

ISSN 2713-3192
DOI 10.15622/ia.2022.21.2
<http://ia.spcras.ru>

ТОМ 21 № 2

ИНФОРМАТИКА И АВТОМАТИЗАЦИЯ

INFORMATICS AND AUTOMATION



СПб ФИЦ РАН

Санкт-Петербург
2022



INFORMATICS AND AUTOMATION

Volume 21 № 2, 2022

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. Merkuryev	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

14-th line V.O., 39, SPIIRAS, 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, 2022

ИНФОРМАТИКА И АВТОМАТИЗАЦИЯ

Том 21 № 2, 2022

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. Каляев	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. Соифер	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. Пшихопов	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, Соединенное 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

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

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», 2022
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

Information Security

F. Mercaldo, F. Martinelli, A. Santone MODEL CHECKING FOR REAL-TIME ATTACK DETECTION IN WATER DISTRIBUTION SYSTEMS	219
--	-----

P. Bui, M. Le, B. Hoang, N. Ngoc, H. Pham DATA PARTITIONING AND ASYNCHRONOUS PROCESSING TO IMPROVE THE EMBEDDED SOFTWARE PERFORMANCE ON MULTICORE PROCESSORS	243
---	-----

I. Lubkin, V. Zolotarev COMPREHENSIVE DEFENSE SYSTEM AGAINST VULNERABILITIES BASED ON RETURN-ORIENTED PROGRAMMING	275
---	-----

Mathematical Modeling, Numerical Methods

M. Abramov, E. Tsukanova, A. Tulupyev, A. Korepanova, S. Aleksanin IDENTIFICATION OF DETERIORATION CAUSED BY AHF, MADS OR CE BY RR AND QT DATA CLASSIFICATION	311
---	-----

G. Algazin, D. Algazina MODELING THE DYNAMICS OF COLLECTIVE BEHAVIOR IN A REFLEXIVE GAME WITH AN ARBITRARY NUMBER OF LEADERS	339
--	-----

Artificial Intelligence, Knowledge and Data Engineering

M. Bobyr, A. Arkhipov, S. Gorbachev, J. Cao, S. Bhattacharyya FUZZY LOGIC APPROACHES IN THE TASK OF OBJECT EDGE DETECTION	376
---	-----

K. Dubrovin, A. Stepanov, A. Verkhoturov, T. Aseeva CROP IDENTIFICATION USING RADAR IMAGES	405
---	-----

A. Shaura, A. Zlobina, I. Zhurbin, A. Bazhenova ANALYSIS OF MULTI-TEMPORAL MULTISPECTRAL AERIAL PHOTOGRAPHY DATA TO DETECT THE BOUNDARIES OF HISTORICAL ANTHROPOGENIC IMPACT	427
---	-----

СОДЕРЖАНИЕ

Информационная безопасность

- Ф. Меркальдо, Ф. Мартинелли, А. Сантоне
ПРОВЕРКА МОДЕЛИ ДЛЯ ОБНАРУЖЕНИЯ АТАК В РЕАЛЬНОМ
ВРЕМЕНИ В СИСТЕМАХ РАСПРЕДЕЛЕНИЯ ВОДЫ 219

- Ф.Х. Буй, М.К. Ле, Б.Т. Хоанг, Н.Б. Нгок, Х.В. Фам
РАЗДЕЛЕНИЕ ДАННЫХ И АСИНХРОННАЯ ОБРАБОТКА ДЛЯ
ПОВЫШЕНИЯ ПРОИЗВОДИТЕЛЬНОСТИ ВСТРОЕННОГО
ПРОГРАММНОГО ОБЕСПЕЧЕНИЯ НА МНОГОКОДНЫХ
ПРОЦЕССОРАХ 243

- И.А. Лубкин, В.В. Золотарев
КОМПЛЕКСНАЯ СИСТЕМА ЗАЩИТЫ ОТ УЯЗВИМОСТЕЙ,
ОСНОВАННЫХ НА ВОЗВРАТНО-ОРИЕНТИРОВАННОМ
ПРОГРАММИРОВАНИИ 275

Математическое моделирование и прикладная математика

- М.В. Абрамов, Е.И. Цуканова, А.Л. Тулупьев, А.А. Корепанова,
С.С. Алексанин
ИДЕНТИФИКАЦИЯ КЛИНИЧЕСКОГО УХУДШЕНИЯ В РЕЗУЛЬТАТЕ
РАЗВИТИЯ ОСН, СПОН ИЛИ ОГМ ПОСРЕДСТВОМ КЛАССИФИКАЦИИ
НА ОСНОВЕ ДАННЫХ ОБ ИНТЕРВАЛАХ RR И QT 311

- Г.И. Алгазин, Д.Г. Алгазина
МОДЕЛИРОВАНИЕ ДИНАМИКИ КОЛЛЕКТИВНОГО ПОВЕДЕНИЯ В
РЕФЛЕКСИВНОЙ ИГРЕ С ПРОИЗВОЛЬНЫМ ЧИСЛОМ ЛИДЕРОВ 339

Искусственный интеллект, инженерия данных и знаний

- М.В. Бобырь, А.Е. Архипов, С.В. Горбачев, Ц. Цао, С. Бхаттачарья
НЕЧЕТКО-ЛОГИЧЕСКИЕ ПОДХОДЫ В ЗАДАЧЕ ДЕТЕКТИРОВАНИЯ
ГРАНИЦ ОБЪЕКТОВ 376

- К.Н. Дубровин, А.С. Степанов, А.Л. Верхотуров, Т.А. Асеева
ИДЕНТИФИКАЦИЯ СЕЛЬСКОХОЗЯЙСТВЕННЫХ КУЛЬТУР С
ИСПОЛЬЗОВАНИЕМ РАДАРНЫХ ИЗОБРАЖЕНИЙ 405

- А.С. Шаура, А.Г. Злобина, И.В. Журбин, А.И. Баженова
АНАЛИЗ ДАННЫХ РАЗНОВРЕМЕННОЙ МУЛЬТИСПЕКТРАЛЬНОЙ
АЭРОФОТОСЪЕМКИ ДЛЯ ОБНАРУЖЕНИЯ ГРАНИЦ ИСТОРИЧЕСКОГО
АНТРОПОГЕННОГО ВОЗДЕЙСТВИЯ 427

F. Mercaldo, F. Martinelli, A. Santone
**MODEL CHECKING FOR REAL-TIME ATTACK DETECTION IN
WATER DISTRIBUTION SYSTEMS**

F. Mercaldo, F. Martinelli, A. Santone **Model Checking for Real-Time Attack Detection in Water Distribution Systems.**

Abstract. Water distribution systems represent critical infrastructures. These architectures are really critical, and irregular behaviour can be reflected in human safety. As a matter of fact, an attacker obtaining control of such an architecture is able to perpetrate a plethora of damages, both to the infrastructure and people. In this paper, we propose an approach to identify irregular behaviours focused on water distribution systems. The designed approach considers a formal verification environment. The logs retrieved from water distribution systems are parsed into a formal model and, by exploiting timed temporal logic, we characterize the behaviour of a water distribution system while an attack is happening. The evaluation, referred to a water distribution system, confirmed the effectiveness of the designed approach in the identification of three different irregular behaviours.

Keywords: critical infrastructure, SCADA, formal verification environment, formal methods, timed automaton, safety, security.

1. Introduction. The networks of many critical infrastructures in countries depend strategically on SCADA systems. To provide some examples, energy generation and transport, gas and oil pipelines, communication systems and aqueducts are now largely managed through industrial automation technologies. [1]. The acronym SCADA (Supervisory Control And Data Acquisition) indicates a computer system for electronic monitoring and control of industrial systems. In a nutshell, with the term SCADA, we are referring to a system for the critical infrastructure management consisting of computers and networked data [2]. Typically, these architectures exploit peripheral devices, for instance, programmable logic controller [3]. The attendant is able to control the critical architectures through the critical infrastructure management system [4]. Sensors and computers cooperate with the aim to guarantee the service provided by the critical infrastructure: the problem is that sensors and computers expose the critical infrastructure to potential irregular behaviours [5].

An attack targeting a SCADA system can easily generate significant physical damage; for this reason attackers are increasingly interested in SCADA systems [6]. For example, in 2003, the Davis-Besse nuclear power plant and the CSX company in the US were victims of the Slammer and Sobig worms respectively. Another attack, Slammer, caused a denial of service that slowed the network down, while Sobig sent spam via email [7]. At the same time, the Sobig malware has infected a computer at CSX headquarters by blocking among other reporting and delivery systems, thus causing train delays [8].

SCADA attack can also afflict damages in airplane passengers. In fact, in 2004, transport companies, such as British Airways, Railcorp and Delta Airlines, were hit by the Sasser worm, which exploited a buffer overflow vulnerability to spread to other systems [9]. Some aggressive variants may have caused network overcrowding. The consequences were delays of trains and planes, in some cases, cancellation of flights [10].

Another critical infrastructure is represented by the oil companies, which have been attack targets. In fact, in 2009, companies operating in the Oil & gas and petrochemical sectors, such as Exxon, Shell and BP, were affected by the Night Dragon malware, distributed using spearphishing technologies. This malware allowed criminals to take remote control of infected computers [11].

Attacks on SCADA systems can also offer to attackers the opportunity not only to perpetrate harm to the population, but also to extract sensitive information. The case of the Stuxnet worm, appeared in 2010; it was a worm engaged in espionage and reprogramming of industrial systems at the Natanz nuclear plant in Iran. The virus intercepted and modified data within a Programmable Logic Controller (PLC). As a consequence, sensible data were gathered, and a fifth of Iranian nuclear centrifuges were destroyed [12].

The most common attacks are aimed at silencing safety warnings or alarms, muting attacks. The typical attack scenario is when the verification of the pressure inside a joint of a gas pipeline is silenced [13]. Much more complex are the attacks that modify the behavior of a SCADA system, for example, altering the pressure levels that the system of a pipeline considers normal [14].

Moreover, considering the obsolete architecture of many systems, it is possible to search through specific scanning tools, systems exposed on the web knowingly or not.

In the identification of cyber-attacks in critical infrastructure are currently exploited features gathered from network protocols: in this context, high-level features for cyber-attacks identification are not explored. For these reasons, in this paper, we design an approach to identify irregular behaviours targeting SCADA systems (with particular regard to water distribution systems), by considering features that we retrieve from a water distribution systems.

We consider a formal verification environment [15–18] to check if properties are satisfied on the modeled water distribution system. Clearly, when a property representing an irregular behaviour is verified, the irregular behaviour is in progress on the considered model.

This paper represents an extension of a preliminary work [19] presented at the 28th International Conference on Enabling Technologies: Infrastructure

for Collaborative Enterprises (WETICE). With respect to the paper in [19] below, we depict the contributions of this paper:

- we propose a property to identify the reply irregular behaviour;
- we extend the experiment to validate better the *underflow* and the *overflow* properties presented in [19];
- we introduce a new property for the identification of a new anomaly, i.e., the reply irregular behaviour.

The work proceeds as follows: the next section presents background concepts about model checking timed automata, Section 3 describes the designed approach, Section 4 discusses the performed experiment aimed to evaluate the approach and, finally, in Section 5 conclusions and future research directions are drawn.

2. Background. In this section, fundamental concepts related to the formal methods technique adopted by the proposed method are provided, i.e., the model checking timed automata.

The toolchain to apply the model-checking technique is composed of a *formal model* and a *temporal logic*: both of them are described in this section.

We consider timed automata, proposed by Alur and Dill [20, 21], as a formal verification technique for real-time systems as, for instance, the water distribution systems. In a nutshell, a timed automaton is composed of a classical finite automaton able to manage clocks, evolving continuously and synchronously with respect to absolute time. Each transition of a timed automaton is labelled by a guard or constraint over clock values, indicating the time in which the transition can be fired, and a set of clocks to be reset when the transition is fired. Each location is constrained by an invariant, which restricts the possible values of the clocks for being in the state, which can then enforce a transition to be taken [22–24].

In this paper, we model a sequence of logs gathered from a water distribution systems as a network of timed automata, i.e., a finite-state machine extended with clock variables [21]. The model is extended with bounded discrete variables that are part of the state. The state of the system is defined by the location of all automata, the clock values and the values of the discrete variables. Automaton may fire an edge (i.e., perform a transition) separately or synchronise with another automaton with the aim to lead a new state.

Below we provide the definition of the syntax and semantics for the basic timed automata. We exploit the following notation: C is a set of clocks and $B(C)$ is the set of conjunctions over simple conditions of the form $x \bowtie c$ or $x - y \bowtie c$, where $x, y \in C$, $c \in \mathbb{N}$ and $\bowtie \in \{ <, \leq, =, >, \geq \}$. A timed automaton is a finite directed graph annotated with conditions over and resets of non-negative real-valued clocks [21].

Definition 1. Timed Automaton. A *timed automaton* is a tuple (L, l_0, C, A, E, I) where L is a set of locations, $l_0 \in L$ is the initial location, C is the set of clocks, A is a set of actions, co-actions and internal τ -action, $E \subseteq L \times A \times B(C) \times 2^C \times L$ is a set of edges between locations with an action, a guard and a set of clocks to be reset, and $I : L \rightarrow B(C)$ assigns invariants to locations.

In the following, we define the semantics of a timed automaton. A clock valuation is a function $u : C \rightarrow \mathbb{R}_{\geq 0}$ from the set of clocks to the non-negative reals. Let $\mathbb{R}^C$ be the set of all clock valuation. Let $u_0(x) = 0$ for all $x \in C$. We will abuse the notation by considering guards and invariants as sets of clock valuations, writing $u \in I(l)$ to mean that u satisfies $I(l)$.

Definition 2. Semantics of Timed Automaton. Let $= (L, l_0, C, A, E, I)$ be a timed automaton. The semantics is defined as a labelled transition system $\langle S, s_0, \rightarrow \rangle$, where $S \subseteq L \times \mathbb{R}^C$ is the set of states, $s_0 = (l_0, u_0)$ is the initial state, and $\rightarrow \subseteq S \times (\mathbb{R}_{\geq 0} \cup A) \times S$ is the transition relation such that:

- $(l, u)d(l, u + d)$ if $\forall d' : 0 \leq d' \leq d \implies u + d' \in I(l)$, and
- $(l, u)a(l', u')$ if there exists $e = (l, a, g, r, l') \in E$ s.t. $u \in g, u' = [r \mapsto 0]u$, and $u' \in I(l')$,

where for $d \in \mathbb{R}_{\geq 0}$, $u + d$ maps each clock x in C to the value $u(x) + d$, and $[r \mapsto 0]u$ denotes the clock valuation which maps each clock in r to 0 and agrees with u over $C \setminus r$.

Timed automata are often composed into a *network of timed automata* over a common set of clocks and actions, consisting of n timed automata $i = (L_i, l_i^0, C, A, E_i, I_i)$, $1 \leq i \leq n$. A location vector is a vector $\bar{l} = (l_1^0, \dots, l_n^0)$. We compose the invariant functions into a common function over location vectors $I(\bar{l}) = \bigwedge_i I_i(l_i)$. We write $\bar{l} [l'_i/l_i]$ to denote the vector where the i -th element l_i of $\bar{l}$ is replaced by l'_i .

