by the Ministry of Science and Higher Education of the Russian Federation, grant № 2019-1339 and Fund for promotion of innovations, grant № 569ГУЦЭС8-D3/62160.

Е.Д. Кулаков, А.С. Михалев, А.В. Саренков, А.Д. Шуталев,
А.Е. ФЕДОРЕВ

АЛГОРИТМ КОРРЕКТИРОВКИ ПОЛОЖЕНИЯ КУСТОВЫХ ПЛОЩАДОК ПРИ РЕШЕНИИ ЗАДАЧИ РАЗРАБОТКИ НЕФТЯНЫХ МЕСТОРОЖДЕНИЙ

Кулаков Е.Д., Михалев А.С., Саренков А.В., Шуталев А.Д., Федорев А.Е. Алгоритм корректировки положения кустовых площадок при решении задачи разработки нефтяных месторождений.

Аннотация. Данная статья посвящена проблеме автоматизации этапа объединения скважин в кусты, рассматриваемого в рамках процесса проектирования разработки нефтяных месторождений. Решение задачи объединения скважин в кусты заключается в определении наилучшего расположения кустовых площадок и распределения скважин по кустам, при которых будут минимизированы затраты на разработку и обслуживание нефтяного месторождения, а ожидаемый дебит максимизирован. Одним из используемых на сегодняшний день подходов является применение оптимизационных алгоритмов. При этом данная задача влечет за собой учет технологических ограничений при поиске оптимального варианта разработки нефтяного месторождения, обоснованным в том числе действующими в отрасли регламентами, а именно минимальное и максимальное допустимое количество скважин в кусте, а также минимально допустимое расстояние между двумя кустовыми площадками. Использование алгоритмов оптимизации не всегда гарантирует оптимальный результат, при котором соблюдаются все заданные ограничения. В рамках данного исследования предложен алгоритм, который позволяет обрабатывать получаемые проектные решения с целью устранения нарушенных ограничений на этапе оптимизации. Алгоритм последовательно решает следующие проблемы: нарушение ограничений на сверхмалое и сверхбольшое количество скважин в кусте; несоответствие числа кустов с заданным; нарушение ограничения на сверхблизкое расположение кустов. Для исследования эффективности разработанного подхода был проведен вычислительный эксперимент на трех сгенерированных синтетических месторождениях с разной геометрией. В рамках эксперимента сравнивалось качество работы оптимизационного метода и предложенного алгоритма, который является надстройкой к оптимизационному. Сравнение проводилось на различных значениях мощности оптимизации, которое обозначает максимальное количество запусков целевой функции. Оценка качества работы сравниваемых подходов определяется величиной штрафа, которая обозначает степень нарушения значений основных ограничений. Критериями эффективности в данной работе являются: среднее значение, среднеквадратичное отклонение, медиана, минимальное и максимальное значения величины штрафа. За счет использования данного алгоритма величина штрафа для первого и третьего месторождений в среднем уменьшается соответственно до 0.04 и 0.03, а для второго месторождения алгоритм позволил получить проектные решения без нарушения ограничений. По результатам проведенного исследования сделано заключение относительно эффективности применения разработанного подхода при решении задачи разработки нефтяных месторождений.

Ключевые слова: нефтяное месторождение, разработка нефтяных месторождений, нефтяная скважина, кустовое бурение, нарушение проектных ограничений, корректировка положения кустовых площадок.

1. Введение. Согласно прогнозам, на ближайшие 20–30 лет производство нефти в России практически обеспечено уже разведанными запасами [1]. Однако по оценкам экспертов основная часть прироста нефтяных запасов приходится на ресурсы, неэффективные для разработки [2 – 5]. Поэтому в среднесрочном периоде Россия может столкнуться с проблемой поддержания достигнутых объемов добычи. Для разработки новых участков потребуется осваивать экстремальные климатические зоны, шельф и другие участки, удаленные от существующих объектов инфраструктуры, что повышает стоимость их вовлечения в разработку.

Современный энергетический рынок отличается нестабильностью и колебанием цен, связанных в первую очередь с геополитической борьбой лидирующих держав планеты [6 – 8] и неблагоприятной эпидемиологической обстановкой в мире [9]. При этом полноценной альтернативы нефти нет, и она еще долго будет одним из важнейших ресурсов для функционирования и развития человечества.

По запасам нефти и газа Россия опережает другие страны, однако отстает по объему геологоразведочных работ и техническому оснащению. С учётом меняющейся структуры запасов для сохранения рентабельности нефтедобывающие компании должны изыскивать возможности сокращения издержек и повышения своей эффективности. Ускорить процесс добычи нефти поможет только инвестирование средств в масштабную геологоразведку и полная цифровизация нефтедобывающей отрасли [10 – 14].

Под процессом цифровизации в данном случае имеется в виду автоматизация всех составляющих этапов разработки и освоения нефтяных месторождений. Внедрение цифровых решений в нефтедобывающую отрасль позволит значительно повысить эффективность производственных процессов в компании за счет оптимизации и анализа данных, а также сократить влияние человеческого фактора на принимаемые проектные решения.

Потенциальный прирост извлекаемых запасов нефти на территории России может составить 6,8 миллиардов тонн за счет технологического развития отрасли [15]. Цифровизация нефтедобывающей отрасли может позволить нарастить добычу к 2035 году до 607 миллионов тонн с учетом экономических и инфраструктурных ограничений [16].

На реализацию проектных решений по обустройству нефтяных месторождений может приходиться до 80% и более общего объема капитальных вложений. Очевидно, что в таком случае

при проектировании требуется уделять внимание поиску оптимальной схемы разбуривания месторождения.

В связи с этим на этапе цифровизации нефтяной отрасли первоочередной становится задача автоматизации этапа планирования нефтяных месторождений за счет внедрения технологий интеллектуального анализа данных, машинного обучения и методов оптимизации [17 – 19].

В рамках данной научной работы рассматривается наиболее распространённый сценарий разработки месторождения УВС – с использованием кустового бурения. Далее будет представлено описание кустового бурения.

2. Описание кустового бурения. Решение задачи разработки нефтяных месторождений заключается в том числе в выборе наиболее эффективного метода разбуривания. Среди существующих методов наиболее эффективным на сегодняшний день является кустовое бурение [20 – 24].

Кустовое бурение представляет собой сооружение скважин (в основном наклонно-направленных), устья которых группируются на близком расстоянии друг от друга на общей ограниченной площадке (основании).

Применение метода кустового бурения имеет следующие преимущества [25 – 27]:

- сокращение материальных и трудовых затрат на инженерное обустройство площадок под скважины;
- уменьшение объемов перевозок;
- повышение рентабельности разработки;
- упрощение процессов добычи и ее автоматизации;
- уменьшение техногенного воздействия на окружающую природную среду.

Используют метод кустового бурения в таких природно-климатических условиях, при которых строительство и обслуживание на значительном удалении друг от друга нефтяных сооружений, коммуникаций и дорог будет стоить сверхбольших капиталовложений [28]. Так, например, кустовое бурение используется в условиях моря на территории шельфовых месторождений. Также указанный выше метод используют в горной, лесной и болотистой местностях.

Обычно при применении метода кустового бурения сталкиваются со следующими основными проблемами:

- определение количества кустовых площадок;
- определение расположения кустовых площадок и распределения скважин между ними.

Число скважин в кусте может достигать 24 штук. Исходя из минимального и максимального количества скважин в кусте и общего количества скважин на территории месторождения, эксперт определяет допустимый диапазон количества кустовых площадок. Но зачастую при решении задачи расположения кустовых площадок и распределения скважин между ними приходится сталкиваться с такими схемами разработки месторождения, при которых их реализация либо будет стоить в разы больше капиталовложений, либо будет вообще невозможна. В связи с этим появляется необходимость определения оптимального диапазона количества кустовых площадок для различных схем разработки месторождений.

При проектировании нефтяных месторождений приходится формировать значительное число исходных проектных вариантов, из которых впоследствии выбирается оптимальный. Поэтому в таких случаях требуется дополнять знания, опыт и интуицию проектировщиков большим объемом геолого-промысловой информации и значительным объемом расчетных данных, что приводит к необходимости применения средств автоматизированного проектирования.

В следующем разделе будет приведен литературный обзор существующих решений задачи кустования скважин.

3. Литературный обзор. На сегодняшний день задаче автоматизации размещения кустовых площадок и распределения скважин между ними посвящено мало работ.

В статье [29] предложены две модели для решения поставленной задачи, которые отличаются наборами исходных данных. В первой модели задача размещения кустовых площадок сводится к задаче целочисленной оптимизации. Во второй модели алгоритм поиска наилучшего расположения кустовых площадок представляет собой итерационную процедуру. При этом для решения задачи предлагается использовать метод, основанный на «лангранжевой релаксации» [30*]. В основе обеих моделей заложена идея минимизации суммарного квадратичного отклонения скважин от кустовых площадок, представленная в дискретном виде.

В работе [30] предлагается методика кустования скважин для различных вариантов систем разработки. Задачу определения положения кустовых площадок предлагается решать, как задачу глобальной оптимизации. В качестве критериев оптимальности авторы статьи используют минимизацию суммарной проходки при бурении и технико-экономический показатель, который учитывает потенциал добывающих скважин и стоимостные характеристики объектов.

Авторы работы [31] предлагают для решения поставленной задачи метод оптимизации на основе интеллектуальной последовательной выборки, согласно которой отбор элементов генеральной совокупности проводится последовательно, на каждом этапе собирается и анализируется информация и принимается решение о дополнительном отборе элементов генеральной совокупности. Предыдущие вычисления алгоритма учитываются при выборе следующей выборки, что в свою очередь позволяет обеспечить более высокую сходимость по сравнению с аналогами. В качестве целевой функции использовался чистый дисконтированный доход за 5 лет.

Авторы статьи [24] предлагают подход определения оптимального распределения скважин между кустами на основе перебора группировок всех скважин. При отборе наилучшего варианта кустования скважин ключевыми показателями являются:

- объем добычи нефти;
- накопленный объем добычи нефти;
- чистый дисконтированный доход.

В статье [32] задача кустования скважин сводится к задаче кластеризации. Предложенный подход представляет собой сочетание классического алгоритма кластеризации и метода оптимизации. В экспериментах в качестве метода оптимизации использовались генетический алгоритм и метод роя частиц. В данном подходе также применяется метод поиска ассоциативных правил для выявления закономерностей между скважинами для выявления наиболее сильных кустов.

Авторы статьи [33] используют модифицированный алгоритм кластеризации k -средних для нахождения оптимальной схемы кустования скважин. В качестве основного критерия оптимальности для алгоритма кластеризации выступает минимизация суммарной проходки наклонных и горизонтальных скважин с учетом сложности их траекторий.

Рассмотренные подходы для решения поставленной задачи обычно заключаются в сведении её к оптимизационному виду. При выборе оптимального размещения кустовых площадок учитываются следующие возможные показатели:

- суммарное квадратичное отклонение скважин от кустовых площадок;
- суммарная проходка при бурении;
- чистый дисконтированный доход;
- объем добычи нефти.

Рассмотренные показатели являются только технико-экономическими, при этом решение данного класса задач требует необходимости учитывать конструкционные особенности. В рассматриваемой статье предлагается подход, учитывающий также и основные технологические ограничения к поставленной задаче.

В следующем разделе будет приведена постановка задачи поиска оптимальной схемы размещения кустовых площадок и распределения скважин между ними при заданных ограничениях.

4. Постановка задачи. Задача кустования нефтяных скважин является частным случаем задачи группирования объектов при наличии ограничений [34]. При этом целевая функция представляет собой сумму расстояний между скважинами и основаниями кустов, к которым они принадлежат:

$$I = \sum_{k=1}^K \sum_{i=1}^{|S_k|} \left((x_i, y_i) - (x_c^k, y_c^k) \right)^2, \quad (1)$$

где K – количество кустов;

s_k – k -ый куст, содержащий $|s_k|$ скважин;

(x_i, y_i) – координаты i -ой скважин;

(x_c^k, y_c^k) – координаты кустовой площадки k -ого куста.

К основным ограничениям данной задачи, которые определяются сводом правил обустройства нефтяных месторождений [35], относятся:

- минимальное и максимальное допустимое количество скважин в кусте;
- максимальная величина отхода;
- минимально допустимое расстояние между двумя кустовыми площадками.

На данный момент, указанные выше ограничения, учитываются при помощи метода штрафных функций [36, 37], который преобразует исходную задачу с ограничениями в задачу безусловной оптимизации. Нарушение какого-либо из ограничений приводит к созданию штрафа, который добавляется к значению целевой функции.

В рамках данного исследования значение штрафа, обозначающая степень нарушения значений основных ограничений будет иметь следующий вид:

$$\varphi = \alpha \cdot \frac{V_{min} + V_{max} + V_{between} + V_{dist}}{K}, \quad (2)$$

где K – количество кустов;

V_{min} – количество кустов, у которых число скважин меньше, чем минимально-допустимое;

V_{max} – количество кустов, у которых число скважин больше, чем максимально-допустимое;

$V_{between}$ – количество кустов, нарушающих ограничение на минимально допустимое расстояние между двумя кустовыми площадками;

V_{dist} – количество кустов, имеющие скважины, нарушающих ограничение на максимальную величину отхода;

α – положительная величина, $\alpha \geq 1$.

Данная оптимизационная задача относится к категории NP трудных задач [38]. Основная проблема таких задач заключается в том, что они не могут быть решены за полиномиальное время. Следовательно, для крупных месторождений решение задачи кустования с помощью оптимизационных алгоритмов не гарантирует получение результата без нарушения указанных ограничений. Таким образом, возникает необходимость модификации существующих возможных подходов к решению данной задачи. В рамках данной статьи предлагается подход к обработке результатов работы оптимизационных методов на основе итеративного алгоритма устранения нарушения заданных ограничений.

5. Предлагаемое решение. Алгоритм выполняет корректировку взаимного расположения кустовых площадок на области месторождения с целью устранения возникших нарушений проектных ограничений, связанных с расположением кустовых площадок и скважин внутри них. Данный подход работает с проблемными кустами и их соседями. Под проблемным понимается куст, у которого нарушается хотя бы одно ограничение из следующего списка:

- минимальное и максимальное допустимое количество скважин в кусте;
- минимально допустимое расстояние между двумя кустовыми площадками.

Данный подход не обрабатывает случаи нарушения ограничения на величину отхода, поскольку выполнение остальных ограничений и правильно заданные пользователем параметры системы разработки автоматически исключают нарушение данного ограничения.

Для работы предлагаемого решения необходим следующий набор параметров:

- информация о расположении скважин;
- значения пользовательских ограничений;
- текущее положение кустовых площадок;
- информация по принадлежности скважин к кустовым площадкам;
- матрица расстояний между кустовыми площадками;
- матрица смежности кустовых площадок.

На основе данных параметров процесс устранения нарушения заданных ограничений может быть представлен в виде алгоритма корректировки положения проблемных кустовых площадок и их соседей. Псевдокод алгоритма представлен в алгоритме 1.

Алгоритм 1. Алгоритм корректировки положения кустовых площадок и их соседей

Вход: информация о расположении скважин *wells*, пользовательские ограничения *rest*, текущее положение кустовых площадок *bushes*, информация по принадлежности скважин к кустовым площадкам *clusters*, матрица расстояний между кустовыми площадками *distance*, матрица смежности кустовых площадок *adj*, максимальное количество итераций *max_iter*

Выход: положение кустовых площадок *bushes*, удовлетворяющее пользовательским ограничениям

Function EliminateClusterViolations(*wells, bushes, clusters, distance, adj, max_iter*)

iteration $\leftarrow$ 0

▷ *current_best* – значение метрики для заданного *bushes*

current_best $\leftarrow$ TargetFunction(*bushes, rest, clusters*)

▷ EstimateBestResult – процедура определения лучшего расположения кустовых площадок исходя из значения метрики

best $\leftarrow$ EstimateBestResult(*current_best, None*)

▷ ChangePosition – процедура переопределения положения кустовых площадок исходя из центра масс скважин в каждом кусте

bushes $\leftarrow$ ChangePosition(*wells, bushes*)

▷ *deque* – очередь проблемных кустов, у которых нарушения по количеству скважин

while True **do**

while len(*deque*) $\neq$ 0 **do**

▷ *intruder* – первый в очереди проблемный куст

▷ *intruder.size* – количество скважин в кусте *intruder*

▷ *rest.min* – минимально-допустимое количество скважин в кусте

if *intruder.size* < *rest.min* **then**

```

▷ EliminateMinCountWells – метод устранения нарушения на сверхмалое
число скважин.
    bushes ← EliminateMinCountWells(wells, rest, bushes, clusters,
adj, intruder)
▷ rest.max – максимально-допустимое количество скважин в кусте
    else if intruder.size > rest.max then
▷ EliminateMaxCountWells – метод устранения нарушения на сверхбольшое
число скважин
    bushes ← EliminateMaxCountWells(wells, rest, bushes, clusters,
adj, intruder)
    bushes ← ChangePosition(wells, bushes)
    current_best ← TargetFunction(bushes, rest, clusters)
    best ← EstimateBestResult(current_best, best)
▷ RecoveryCountClusters – процедура приведения числа кустов к
исходному
    bushes ← RecoveryCountClusters(bushes, rest, clusters)
▷ EliminateExtraNearbyBushes – метод устранения нарушения
сверхблизости кустов
    bushes ← EliminateExtraNearbyBushes(bushes, rest, adj, clusters)
    current_best ← TargetFunction(bushes, rest, clusters)
    best ← EstimateBestResult(current_best, best)
▷ best.count_violations – общее число нарушений у расположения
кустовых площадок
    if best.count_violations = 0 then
        break
    if iteration ≥ max_iter then
        break
▷ определение новой очереди проблемных кустов deque исходя из bushes
    iteration ← iteration + 1
return best

```

В качестве вспомогательных методов выступают:

- процедура расчета целевой функции TargetFunction, которая была описана в формуле (1);
- процедура оценки оптимальности полученного решения EstimateBestResult исходя из значения целевой функции;
- функция переопределения положения кустовых площадок ChangePosition, использующая расчет центра масс скважин;
- алгоритмы устранения нарушения на количество скважин в кусте: EliminateMinCountWells (алгоритм 2), EliminateMaxCountWells (алгоритм 3);
- процедура приведения количества кустовых площадок к заданному числу RecoveryCountWellPads (алгоритм 4);

– алгоритм устранения нарушения на сверхблизость кустовых площадок *EliminateNearbyWellPads* (алгоритм 5).

Алгоритм 2. Алгоритм устранения нарушения на сверхмалое количество скважин в кусте

Вход: информация о расположении скважин *wells*, пользовательские ограничения *rest*, текущее положение кустовых площадок *bushes*, информация по принадлежности скважин к кустовым площадкам *clusters*, матрица смежности кустовых площадок *adj*, номер проблемного куста *intruder* **Выход:** новое положение кустовых площадок *bushes*, в котором *intruder* не имеет нарушения

Function *EliminateMinCountWells(wells, rest, bushes, clusters, adj, intruder)*

```

▷ intruder.size – количество скважин в кусте intruder
  ▷ count_add – количество скважин, которые надо добавить к intruder
    count_add ← rest.min – intruder.size
  ▷ neighbours – список соседей куста intruder, у которых нет нарушения
    сверхмалого числа скважин
  ▷ count_donor – количество скважин у соседей intruder, которые
    возможно передать ему без нарушения ограничений
    count_donor ← – len(neighbours) * rest.min
  for i ← 0 to len(neighbours) – 1 do
    count_donor ← count_donor + neighbours[i].size
  if count_donor ≥ count_add then
    ▷ проблема решается изъятием скважин у соседей
    ▷ candidate_wells – список скважин кустов из neighbours
    for i ← 0 to len(candidate_wells) – 1 do
      if count_add = 0 then
        break
      ▷ index_bush – куст для candidate_wells[i] скважины
      if index_bush.size – rest.min ≥ 1 then
        ▷ скважина candidate_wells[i] принадлежит кусту intruder
        count_add ← count_add – 1
      ▷ переопределение положения куста intruder на основе центра масс
    скважин
  else
    ▷ проблема решается объединением куста нарушителя с одним из
    его соседей
    ▷ neighbours – список соседей куста intruder, у которых нет
    нарушения сверхбольшого числа скважин
  ▷ index_bush – номер куста из neighbours, от которого минимальное
  расстояние до intruder
    ▷ объединение кустов index_bush и intruder
    ▷ определение положения нового куста на основе центра масс
    скважин
  return bushes

```

Алгоритм 3. Алгоритм устранения нарушения на сверхбольшое количество скважин в кусте

Вход: информация о расположении скважин *wells*, пользовательские ограничения *rest*, текущее положение кустовых площадок *bushes*, информация по принадлежности скважин к кустовым площадкам *clusters*, матрица смежности кустовых площадок *adj*, номер проблемного куста *intruder* **Выход:** новое положение кустовых площадок *bushes*, в котором *intruder* не имеет нарушения

Function EliminateMaxCountWells (*wells*, *rest*, *bushes*, *clusters*, *adj*, *intruder*)

```

    ▷ intruder.size – количество скважин в кусте intruder
    ▷ count_sub – количество скважин, которые надо добавить к intruder
    count_sub ← intruder.size – rest.max
    ▷ проблема решается раздачей скважин куста intruder его соседям
    ▷ recipient_wells – отсортированный список скважин куста intruder по
    возрастанию расстояния до кустовой площадки
    for i ← 0 to len(recipient_wells) – 1 do
        if count_sub = 0 then
            break
        ▷ neighbours – список соседей куста intruder, у которых нет нарушения
        сверхбольшого числа скважин
        ▷ recipient_wells[i].distance – расстояние от скважины до кустовой площадки
        intruder
        ▷ candidate_owner – ближайший из списка соседей neighbours к скважине
        recipient_wells[i]
        ▷ distance – расстояние от candidate_owner до recipient_wells[i]
        if distance < 1.1 * recipient_wells[i].distance then
            ▷ скважина recipient_wells[i] принадлежит кусту candidate_owner
            count_sub ← count_sub – 1
        if intruder.size > rest.max then
            ▷ проблема решается разделением куста intruder на два отдельных куста
            ▷ C – центр масс куста
            ▷ D – наиболее удаленная скважина от C
            ▷ Расчет уравнения прямой CD
            ▷ Расчет уравнения перпендикуляра прямой CD, проходящего в точке C
            ▷ На основании уравнения перпендикуляра к CD разделение скважин на две
            группы.
            ▷ В каждой группе расчет положения кустовой площадки на основе центра
            масс скважин
    return bushes

```

После выполнении алгоритмов устранения нарушения на сверхбольшое и сверхмалое количество скважин в кусте не исключено, что количество кустов на месторождении может отличаться от заданного. На этот случай, если для пользователя очень важным является установленное значение количества кустов, была разработана процедура приведения количества кустовых площадок к заданному. Предлагаемый подход решает проблему несоответствия текущего числа кустов и заданного. Его псевдокод представлен в алгоритме 4.