In the following, we define the semantics of a network of timed automata.

Definition 3. Semantics of a network of Timed Automata: Let $= (L, l_i^0, C, A, E_i, I_i)$ be a network of n timed automata. Let $\bar{l}_0 = (l_1^0, \dots, l_n^0)$ be the initial location vector. The semantics is defined as a transition system $\langle S, s_0, \rightarrow \rangle$, where $S = (L_1 \times \dots \times L_n) \times \mathbb{R}^C$ is the set of states, $s_0 = (\bar{l}_0, u_0)$

is the initial state, and $\rightarrow \subseteq S \times S$ is the transition relation defined by:

- $(\bar{l}, u) \rightarrow (\bar{l}, u + d)$ if $\forall d' : 0 \leq d' \leq d \implies u + d' \in I(\bar{l})$, and
- $(\bar{l}, u) \rightarrow (\bar{l}[l'_i/l_i], u')$ if there exists $l_i \tau gr l_i^i$ s.t.
 $u \in g, u' = [r \mapsto 0]u$ and $u' \in I(l')$.
- $(\bar{l}, u) \rightarrow (\bar{l}[l'_j/l_j, l'_i/l_i], u')$ if there exists $l_i c? g_i r_i l_i^i$ and
 $l_j c! g_j r_j l_j^j$ s.t. $u \in (g_i \wedge g_j)$,
 $u' = [r_i \cup r_j \mapsto 0]u$ and $u' \in I(l')$.

Modelling languages usually extend timed automata with the following additional features:

– *binary synchronisation*: channels are declared as `chan c`. An edge labelled with `c!` synchronises with another labelled `c?`. A synchronisation pair is chosen non-deterministically if several combinations are enabled;

– *broadcast channels*: are declared as `broadcast chan c`. In a broadcast synchronisation one sender `c!` can synchronise with an arbitrary number of receivers `c?`. Any receiver that can synchronise in the current state must do so. If there are no receivers, then the sender can still execute the `c!` action, i.e. broadcast sending is never blocking;

– *initialisers*: are used to initialise integer variables and arrays of integer variables. For instance, `int i := 2`; or `int i[3] := {1, 2, 3}`.

Once defined the formal model, we need to verify the model with regard to a requirement specification. Similarly to the model, the requirement specification must be expressed in a formally well-defined language. Several such logics exist in the scientific literature. In this paper, we consider the Timed Computational Temporal Logic (TCTL) [21, 25], which extends the classical untimed branching-time logic CTL [26] with time constraints on modalities.

The TCTL syntax is defined by the following grammar:

$$\phi ::= a \mid \neg\phi \mid \phi \vee \phi \mid \mathbf{E}\phi \mathbf{U}_I \phi \mid \mathbf{A}\phi \mathbf{U}_I \phi,$$

where $a \in AP$ (we denote with AP a set of atomic propositions), and I is an interval of $\mathbb{R}_+$ with integral bounds.

There are two possible semantics for TCTL, one which is said *continuous*, and the other one which is more discrete and is said *pointwise*. We consider the second one, i.e., the *pointwise* semantic (Table 1).

Where $\varrho[\pi]$ is the state of ϱ at position π , and duration $\varrho \leq \pi$ is the prefix of ϱ ending at position π , and $\text{duration}(\varrho \leq \pi)$ is the sum of all delays along ϱ up to position π .

Table 1. Timed temporal logic semantics

$(l, u) \models a \iff a \in (l)$
$(l, u) \models a \neg \phi(l, v) \not\models \phi$
$(l, u) \models \phi \vee \psi \iff (l, u) \models \phi \text{ or } (l, u) \models \psi$
$(l, v) \models \mathbf{E}\phi \mathbf{U}_I \psi \iff \text{there is an infinite run } \varrho \text{ in } \text{from } (l, u) \text{ such that } \varrho \models \phi \mathbf{U}_I \psi$
$(l, u) \models \mathbf{A}\phi \mathbf{U}_I \psi \iff \text{any infinite run } \varrho \text{ in } \text{from } (l, u) \text{ is such that } \varrho \models \phi \mathbf{U}_I \psi$
$\varrho \models \phi \mathbf{U}_I \psi \iff \text{there exists a position } \pi > 0 \text{ along } \varrho \text{ such that } \text{varrho}[\pi] \models \psi, \text{ for every position } 0 < \pi' < \pi, \varrho[\pi'] \models \psi, \text{ and } \text{duration}(\varrho \leq_\pi) \in I.$

In the pointwise semantics, a position in a run:

$$\varrho = s_0 \tau_1, e_1 s_1 \tau_2, e_2 s_2 \dots s_{n-1} \tau_n, e_n s_n,$$

is an integer i and the corresponding state s_i . In this semantics, formulas are checked only right after a discrete action has been done. Sometimes, the pointwise semantics is given in terms of actions and timed words, but it does not change anything.

As usually in CTL, TCTL: $\models a \vee \neg a$ standing for true, $\models \neg$ standing for false, the implication $\phi \rightsquigarrow \psi \equiv (\neg \phi \vee \psi)$, the eventual operator $F_I \phi \equiv \text{tt } \mathbf{U}_I \phi$ and the global operator $G_I \phi \equiv \neg (F_I \neg \phi)$.

Formulae can be classified into *reachability*, *safety* and *liveness*. Figures 1, 2, 3 and 4 shows examples of different path formulae.

Reachability properties: they are looking if whether a given state formula, φ , possibly can be satisfied by any reachable state. Reachability properties are often used while designing a model to perform sanity checks. We express that some state satisfying φ should be reachable using the path formula $\mathbf{E} \varphi$.

Safety properties are expressed in the following form: “something bad will never happen”. These properties are usually formulated positively, for example, something good is invariantly true. For instance, let φ a state formula, we express that φ should be true in all states that are reachable with the path formulae $\mathbf{A} [] \varphi$.

Liveness properties are of the following form: something will eventually happen. Liveness is expressed with the path formula $\mathbf{A} \varphi$ meaning φ is eventually satisfied. A useful form is the leads to or response property, written $\psi \rightsquigarrow \varphi$ which is read as whenever ψ is satisfied, then eventually φ will be satisfied.

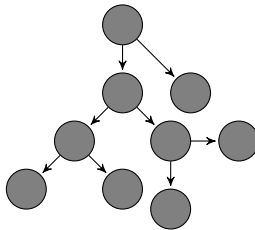


Fig. 1. $A \models \varphi$

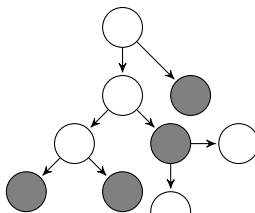


Fig. 2. $A \not\models \varphi$

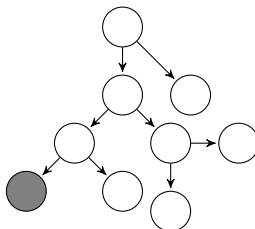


Fig. 3. $E \models \varphi$

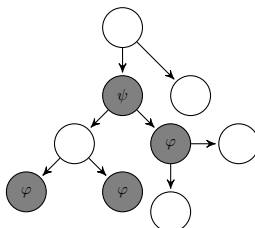


Fig. 4. $\psi \rightsquigarrow \varphi$

Once the model and temporal logic properties are defined, we need something enabling us to check whether the timed automata network (i.e., the formal model) satisfies the defined properties. To this aim, we consider formal verification, a system process exploiting mathematical reasoning to verify whether a system under analysis (i.e., the model) satisfies some requirements (i.e., the timed temporal properties).

Several verification techniques have been proposed in the last years. In this paper, we resort to model checking [24, 27].

In the model checking technique the properties are formulated in temporal logic: each property is evaluated against the system. The model checker accepts as input a model and a property, it returns “true” whether the system satisfies the formula and “false” otherwise. The performed check is an exhaustive state space search that is guaranteed to terminate since the model is finite. In this paper, we consider as model checker UPPAAL¹ [22, 23, 28], an integrated tool environment for modeling, validation and verification of real-time systems modeled as timed automata networks. Thus, the syntax of UPPAAL expression is given in Table 2.

Table 2. Syntax of expressions in BNF

<i>Expression</i>	$\rightarrow ID \mid NAT$
	$\mid Expression \ '[Expression]'$
	$\mid ('Expression')$
	$\mid Expression \ ' + +' \mid ' + +' Expression$
	$\mid Expression \ ' - -' \mid ' - -' Expression$
	$\mid Expression \ AssignOp \ Expression$
	$\mid UnaryOp \ Expression$
	$\mid Expression \ BinaryOp \ Expression$
	$\mid Expression \ '?' \ Expression \ ' :! \ Expression$
	$\mid Expression \ ' : ID$
<i>UnaryOp</i>	$\rightarrow ' - ' \mid ' ! ' \mid ' not '$
<i>BinaryOp</i>	$\rightarrow ' < ' \mid ' \leq ' \mid ' = ' \mid ' ! = ' \mid ' \geq ' \mid ' > '$
	$\mid ' + ' \mid ' - ' \mid ' * ' \mid ' / ' \mid ' \% ' \mid ' \& '$
	$\mid ' ' \mid ' \sim ' \mid ' \ll ' \mid ' \gg ' \mid ' \% \% ' \mid ' '$
	$\mid ' < ? ' \mid ' > ? ' \mid ' \wedge ' \mid ' \vee ' \mid ' \rightsquigarrow '$
<i>AssignOp</i>	$\rightarrow ' : = ' \mid ' + = ' \mid ' - = ' \mid ' * = ' \mid ' / = ' \mid ' \% = '$
	$\mid ' = ' \mid ' \% = ' \mid ' ^ = ' \mid ' \ll = ' \mid ' \gg = '$

¹<http://www.uppaal.org/>

Expressions are used with the following labels:

- *guard*: a particular expression satisfying the following conditions: (i) it is side-effect free; (ii) it evaluates to a boolean; (iii) only clocks, integer variables, and constants are referenced (or arrays of these types); (iv) clocks and clock differences are only compared to integer expressions; (v) guards overlocks are essentially conjunctions (disjunctions are allowed over integer conditions);
- *synchronisation*: a synchronisation label is either on the form Expression! or Expression? or is an empty label. The expression must be side-effect free, evaluate to a channel, and only refer to integers, constants and channels;
- *assignment*: an assignment label is a comma-separated list of expressions with a side-effect; expressions must only refer to clocks, integer variables, and constants and only assign integer values to clocks. *invariant*: an expression that satisfies the following conditions: it is a side-effect free; only clock, integer variables, and constants are referenced; it is a conjunction of conditions of the form $x < e$ or $x \leq e$ where x is a clock reference and e evaluates to an integer.

3. The Designed approach. In the designed approach, we consider the cistern level gauging as features. In details, if the water distribution system is formed by two cisterns, i.e., two features describing the water gauging in the two cisterns are considered.

The designed approach consists of two main steps: the *Formal Model Creation* (Figure 5) and the *Formal Model Verification* (Figure 8).

From the water distribution system under analysis and the technical report, a technician marks the specific day log as *irregular*. The features gathered from the cistern levels and the label to mark a specific trace as *irregular* are the input for the *one day log* stored in CSV files, containing the cistern gauging for one day at a fixed range time (1 hour).

The *Discretisation* is aimed to discretise each cistern level feature. The numeric values are split into 3 different ranges [29]. Furthermore, each discretised feature previously gathered is converted into a timed model. The discretisation process is aimed to convert continuous values into discrete ones. A plethora of methods were proposed by the research community for numeric values discretisation: in our approach, we exploit one discussed by researchers in [29]. The idea behind this method divides the numeric features into three intervals by exploiting the equal-width partitioning that basically divides the values of a certain numeric value into three equal-size intervals. In our case, the cistern level continuous values are divided into one of three different classes (*Up*, *Basal*, and *Low*).

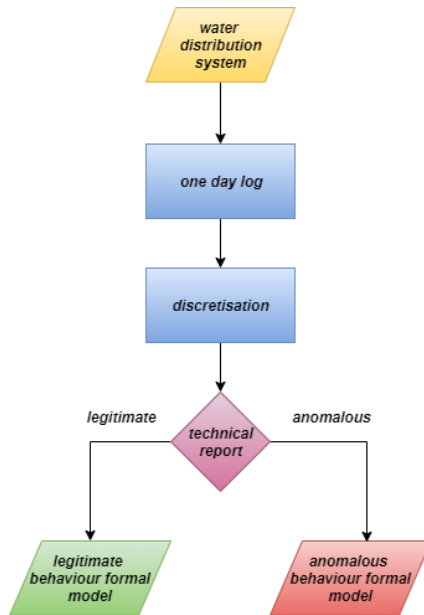


Fig. 5. Formal model generation

We compute the interval width in the following way: $W = (Max - Min)/3$, where Max and Min indicate the maximum and the minimum values. The partitioning is applied to all features (i.e., the cisterns). Moreover, once obtained the discrete values from the discretisation process, from each feature is generated a timed automaton (i.e., *formal model* in Figure 5).

With the aim to better explain how the timed automaton is built, let us consider the following example. In particular, in Table 3 we represent a fragment of a discretised feature.

With the first column (i.e., *Time* in Table 3) we are referring to the interval time (in the fragment $1 \leq t \leq 6$), while with the F_1 and F_2 columns, we are referring to the two considered features (i.e., the cistern levels). By considering the fragment in Table 3, in the t_3 time interval, the F_1 reaches the *Up* value, while the F_2 feature reaches the *Low* value. Once obtained the discretised values for the features (i.e., the cistern values), it is possible to generate the formal model. In detail, we build a timed automata network: in a nutshell, for each discretised feature, we generated an automaton.

Table 3. A fragment related to the feature discretisation

<i>Time</i>	F_1	F_2
t_1	<i>Up</i>	<i>Up</i>
t_2	<i>Up</i>	<i>Low</i>
t_3	<i>Up</i>	<i>Low</i>
t_4	<i>Basal</i>	<i>Up</i>
t_5	<i>Low</i>	<i>Basal</i>
t_6	<i>Up</i>	<i>Basal</i>

By considering the discretised fragment, we shown in Table 3 the timed automata network that is built in the following way: if the same discrete value is repeated in a consecutive time interval, the automaton related to the feature under analysis exhibits a loop: for instance, the automaton built from the F_1 discretised feature contains a loop related to the t_1 , t_2 and t_3 time intervals (because the *Up* discretised values are repeated three times). Moreover, the resulting automaton for the F_2 feature shows two loops: the first loop is related to the t_2 and t_3 time intervals (*Low* is the repeated discrete value), while the second loop is related to the t_5 and t_6 time intervals, and in this last case the repeated discrete value is *Basal*).

In order to exit from the loops, a guard is exploited. Differently, with regard to the entering condition, an invariant is considered. Moreover, for each automaton, we consider a number of clocks equal to two: the first clock (x) to ensure the entering condition into the loop and the second one (y) to ensure the exit condition. Moreover, each automaton is responsible for storing the count related to the respective *Up*, *Basal* and *Low* values.