Алгоритм 4. Процедура приведения количества кустовых площадок к заданному

Вход: заданное число кустовых площадок $original_N$, текущее число кустовых площадок $current_N$, текущее положение кустовых площадок $bushes$, матрица смежности кустовых площадок adj

Выход: новое положение кустовых площадок $bushes$, соответствующее заданному количеству кустов

Function RecoveryCountWellPads ($original_N$, $current_N$, $bushes$, adj)

$count \leftarrow original_N - current_N$

if $count > 0$ **then**

▷ проблема решается добавлением новых кустов

▷ $graph$ – неориентированный граф на основе матрицы смежности adj

▷ Поиск в $graph$ циклов, каждый из которых имеет минимум 4 вершины и не содержит вершин внутри себя

▷ $cycles$ – список циклов, найденных в $graph$

if $len(cycles) < count$ **then**

▷ для каждого цикла вычисляется его центр масс координат его вершин

▷ $mas_centers$ – пустой список, размерности $len(cycles)$

for $i \leftarrow 0$ **to** $len(cycles) - 1$ **do**

▷ $position$ – центр масс координат вершин $cycles[i]$

$mas_centers[i] \leftarrow position$

▷ $mas_centers$ – список, каждый элемент которого центр масс соответствующего цикла из $cycles$

for $i \leftarrow 0$ **to** $count - len(cycles) - 1$ **do**

▷ $position$ – случайно сгенерированное положение куста на месторождении

▷ к $mas_centers$ добавляется $position$

else

▷ для каждого цикла вычисляется показатель целесообразности добавления в центре него новой кустовой площадки по формуле 2

▷ сортировка списка $cycles$ по убывания показателя целесообразности (формула 2)

▷ $mas_centers$ – пустой список, размерности $count$

for $i \leftarrow 0$ **to** $count - 1$ **do**

▷ $position$ – центр масс координат вершин $cycles[i]$

$mas_centers[i] \leftarrow position$

▷ UnionLists – функция объединения двух или более списков

$bushes \leftarrow UnionLists(bushes, mas_centers)$

else

▷ проблема решается удалением лишних кустов

▷ для каждого куста вычисляется показатель значимости кустовой площадки по формуле 3

▷ сортировка списка $bushes$ по возрастанию показателя значимости (формула 3)

for $i \leftarrow 0$ **to** $1 - count$ **do**

▷ из $bushes$ удаляем первый элемент

return $bushes$

В рамках представления процедуры приведения количества кустовых площадок к заданному были предложены критерии отбора циклов вершин и кустовых площадок.

Показатель целесообразности добавления в центре его новой кустовой площадки имеет следующий вид:

$$metrics_{cycle} = \frac{1}{n+1} \sum_{i=1}^n size_i, \quad (3)$$

где n – число кустов в цикле;

$size_i$ – количество скважин в i -ом кусте цикла.

Показатель значимости кустовой площадки для нефтяного месторождения в рамках отсеивания лишних кустов определяется следующим образом:

$$metrics_{bush} = \frac{1}{n} \left(\sum_{i=1}^n size_i + size_j \right), \quad (4)$$

где j – номер куста;

n – число соседей j -го куста;

$size_j$ – количество скважин в j -ом кусте цикла.

Следующая компонента предлагаемого подхода решает проблему сверхблизости кустов. Данная проблема решается двумя способами:

- процедура взаимного отталкивания проблемных кустов до приемлемого расстояния между ними;
- процедура определения нового положения проблемных кустов.

Предлагаемый алгоритм вначале опирается на первую процедуру и в случае нарушения сверхблизости кустов с соседями запускает вторую процедуру. Во второй процедуре используется расстояние Хаусдорфа [39]. Псевдокод данной компоненты представлен в алгоритме 5.

Алгоритм устранения нарушения на сверхблизость кустовых площадок строится на использовании области для размещения пары кустов без нарушения ограничений. Данная область представляет собой геометрическое место точек, в котором расстояние от кустов из списка соседей текущей пары проблемных кустов для каждой точки будет больше, чем минимально-допустимое расстояние между парой кустовых площадок *distance*. Зачастую без дополнительных

ограничений такая область будет не применима для дальнейшего использования в силу своей неограниченности.

Для пресечения описанного выше случая используются следующие дополнительные ограничения:

- любая точка должна быть в пределах области месторождения;

- пусть два куста P и Q из списка соседей являются соседями между собой, а расстояние PQ больше, чем $2 * distance$. C – середина отрезка между кустовыми площадками текущей пары проблемных кустов. Тогда, при разделении плоскости прямой PQ , любая точка из ГМТ должна быть в той же полуплоскости, что и вершина C .

Алгоритм 5. Алгоритм устранения нарушения на сверхблизость кустовых площадок

Вход: информация о расположении скважин *wells*, пользовательские ограничения *rest*, текущее положение кустовых площадок *bushes*, информация по принадлежности скважин к кустовым площадкам *clusters*, матрица смежности кустовых площадок *adj*

Выход: новое положение кустовых площадок *bushes*, без нарушения на сверхблизость кустовых площадок

Function EliminateNearbyWellPads (*wells, rest, bushes, clusters, adj*)

▷ *deque* – очередь проблемных пар кустов, у которых нарушения на сверхблизость

while len(*deque*) ≠ 0 **do**

▷ *tuple* – текущая пара проблемных кустов

▷ A и B – координаты кустов из *tuple*

$C \leftarrow (A+B)/2$

▷ e – единичный вектор прямой AB

▷ расчет положения точек D и E на прямой AB

▷ *rest.distance* – минимально-допустимое расстояние между парой кустовых площадок

$D \leftarrow C + e * rest.distance/2$

$E \leftarrow C - e * rest.distance/2$

▷ *neighbours* – все соседние кусты для A и B , исходя из матрицы смежности *adj*

▷ расчет все расстояний *mas_distances* от D и E до кустов из *neighbours*

▷ All – оператор and для булевой последовательности

if All(*mas_distance* ≥ *rest.distance*) **then**

▷ точки D и E – новые положения кустов A и B

else

▷ расчет области *region*, в которой можно разместить кустовые площадки из текущей пары без нарушения ограничения

▷ *points* – список точек границы *region*

```

for  $i \leftarrow 0$  to  $\text{len}(\text{points}) - 1$  do
    ▷ расчет расстояния Хаусдорфа  $\text{distance}$  между  $\text{points}[i]$  и
    границей  $\text{region}$ 
    if  $\text{distance} \geq \text{rest.distance}$  then
         $R \leftarrow \text{points}[i]$ 
        ▷ точка  $T$  – точка на границе  $\text{region}$ ,  $|RT| = \text{distance}$ 
        break
     $S \leftarrow (R+T)/2$ 
    ▷  $e$  – единичный вектор прямой  $RT$ 
    ▷ расчет положения точек  $X$  и  $Y$  на прямой  $RT$ 
     $X \leftarrow S + e * \text{rest.distance}/2$ 
     $Y \leftarrow S - e * \text{rest.distance}/2$ 
    ▷ точки  $X$  и  $Y$  – новые положения кустов  $A$  и  $B$ 
return  $\text{bushes}$ 
    
```

Пример такой области с дополнительными ограничениями представлен на рисунке 1.

После завершения описания предложенного подхода отдельных пояснений требует способ нахождения матрицы смежности кустовых площадок. Итеративная схема используемой в работе процедуры расчета матрицы смежности кустовых площадок представлена ниже.

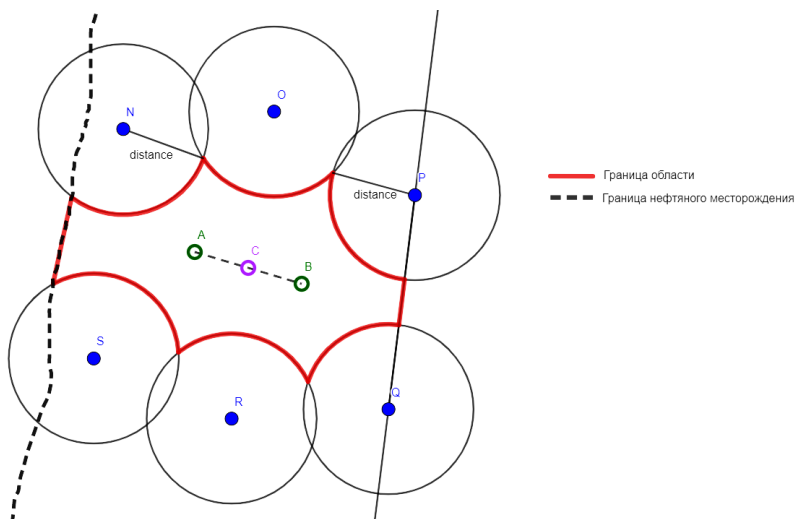


Рис. 1. Пример области, в которую можно разместить пару кустовых площадок

Шаг 1: Расчет матрицы расстояний между кустовыми площадками *distance*, у которой каждый элемент рассчитывается следующим образом:

$$distance_{i,j} = \left| (x_c^i, y_c^i) - (x_c^j, y_c^j) \right|, \quad i, j \in [1, K], \quad (5)$$

где K – количество кустов;

(x_c^i, y_c^i) – координаты кустовой площадки i -ого куста.

Шаг 2: Расчет расстояния между кустовыми площадками и их самыми удаленными скважинами (Далее – радиус зоны влияния куста) *radius*:

$$radius_i = \max_l \left| (x_c^i, y_c^i) - (x_l, y_l) \right|, \quad i \in [1, K], \quad (6)$$

где K – количество кустов;

(x_c^i, y_c^i) – координаты кустовой площадки i -ого куста;

(x_l, y_l) – координаты l -ой скважины, $l \in [1, |S_i|]$;

S_i – i -ый куст, содержащий $|S_i|$ скважин.

Шаг 3: Обход каждой пары кустов. Если *condition* (7) больше или равен 0, то пара кустов i и j являются соседями, в ином случае пара кустов i и j не являются соседями.

$$condition = radius_i + radius_j - \lambda \cdot distance_{i,j}, \quad (7)$$

где $radius_i$ – радиус зоны влияния i -го куста;

$distance_{i,j}$ – расстояние между i -ой и j -ой кустовыми площадками;

λ – поправочный коэффициент, $\lambda \in (0, 1)$ (В рамках исследования авторами использовался λ равным 0.83).

В целом после завершения Шага 3 процедура находит матрицу смежности кустовых площадок. Но такая матрица смежности зачастую не удовлетворяет условию планарности (рисунок 2), которое необходимо для адекватной работы следующего этапа автоматизации процесса проектирования разработки нефтяного месторождения, а именно планирования ввода в эксплуатацию месторождения. Поэтому далее описана

часть процедуры, посвященная приведению графа, построенного на основе матрицы смежности, к планарному.

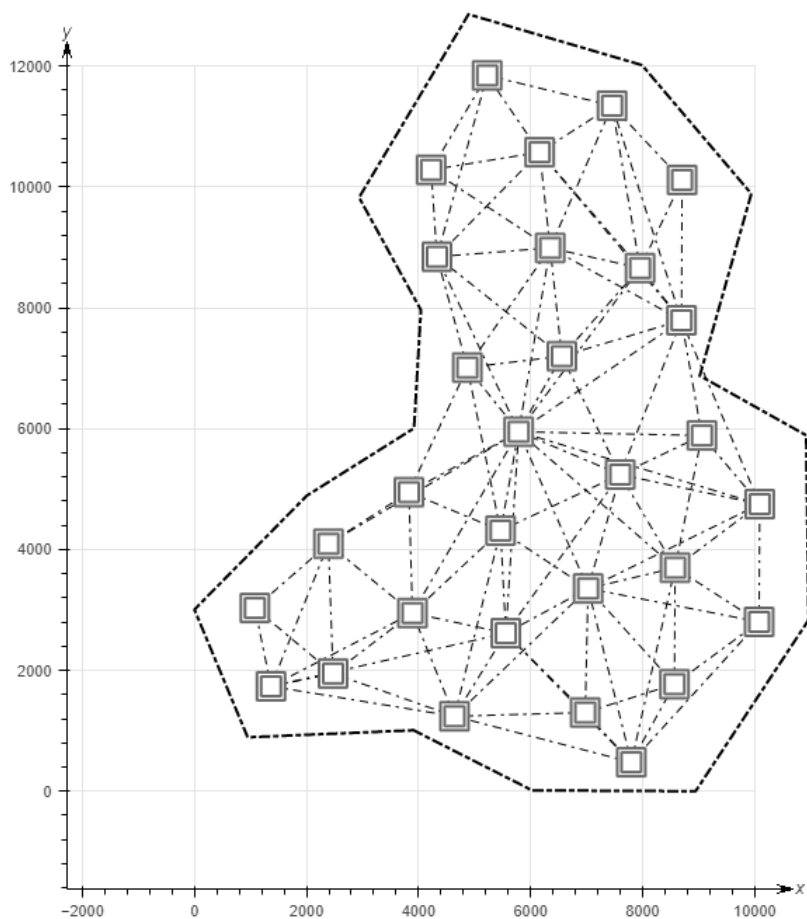


Рис. 2. Пример работы процедуры нахождения матрицы смежности

Шаг 4: Обход по всем кустам. Для i -ого куста *neighbours* – список его соседей.

Шаг 5: Обход по всем соседям i -ого куста. *common_neighbours* – список общих соседей i -ого и *neighbours_j* кустов. Далее i -ый куст будет именоваться A , а *neighbours_j* куст – B .

Шаг 6: Расчет расстояния между кустами A и B .

Шаг 7: Обход всех общих соседей кустов A и B . Текущий общий сосед кустов A и B будет именоваться C . Расчет расстояний AC и BC .

Шаг 8: Если угол напротив стороны AB больше или равен 90 градусов (косинус угла напротив стороны AB value (8) меньше или равен 0), то для пары i -ого и $neighbours_j$ кустов завершается обход их общих соседей. Переход к шагу 9. Иначе переход к шагу 5.

$$\text{value} = \frac{BC^2 + AC^2 - AB^2}{2 \cdot BC \cdot AC}. \quad (8)$$

Шаг 9: Если у куста A или куста B нет скважин, то переход к шагу 5.

Шаг 10: Расчет полигональных структур вокруг скважин куста A и скважин куста B .

Шаг 11: Пусть будет построен отрезок между кустами A и B . Тогда точка P – пересечение отрезка AB и полигональной структуры куста A , а точка Q – пересечение отрезка AB и полигональной структуры куста B .

Шаг 12: Расчет расстояния между кустом A и точкой P . Расчет расстояния между точкой Q и кустом B .

Шаг 13: Расчет отношения $ratio$ между длинами сторон PQ и AB .

Шаг 14: Если длина $AP < radius_i$ и длина $QB < radius_j$, то переход к шагу 14. Иначе переход к шагу 15.

Шаг 15: Точки P и Q находятся внутри соответственно зон влияния кустов A и B . Если $ratio \geq 0.15$, то необходимо удалить связь между i -ым и $neighbours_j$ кустами. Переход к шагу 4.

Шаг 16: Как минимум одна из точек P и Q находятся за пределами зоны влияния своего куста. Если $ratio \geq 0.13$, то необходимо удалить связь между i -ым и $neighbours_j$ кустами. Переход к шагу 4.

Используемые пороговые значения для коэффициента $ratio$ в шагах 15 и 16 процедуры нахождения матрицы смежности кустовых площадок были установлены опытным путем в процессе настройки работы процедуры.

6. Численные эксперименты. Для исследования разработанного алгоритма корректировки положения кустовых площадок была разработана следующая схема:

– Подготовка тестовых нефтяных месторождений для проведения исследования;

– Размещение кустовых площадок с использованием оптимизационного алгоритма на тестовых месторождениях. Фиксация результатов размещения;

– Запуск разработанного алгоритма на основе каждого сохраненного результата размещения кустовых площадок при помощи оптимизационного алгоритма. Фиксация результатов корректировки положения кустов.

В рамках создания тестовых нефтяных месторождений будут рассмотрены разные геометрические фигуры, потенциально охватывающие большую часть спектра возможных конфигураций залежей углеводородного сырья. Для подготовки тестовых нефтяных месторождений был определен список характерных признаков, комбинации которых формируют возможные конфигурации залежей:

- неправильная форма контура;
- вытянутая форма контура;
- залежь с обширной зоной отсутствия коллектора;
- несколько несвязанных блоков залежи.

Всего для эксперимента было подготовлено три синтетических тестовых нефтяных месторождения, ключевые характеристики которых представлены в таблице 1. Внешний вид тестовых месторождений продемонстрирован на рисунках 3 – 5.

Особенностью первого месторождения является наличие 5 разведочных ранее пробуренных скважин, которые не подлежат объединению в кусты и, соответственно, не участвуют в процессе кустования. Особенностью же последнего месторождения является горизонтальный тип скважин. Длина ствола и азимут составляют 450 м и 65 градусов соответственно.

Таблица 1. Характеристики тестовых месторождений

Характеристики	Номер тестового месторождения		
	1	2	3
Площадь (км ²)	172.03	76.04	368.92
Количество залежей	2	4	1
Количество добывающих скважин	352	322	386
Количество нагнетательных скважин	169	303	373
Общее количество скважин	521	625	759
Тип сетки скважин	обращенная семиточечная	обращенная пятиточечная	шахматная рядная
Тип скважин	вертикальный	вертикальный	горизонтальный
Расстояние между скважинами (м)	600	500	700

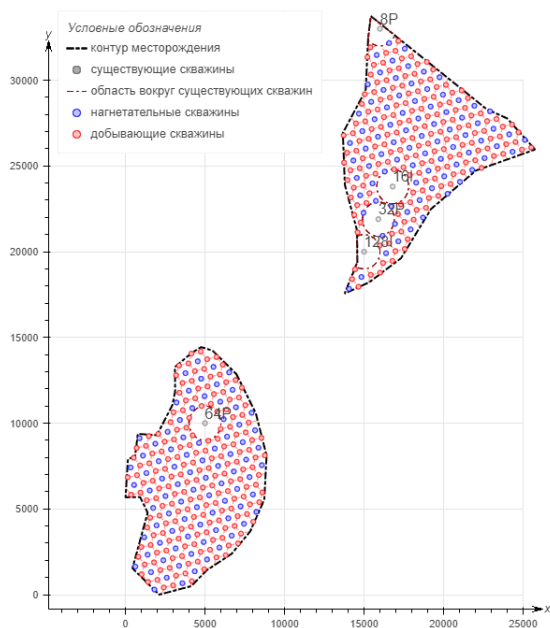


Рис. 3. Схема представления первого синтетического месторождения

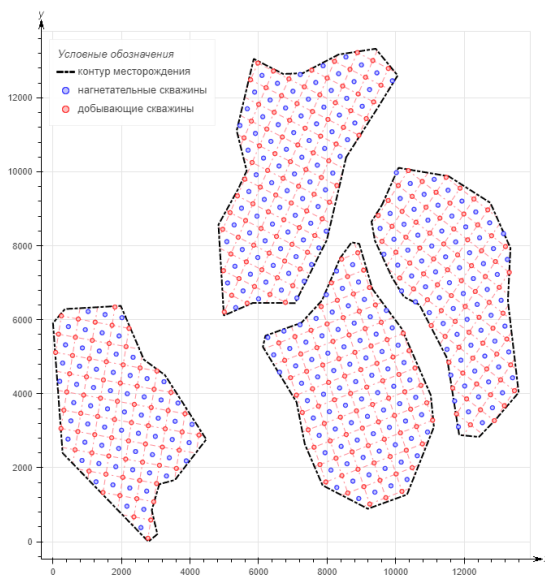


Рис. 4. Схема представления второго синтетического месторождения

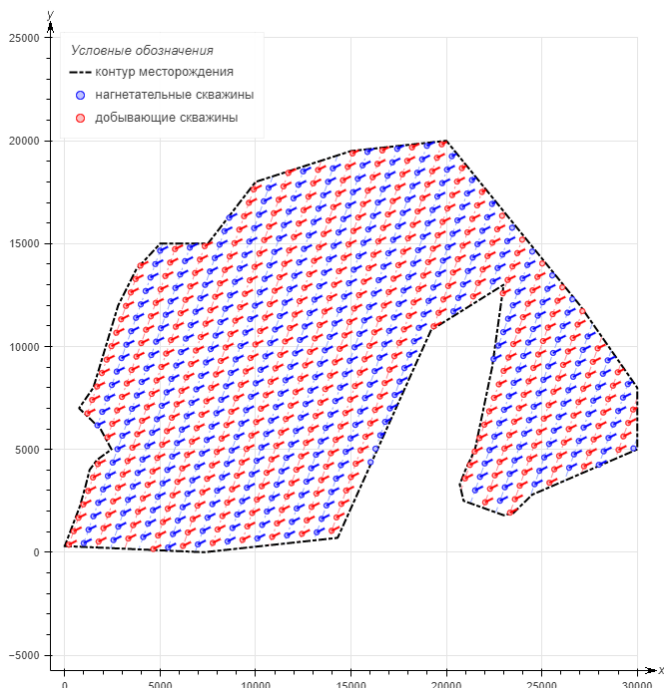


Рис. 5. Схема представления третьего синтетического месторождения

Следующим этапом необходимо запустить задачу размещения кустовых площадок и скважин между ними при помощи глобального метода оптимизации. Для этого необходимо определиться со значениями основных ограничений, а также с методом глобальной оптимизации. В качестве алгоритма глобальной оптимизации был выбран метод роя частиц, который успешно применяется при решении прикладных задач [40 – 42].