The values for the F_x automaton are labelled with a subscript $x \in \{1, 2\}$ (we consider two features). The s channel permits the automata synchronisation and, for this reason, is not stored locally. The aim of the channels is to ensure the automata network progression. This mechanism basically is a hand-shaking synchronization: two processes take a transition at the same time, the first one will have an $s!$, while the other an $s?$, in order to ensure the synchronization. As a matter of fact, one sender event $s!$ is able to synchronise with a number of $s?$ receiver events. We resort to channels in order to avert incoherence from the discretised feature values and the interval times. For instance, the automaton related to the first feature is indicated with F_1 and its variables are u_1 , b_1 and l_1 respectively for *Up*, *Basal* and *Low* values. As a consequence, the automaton for the second variable is F_1 , and its variables are u_2 , b_2 and l_2 , respectively for *Up*, *Basal* and *Low* values of this second automaton.

Figures 6 and 7 indicate the model, respectively, gathered from the F_1 and F_2 discretisation. In both the figures, we indicate with u the *Up* value, with b the *Basal* value, and with l the *Low* value, all the three values are present in Table 3.

Below we provide a brief explanation about the models represented in Figures 6 and 7: the F_1 automaton is iterating in the loop (node 1 in Figure 6) for three-time intervals (i.e., $y_1 < 3$), while the F_2 automaton after the increment of the u_1 variable (node 1 in Figure 7) is iterating for two-time intervals (i.e., $y_2 < 2$). Subsequently, the F_1 automaton does not exhibit any loops, while the F_2 automaton is iterating for two-time intervals in the loop in node 3 in Figure 7, and then it continues with the last node.

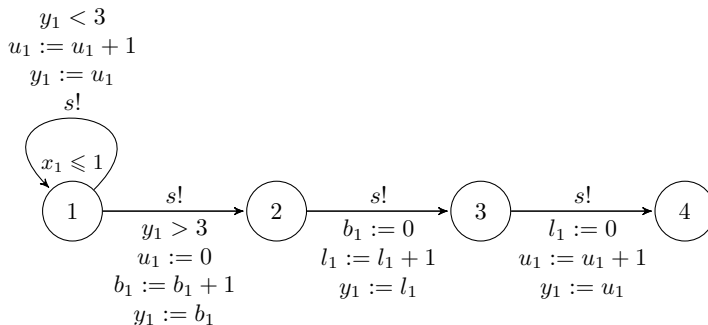


Fig. 6. The F_1 model

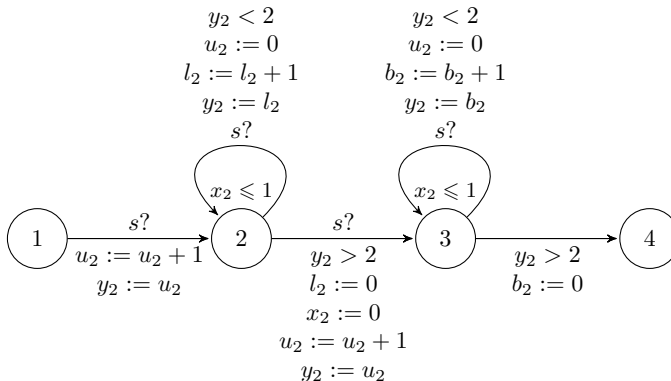


Fig. 7. The F_2 model

Once generated the formal model, the *Formal Model Verification* step (Figure 8) is aimed at checking if the modeled water distribution system is violated.

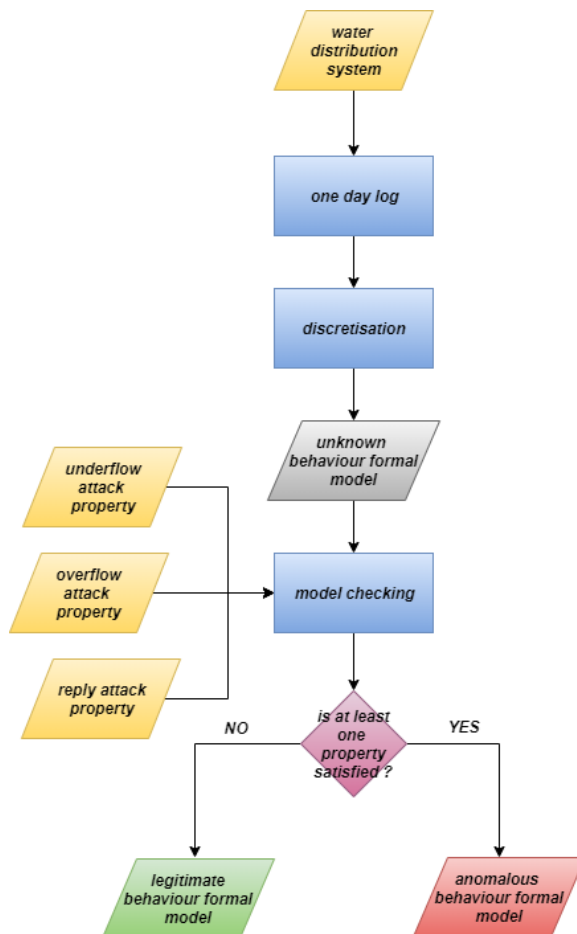


Fig. 8. Formal model verification

The *Formal Model Verification* receives as input a formal model (*formal model*) built in the previous step and a set of properties. The set of logic properties is checked against the formal model generated (*formal verification environment*) using the UPPAAL (an acronym based on a combination of UPPsala and Aalborg universities) formal verification environment². UPPAAL represents an integrated tool environment for modeling, validation and verification of real-time systems modeled as networks of timed automata.

4. Experimental evaluation. From the physical layout point of view, the following generic scenario is considered: the analysed SCADA water distribution system is composed by a generic water distribution system operator where, recently, has introduced a new technology to enable remote data collection from sensors and remote control of actuators. After the technology was introduced, anomalous levels in two tanks were observed; for instance, water overflow in one tank occurred. By looking for the causes, experts domain suspect potential cyberattacks. In particular, they consider some kind of malicious behaviours [30, 31] aimed to activate and deactivate the actuators.

The dataset³ exploited in the experimentation of the designed approach contains one-day logs from a water distribution system composed from the water levels of the two cisterns (i.e., *ct1* and *ct2*).

We have hourly gauging of the water under *regular* operating conditions, and when an *overflow* (*OF*), an *underflow* (*UF*) or a *reply* (*R*) irregular behaviours targeting, respectively, the *ct1* and *ct2* cisterns are happening. Logs referred to 30 days of gauging are considered: ten marked with the *OF* irregular behaviour on *ct1*, ten with the *UF* irregular behaviour on *ct2* and the last ten days with *reply* (*R*) irregular behaviour on both the *ct1* and *ct2* cisterns (irregular behaviours have happened each day), while the remaining ten days without irregular behaviours [5, 32, 33]. For each day log a formal model is built. In total, we consider 40 days.

To summarise, in the experiment, we consider 40 days: 30 days related to irregular behaviours and 10 days with legitimate behaviours. The 30 days contain the following irregular behaviours: 10 days with overflow attack each day, 10 days with underflow attack each day and 10 days with reply attack each day.

4.1. The Properties. Table 4 indicates the *OF*, *UF* and *R* irregular behaviour identification properties.

²<http://www.uppaal.org/>

³<http://www.batadal.net>

Table 4. Timed temporal logic property for OF , UF and R water distribution irregular behaviour identification

$$\begin{aligned}
 E\varphi, \text{ where } \varphi &= h_1 \geq 11 \\
 E\chi, \text{ where } \chi &= l_2 \geq 9 \\
 E\psi, \text{ where } \psi &= b_1 \geq 12 \wedge b_2 \geq 12 \\
 E\varphi \vee \chi \vee \psi
 \end{aligned}$$

The φ property, referred to the OF irregular behaviour, is able to verify if an OF is in progress (i.e., when the $ct1$ level exhibits an *up* value at least 11 times). The UF irregular behaviour, described by the χ property, is able to verify if an UF is in progress (i.e., when the $ct2$ level exhibits a *low* value at least 9 times). The last property, i.e., ψ (for the R irregular behaviour identification), is aimed to verify if both the $ct1$ and the $ct2$ cisterns present the same value for more than 12 times. We want to verify if the φ , the χ and the ψ properties, possibly, can be satisfied by a reachable state. We express that some state satisfying φ should be reachable using the path property $E\varphi$. Similar considerations can be done for the χ and ψ properties. Furthermore, provide the property expressing that can happen an OF , or UF or a R irregular behaviours (i.e., $E\varphi \vee \chi \vee \psi$).

The properties were formulated with the help of domain experts. As a matter of fact, the properties are aimed to explain the domain experts knowledge. In particular, the domain experts suggest that whether $ct1$ exhibits for a number of times equal or greater than 11 an increasing value, this is reflecting of an overflow attack in progress, while if $ct2$ exhibits a value that is decreasing for a number of times equal or greater than 9, this is reflecting in an underflow attack in progress. With regard to the reply attack, expert domains suggested that whether both $ct1$ and $ct2$ are showing basal values for a number of times equal or greater than 12 this is symptomatic for this kind of attack in progress.

4.2. The Experiment. The property verification outcomes are indicated in Table 5.

With L^x we mark the log of the water cisterns for a single log, where $L \in \{\text{irregular}, \text{regular}\}$, x identifies the log under analysis. The dataset comprises 40-day log, 10 exhibiting normal behaviour, while the remaining 30 are afflicted by OF , UF and R irregular behaviours. The column (φ) identifies the models verified as *true* (✓ in Table 5) or *false* (✗ in Table 5) to the OF

property, while the column (χ) identifies the models verified as *true* or *false* to the *UF* property.

From the formal verification environment output indicated in Table 5 we can state the φ and the χ are able to identify all the models afflicted by the *OF* irregular behaviour (i.e., *irregular*_o¹, *irregular*_o², *irregular*_o³, *irregular*_o⁴, *irregular*_o⁵, *irregular*_o⁶, *irregular*_o⁷, *irregular*_o⁸, *irregular*_o⁹ and *irregular*_o¹⁰). The χ property is able to identify all the models afflicted by the *UF* irregular behaviour (i.e., *irregular*_u¹, *irregular*_u², *irregular*_u³, *irregular*_u⁴, *irregular*_u⁵, *irregular*_u⁶, *irregular*_u⁷, *irregular*_u⁸, *irregular*_u⁹ and *irregular*_u¹⁰). The ψ property is able to identify all the models afflicted by the *R* irregular behaviour (i.e., *irregular*_r¹, *irregular*_r², *irregular*_r³, *irregular*_r⁴, *irregular*_r⁵, *irregular*_r⁶, *irregular*_r⁷, *irregular*_r⁸, *irregular*_r⁹ and *irregular*_r¹⁰). Furthermore, all the models without irregular behaviour (i.e., *regular*¹, *regular*², *regular*³, *regular*⁴, *regular*⁵, *regular*⁶, *regular*⁷, *regular*⁸, *regular*⁹ and *regular*¹⁰) are marked as *false* by φ , χ and ψ properties.

Moreover, we consider four different metrics to evaluate the performance of the proposed approach: Precision, Recall, F-Measure and Accuracy.

The precision has been computed as the proportion of the observations that truly belong to investigated logs among all those which were assigned to the specific attack. It is the ratio of the number of relevant records retrieved to the total number of irrelevant and relevant records retrieved:

$$lclPrecision = \frac{tp}{tp + fp},$$

where *tp* indicates the number of true positives (for instance, whether we are evaluating the φ formula, this value represents the number of one-day logs whose related *overflow attack* model is correctly labelled as *true* by the formal verification environment) and *fp* indicates the number of false positives (for instance, whether we are evaluating the φ formula, this value represents the number of models whose related *legitimate* model is wrongly labelled as *true* by the formal verification environment).

Table 5. Property verification

Performances Model	φ	χ	ψ
<i>irregular</i> ¹	✓	✗	✗
<i>irregular</i> ²	✓	✗	✗
<i>irregular</i> ³	✓	✗	✗
<i>irregular</i> ⁴	✓	✗	✗
<i>irregular</i> ⁵	✓	✗	✗
<i>irregular</i> ⁶	✓	✗	✗
<i>irregular</i> ⁷	✓	✗	✗
<i>irregular</i> ⁸	✓	✗	✗
<i>irregular</i> ⁹	✓	✗	✗
<i>irregular</i> ¹⁰	✓	✗	✗
<i>irregular</i> ₁ ¹	✗	✓	✗
<i>irregular</i> ₂ ¹	✗	✓	✗
<i>irregular</i> ₃ ¹	✗	✓	✗
<i>irregular</i> ₄ ¹	✗	✓	✗
<i>irregular</i> ₅ ¹	✗	✓	✗
<i>irregular</i> ₆ ¹	✗	✓	✗
<i>irregular</i> ₇ ¹	✗	✓	✗
<i>irregular</i> ₈ ¹	✗	✓	✗
<i>irregular</i> ₉ ¹	✗	✓	✗
<i>irregular</i> ₁₀ ¹	✗	✓	✗
<i>irregular</i> ₁ ²	✗	✗	✓
<i>irregular</i> ₂ ²	✗	✗	✓
<i>irregular</i> ₃ ²	✗	✗	✓
<i>irregular</i> ₄ ²	✗	✗	✓
<i>irregular</i> ₅ ²	✗	✗	✓
<i>irregular</i> ₆ ²	✗	✗	✓
<i>irregular</i> ₇ ²	✗	✗	✓
<i>irregular</i> ₈ ²	✗	✗	✓
<i>irregular</i> ₉ ²	✗	✗	✓
<i>irregular</i> ₁₀ ²	✗	✗	✓
<i>regular</i> ₁ ¹	✗	✗	✗
<i>regular</i> ₂ ¹	✗	✗	✗
<i>regular</i> ₃ ¹	✗	✗	✗
<i>regular</i> ₄ ¹	✗	✗	✗
<i>regular</i> ₅ ¹	✗	✗	✗
<i>regular</i> ₆ ¹	✗	✗	✗
<i>regular</i> ₇ ¹	✗	✗	✗
<i>regular</i> ₈ ¹	✗	✗	✗
<i>regular</i> ₉ ¹	✗	✗	✗
<i>regular</i> ₁₀ ¹	✗	✗	✗

The recall has been computed as the proportion of attacks that were assigned to a given class among all the attacks that truly belong to the class. It is the ratio of the number of relevant records retrieved to the total number of relevant records:

$$lclRecall = \frac{tp}{tp + fn},$$

where tp indicates the number of true positives and fn indicates the number of false negatives (for instance, whether we are evaluating the φ formula, this value represents the number of one-day log whose related *overflow attack* model is wrongly labelled as *false* by the formal verification environment).

The F-Measure is a measure of a test's accuracy. This score can be interpreted as a weighted average of the precision and recall:

$$lclF-Measure = 2 * \frac{Precision * Recall}{Precision + Recall},$$

where *Precision* and *Recall* are obtained by following the formulae above explained.

The Accuracy is the fraction of the model correctly identified, and it is computed as the sum of true positives and negatives divided all the evaluated models:

$$lclAccuracy = \frac{tp + tn}{tp + fn + fp + tn},$$

where tn indicates the number of true negatives (for instance, whether we are evaluating the φ formula, this value represents the number of one-day logs whose related *legitimate* model is correctly labelled as *false* by the formal verification environment).