Чтобы ограничить использование вычислительных ресурсов при решении задач оптимизации требуется ввести понятия мощности оптимизации N_{max} . Данное понятие обозначает максимальное количество запусков целевой функции. Для популяционных алгоритмов мощность оптимизации можно определить следующим образом:

$$N_{max} = NP \cdot Iter_{max}, \quad (9)$$

где NP – размер популяции;

$Iter_{max}$ – количество поколений.

В рамках данной задачи, целевая функция описана в формуле (2). В качестве входных параметров выступают координаты по осям X и Y размещаемых кустовых площадок. Следовательно, размер популяции оптимизационного алгоритма для решения поставленной задачи будет равен удвоенному числу кустовых площадок.

Отношение размера популяции и количества поколений в рамках данного исследования будет величиной постоянной и будет равно один к двум. Таким образом, выбрав значение мощности оптимизации для размещения m кустовых площадок, автоматически определится размер популяции и количество поколений.

Далее требуется определить значения основных ограничений, которые представлены ниже:

- минимальное допустимое количество скважин в кусте – 10;
- максимальное допустимое количество скважин в кусте – 24;
- максимальная величина отхода – 2500 м.;
- минимально допустимое расстояние между двумя кустовыми площадками – 1000 м.

Сравнение результатов работы оптимизационного алгоритма и разработанного алгоритма корректировки положения кустовых площадок после выполнения оптимизационного алгоритма будет проводиться на основе анализа величины штрафа, которая обозначает степень нарушения значений основных ограничений.

Стоит отметить, что для сравнения берется величина штрафа до умножения на α . Такая величина состоит из среднего арифметического числа кустов по каждому из нарушений ограничений. Но, при этом один и тот же куст не может иметь одновременно нарушения на сверхмалое и на сверхбольшое число скважин. Следовательно, допустимый диапазон величины штрафа будет равен $[0, 3]$.

Для каждого тестового месторождения были определены количество кустовых площадок, которые необходимо разместить, и перечень принимаемых значений мощности оптимизации. Данные показатели представлены в таблице 2. Для каждого выбранного значения мощности оптимизации было определено число реализаций равное 50.

Таблица 2. Количественные показатели для оптимизационной задачи

Показатели	Номер тестового месторождения		
	1	2	3
Количество кустовых площадок	25	30	35
Минимальное принимаемое значение для мощности оптимизации	4000	4800	5600
Максимальное принимаемое значение для мощности оптимизации	84000	100800	117600
Шаг для мощности оптимизации	10000	12000	14000

Для проведения исследований в рамках публикации использовался программный продукт, разработанный авторами [43]. Эксперимент проводился на ЭВМ с ОП Windows со следующими характеристиками:

- ОЗУ: 32 Гб;
- количество ядер: 8;
- тактовая частота: 3700 Ггц.

Результаты экспериментов представлены в таблицах 3 – 5. Визуализация данных в таблицах представлена на рисунках 6 – 8.

Таблица 3. Значения штрафа для первого тестового месторождения

	Значения мощности оптимизации								
	4000	14000	24000	34000	44000	54000	64000	74000	84000
метод оптимизации									
min	1	0,6	0,44	0,48	0,44	0,28	0,48	0,48	0,4
median	1,28	0,9	0,78	0,68	0,72	0,7	0,76	0,64	0,72
mean	1,26	0,91	0,77	0,72	0,75	0,68	0,75	0,68	0,73
max	1,56	1,2	1,04	1,08	1,08	0,96	1,04	1,08	1,08
sd	0,12	0,15	0,14	0,14	0,14	0,16	0,15	0,14	0,14
метод оптимизации с алгоритмом корректировки положения кустовых площадок									
min	0	0	0	0	0	0	0	0	0
median	0,04	0,04	0,04	0,04	0,04	0,04	0,04	0	0,04
mean	0,04	0,04	0,04	0,03	0,03	0,04	0,04	0,03	0,03
max	0,2	0,2	0,16	0,24	0,12	0,2	0,16	0,12	0,16
sd	0,05	0,04	0,04	0,05	0,03	0,05	0,05	0,04	0,04

Таблица 4. Значения штрафа для второго тестового месторождения

	Значения мощности оптимизации								
	4800	16800	28800	40800	52800	64800	76800	88800	100800
метод оптимизации									
min	0,5	0,23	0,07	0,1	0,07	0,07	0,1	0,07	0,07
median	0,7	0,4	0,27	0,23	0,2	0,2	0,2	0,2	0,17
mean	0,7	0,39	0,27	0,22	0,19	0,21	0,2	0,2	0,16
max	0,9	0,57	0,4	0,37	0,4	0,37	0,33	0,43	0,37
sd	0,08	0,08	0,08	0,05	0,07	0,07	0,06	0,07	0,07
метод оптимизации с алгоритмом корректировки положения кустовых площадок									
min	0	0	0	0	0	0	0	0	0
median	0	0	0	0	0	0	0	0	0
mean	0,01	0,01	0,01	0,01	0,01	0,02	0,01	0,01	0,01
max	0,07	0,13	0,1	0,13	0,1	0,17	0,07	0,17	0,07
sd	0,02	0,03	0,02	0,03	0,02	0,04	0,02	0,03	0,02

Таблица 5. Значения штрафа для третьего тестового месторождения

	Значения мощности оптимизации								
	5600	19600	33600	47600	61600	75600	89600	103600	117600
метод оптимизации									
min	0,83	0,37	0,26	0,23	0,17	0,11	0,14	0,14	0,14
median	0,97	0,66	0,43	0,37	0,34	0,29	0,29	0,29	0,29
mean	0,97	0,65	0,44	0,37	0,33	0,31	0,28	0,3	0,29
max	1,14	0,91	0,74	0,6	0,51	0,49	0,46	0,43	0,54
sd	0,08	0,11	0,09	0,07	0,07	0,08	0,07	0,07	0,08
метод оптимизации с алгоритмом корректировки положения кустовых площадок									
min	0	0	0	0	0	0	0	0	0
median	0,03	0,03	0	0,03	0,03	0,03	0,03	0,03	0,03
mean	0,03	0,03	0,02	0,02	0,03	0,03	0,03	0,03	0,02
max	0,11	0,14	0,14	0,11	0,17	0,17	0,11	0,2	0,09
sd	0,03	0,03	0,03	0,03	0,04	0,04	0,03	0,04	0,02

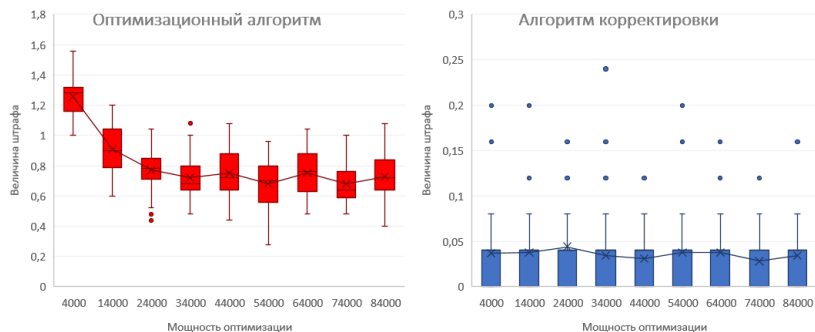


Рис. 6. Зависимость величины штрафа от мощности оптимизации (1-ое месторождение)

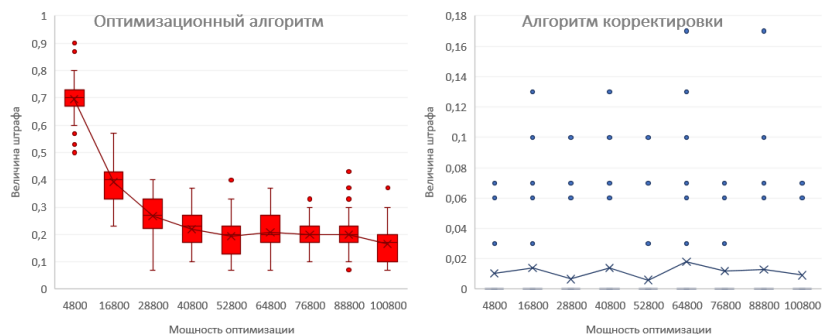


Рис. 7. Зависимость величины штрафа от мощности оптимизации (2-ое месторождение)

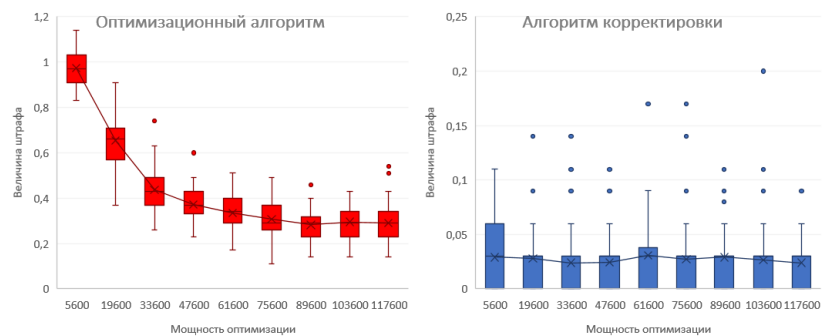


Рис. 8. Зависимость величины штрафа от мощности оптимизации (3-ое месторождение)

7. Обсуждение и выводы. В рамках решения задачи размещения кустовых площадок на нефтяном месторождении было проведено сравнение работоспособности оптимизационного алгоритма метода роя частиц и надстройки над ним, представляющее собой алгоритм корректировки положения кустовых площадок. Для проведения эксперимента были подготовлены три тестовых синтетических нефтяных месторождения. Сравняемым показателем являлась величина штрафа за нарушение основных ограничений задачи.