By the metrics computation, we reach a value equal to 1 for all metrics we have considered (i.e., Precision, Recall, F-Measure and Accuracy). This is symptomatic that the designed approach is able to correctly identify the several irregular behaviours we considered in the experiment with no false positive.

5. Conclusion and future work. The risks of an attack on SCADA systems are concrete and real: these are very sensitive objectives given the importance of the processes they govern and the impact that a disservice can cause on the community. Precisely, because of their structure distributed over the territory, hitting a SCADA system can really affect hundreds or thousands of other sites or plants, which in turn become unmanageable and therefore dangerous. There are countless organizations and goals that move interests in carrying out or making increasingly sophisticated actions and attack strategies, ranging from fraud and physical to creating disservices.

For this reason, to mitigate the irregular behaviours in critical infrastructure and, in particular, in SCADA water distribution systems, in this paper, we design a timed automata-based approach to identify OF , UF and R irregular behaviours on the water distribution system.

We propose formal model logs obtained from SCADA systems in terms of a timed automata network by exploiting the UPPAAL formal verification environment.

An accuracy equal to 1 is reached, symptomatic of the effectiveness of the proposed approach in attack detection in water distribution systems.

The proposed method can be adopted for real-time irregular behaviour detection in critical systems while the monitored SCADA system is under attack. As a matter of fact, in critical contexts, it is of fundamental importance to identify a threat when the latter is in progress in order to minimize data and put the system in a position to work safely. Moreover, considering that the proposed approach exploits formal methods and temporal logic formula, it is possible to understand the reason why a certain property is resulting false using the counterexample. Although the principal aim of the counterexample is to assist the designer in finding the source of the error in complex systems design, the counterexample can be exploited for other purposes. For instance, in the context of the proposed method can be used to understand how many time intervals the property will become true, thus providing a sort of probability of risk that irregular behavior may occur. This aspect can be of interest because it can be considered to forecast future irregular behaviours, thus taking measures before they can possibly take place.

As future work, we plan to evaluate the proposed of other SCADA critical systems (for instance, relating to oil/gas or power management and distribution) with the aim to validate the proposed method in another critical environment. Moreover, we plan to evaluate the proposed method with systems with much faster dynamics (for instance, power distribution/generation), i.e., with significantly larger models.

References

1. S. Cheruvu, A. Kumar, N. Smith, and D.M. Wheeler, "Conceptualizing the secure internet of things," in *Demystifying Internet of Things Security*, pp. 1–21, Springer, 2020.
2. K. Jia, J. Xiao, S. Fan, and G. He, "A mqtt/mqtt-sn-based user energy management system for automated residential demand response: Formal verification and cyber-physical performance evaluation," *Applied Sciences*, vol. 8, no. 7, p. 1035, 2018.
3. S.A. Boyer, SCADA: supervisory control and data acquisition. International Society of Automation, 2009.
4. B. Miller and D.C. Rowe, "A survey scada of and critical infrastructure incidents.," *RIIT*, vol. 12, pp. 51–56, 2012.
5. R. Taormina, S. Galelli, N.O. Tippenhauer, E. Salomons, A. Ostfeld, D.G. Eliades, M. Aghashahi, R. Sundararajan, M. Pourahmadi, M.K. Banks, et al., "Battle of the attack

- detection algorithms: Disclosing cyber attacks on water distribution networks”, *Journal of Water Resources Planning and Management*, vol. 144, no. 8, p. 04018048, 2018.
6. P.K. Hajoary and K. Akhilesh, “Role of government in tackling cyber security threat,” in *Smart Technologies*, pp. 79–96, Springer, 2020.
7. R. Meyur, “A bayesian attack tree based approach to assess cyber-physical security of power system,” in *2020 IEEE Texas Power and Energy Conference (TPEC)*, pp. 1–6, IEEE, 2020.
8. I.N. Fovino, A. Carcano, M. Masera, and A. Trombetta, “An experimental investigation of malware attacks on scada systems,” *International Journal of Critical Infrastructure Protection*, vol. 2, no. 4, pp. 139–145, 2009.
9. A. Carcano, I.N. Fovino, M. Masera, and A. Trombetta, “Scada malware, a proof of concept,” in *International Workshop on Critical Information Infrastructures Security*, pp. 211–222, Springer, 2008.
10. T. Alladi, V. Chamola, and S. Zeadally, “Industrial control systems: Cyberattack trends and countermeasures,” *Computer Communications*, 2020.
11. F. Daryabar, A. Dehghantanha, N.I. Udzir, S. bin Shamsuddin, et al., “Towards secure model for scada systems,” in *Proceedings Title: 2012 International Conference on Cyber Security, Cyber Warfare and Digital Forensic (CyberSec)*, pp. 60–64, IEEE, 2012.
12. T. Wu, J.F.P. Disso, K. Jones, and A. Campos, “Towards a scada forensics architecture,” in *1st International Symposium for ICS SCADA Cyber Security Research 2013 (ICS-CSR 2013) 1*, pp. 12–21, 2013.
13. D. Upadhyay and S. Sampalli, “Scada (supervisory control and data acquisition) systems: Vulnerability assessment and security recommendations,” *Computers Security*, vol. 89, p. 101666, 2020.
14. M. Gaiceanu, M. Stanculescu, P.C. Andrei, V. Solcanu, T. Gaiceanu, and H. Andrei, “Intrusion detection on ics and scada networks,” in *Recent Developments on Industrial Control Systems Resilience*, pp. 197–262, Springer, 2020.
15. M.G. Cimino, N. De Francesco, F. Mercaldo, A. Santone, and G. Vaglini, “Model checking for malicious family detection and phylogenetic analysis in mobile environment,” *Computers Security*, vol. 90, p. 101691, 2020.
16. L. Brunese, F. Mercaldo, A. Reginelli, and A. Santone, “Formal methods for prostate cancer gleason score and treatment prediction using radiomic biomarkers,” *Magnetic resonance imaging*, vol. 66, pp. 165–175, 2020.
17. L. Brunese, F. Mercaldo, A. Reginelli, and A. Santone, “An ensemble learning approach for brain cancer detection exploiting radiomic features,” *Computer methods and programs in biomedicine*, vol. 185, p. 105134.
18. L. Brunese, F. Mercaldo, A. Reginelli, and A. Santone, “Prostate gleason score detection and cancer treatment through real-time formal verification,” *IEEE Access*, vol. 7, pp. 186236–186246, 2019.
19. F. Mercaldo, F. Martinelli, and A. Santone, “Real-time scada attack detection by means of formal methods,” in *2019 IEEE 28th International Conference on Enabling Technologies: Infrastructure for Collaborative Enterprises (WETICE)*, pp. 231–236, IEEE, 2019.
20. R. Alur and D. Dill, “Automata for modeling real-time systems,” in *International Colloquium on Automata, Languages, and Programming*, pp. 322–335, Springer, 1990.
21. R. Alur and D.L. Dill, “A theory of timed automata,” *Theoretical computer science*, vol. 126, no. 2, pp. 183–235, 1994.
22. G. Behrmann, K.G. Larsen, O. Moller, A. David, P. Pettersson, and W. Yi, “Uppaal-present and future,” in *Proceedings of the 40th IEEE Conference on Decision and Control (Cat. No. 01CH37228)*, vol. 3, pp. 2881–2886, IEEE, 2001.

23. G. Behrmann, A. David, and K.G. Larsen, "A tutorial on uppaal 4.0," Department of computer science, Aalborg university, 2006.
24. G. Behrmann, A. David, and K.G. Larsen, "A tutorial on uppaal," in Formal methods for the design of real-time systems, pp. 200–236, Springer, 2004.
25. P. Bouyer, "Model-checking timed temporal logics," Electronic Notes in Theoretical Computer Science, vol. 231, pp. 323–341, 2009.
26. E.M. Clarke, E.A. Emerson, and A.P. Sistla, "Automatic verification of finite-state concurrent systems using temporal logic specifications," ACM Transactions on Programming Languages and Systems (TOPLAS), vol. 8, no. 2, pp. 244–263, 1986.
27. N.D. Francesco, G. Lettieri, A. Santone, and G. Vaglini, "Heuristic search for equivalence checking," Software and System Modeling, vol. 15, no. 2, pp. 513–530, 2016.
28. H.E. Jensen, K.G. Larsen, and A. Skou, "Scaling up uppaal," in International Symposium on Formal Techniques in Real-Time and Fault-Tolerant Systems, pp. 19–30, Springer, 2000.
29. J. Dougherty, R. Kohavi, and M. Sahami, "Supervised and unsupervised discretization of continuous features," in Machine Learning Proceedings 1995, pp. 194–202, Elsevier, 1995.
30. F. Mercaldo and A. Santone, "Deep learning for image-based mobile malware detection," Journal of Computer Virology and Hacking Techniques, pp. 1–15, 2020.
31. A. De Lorenzo, F. Martinelli, E. Medvet, F. Mercaldo, and A. Santone, "Visualizing the outcome of dynamic analysis of android malware with vizmal," Journal of Information Security and Applications, vol. 50, p. 102423, 2020.
32. R. Taormina, S. Galelli, N.O. Tippenhauer, A. Ostfeld, and E. Salomons, "Assessing the effect of cyber-physical attacks on water distribution systems," in World Environmental and Water Resources Congress 2016, pp. 436–442, 2016.
33. E. Salomons and A. Ostfeld, "Simulation of cyber-physical attacks on water distribution systems with epanet," in Proceedings of the Singapore Cyber-Security Conference (SGCRC) 2016: Cyber-Security by Design, vol. 14, p. 123, 2016.

Mercaldo Francesco — Ph.D., Researcher, University of Molise. Research interests: artificial intelligence, mobile computing, Android (operating system), invasive software, security of data, biomedical MRI, medical image processing, pattern classification. The number of publications — 61. francesco.mercaldo@iit.cnr.it; 1, Via Francesco De Sanctis St., 86100, Campobasso, Italy; office phone: +3908744041, 171.

Martinelli Fabio — Ph.D., Dr.Sci., Research director, institute for informatics and telematics, National Research Council of Italy. Research interests: formal analysis of network and system security, security protocols, secure service, information flow analysis, PKI and trust management, access and usage control, web and GRID services, mobile devices, formal analysis of concurrent, distributed, mobile systems, program synthesis. The number of publications — 166. fabio.martinelli@iit.cnr.it; 1, Via Giuseppe Moruzzi St., 56124, Pisa, Italy; office phone: +390503153425, 443.

Santone Antonella — Ph.D., Associate professor, Department of engineering, University of Molise. Research interests: formal description techniques, temporal logic, concurrent and distributed systems modeling, heuristic search, formal methods for systems biology and for security engineering. The number of publications — 138. antonella.santone@unimol.it; 1, Via Francesco De Sanctis St., 86100, Campobasso, Italy; office phone: +3908744041, 171.

Ф. МЕРКАЛЬДО, Ф. МАРТИНЕЛЛИ, А. САНТОНЕ
**ПРОВЕРКА МОДЕЛИ ДЛЯ ОБНАРУЖЕНИЯ В РЕАЛЬНОМ
ВРЕМЕНИ АТАК В СИСТЕМАХ РАСПРЕДЕЛЕНИЯ ВОДЫ**

Меркальдо Ф., Мартинелли Ф., Сантоне А. Проверка модели для обнаружения в реальном времени атак в системах распределения воды.

Аннотация. Системы распределения воды представляют собой критическую инфраструктуру. Эти архитектуры очень важны, и нестандартное поведение может отразиться на безопасности человека. Фактически, злоумышленник, получивший контроль над такой архитектурой, может нанести множество повреждений как инфраструктуре, так и людям. В этой статье мы предлагаем подход к выявлению нестандартного поведения, ориентированного на системы распределения воды. Разработанный подход рассматривает формальную среду проверки. Журналы, полученные из систем распределения воды, анализируются в формальную модель, и, используя временную логику, мы характеризуем поведение системы распределения воды во время атаки. Оценка, относящаяся к системе распределения воды, подтвердила эффективность разработанного подхода при выявлении трех различных нестандартных режимов работы.

Ключевые слова: критическая инфраструктура, SCADA, формальная среда верификации, формальные методы, таймер, безопасность, охрана.

Меркальдо Франческо — Ph.D., научный сотрудник, Университет Молизе. Область научных интересов: искусственный интеллект, мобильные вычисления, Android (операционная система), инвазивное программное обеспечение, безопасность данных, биомедицинская МРТ, обработка медицинских изображений, классификация образов. Число научных публикаций — 61. francesco.mercaldo@iit.cnr.it; Виа Франческо де Санктис, 1, 86100, Кампобассо, Италия; р.т.: +3908744041, 171.

Мартинелли Фабио — Ph.D., Dr.Sci., директор по исследованиям, институт информатики и телематики, Национальный исследовательский совет Италии. Область научных интересов: формальный анализ сетевой и системной безопасности, протоколы безопасности, безопасный сервис, анализ информационных потоков, PKI и доверительное управление, контроль доступа и использования, веб- и GRID-сервисы, мобильные устройства, формальный анализ параллельных, распределенных, мобильных систем, программный синтез. Число научных публикаций — 166. fabio.martinelli@iit.cnr.it; Виа Джузеппе Моруцци, 1, 56124, Пиза, Италия; р.т.: +390503153425, 443.

Сантоне Антонелла — Ph.D., доцент, инженерный факультет, Университет Молизе. Область научных интересов: методы формального описания, темпоральная логика, моделирование параллельных и распределенных систем, эвристический поиск, формальные методы системной биологии и техники безопасности. Число научных публикаций — 138. antonella.santone@unimol.it; Виа Франческо де Санктис, 1, 86100, Кампобассо, Италия; р.т.: +3908744041, 171.

Литература

1. S. Cheruvu, A. Kumar, N. Smith, and D.M. Wheeler, "Conceptualizing the secure internet of things," in *Demystifying Internet of Things Security*, pp. 1–21, Springer, 2020.
2. K. Jia, J. Xiao, S. Fan, and G. He, "A mqtt/mqtt-sn-based user energy management system for automated residential demand response: Formal verification and cyber-physical performance evaluation," *Applied Sciences*, vol. 8, no. 7, p. 1035, 2018.
3. S.A. Boyer, SCADA: supervisory control and data acquisition. International Society of Automation, 2009.
4. B. Miller and D.C. Rowe, "A survey scada of and critical infrastructure incidents.," *RIIT*, vol. 12, pp. 51–56, 2012.
5. R. Taormina, S. Galelli, N.O. Tippenhauer, E. Salomons, A. Ostfeld, D.G. Eliades, M. Aghashahi, R. Sundararajan, M. Pourahmadi, M.K. Banks, et al., "Battle of the attack detection algorithms: Disclosing cyber attacks on water distribution networks", *Journal of Water Resources Planning and Management*, vol. 144, no. 8, p. 04018048, 2018.
6. P.K. Hajoary and K. Akhilesh, "Role of government in tackling cyber security threat," in *Smart Technologies*, pp. 79–96, Springer, 2020.
7. R. Meyur, "A bayesian attack tree based approach to assess cyber-physical security of power system," in *2020 IEEE Texas Power and Energy Conference (TPEC)*, pp. 1–6, IEEE, 2020.
8. I.N. Fovino, A. Carcano, M. Masera, and A. Trombetta, "An experimental investigation of malware attacks on scada systems," *International Journal of Critical Infrastructure Protection*, vol. 2, no. 4, pp. 139–145, 2009.
9. A. Carcano, I.N. Fovino, M. Masera, and A. Trombetta, "Scada malware, a proof of concept," in *International Workshop on Critical Information Infrastructures Security*, pp. 211–222, Springer, 2008.
10. T. Alladi, V. Chamola, and S. Zeadally, "Industrial control systems: Cyberattack trends and countermeasures," *Computer Communications*, 2020.
11. F. Daryabar, A. Dehghantanha, N.I. Udzir, S. bin Shamsuddin, et al., "Towards secure model for scada systems," in *Proceedings Title: 2012 International Conference on Cyber Security, Cyber Warfare and Digital Forensic (CyberSec)*, pp. 60–64, IEEE, 2012.
12. T. Wu, J.F.P. Disso, K. Jones, and A. Campos, "Towards a scada forensics architecture," in *1st International Symposium for ICS SCADA Cyber Security Research 2013 (ICS-CSR 2013)* 1, pp. 12–21, 2013.
13. D. Upadhyay and S. Sampalli, "Scada (supervisory control and data acquisition) systems: Vulnerability assessment and security recommendations," *Computers Security*, vol. 89, p. 101666, 2020.
14. M. Gaiceanu, M. Stanculescu, P.C. Andrei, V. Solcanu, T. Gaiceanu, and H. Andrei, "Intrusion detection on ics and scada networks," in *Recent Developments on Industrial Control Systems Resilience*, pp. 197–262, Springer, 2020.
15. M.G. Cimino, N. De Francesco, F. Mercaldo, A. Santone, and G. Vaglini, "Model checking for malicious family detection and phylogenetic analysis in mobile environment," *Computers Security*, vol. 90, p. 101691, 2020.
16. L. Brunese, F. Mercaldo, A. Reginelli, and A. Santone, "Formal methods for prostate cancer gleason score and treatment prediction using radiomic biomarkers," *Magnetic resonance imaging*, vol. 66, pp. 165–175, 2020.
17. L. Brunese, F. Mercaldo, A. Reginelli, and A. Santone, "An ensemble learning approach for brain cancer detection exploiting radiomic features," *Computer methods and programs in biomedicine*, vol. 185, p. 105134.

18. L. Brunese, F. Mercaldo, A. Reginelli, and A. Santone, "Prostate gleason score detection and cancer treatment through real-time formal verification," *IEEE Access*, vol. 7, pp. 186236–186246, 2019.
19. F. Mercaldo, F. Martinelli, and A. Santone, "Real-time scada attack detection by means of formal methods," in *2019 IEEE 28th International Conference on Enabling Technologies: Infrastructure for Collaborative Enterprises (WETICE)*, pp. 231–236, IEEE, 2019.
20. R. Alur and D. Dill, "Automata for modeling real-time systems," in *International Colloquium on Automata, Languages, and Programming*, pp. 322–335, Springer, 1990.
21. R. Alur and D.L. Dill, "A theory of timed automata," *Theoretical computer science*, vol. 126, no. 2, pp. 183–235, 1994.
22. G. Behrmann, K.G. Larsen, O. Moller, A. David, P. Pettersson, and W. Yi, "Uppaal-present and future," in *Proceedings of the 40th IEEE Conference on Decision and Control (Cat. No. 01CH37228)*, vol. 3, pp. 2881–2886, IEEE, 2001.
23. G. Behrmann, A. David, and K.G. Larsen, "A tutorial on uppaal 4.0," *Department of computer science, Aalborg university*, 2006.
24. G. Behrmann, A. David, and K.G. Larsen, "A tutorial on uppaal," in *Formal methods for the design of real-time systems*, pp. 200–236, Springer, 2004.
25. P. Bouyer, "Model-checking timed temporal logics," *Electronic Notes in Theoretical Computer Science*, vol. 231, pp. 323–341, 2009.
26. E.M. Clarke, E.A. Emerson, and A.P. Sistla, "Automatic verification of finite-state concurrent systems using temporal logic specifications," *ACM Transactions on Programming Languages and Systems (TOPLAS)*, vol. 8, no. 2, pp. 244–263, 1986.
27. N.D. Francesco, G. Lettieri, A. Santone, and G. Vaglini, "Heuristic search for equivalence checking," *Software and System Modeling*, vol. 15, no. 2, pp. 513–530, 2016.
28. H.E. Jensen, K.G. Larsen, and A. Skou, "Scaling up uppaal," in *International Symposium on Formal Techniques in Real-Time and Fault-Tolerant Systems*, pp. 19–30, Springer, 2000.
29. J. Dougherty, R. Kohavi, and M. Sahami, "Supervised and unsupervised discretization of continuous features," in *Machine Learning Proceedings 1995*, pp. 194–202, Elsevier, 1995.
30. F. Mercaldo and A. Santone, "Deep learning for image-based mobile malware detection," *Journal of Computer Virology and Hacking Techniques*, pp. 1–15, 2020.
31. A. De Lorenzo, F. Martinelli, E. Medvet, F. Mercaldo, and A. Santone, "Visualizing the outcome of dynamic analysis of android malware with vizmal," *Journal of Information Security and Applications*, vol. 50, p. 102423, 2020.
32. R. Taormina, S. Galelli, N.O. Tippenhauer, A. Ostfeld, and E. Salomons, "Assessing the effect of cyber-physical attacks on water distribution systems," in *World Environmental and Water Resources Congress 2016*, pp. 436–442, 2016.
33. E. Salomons and A. Ostfeld, "Simulation of cyber-physical attacks on water distribution systems with epanet," in *Proceedings of the Singapore Cyber-Security Conference (SGCRC) 2016: Cyber-Security by Design*, vol. 14, p. 123, 2016.

P. BUI, M. LE, B. HOANG, N. NGOC, H. PHAM

DATA PARTITIONING AND ASYNCHRONOUS PROCESSING TO IMPROVE THE EMBEDDED SOFTWARE PERFORMANCE ON MULTICORE PROCESSORS

Bui P., Le M, Hoang B., Ngoc N., Pham H. Data Partitioning and Asynchronous Processing to Improve the Embedded Software Performance on Multicore Processors.

Abstract. Nowadays, ensuring information security is extremely inevitable and urgent. We are also witnessing the strong development of embedded systems, IoT. As a result, research to ensure information security for embedded software is being focused. However, studies on optimizing embedded software on multi-core processors to ensure information security and increase the performance of embedded software have not received much attention. The paper proposes and develops the embedded software performance improvement method on multi-core processors based on data partitioning and asynchronous processing. Data are used globally to be retrieved by any threads. The data are divided into different partitions, and the program is also installed according to the multi-threaded model. Each thread handles a partition of the divided data. The size of each data portion is proportional to the processing speed and the cache size of the core in the multi-core processor. Threads run in parallel and do not need synchronization, but it is necessary to share a general global variable to check the executing status of the system. Our research on embedded software is based on data security, so we have tested and assessed the method with several block ciphers like AES, DES, etc., on Raspberry PI3. The average performance improvement rate achieved was 59.09%.

Keywords: embedded software performance improvement, multicore processors, multithread, data partitioning, asynchronous processing.

1. Introduction. Data encryption has been studied with many algorithms such as AES, DES, etc., to ensure data security and is focused on research for sequential processing software systems. It hasn't been being studied on embedded systems. Because encrypting data will increase the cost of data processing time compared to unencrypted data, the problem is to both ensure data security and maximize performance by adapting encryption algorithm based on embedded software configuration.

Today, in the fast development trend of the 4th industrial revolution, information technology systems, especially IoT systems, also flourish in both hardware and programming models. In this trend, multicore processors are widely studied and applied. Not only information technology systems in general, but embedded systems also use more and more multicore processors.

The purpose of using multicore processors is to improve the performance of embedded systems under limited resource conditions. There are many studies that have been done to improve the performance of embedded software. Improvements to embedded software performance can be made on different levels: from processor level to operating system, programming and design levels. However, most of the old programming

models, which have only been programmed in serial/sequential processing, have not yet promoted the parallel processing.

In recent years, there have also been several teams that have performed researches and developed methods to improve embedded software performance on multicore processors. These studies mainly focus on multithread models – the distribution of processing flows for parallel implementation as the optimal approach based on co-design is also studied and applied to multicore processors, such as [1]. In the study, the author proposed and developed a method of co-designing hardware and software for network systems on a chip (NoC), which improved the average performance by 33.1%. However, due to a data conflict between the local data set of the failed transaction and the global end-level caching (LLC), the transaction cancellation rate is still high in software transactional memory, and it is costly to identify and update the cancelled data.

In study [2], the authors built a user space memory scheduler that allocates the ideal memory node for tasks by monitoring the characteristics of the heterogeneous memory architecture to optimize application performance for the NUMA multicore processors. The average performance improvement rate achieved by this method is 25%. The authors focused on the development of the scheduler but used the Princeton Application Repository for Shared-Memory Computers (PARSEC); therefore, it is not yet possible to evaluate the compatibility of a scheduler with the ability to parallel the data on the NUMA configuration.

Study [3] presented a multi-threaded approach to the travelling salesman problem to improve performance. Under certain hardware limitations, the proposed math can take full advantage of multi-core chips and effectively balance out the contradiction between increasing data size and computing efficiency, thus gaining a satisfactory solution. Nevertheless, a parallel threaded approach that depends on the set of input parameters of the problem is not new; instead, it just gives case studies on thread programming to solve the problem.

Study [4] presented a platform used in multithreaded programming to improve performance as OpenMP. The platform can be used to develop both desktop applications and embedded apps. This study analyses the effects of different schedules and segment sizes on the at-gain speed of multi-core platforms that use different shared memory in regular workloads. The results illustrated that different multi-core technology showed different acceleration values, and different multi-core platforms were better than others in terms of speed as the number of cores was increased.

In study [5], the authors presented the method of sharing the data to be processed to points to improve the thoughts of points in the problem.

This study demonstrates that each specific problem has its own method of implementing performance improvement.

Studies presented on hardware configuration-based performance improvement methods are in [6]. The authors proposed a single-command, multi-data model to improve performance for multi-disciplinary chips.

Most of these studies focus on solving only one layer of specific problems and have not fully solved the problem of synchronization, linking data between threads. At the same time, the studies have not looked at data independence and data partitioning for parallel processing.

With big data processing problems and independent data sections, a more efficient model is needed to promote data independence and simplify synchronized time, linking data between threads. Therefore, in this study, we propose a model of data partitioning and asynchronous processing to process data in parallel to improve embedded software performance.

In study [7], the authors proposed a parallel multithreaded pipeline to filter, clean and classify information in the information collection phase. We developed the pipeline so that it can be easily re-applied to any type of heterogeneous aggregation and run efficiently on medium to low resource infrastructures where I/O speed is the main limitation.

In consideration of undesirable elements in synchronous programming, asynchronous programming has emerged as a programming style to overcome the limitations. Usually, in asynchronous programming models, the methods are included in the queue list for implementation, and the order in which the method is implemented is serial yet unidentified. Nevertheless, a lag still exists with this method due to the need for serial / sequential execution; therefore, there have been researches being carried out to solve the parallel asynchronous problems. The idea of asynchronous programming is to divide the execution of the original program into tasks that are running over a short period of time. Furthermore, each task is performed as recursive sequence software that defines new methods to be enforced later when accessing shared memory.

The simplicity of the asynchronous program's combination makes asynchronous programming a preferred option for deploying API or systems that require reactive systems. The use of asynchronous programming for main servers, desktop applications, and embedded systems is increasingly being used and varied with problems [8], such as JavaScript tools of modern web browsers, Grand Central Dispatch in MacOS and iOS, Linux job queues, asynchrony in .NET, and deferred procedure calls in Windows core are all based on asynchronous programming. Even in single processing settings (i.e. without any parallelism), asynchronized frameworks such as Node.js are becoming widely used to design extremely

scalable (web) servers.

Many studies have been carried out to present asynchronous programming models. Nevertheless, some tasks posted in asynchronous programs may have different levels of execution priority, as addressed in study [8]. A major drawback of existing tasks considering execution priorities is the complexity of models that host those priorities. An asynchronous program divides the behaviour of the accumulated program into short-running tasks. Each task basically acts as a recursive sedated program, in addition to accessing the memory shared by all tasks that can post new tasks to the task buffer for later execution. Tasks from each buffer perform in succession: when one task is completed, another is taken from the buffer and run until completed. Programming a reactive system requires the designer simply to ensure that no single task is performed for too long to prevent others more urgent tasks from being executed. Figure 1 illustrates the general architecture of asynchronous programs.

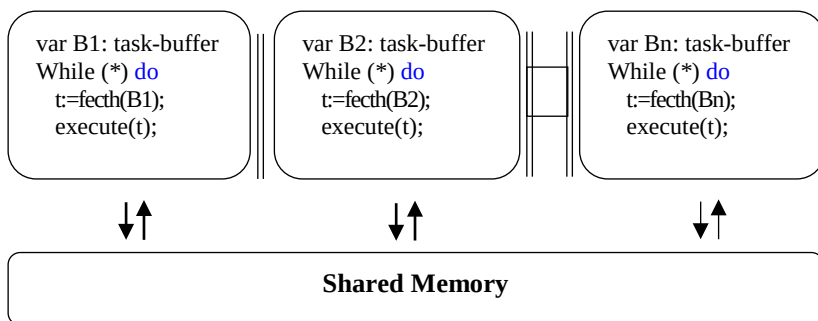


Fig. 1. A general architecture of asynchronous programs with N task buffers. Tasks from the same buffer execute serially, but concurrently with tasks of other buffers [8]

In study [9], the authors and his colleagues proposed an official model of asynchronous event-oriented programs, which narrows the semantic gap between programs and existing models, especially by allowing dynamic creation of tasks, events, buffer tasks and strings simultaneously, and accurately capture the interaction between these quantities. Instead of assigning each calculated task to a blocked dedicated thread for certain conditions, the system maintains simple sets of events that pend tasks, buffers of tasks with enabled events and workflows for performing tasks in the buffer.

Figure 2 presents asynchronous event-driven programs. The pending tasks (drawn as triangles) are moved to their designated task buffers (drawn

as boxes) once their designated events (drawn as circles) are triggered. Threads (drawn as diamonds) execute buffered tasks to completion, such that no two tasks from the same buffer (drawn with the same colour) execute in parallel. In addition to the ability to use events to synchronize between tasks, the task cache itself provides another means of coordination: the system can allow tasks from separate buffers to perform in parallel while ensuring that tasks from the same buffer perform sequentially. Authoring the management of events, tasks, buffers, and flows to the system often increases the efficiency of the system.