Была введена мощность оптимизации, понятие, которое ограничивает вычислительные возможности метода оптимизации. Для каждого тестового месторождения была определена последовательность принимаемых значений этого параметра, и относительно них был проведен запуск рассматриваемых подходов. Полученные результаты были представлены в табличном и графическом форматах.

На приведенных графиках для визуализации результатов была выбрана диаграмма «ящик с усами», которая отлично демонстрирует значения медианы и математического ожидания, минимальные и максимальные значения, а также отображает выбросы. Таким образом, для каждого принимаемого значения мощности оптимизации на графике приводятся соответствующие статические показатели относительно величины штрафа.

Применение оптимизационного подхода для всех тестовых нефтяных месторождений не позволило получить решение, удовлетворяющее пользовательским ограничениям в рамках поставленной задачи. Предложенный в рамках данной статьи алгоритм в большинстве случаев находит оптимальное решение без нарушения заданных ограничений.

Полученные результаты можно связать с тем, что представленное решение является надстройкой над оптимизационным методом. Алгоритм корректировки использует в качестве входной информации расположение кустовых площадок, которое определяется в ходе работы метода оптимизации. В силу случайного характера размещения кустовых площадок метод оптимизации может генерировать решения, при которых устранение нарушений либо невозможно, либо будет требовать больших вычислительных ресурсов. Примеры некоторых таких ситуаций:

- кустовая площадка находится за пределами месторождения и не имеет соседей;

- кустовая площадка и ее соседи имеют нарушения по ограничению максимального числа скважин;

- кустовая площадка и ее соседи имеют нарушения по ограничению минимального числа скважин.

В целом сложившуюся проблему можно решить двумя способами:

- замена оптимизационного подхода на другие решения;
- создание нового самостоятельного алгоритма размещения кустовых площадок.

Новое решение данной задачи может заключаться в представлении её в виде задачи упаковки.

Также возможна замена текущего метода роя частиц другими оптимизационными подходами, позволяющие решать задачи условной оптимизации.

8. Заключение. В рамках данного исследования был предложен алгоритм корректировки положения кустовых площадок при решении задачи разработки нефтяных месторождений.

Предлагаемый подход основан на идеи изменения взаимного расположения кустовых площадок на области месторождения, с целью устранения нарушений проектных ограничений, связанных с планированием месторождения. В качестве входных данных предложенное решение использует результат размещения кустовых площадок, который был получен при помощи метода оптимизации. В целом за счет использования данного подхода при разработке месторождения может быть достигнут результат, при котором количество скважин для каждого из кустов будет в пределах допустимых нормативов, а расстояние между кустами – не меньше заданного.

В рамках численных экспериментов были сформированы синтетические месторождения, которые отличаются друг от друга геометрией, количеством залежей и конфигурацией сетки скважин. На них было проведено размещение кустовых площадок при помощи алгоритма оптимизации «метод роя частиц». Следующим шагом был запущен алгоритм корректировки положения кустовых площадок. Таким образом, выполнялось сравнение штрафного показателя после завершения работы оптимизационного метода и предлагаемого подхода. В ходе анализа полученных результатов были сделаны выводы о том, что предложенный алгоритм в большинстве случаев может справляться с поставленной задачей.

В дальнейшем в рамках решения данной задачи планируется внедрение представленного подхода в сам процесс решения

оптимизационной задачи как дополнительный этап, позволяющий выполнять коррекцию найденных решений под имеющиеся ограничения. Другое возможное направление исследований в рамках решения задачи автоматизации процесса проектирования разработки нефтяных месторождений будет заключаться в формировании новой методологии размещения кустовых площадок без использования глобальных оптимизационных алгоритмов.

Литература

1. Кутузова М. Катастрофа с надеждой на будущее // Нефть России. 2016. № 11–12. С. 28–31.
2. Нефтянка, шаг вперед. Справится ли Россия? URL: <https://teknoblog.ru/2017/01/08/73260>. (дата обращения: 10.10.2022).
3. Баскова М.Л. Анализ развития нефтяной отрасли России // NovaInfo.Ru. 2015. № 33. С. 76–81.
4. Фрай М.Е. Оценка современного состояния нефтяной промышленности России // Вестник Удмуртского университета. Серия «Экономика и право». 2015. № 2. С. 75–85.
5. Эдер Л.В., Филимонова И.В., Проворная И.В., Мамахатов Т.М. Состояние нефтяной промышленности России: добыча, переработка, экспорт // Минеральные ресурсы России. Экономика и управление. 2016. № 6. С. 41–51.
6. Yergin D. The Prize: The Epic Quest for Oil, Money, and Power // Simon & Schuster. 1990. 887 p.
7. Prodromou T., Demirel R. Oil Price shock and cost of capital: Does market liquidity play a role? // Energy Economics. 2022. pp. 1–39.
8. Al Jabri S., Raghavan M., Vespignani J. Oil prices and fiscal policy in an oil-exporter country: Empirical evidence from Oman // Energy Economics. 2022. vol. 111. pp. 1–33.
9. Черная полоса Господству нефти в мире приходит конец. В 2020-м она получила удар, от которого может не оправиться. URL: <https://lenta.ru/articles/2020/12/30/petroleum/>. (дата обращения: 10.10.2022).
10. Воробьев А.Е., Воробьев К.А., Тчаро Х. Цифровизация нефтяной промышленности // Изд-во «Спутник+». 2018. 327 с.
11. Пескова Д.Р., Ходковская Ю.В., Шарафутдинов Р.Б. Цифровизация бизнес-процессов в нефтегазовых компаниях // Евразийский юридический журнал. 2018. № 9. С. 438–444.
12. Тчаро Х., Воробьев А.Е., Воробьев К.А. Цифровизация нефтяной промышленности: базовые подходы и обоснование «интеллектуальных» технологий // Вестник Евразийской науки. 2018. № 2. С. 1–17.
13. Абасова Х.А. Характеристика финансовых рисков и их особенности в нефтяной промышленности // Финансы и кредит. 2013. № 9. С. 61–68.
14. Ващук Д.О., Курочкин Е.С., Беломестнов А.В. Цифровизация нефтегазовой промышленности // Инновационный потенциал цифровой экономики: Сб. научн. статей Междунар. науч.-практ. конф. (г. Курск, 28 октября 2021 г.). Курск: ЮЗГУ, 2021. С. 79–82.
15. Козлова Д. Перспективы и барьеры цифровой трансформации нефтегазового комплекса России // Нефтегазовая вертикаль: Спецвыпуск. Аналитическое обозрение «Вектор». 2018. № 15–16. С. 19–26.

16. Козлова Д.В., Пигарев Д.Ю. Интеллектуальная добыча почему России необходимо изменить подход к государственному стимулированию отрасли. // *Neftegaz.Ru*. 2018. № 7. С. 32–39.
17. Tariq Z., Aljawad M.S., Hasan A., Murtaza M., Mohammed E., El-Husseine A., Alarifi S.A., Mahmoud M., Abdulraheem A. A systematic review of data science and machine learning applications to the oil and gas industry // *Journal of Petroleum Exploration and Production Technology*. 2021. no. 11. pp. 4339–4374.
18. Sircar A., Yadav K., Rayavarapu K., Bist N., Oza H. Application of machine learning and artificial intelligence in oil and gas industry // *Petroleum Research*. 2021. no. 6. pp. 379–391.
19. Hanga K.M., Kovalchuk Y. Machine learning and multi-agent systems in oil and gas industry applications: A survey // *Computer Science Review*. 2019. vol. 34. pp. 1–17.
20. Куст скважин. URL: <https://proektirovanie.gazprom.ru/about/subsidiaries/47/>. (дата обращения: 10.10.2022).
21. Силкина Т.С., Сахимова Э.Э., Лямина Н.Ф., Саушин А.З. Анализ высокотехнологичного метода добычи нефти с использованием уплотняющей сетки при бурении новой скважины в сложной кустовой системе // *Новейшие технологии освоения месторождений углеводородного сырья и обеспечение безопасности экосистем Каспийского шельфа: Материалы XII Междунар. науч.-практ. конф. (г. Астрахань, 3 сентября 2021 г.)*. Астрахань: АГТУ, 2021. С. 62–65.
22. Liu G., Wang G., Zhao Z., Ma F. A new pattern of cluster-layout for deep geothermal reservoirs: Case study from Dezhou geothermal field, China // *Renewable Energy*. 2020. pp. 484–499.
23. Kheirollahi H., Chahardowli M., Simjoo M. A new method of well clustering and association rule mining // *Journal of Petroleum Science and Engineering*. 2022. vol. 214.
24. Abramov A. Optimization of well pad design and drilling – well clustering // *Petroleum Exploration and Development*. 2019. no. 3. pp. 614–620.
25. Арбузов В.Н. Эксплуатация нефтяных и газовых скважин Часть 1 // Изд-во ТПУ, 2011. 200 с.
26. Лошаков Д.С., Васильев С.И., Милосердов Е.Е., Ганиев Д.Ф., Герлинский П.В. Проблемы обустройства кустовых оснований при наличии многолетних мерзлых пород // *Горная промышленность*. 2016. № 6. С. 74–75.
27. Денисов П.Г. Сооружение буровых // *М: Недра*. 1989. 397 с.
28. Богаткина Ю.Г., Еремин Н.А., Лындин В.Н. Особенности бурения и формирования затрат в строительстве нефтегазовых скважин кустовым методом Ханты-Мансийского АО // *Проблемы экономики и управления нефтегазовым комплексом*. 2021. № 10. С. 5–9.
29. Ермолаев А.И., Кувичко А.М., Соловьев В.В. Модели и алгоритмы размещения кустовых площадок и распределения скважин по кустам при разработке нефтяных и газовых месторождений // *Автоматизация, телемеханизация и связь в нефтяной промышленности*. 2011. № 9. С. 29–32.
30. Можиль А.Ф., Третьяков С.В., Дмитриев Д.Е., Гильмутдинова Н.З., Есипов С.В., Карачев А.А. Технико-экономическая оптимизация кустования скважин при интегрированном концептуальном проектировании // *Нефтяное хозяйство*. 2016. № 4. С. 126–129.
31. Robertson E., Iyer N., Klenner R.C.L., Liu G. Optimization of unconventional well-pad area using reservoir simulation and intelligent sequential sampling // *Unconventional Resources Technology Conference*. 2017. pp. 1–12.