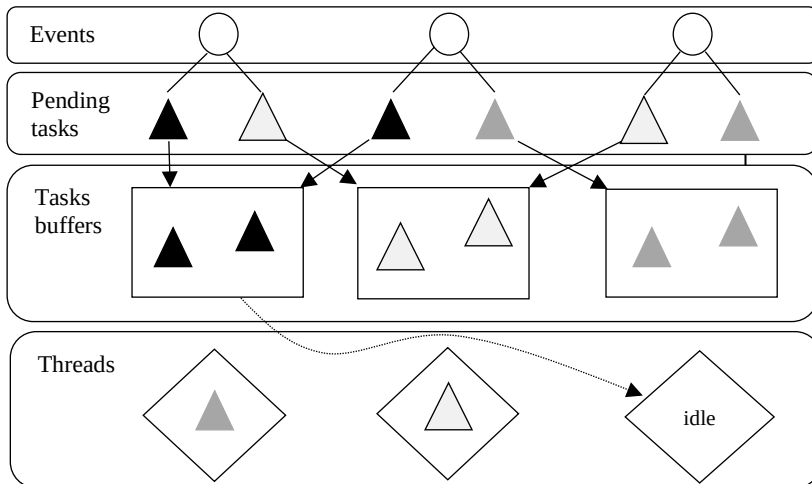


Fig. 2. Asynchronous event-driven programs [9]

Most of the studies only focus on solving a specific class of problems and have not completely solved the problem of synchronization and data linkage between threads. At the same time, studies have not considered the issue of data independence and data division for parallel processing.

With problems related to a big amount of processing data and/or independent pieces of data, a more efficient model is needed to promote the independence of the data and minimize the time for synchronization and linked data between the streams. Therefore, in this study, we propose a model of data division and asynchronous processing for parallel data processing to improve embedded software performance.

The rest of the paper is organized as follows: Section II – Survey, analysis, synthesis of related research; Section III – Presentation on ideas,

process and content of method's development; Section IV – Experiment for method's testing and evaluation; Section V – Conclusion and trends of development.

2. Related work. The problem of optimizing performance based on parallel calculations has been researched and developed by many research groups in different approaches such as CPU hardware level, co-design level, operating system level and application level multi-threading. The CPU level parallelization includes typical optimization techniques such as order scheduling, command pipeline, in-of-order, out-of-order, CPU configuration, setting configuration, etc. [10]. These algorithms and methods are the basis for the direction of research and implementation approach of study [10].

Study [11] analysed cipher-related algorithms and then designed and deployed the AES algorithm in CUDA. The data will be divided into 64K blocks by the CPU and transferred to the cores of the GPU and performed in parallel the sub Bytes, Shift row, Mix column and AddroundKey functions for encryption. Encrypted data were collected on the GPU and transferred upwards to the CPU. This technique improved the performance of the AES encryption process with the help of a GPU. However, the cost was involved in transporting data from the CPU to the GPU and vice versa.

Data partitioning as an important sub-factor in database tools has been studied in the previous work. Polychroniou, O. and Ross, K.A. [12] provided extensive analysis of data partitioning across multiple methods, such as partition type (base, hash or range) and shuffle strategy. It has been proved that for more than 16 partitions, partitions that write in combination with direct memory writing and skipping cache memory would work best. Partition throughputs are reported to be 1.1 billion data sets/s for 8,192 partitions with 64 threaded parallel executions on 32-core servers.

Schuhknecht F.M. et al. in [13] presented a set of experiments performing base-based partitioning. Existing optimizations (combinations, time-consuming storage, etc.) are enabled, and new optimizations (pre-fetch for recording, micro-row layout) are added step by step to observe their impact on the total duration of implementation. The fastest FPGA data partition deployment to date was presented by Wang and his associates [14] with 256 Million Data Sets/s for 8,192 partitions. They improved the existing OpenCL implementation of a partitioned and deployed it on the FPGA. The partitioned data are assumed to be in the DRAM that is directly connected to the FPGA. The resulting partitions are recorded to the same DRAM, and the transfer via PCIe to the server memory is necessary if the partitioned data will be used by the CPU for further operations.

Study [15] shows that the probability of parallel calculation of data depends heavily on its data partitions. The solutions implemented by the status quo of the systems have not yet been optimized. Community-recommended techniques for finding optimal data partitions are not applied directly when relevant to complex user-defined data functions and models. A parallel data program compiled into an execution plan chart (Execution Plan Graph -EPG) is a directional AC chart with multiple stages. Therefore, data partitioning affects many aspects of how a job is run in a cluster, including parallelism, workloads for each vertex, and network traffic between vertices. In this study, the authors proposed the system architecture as Figure 3 to partition the data.

Figure 3 shows the architecture of the system. First, the system compiles a certain data parallel program into a work plan chart (EPG) with the original data partitions (e.g., provided by the user). The code analysis module takes this EPG and the code for each vertex in the EPG as input to infer the calculated complexity of the program per-vertex and important data features. This step is important because it not only provides information about the relationship between the input data size versus the cost of calculation and I/O but also guides the process of data analysis, for example, providing suggestions for strategically sampling data and estimating data statistics. This information can then be used to identify image recordings to process and distribute them more evenly. The data Analysis module linearly scans data to create compact data expressions.

The Modelling/Estimation Module uses code and data analysis results to estimate the runtime cost of each vertex, including CPU time, output data size, and network traffic.

The authors in study [15] gave the example of creating a partitioned graph, as shown in Figure 4. The two root nodes represent two partitions of the sampled input data. The cost optimization module inserts an additional partitioning stage into the existing EPG in search of an optimized partitioning scheme. First, the two inputs are divided into 8 partitions (for example, $h1(k) \bmod 8$ hash partition) and EPG is updated accordingly. The Cost Estimation module then determines the critical path up to the current period in the updated EPG, including the vertex associated with Partition 5. To reduce costs, it divides Partition 5 into the other two partitions (for example, the $h2(k) \bmod 2$ hash partition). Meanwhile, Partitions 0, 2 and 4 all have small costs to reduce I/O, the cost of starting peaks and therefore the underlying overall cost. The process of estimating costs and optimizing by ingesting and separating this recursive data is repeated until it converges. Each iteration is a greedy step towards minimizing overall costs. In

summary, study [15] has proposed a system architecture to find the optimal for data partitioning based on greedy algorithms.

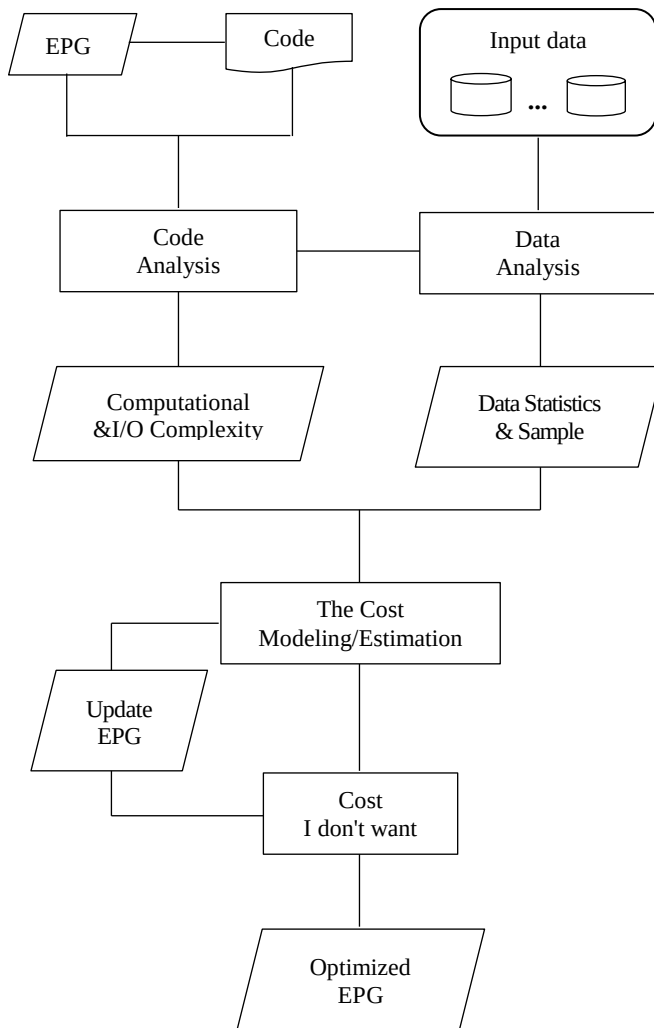


Fig. 3. System architecture, from study [15]

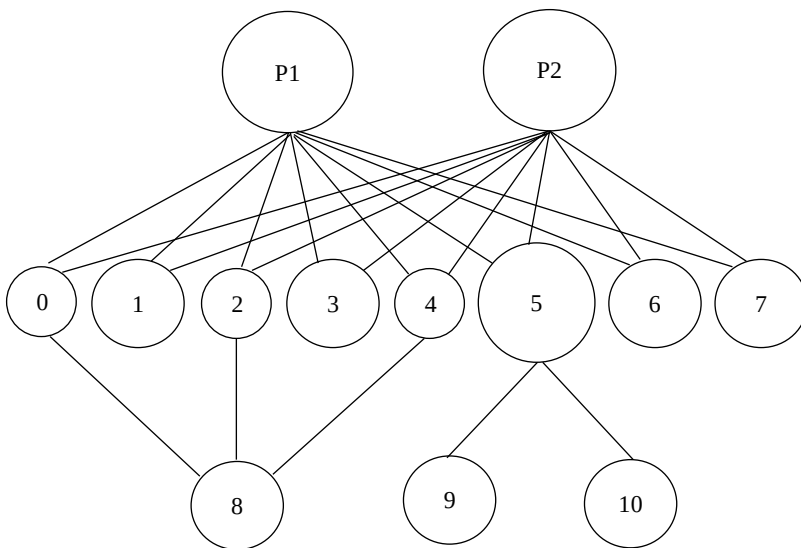


Fig. 4. Example of generating an optimized data partitioning scheme represented by the partition graph

The article [16] proposes a data set manipulation method to divide the training data set into many groups with different characteristics to train each classifier. The author partitioned the training data into groups of the same distance. If each individual SVM is trained using each of these groups, the trained classifiers have increased diversity. The test results showed an average accuracy of 70% by repeating the data set for a single letter four times.

In study [17], the authors et al. proposed an automatic local data partitioning algorithm, which can automatically recognize the local maxima of the data density from experimental observations and use them to serve as the focal point to form data partitions. This method is free of user's parameters and prior assumptions. Numerical results based on the benchmark dataset proved the validity of the proposed algorithm and demonstrated its high performance and computational efficiency compared to modern clustering algorithms.

Kara K. et al. [18] explored the use of FPGA to speed up data partitioning. The paper studies new hybrid architectures in which FPGA is placed as a co-processor located on one socket and has consistent access to the same memory as the CPU located on the other socket. Such an architecture reduces the cost of data transfer between the CPU and the

FPGA, allowing the execution of the combined algorithm in which the partition occurs on the FPGA and the construction and exploration stages of a connection that occurs on the CPU. The FPGA-only partitioning and execution reduce the cost of data transfer to the CPU and vice versa, but storing and performing on big data are matters for further study of the paper.

Study [19] considered the issue of optimal distribution of workloads for the parallel implementation of data between the processing components of non-click calculation systems. Zhong Z. et al. presented a solution that uses functional performance models (FPM) of processing agents and FPM-based data partitioning algorithms. The effectiveness of this method is proven by experiments with parallel matrix multiplication. The FPM-based data partitioning has proven to be fully and efficiently used to balance the workloads of many data parallel applications across modern hem helmless disk forms. However, more effort is still needed to improve this approach in some respects. For example, only the calculated performance of processing units is used to partition data so far, while the cost of communicating between processes is not considered. In some dense matrix applications on highly haemolysed disk forms, the performance of a processing unit may depend on the shape of the matrix block assigned to it. In that case, the multi-dimensional performance model has more than one parameter that may be required to describe the performance accurately. For large complex applications where the computational kernel cannot be easily separated from the application or there is more than one computational kernel, it may take more effort to balance the workload with this approach.

In study [20], the authors compared and evaluated three algorithms AES, DES, and 3DES by nine factors: key length, passcode type, block size, developed, code break resistance, security, ability key, ASCII printable character key, and the time it takes to check all possible keys at 50 billion seconds. The results have proven that AES algorithms are better than DES and 3DES, as shown in Table 1. However, this study did not evaluate the performance of data encryption processing of these three methods.

Study [21] compared the results of the AES encoding techniques to increase encryption efficiency on FPGA integrated circuits. Five techniques were outlined in the article.

The CB-KB-S technique: Both encryption and key creation are in the RAM Block. Processing is carried out in a singly way, first creating an extended key then the encryption process. The speed is slow due to the implementation of the order; however, it gives accurate results, therefore this technique is applied to problems that require high accuracy.

Table 1. Raspberry PI 3 embedded computer configuration

Factors	AES	3DES	DES
Key Length	128, 192, Or 256 bits	(kl, k2 and k3) 168bits (kl and k2 is same) 112 bits	56 bits
Cipher Type	Symmetric block cipher	Symmetric block cipher	Symmetric block cipher
Block Size	128, 192, or 256 bits	64 bits	64 bits
Developed	2000	1978	1977
Cryptanalysis resistance	Strong against differential, truncated differential, linear, interpolation and square attacks	Vulnerable to differential, Brute force attacker could be analyzed plaint text using differential cryptanalysis.	Vulnerable to differential and linear cryptanalysis; weak substitution tables
Security	Considered secure	One only weak which is Exit in DES.	Proven inadequate
Possible Keys	2128, 2192, or 2256	2112 or 2168	256
Possible ASCII printable character key	9516, 9524, or 9532	9514 or 9521	957
Time required to check all possible keys at 50 billion keys per second**	For a 128-bit key: 5x 1021years	For a 112-bit key: 800Days	for a 56-bit key: 400Days

The CB-KB-P technique: Both encryption and key creation are in the RAM Block but are used in a private key writing module. This module will run in parallel with the encryption process instead of performing a sequential key creation according to the CB-KB-S algorithm. With this technique, the application of parallel processing of encryption blocks increases performance due to reduced latency during the sequential process. Still, it reduces security due to the extended key born before the encryption process.

The CB-KC-S technique: With this technique, the encryption process is carried out at Block RAM, and the extension key is done in Logic Block. These two modules are done in a sequence of the extended key that is performed and then encrypted. This technique will reduce the processing in

a way that conducts two tasks in two different places. On the other hand, data exchange between the parties also causes certain latency in encryption.

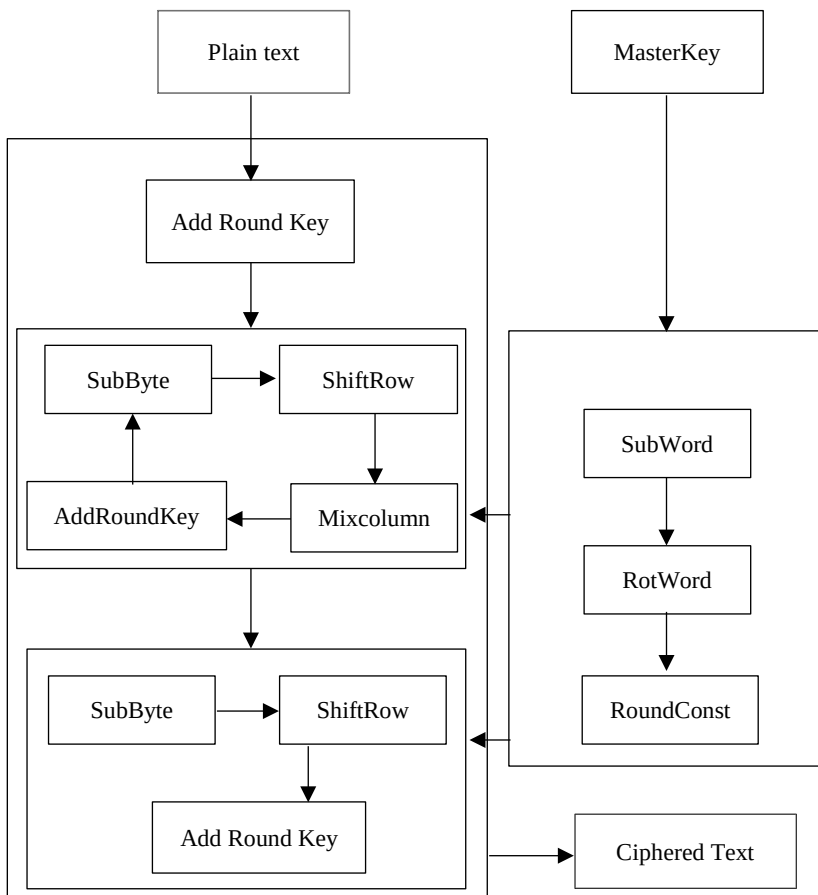


Fig. 5. Standard overview of AES algorithm

The CB-KC-P technique: like the CB-KC-S technique, the encryption process in this technique will be carried out at block RAM, and the extension key is done in Logic Block. The difference with the CB-KC-S technique is the fact that this technique conducts parallel unlocking and encryption. Since this process takes place in parallel, the performance of the technique is considered very good, but the safety of this technique, as well

as the CB-KB-P technique, is not high because there is no constraint on the encoding result with extended key birth.

The CC-KC-S technique: this technique brings both encryption and extended key being executed in Logic Block. Because of the same implementation in Logic Block, there is no longer lag due to data exchange between Block RAM and Logic Block, so the efficiency of this technique achieved is the highest.

Most of the applications are written using a synchronous batch programming model. Therefore, they are not optimal for low-latency or asynchronous communication algorithms. The study [22] proposes constructs for asynchronous multi-GPU programming and describes their implementation in a thin runtime environment, performing common math operations and distributed task lists. The authors et al. demonstrated that this approach achieves performance gains and exhibits strong scalability for heterogeneous systems, yielding more than 7x speedup for some algorithms.

In study [23], asynchronous programming was built based on a multithreaded basis due to the fact that APIs have been deployed using multithreaded programs with shared memory. The execution of software flows is not affected by the number of processors in the system, which has been proven by the fact that threads are enforced as recursive sedation software that runs simultaneously with alternating write and read commands. It is likely that being able to walk away, in this case, might cause the complexity of simultaneous programming models. This study gave the authors the idea of using shared variables for this paper's asymable parallel programming solution.

The reduction is then as follows. First, the pre-emptive priority scheduler is made explicit by adding to the program the code shown in Figure 6. The hyper period procedure executes each thread one time, choosing non-deterministically a sleeping thread to execute via the choose operation, which returns an index that satisfies the supplied guard. An infinite cycle of hyper periods is simulated by invoking hyper period in a non-terminating loop. During each hyper period, the scheduler has two tasks: (i) it must ensure that each thread i is awoken so that i can execute its task, and (ii) the wake-ups should happen non-deterministically. The first task is handled by defining a Boolean array of size n , where each entry in the array denotes whether a thread t is sleeping or not. (In Figure 6, the array is named Sleeping.) The scheduler loops until all threads have been awoken and completed their periodic task.

The second task is handled by performing a source-to-source transformation on the code of each thread so that it non-deterministically invokes Schedule before each statement st . That is if a thread is comprised

of program statements $st1, \dots, stk$, then the transformed program will have program statements $st0\ 1, \dots, st0\ k$, where each $st0$ is defined as: $st0, Schedule(); st$. In the definition of *Schedule*, the function *nondet* non-deterministically returns *true* or *false*. When *Schedule* is invoked, the code of a higher-priority thread $ti\ 0$ than the thread ti whose code is currently executing may be invoked, which corresponds to ti being pre-empted by $ti\ 0$. Before executing a thread ti by invoking $Threads[i].entry()$, the flag $Sleeping[i]$ is set to false to ensure that ti is executed exactly once per hyperperiod H .

<pre>// Sleeping flags Sleeping[n] = {true,...,true}; // Thread priorities Priorities[n] = ...; // Thread entry points Threads[n] = ...; // 0 => choose any thread Prio = 0; void Hyperperiod() { while (Vi Sleeping[i]) { j = choose j: Sleeping[j]; Sleeping[j] = false; Threads[j].entry(); } }</pre>	<pre>// Wake-up higher-priority thread void Schedule() { // Save current priority int prevPrio = Prio; for i in (1..n) { if (Priorities[i] <= Prio) continue; if (nondet() && Sleeping[i]) { Prio = i; Sleeping[i]=false; Threads[i].entry(); break; } } Prio = prevPrio; }</pre>
--	--

Fig. 6. Pseudo-code to execute one hyperperiod [23]

The study [24] proposes a light-weighted asynchronous processing solution that is based on the idea of Pareto principle about coloring algorithms. It is to separate the vertices processing based on the color distribution by partitioning the vertices to achieve the maximum parallelism and to reduce data transmission costs while processing the partitions.

Lian X. et al. have performed research [25] on two common asynchronous parallel implementations for Stochastic Gradient on computer networks and shared memory systems. Two algorithms (ASYSG-CON and ASYSG-INCON) were used to describe the two above implementations. This study was able to explain their converging and accelerating properties, mainly due to the heterogeneity of the majority of deep learning formulas and asym asymable parallel mechanisms.

In study [26], the authors modeled the dependence of the output on tens of thousands of causal events that recursively use a new time-coding

scheme for real-time processing of event streams. In tests using real data, the study proved useful in real-world applications.

Based on the results of the study reviewed above, we can see that the problem of optimizing embedded software on multi-core processors also focuses a lot on multi-threaded models to parallelize the processing flows in the program. With independent data problems, it is necessary to process a large amount of data that requires a more efficient parallel model and minimizes synchronized time, linking data between threads. This class of problems can be solved by dividing data into sub-sets for asynchronous processing. This method will be developed, tested and evaluated in the following sections.

3. Ideas of research and experimental processes

3.1. Ideas. The main idea of the proposed method is to build a global data partitioning model based on the configuration of processor builders and asynchronous data processing between threads for performance improvements. Global data are used so that every thread is accessible, and the data division creates independent data to execute concurrent processing on processor cores. An asynchronous data processing model is applied to eliminate intermediate data processing time and link data between threads.

3.2. Research and experimental process. The process of developing and experimenting with the method described as shown in Figure 7 is carried out in 4 steps as follows:

1. Analyse configuration of a processor: From the information of the processor, we analyse the processor's configuration information to give the processor's configuring the filling.
2. Calculate partitioning rates: Based on the configuration of the chromium 4th of the processor, we calculate, determine the data partitioning ratios.
3. Data partitioning: To select the rate type and global data, we divide it into data parts.
4. Multi-thread programming: From the processor configuration information and data sections, we determine the number of data and set the processing flows on each divided data area. Map the threads to the corresponding characters.

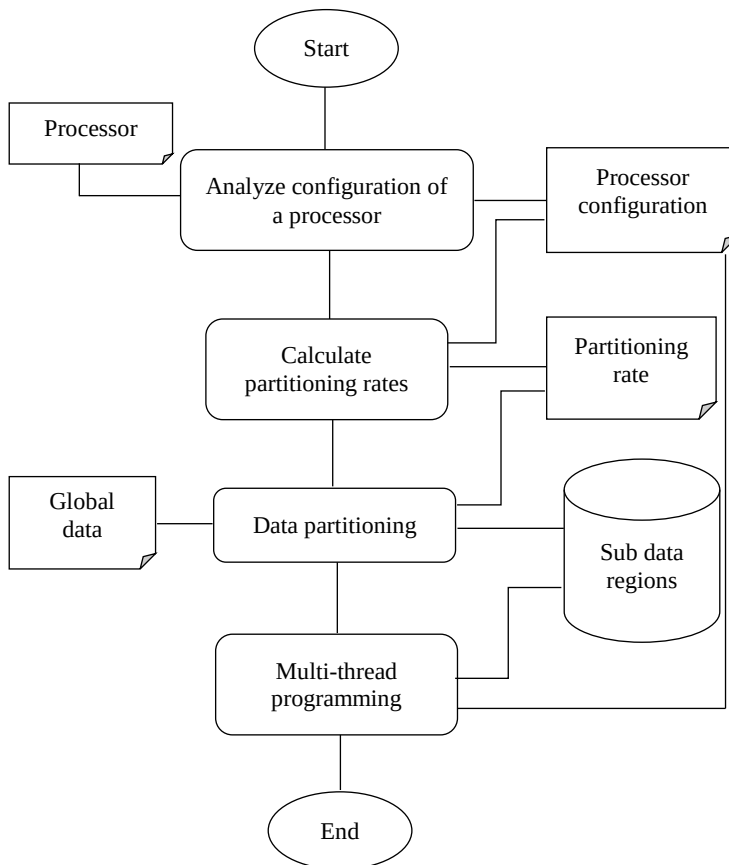


Fig. 7. Research and experimental process

3.3. Data partitioning. Data partitioning is a data delivery technique to improve data processing performance. Processing performance can be improved in one of two ways. Firstly, it is possible in some cases to pre-determine that a partition does not need to be accessed for processing depending on how the data are partitioned. Secondly, when data are partitioned, in some cases, parallelism can be achieved in data processing because different partitions can be accessed and processed in parallel. In this paper, we focus on data partitioning towards parallel processing on multi-core processor cores, since depending on the properties of the processor, the partitioning is based on the processing capacity, the size of the human cache.

Definition 1 – ratio by speed. The speed ratio is the ratio between the processing capacity of the current core in the set of microprocessors. This ratio is used to divide data by the speed of the core. The speed ratio is denoted as r_s , described in Equation (1).

$$r_s = \frac{s_i}{\sum_{i=1}^N s_i}, \quad (1)$$

where:

- s_i is the speed of the i^{th} core;
- N is the number of cores in a processor.

Definition 2 – ratio by buffer size. Caching size ratio is the rate at which data are stored on the cache of the current core, the set of microprocessors. This ratio is used to partition data according to the buffer size of the core. The buffer size ratio is denoted as r_c , described in Equation (2):

$$r_c = \frac{c_i}{\sum_{i=1}^N c_i}, \quad (2)$$

where:

- c_i is the cache size of the i^{th} core;
- N is the number of cores in a processor.

Definition 3 – aggregate ratio. The aggregate ratio is the ratio of the combination of the speed ratio and the core cache size ratio, currently the set of microprocessors. This scale is used to partition the data according to the aggregate ratio of the core. The aggregate ratio is denoted as r , described in Equation (3):

$$r = \alpha \times r_s + \beta \times r_c, \quad (3)$$

where:

- α and β are positive coefficients;
- $\alpha + \beta = 1$.

From the above definitions, we offer three ways to divide the original data set into sub-set. Call D the size of the data to be processed, and D_i is the size of the component data partition i being calculated by Equations (4), (10) and (11) depending on the division.

Split data proportional to the speed of core/processor. The size of the data divided by the speed ratio of the i^{th} component is determined based on the speed ratio of the i^{th} core and the size of the data to be processed:

$$Ds_i = r_s \times D = \frac{s_i}{\sum_{i=1}^N s_i} \times D. \quad (4)$$

Proposition. Suppose we have a processor with N cores, each core has the same speed s_i , has the same cache capacity, and the performance r_s is calculated by definition 1. There is data block D divided by cores that execute concurrently.

According to the normal way of data division, there are N cores, D data block will be divided into N parts then, each core will perform D/N data processing. Since the cores execute concurrently, the time it takes to process the data block D needs to run out t_{max} , where t_{max} is the time of the lowest-speed core s_{min} or, in other words, the time that the least-performing core needs to process the assigned data portion.

$$t_{max} = \frac{D}{N} \times \frac{1}{s_{min}}. \quad (5)$$

We improved the data division to increase the performance of the above D block processing by dividing the volume for each core based on the performance of each core. Then the data is divided according to Equations (4). Since the cores execute concurrently and the data of each core is divided according to the performance of each core, the processing time of each core is equivalent and is called t .

$$t = \frac{D \times s_i}{\sum_{i=1}^N s_i} \times \frac{1}{s_i}. \quad (6)$$

Need to prove that $t < t_{max}$.

From (5),

$$t_{max} = \frac{D}{N \times s_{min}} = \frac{D}{\sum_{i=1}^N s_{min}}. \quad (7)$$

From (6),

$$t = \frac{D}{\sum_{i=1}^N s_i}. \quad (8)$$

Since s_{min} is always less than or equal to $s_i \forall i \in 1, \dots, N$, according to (7) and (8), $t < t_{max}$.

So with the performance of different cores and the data block being large enough, our data division is better than the normal data division.

Since the processing speed of the core is directly proportional to the cache capacity of the core, the above statement is true for the data division Equations (9) and (10).

Split data proportional to buffer size. The size of the data divided by the caching size ratio of the i^{th} component is determined based on the caching size ratio of the i^{th} core and the size of the data to be processed:

$$Dc_i = r_c \times D = \frac{c_i}{\sum_{i=1}^N c_i} \times D. \quad (9)$$

Split data using a combination of two buffer sizes and speed parameters. The data size of the i^{th} component is determined based on the aggregate ratio of the i^{th} core and the size of the data to be processed:

$$D_i = r \times D = (\alpha \times r_s + \beta \times r_c) \times D. \quad (10)$$

3.4. Building a multi-threaded model for asynchronization optimization. When solving the problem in a multi-threaded model, in addition to execution time context transfer time, the system also takes more time to synchronize the flows. Thread synchronized time has a great influence on system performance. When dividing data, if data processed by threads depend on each other, it is imperative to synchronize the flows. However, if the data parts are processed by independent threads, there is no need for synchronizing; just determine when all flows are done. The challenge is how to solve the problem of asynchrony with independent data. To solve this problem, we propose organizations of input, output and counting variables in the global memory area. Because it is a global memory, all threads are retrieved. At that time, the input data will be divided into regions with dimensions determined by formulas (4), (9), (10), depending on the divisions. The output data also are in global memory areas. Each thread stores the input data in the corresponding memory using this global memory. The processing process is shown in Figure 8, including the following steps:

- Init(): Function that implements two alarms of data area D as a global byte array.
- Division(D): Divides array D into D_i sub-arrays and calculates

the number of threads to perform, assigning global variables using the number of threads to perform.

- Multithread(): Execution of threads performed on 1 Array of D_i . When a thread is done, the global variable decreases by 1.
- Check (global): Check the global = 0 variable; all threads are finished, no synchronizing and ending is required.