32. Kheirollahi H., Chahardowli M., Simjoo M. A new method of well clustering and association rule mining // *Journal of Petroleum Science and Engineering*. 2022. vol. 214. DOI: 10.15530/jurtec-2017-2673695.
33. Шатровский А.Г., Чинаров А.С., Салихов М.Р. Группирование проектных скважин для размещения кустовых площадок на примере многопластового месторождения // *ПРОНЕФТЬ. Профессионально о нефти*. 2020. №3. С. 44–49.
34. Kulakov E.D., Mikhalev A.S., Kuznetsov A.S., Sarenkov A.V., Shutalev A.D., Gorokhov A.P., Fedoreev A.E. Planning development automation of oil fields // *AIP Conference Proceedings* 2402. 2021. pp. 1–13.
35. Обустройство нефтяных и газовых месторождений. Требования пожарной безопасности. URL: <https://docs.cntd.ru/document/1200122146>. (дата обращения: 10.10.2022).
36. Cherif L., Merikhi B. A penalty method for nonlinear programming // *RAIRO - Operations Research*. 2019. vol. 53. no. 1. pp. 29–38.
37. Mnif M., Pham H. Stochastic optimization under constraints // *Stochastic Processes and their Applications*. 2001. vol. 93. pp. 149–180.
38. Horst R., Pardalos P.M., Thoai N.V. Introduction to Global Optimization // *Kluwer Academic Publishers*. 2000. 354 p.
39. Schuster H.G., Just W. Deterministic Chaos // *John Wiley & Sons*. 2006. pp. 312.
40. Pervaiz S., Ul-Qayyum Z., Bangyal W., Gao L., Ahmad J. A systematic literature review on particle swarm optimization techniques for medical diseases detection // *Comput Math Methods Med*. 2021. pp. 1–10.
41. Jiao J., Ghoreshi S., Moradi Z., Oslub K. Coupled particle swarm optimization method with genetic algorithm for the static-dynamic performance of the magneto-electro-elastic nanosystem // *Engineering with Computers*. 2022. no. 38. pp. 2499–2513.
42. Chen H., Fan D., Huang W., Huang J., Cao C., Yang L., He Y., Zeng L. Particle swarm optimization algorithm with mutation operator for particle filter noise reduction in mechanical fault diagnosis // *International Journal of Pattern Recognition and Artificial Intelligence*. 2020. vol. 34. no. 10. pp. 1–8.
43. Патент РФ RU2021682129 Программа "Smart Oil Planning v1.0" 30/12/2021.
44. Lemarechal C. Lagrangian Relaxation // *Lecture Notes in Computer Science*. 2001. vol. 2241. pp. 112–156.

Кулаков Егор Дмитриевич — аспирант, институт космических и информационных технологий, ФГАОУ ВО «Сибирский федеральный университет». Область научных интересов: анализ данных, машинное обучение, математическое моделирование, методы оптимизации. Число научных публикаций — 9. eg.2015j@yandex.ru; проспект Свободный, 79, 660041, Красноярск, Россия; р.т.: +7(391)244-8625.

Михалев Антон Сергеевич — старший преподаватель, институт космических и информационных технологий, ФГАОУ ВО «Сибирский федеральный университет». Область научных интересов: искусственный интеллект, анализ данных, машинное обучение, компьютерное зрение, глубокое обучение, математическое моделирование, методы оптимизации. Число научных публикаций — 56. asmikhalev@sfu-kras.ru; проспект Свободный, 79, 660041, Красноярск, Россия; р.т.: +7(391)244-8625.

Саренков Александр Валерьевич — директор по разработке, ООО "Геомикс". Область научных интересов: цифровая трансформация нефтяной и горнорудной отрасли, разработка наукоемкого программного обеспечения. Число научных публикаций — 8. a.sarenkov@rt-eg.ru; улица Академика Королева, 13/1, 129515, Москва, Россия; р.т.: +7(800)201-6596.

Шуталев Артем Дмитриевич — заведующий сектором проектирования разработки, ООО "РН-КрасноярскНИПИнефть". Область научных интересов: проектирование разработки, моделирование нефтяных и газовых месторождений, проектное управление. Число научных публикаций — 7. ShutalevAD@gmail.com; улица 9 Мая, 65д, 660098, Красноярск, Россия; р.т.: +7(391)200-8830.

Федореев Артем Евгеньевич — главный специалист - владелец продукта, ООО "РН-КрасноярскНИПИнефть". Область научных интересов: цифровизация процессов на основе технологий RPA/OCR, проектное управление, разработка концепций и формирование MVP проектов, реинжиниринг и оптимизация процессов. Число научных публикаций — 6. FedoreevAE@knipi.rosneft.ru; улица 9 Мая, 65д, 660098, Красноярск, Россия; р.т.: +7(391)200-8830.

E. KULAKOV, A. MIKHALEV, A. SARENKOV, A. SHUTALEV, A. FEDOREEV
**POSITION CORRECTION ALGORITHM OF WELL PADS WHEN
SOLVING THE PROBLEM OF DEVELOPING OIL FIELDS**

Kulakov E., Mikhalev A., Sarenkov A., Shutalev A., Fedoreev A. **Position Correction Algorithm of Well Pads When Solving the Problem of Developing Oil Fields.**

Abstract. This article is devoted to the problem of automation of the stage of combining wells into clusters, considered as part of the process of designing the development of oil fields. The solution to the problem of combining wells into clusters is to determine the best location of well pads and the distribution of wells into clusters, in which the costs of developing and maintaining an oil field will be minimized, and the expected flow rate will be maximized. One of the currently used approaches to solving this problem is the use of optimization algorithms. At the same time, this task entails taking into account technological limitations when searching for the optimal option for the development of an oil field, justified, among other things, by the regulations in force in the industry, namely, the minimum and maximum allowable number of wells in a pad, as well as the minimum allowable distance between two well pads. The use of optimization algorithms does not always guarantee an optimal result, in which all specified constraints are met. Within the framework of this study, an algorithm is proposed that allows us to work out the resulting design solutions in order to eliminate the violated restrictions at the optimization stage. The algorithm consistently solves the following problems: violation of restrictions on the ultra-small and ultra-large number of wells in a pad; discrepancy between the number of pads with a given one; violation of the restriction of the ultra-close arrangement of pads. To study the effectiveness of the developed approach, a computational experiment was conducted on three generated synthetic oil fields with different geometries. As part of the experiment, the quality of the optimization method and the proposed algorithm, which is a raise to the optimization method, were compared. The comparison was carried out on different values of optimization power, which denotes the maximum number of runs of the target function. The evaluation of the quality of the work of the compared approaches is determined by the amount of the fine, which indicates the degree of violation of the values of the main restrictions. The efficiency criteria in this work are: the average value, the standard deviation, the median, and the minimum and maximum values of the penalty. Due to the use of this algorithm, the value of the penalty for the first and third oil fields is reduced on average to 0.04 and 0.03 respectively, and for the second oil field, the algorithm allowed to obtain design solutions without violating restrictions. Based on the results of the study, a conclusion was made regarding the effectiveness of the developed approach in solving the problem of oil field development.

Keywords: oil field, development of oil fields, oil well, cluster drilling, violation of design restrictions, adjustment of the position of cluster pads.

References

1. Kutuzova M. [Disaster with hope for the future]. *Neft' Rossii – Oil of Russia*. 2016. no. 11–12. pp. 28–31. (In Russ.).
2. Neft'yanka, shag vpered. *Spravit'sya li Rossiya? [Oil industry, step forward. Will Russia cope?]* Available at: <https://teknoblog.ru/2017/01/08/73260>. (accessed: 10.10.2022). (In Russ.).
3. Baskova M.L. [Analysis of the development of the oil industry in Russia] *NovaInfo.Ru*. 2015. no. 33. pp. 76–81. (In Russ.).

4. Fraj M.E. [Assessment of the current state of the oil industry in Russia] Vestnik Udmurtskogo universiteta. Seriya «Ekonomika i pravo» – Bulletin of the Udmurt University. Series "Economics and Law". 2015. no. 2. pp. 75–85. (In Russ.).
5. Eder L.V., Filimonova I.V., Provornaya I.V., Mamahatov T.M. [The state of the oil industry in Russia: production, processing, export]. Mineral'nye resursy Rossii. Ekonomika i upravlenie. – Mineral resources of Russia. Economics and Management. 2016. no. 6. pp. 41–51. (In Russ.).
6. Yergin D. The Prize: The Epic Quest for Oil, Money, and Power. Simon & Schuster. 1990. 887 p.
7. Prodromou T., Demirer R. Oil Price shock and cost of capital: Does market liquidity play a role? Energy Economics. 2022. pp. 1–39.
8. Al Jabri S., Raghavan M., Vespignami J. Oil prices and fiscal policy in an oil-exporter country: Empirical evidence from Oman. Energy Economics. 2022. vol. 111. pp. 1–33.
9. Chernaya polosa Gospodstvu nefti v mire prihodit konec. V 2020-m ona poluchila udar, ot kotorogo mozhet ne opravitsya. [Black bar The dominance of oil in the world is coming to an end. In 2020, she received a blow from which she may never recover.] Available at: <https://lenta.ru/articles/2020/12/30/petroleum/>. (accessed: 10.10.2022). (In Russ.).
10. Vorob'ev A.E., Vorob'ev K.A., Tcharo H. Cifrovizaciya neftyanoj promyshlennosti [Digitalization of the oil industry]. Izdatel'stvo «Sputnik+». 2018. 327 p. (In Russ.).
11. Peskova D.R., Hodkovskaya Yu.V., Sharafutdinov R.B. [Digitalization of business processes in oil and gas companies]. Evrazijskij yuridicheskij zhurnal – Eurasian Law Journal. 2018. no. 9. pp. 438–444. (In Russ.).
12. Tcharo H., Vorob'ev A.E., Vorob'ev K.A. [Digitalization of the oil industry: basic approaches and rationale for "intelligent" technologies]. Vestnik Evrazijskoj nauki – Bulletin of Eurasian Science. 2018. no. 2. pp. 1–17. (In Russ.).
13. Abasova H.A. [Characteristics of financial risks and their features in the oil industry]. Finansy i kredit – Finance and credit. 2013. no. 9. pp. 61–68. (In Russ.).
14. Vashchuk D.O., Kurochkin E.S., Belomestnov A.V. [Digitalization of the oil and gas industry]. Innovacionnyj potencial cifrovoj ekonomiki: Sb. nauchn. statej Mezhdunar. nauch.-prakt. konf. [Innovative potential of the digital economy: Collected papers]. Kursk: YuZGU. 2021. pp. 79–82. (In Russ.).
15. Kozlova D. [Prospects and barriers for digital transformation of the Russian oil and gas complex]. Neftegazovaya vertikal': Specvypusk. Analiticheskoe obozrenie «Vektor» – Oil and gas vertical: Special issue. Analytical review "Vector". 2018. no. 15–16. pp. 19–26. (In Russ.).
16. Kozlova D.V., Pigarev D.Yu. [Intelligent mining why Russia needs to change the approach to state stimulation of the industry]. Neftegaz.Ru. 2018. no. 7. pp. 32–39. (In Russ.).
17. Tariq Z., Aljawad M.S., Hasan A., Murtaza M., Mohammed E., El-Husseine A., Alarifi S.A., Mahmoud M., Abdurraheem A. A systematic review of data science and machine learning applications to the oil and gas industry. Journal of Petroleum Exploration and Production Technology. 2021. no. 11. pp. 4339–4374.
18. Sircar A., Yadav K., Rayavarapu K., Bist N., Oza H. Application of machine learning and artificial intelligence in oil and gas industry. Petroleum Research. 2021. no. 6. pp. 379–391.
19. Hanga K.M., Kovalchuk Y. Machine learning and multi-agent systems in oil and gas industry applications: A survey. Computer Science Review. 2019. vol. 34. pp. 1–17.
20. Kust skvazhin [Well pad]. Available at: <https://proektirovanie.gazprom.ru/about/subsidiaries/47/>. (accessed: 10.10.2022). (In Russ.).

21. Silkina T.S., Sahimova E.E., Lyamina N.F., Saushin A.Z. [Analysis of a high-tech oil recovery method using sealing mesh when drilling a new well in a complex cluster system]. *Novejshie tekhnologii osvoeniya mestorozhdenij uglevodorodnogo syr'ya i obespechenie bezopasnosti ekosistem Kaspijskogo shel'fa: Materialy XII Mezhdunar. nauch.-prakt. konf.* [The latest technologies for the development of hydrocarbon deposits and ensuring the safety of the ecosystems of the Caspian shelf: Materials of the XII International scientific and practical conference] Astrahan': AGTU, 2021. pp. 62–65.
22. Liu G., Wang G., Zhao Z., Ma F. A new pattern of cluster-layout for deep geothermal reservoirs: Case study from Dezhou geothermal field, China. *Renewable Energy*. 2020. pp. 484–499.
23. Kheirollahi H., Chahardowli M., Simjoo M. A new method of well clustering and association rule mining. *Journal of Petroleum Science and Engineering*. 2022. vol. 214.
24. Abramov A. Optimization of well pad design and drilling – well clustering. *Petroleum Exploration and Development*. 2019. no. 3. pp. 614–620.
25. Arbuzov V.N. *Ekspluatatsiya neftyanyh i gazovyh skvazhin Chast' 1* [Operation of oil and gas wells Part 1]. Izdatel'stvo TPU. 2011. 200 p. (In Russ.).
26. Loshakov D.S., Vasil'ev S.I., Miloserdov E.E., Ganiev D.F., Gerlinskij P.V. [Problems of arranging well pads in the presence of permafrost]. *Gornaya promyshlennost' – Mining industry*. 2016. no. 6. pp. 74–75. (In Russ.).
27. Denisov P.G. *Sooruzhenie burov [Construction of drilling]*. M: Nedra. 1989. 397 p. (In Russ.).
28. Bogatkina Yu.G., Eremin N.A., Lyndin V.N. [Features of drilling and formation of costs in the construction of oil and gas wells by the cluster method of Khanty-Mansi AA]. *Problemy ekonomiki i upravleniya neftegazovym kompleksom – Problems of economics and management of the oil and gas complex*. 2021. no. 10. pp. 5–9. (In Russ.).
29. Ermolaev A.I., Kuvichko A.M., Solov'ev V.V. [Models and algorithms for well pad placement and distribution of wells by pads in the development of oil and gas fields]. *Avtomatizatsiya, telemekhanizatsiya i svjaz' v nefjtanoj promyshlennosti – Automation, telemechanization and communications in the oil industry*. 2011. no. 9. pp. 29–32. (In Russ.).
30. Mozhchil' A.F., Tret'jakov S.V., Dmitriev D.E., Gil'mutdinova N.Z., Esipov S.V., Karachev A.A. [Technical and economic optimization of well padding with integrated conceptual design]. *Neftjanoe hozjajstvo – Oil industry*. 2016. no. 4. pp. 126–129. (In Russ.).
31. Robertson E., Iyer N., Klenner R.C.L., Liu G. Optimization of unconventional well-pad area using reservoir simulation and intelligent sequential sampling. *Unconventional Resources Technology Conference*. 2017. pp. 1–12.
32. Kheirollahi H., Chahardowli M., Simjoo M. A new method of well clustering and association rule mining. *Journal of Petroleum Science and Engineering*. 2022. vol. 214. DOI: 10.15530/urtec-2017-2673695.
33. Shatrovskij A.G., Chinarov A.S., Salihov M.R. [Grouping of project wells for placement of well pads on the example of a multi-layer field]. *PRONEFT'. Professional'no o nef'ti – PRONEFT'. Professionally about oil*. 2020. no. 3. pp. 44–49. (In Russ.).
34. Kulakov E.D., Mikhalev A.S., Kuznetsov A.S., Sarenkov A.V., Shutalev. A.D., Gorokhov A.P., Fedoreev A.E. Planning development automation of oil fields. *AIP Conference Proceedings* 2402. 2021. pp. 1–13.

35. Obustroystvo neftyanyh i gazovyh mestorozhdenij. Trebovaniya pozharnoj bezopasnosti [Arrangement of oil and gas fields. Fire safety requirements]. Available at: <https://docs.cntd.ru/document/1200122146>. (accessed: 10.10.2022). (In Russ.).
36. Cherif L.B., Merikhi B. A penalty method for nonlinear programming. *RAIRO - Operations Research*. 2019. vol. 53. no. 1. pp. 29–38.
37. Mnif M., Pham H. Stochastic optimization under constraints. *Stochastic Processes and their Applications*. 2001. vol. 93. pp. 149–180.
38. Horst R., Pardalos P.M., Thoai N.V. *Introduction to Global Optimization*. Kluwer Academic Publishers. 2000. 354 p.
39. Schuster H.G., Just W. *Deterministic Chaos*. John Wiley & Sons. 2006. pp. 312.
40. Pervaiz S., Ul-Qayyum Z., Bangyal W.H., Gao L., Ahmad J. A systematic literature review on particle swarm optimization techniques for medical diseases detection. *Comput Math Methods Med*. 2021. pp. 1–10.
41. Jiao J., Ghoresli S., Moradi Z., Oslub K. Coupled particle swarm optimization method with genetic algorithm for the static-dynamic performance of the magneto-electro-elastic nanosystem. *Engineering with Computers*. 2022. no. 38. pp. 2499–2513.
42. Chen H., Fan D.L., Huang W., Huang J., Cao C., Yang L., He Y., Zeng L. Particle swarm optimization algorithm with mutation operator for particle filter noise reduction in mechanical fault diagnosis. *International Journal of Pattern Recognition and Artificial Intelligence*. 2020. vol. 34. no. 10. pp. 1–8.
43. RU2021682129 [Program "Smart Oil Planning v1.0"] 30/12/2021. (In Russ.).
44. Lemarechal C. *Lagrangian Relaxation*. Lecture Notes in Computer Science. 2001. vol. 2241. pp. 112–156.

Kulakov Egor — Postgraduate student, Institute of space and information technologies, Siberian Federal University. Research interests: data analysis, machine learning, mathematical modeling, optimization methods. The number of publications — 9. eg.2015j@yandex.ru; 79, Svobodny Av., 660041, Krasnoyarsk, Russia; office phone: +7(391)244-8625.

Mikhalev Anton — Senior lecturer, Institute of space and information technologies, Siberian Federal University. Research interests: artificial intelligence, data analysis, machine learning, computer vision, deep learning, mathematical modeling, optimization methods. The number of publications — 56. asmikhalev@sfu-kras.ru; 79, Svobodny Av., 660041, Krasnoyarsk, Russia; office phone: +7(391)244-8625.

Sarenkov Aleksandr — Director of development, LLC "Geomix". Research interests: digital transformation of the oil and mining industry, development of high-tech soft-ware. The number of publications — 8. a.sarenkov@rt-eg.ru; 13/1, Akademika Koroleva St., 129515, Moscow, Russia; office phone: +7(800)201-6596.

Shutalev Artem — Head of the field development sector, LLC RN-KrasnoyarskNIPIneft. Research interests: reservoir engineering and simulation, project management. The number of publications — 7. ShutalevAD@gmail.com; 65д, 9 May St., 660098, Krasnoyarsk, Russia; office phone: +7(391)200-8830.

Fedoreev Artem — Main specialist, product owner, LLC RN-KrasnoyarskNIPIneft. Research interests: digitalization of processes based on technologies RPA/OCR, project management, concept development and MPV project formation, reengineering and process optimization. The number of publications — 6. FedoreevAE@knipi.rosneft.ru; 65д, 9 May St., 660098, Krasnoyarsk, Russia; office phone: +7(391)200-8830.

Руководство для авторов

Взаимодействие автора с редакцией осуществляется через личный кабинет на сайте журнала «Информатика и автоматизация» <http://ia.spcras.ru/>. При регистрации авторам рекомендуется заполнить все предложенные поля данных. Подготовка статьи ведется с помощью текстовых редакторов MS Word 2007 и выше или LaTeX. Объем основного текста (до раздела Литература) - от 20 до 30 страниц включительно. Переносы разрешены. Номера страниц не проставляются. Основная часть текста статьи разбивается на разделы, среди которых являются обязательными: введение, хотя бы один «содержательный» раздел и заключение. Допускается также мотивированное содержанием и структурой материала выделение подразделов. В основную часть опускается помещать рисунки, таблицы, листинги и формулы. Правила их оформления подробно рассмотрены на нашем сайте в разделе «Руководство для авторов».

Author guidelines

Interaction between each potential author and the Editorial board is realized through the personal account on the website of the journal "Informatics and Automation" <http://ia.spcras.ru/>. At the registration the authors are requested to fill out all data fields in the proposed form. The submissions should be prepared using MS Word 2007, LaTeX. The text of the paper in the main part should not exceed 30 pages. Pages are not numbered; hyphenations are allowed. Certain figures, tables, listings and formulas are allowed in the main section, and their typography is considered in more detail at the journal web.

Signed to print 17.03.2023. Passed for print 01.04.2023.

Printed in Publishing center GUAP.

Address: 67 litera A, B. Morskaya, St. Petersburg, 190000, Russia

Founder and Publisher: SPC RAS.

Address: 39 litera A, 14th Line V.O., St. Petersburg, 199178, Russia.

The journal is registered in the Federal Service for Supervision of Communications,
Information Technology, and Mass Media,

Registration Certificate (registration number) ПИ № ФС77-79228 dated September 25, 2020

Subscription Index П5513, Russian Post Catalog

Подписано к печати 17.03.2023. Дата выхода в свет 01.04.2023.

Формат 60×90 1/16. Усл. печ. л. 15,5. Заказ № 103. Тираж 300 экз., цена свободная.

Отпечатано в Редакционно-издательском центре ГУАП.

Адрес типографии: Б. Морская, д. 67, лит. А, г. Санкт-Петербург, 190000, Россия

Учредитель и издатель: СПб ФИЦ РАН.

Адрес учредителя и издателя: 14-я линия В.О., д. 39, лит. А, г. Санкт-Петербург, 199178, Россия

Журнал зарегистрирован Федеральной службой по надзору в сфере связи,
информационных технологий и массовых коммуникаций,
свидетельство о регистрации (регистрационный номер) ПИ № ФС77-79228 от 25 сентября 2020 г.

Подписной индекс П5513 по каталогу «Почта России»