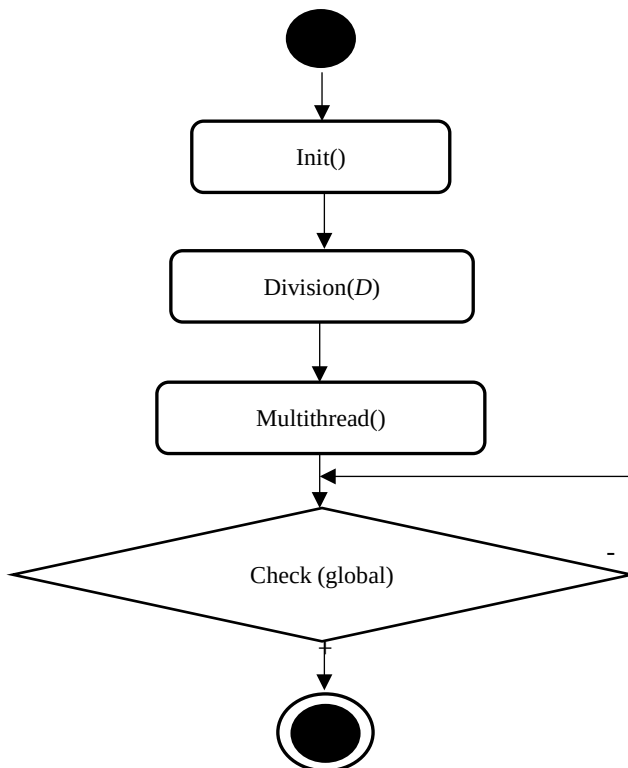


Fig. 8. Processing

4. Experiments

4.1. Experimental description. To evaluate the proposed method, we conduct experiments on the embedded Raspberry PI 3 computer – configured as in Table 2 on the embedded computer. We implement the encryption algorithms AES, DES, 3DES, etc., installed according to the mono-threaded model and multi-threaded model with the data divided according to formulas (4), (9) and (10).

Table 2. Raspberry PI 3 Embedded computer configuration

Ingredients		Values
Hardware	CPU	64 bit quad-core CPUARM Cortex A53 processor, 1.2 GHz speed 4 processors
	RAM	1 GB
	Memory card	32 GB
	Peripheral	4 USB ports HDMI port, full HDMI support Ethernet port (or LAN port)
Software	OS	Raspbian
	Application	Python Interpreter

4.2. Experimental system model. The experimental system model shown in Figure 9 includes:

1. **Camera:** Video obtained from the camera will go to the embedded Raspberry computer.
2. **Raspberry embedded computer:** The original video received from the camera is encrypted and transferred the code to storage on the IoT cloud server.
3. **IoT cloud server:** Provides API for execution on the mobile and the computer via a web browser.
4. **Smartphone:** Enforces android app with the server API to decrypt video transmitted from the server.
5. **PC:** Enforce on a web browser with API from the server to decrypt video transmitted from the server.

To evaluate the method, on the Raspberry, we install the AES algorithm, DES, Triple DES (3DES) to encode the video according to the single-threaded model and according to the multi-threaded model based on the proposed method. The Raspberry has a 4-ed processor with a symmetrical architecture, so according to the method developed in previous sections, we divide the data into 4 partitions of the same size; each partition is enforced in one thread. Experimentally, the program is divided into 4 threads corresponding to 4 characters. Each thread handles a corresponding data area.

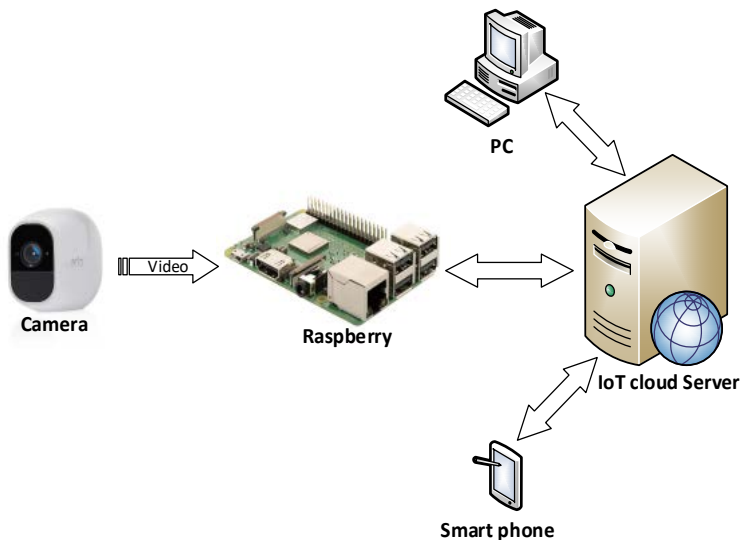


Fig. 9. Experimental model

Table 3. AES encryption algorithm executing time on Raspberry PI 3

Data size (MB)	Execution Time (ms)		Improvement rate (%)
	Single threading	Multithreading	
1	64.23	35.80	44.26
2	118.59	61.44	48.19
4	230.39	112.72	51.07
8	455.22	216.37	52.47
16	912.86	416.25	54.40
32	1851.16	843.74	54.42
64	3798.15	1629.91	57.09
Average			51.78

Table 4. DES encryption algorithm executing time on Raspberry PI 3

Data size (MB)	Execution Time (ms)		Improvement rate (%)
	Single threading	Multithreading	
1	90.53	44.19	51.19
2	177.97	76.67	56.92
4	347.27	140.28	59.60
8	699.05	271.78	61.12
16	1398.91	541.06	61.32
32	2769.94	1082.18	60.93
64	5627.24	2136.39	62.03
Average			57.59

Table 5. 3 DES encryption algorithm executing time on Raspberry PI 3

Data size (MB)	Execution Time (ms)		Improvement rate (%)
	Single threading	Multithreading	
1	201.91	72.31	64.18
2	401.41	135.38	66.27
4	817.61	264.61	67.63
8	1617.32	520.43	67.82
16	3220.44	1169.88	63.67
32	6431.65	2025.93	68.05
64	12812	4069.99	68.23
Average			66.55

4.3. Evaluation of experimental results. The experimental results are aggregated in Table 6 and are described in Figure 10. The experimental results show an average performance improvement rate of 59.09%. In particular, the experimental result with the AES algorithm is 51.78%, DES is 57.59%, and Triple DES is 66.55%. The average performance improvement rate increases as data size grows. When the data size is small, the performance improvement rate is low when using multithreads due to

the loss of CPU rotation time between threads and intermediate data processing time. However, when the processing data size is large, the rate of performance improvements is high and gradually progresses to the proportionality to the number of threads. In the proposed model, since it does not take time to process data synchronizing between threads, it only takes time to divide the original data, so it works well when the data size is large. Although it is not only in terms of performance, this model also increases the size of the program's occupied memory.

Table 6. Improvement rate comparison between AES, DES and 3DES

Data size (MB)	Improvement rate (%)		
	AES	DES	3DES
1	44.26	51.19	64.18
2	48.19	56.92	66.27
4	51.07	59.6	67.63
8	52.47	61.12	67.82
16	54.4	61.32	63.67
32	54.42	60.93	68.05
64	57.09	62.03	68.23
Average	51.78	57.59	66.55

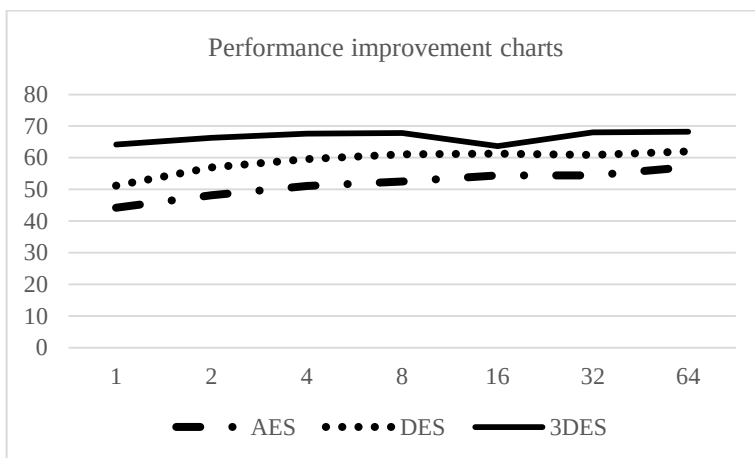


Fig. 10. Performance improvement charts

During the experiment, we also compared and evaluated the performance of data partition and asynchronous processing between the three AES, DES, and 3DES encryption algorithms on the same data set size 1, 2, 4, 8, 16, 64 and as shown in Table 7 and in Figure 11. Out of the three algorithms, the execution time of AES is the least and the greater the data size, the bigger the execution time discrepancy between the three algorithms.

Table 7. Performance comparison between AES, DES and 3DES

Data size (MB)	Execution Time (ms)		
	AES	DES	3DES
1	35.80	44.19	72.31
2	61.44	76.67	135.38
4	112.72	140.28	264.61
8	216.37	271.78	520.43
16	416.25	541.06	1169.88
32	843.74	1082.18	2025.93
64	1629.91	2136.39	4069.99

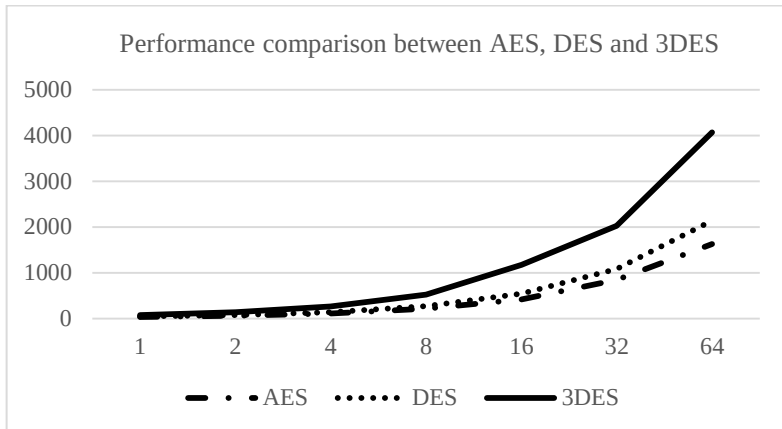


Fig. 11. Performance comparison between AES, DES and 3DES

The graph in Figure 12 compares the rate of performance improvement in the proposed method with other methods [1, 2, 3]

([2] is 25%, [3] – 21%, [1] – 33.1 %. The recommended method's performance improvement rate is also better than the average improvement rate as aggregated in studies [1, 2, 3].

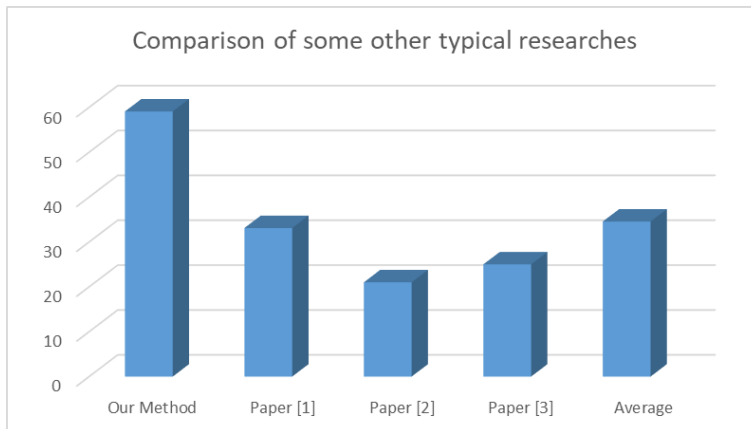


Fig. 12. Comparison of some other typical studies

5. Conclusion and future work. The paper proposes and develops a method to improve embedded software performance on multi-core processors based on data partitioning and asynchronous processing. The issue of performance improvement for embedded software on multi-core processors is of high practical significance. Especially, the problems need to be handled and ensured information security because this type of data takes time to process for data safety. Therefore, the performance improvement based on data partitioning and asynchronous processing of the paper meets the need of improving the speed of embedded software for data with independent divisibility such as block cipher.

Data are divided into sections proportional to the number of cores, the speed of the core and the cache size. Data are declared globally to be shared for all threads, so it does not take time to sync and link data. Each part of the data is mapped to the corresponding execution threads. Threads are performed asynchronously. The proposed method works well, especially when the data size is large. A high-performance improvement rate compared to previous studies has been shown.

Despite the positive results, the proposed method still has some limitations such as large memory appropriation; ineffective when the data size is small; the problem of mapping threads to the corresponding cores; the new paper only calculates the data division based on performance at the

beginning of execution, and during execution, the performance of each CPU core changes is not taken into account.

As future work, we will continue to improve the method in several aspects: expanding the problem to process data in sync; map the execution threads to the corresponding cores; and monitor and change the data division according to performance at each time when each CPU core has finished processing data.

References

1. Yao, Y. Power and Performance Optimization for Network-on-Chip based Many-Core Processors. PhD thesis. KTH. School of Electrical Engineering and Computer Science (EECS). 2019.
2. Lim, G. and Suh, S.-B. User-Level Memory Scheduler for Optimizing Application Performance in NUMA-Based Multicore Systems. IEEE 5th International Conference on Software Engineering and Service Science. 2014. 10.1109/ICSESS.2014.6933553.
3. Wei, X., Ma, L., Zhang, H. & Liu, Y. Multi-core, Multi-thread based Optimization Algorithm for Large-scale Traveling Salesman Problem. Alexandria Engineering Journal 60, 2021, pp. 189-197.
4. Khalib, Z.I.A., Ng, H.Q. Elshaikh, M., and Othman, M.N., Optimizing Speedup on Multicore Platform with OpenMP Schedule Clause and Chunk Size, IOP Conference Series. 2020. Materials Science and Engineering 767, 012037.
5. Lingampalli, S., Mirza, F., Raman, T. and Agonafer, D. Performance Optimization of Multi-core Processors using Core Hopping - Thermal and Structural. Proc. of the 28th Annual IEEE Semiconductor Thermal Measurement and Management Symposium (SEMI-THERM). 2012. pp. 112-117.
6. Gunther, N.J., Subramanyam, S., and Parvu, S. A Methodology for Optimizing Multithreaded System Scalability on Multi-cores. Programming Multi-core and Many-core Computing Systems. 2011. CoRR abs/1105.4301.
7. Rengasamy, V., Fu, T.-Y., Lee, W.-C., and Madduri, K. Optimizing Word2Vec Performance on Multicore Systems. Proceedings of the Seventh Workshop on Irregular Applications. 2017. Architectures and Algorithms. Association for Computing Machinery. New York. NY. USA.
8. Wipe, E., Miller, J.E., Choi, I., Yeung, D. Amarasinghe, S.P., and Agarwal, A. Multicore Performance Optimization Using Partner Cores. 2011. in Michael McCool & Mendel Rosenblum. 'HotPar'. USENIX Association.
9. Zhou, Y., He, F., Hou, N., and Qiu, Y. Parallel Ant Colony Optimization on Multi-core SIMD CPUs. Future Generation Computer Systems 79. 2018. pp. 473-487.
10. Emmi, M., Lal, A., and Qadeer, S. Asynchronous Programs with Prioritized Task-buffers. SIGSOFT FSE. 2012. 48.
11. Emmi, M., Ganty, P., Majumdar, R., Rosa-Velardo, F. Analysis of Asynchronous Programs with Event-Based Synchronization ESOP 2015: Programming Languages and Systems. 2015. pp. 535-559.
12. Kornaros, G. Multi-Core Embedded Systems. CRC Press. 2010. Inc., USA.
13. Bodake, V. and Gawande, R.M. A Review on An Encryption Engines For Multi Core Processor Systems. IOSR Journal of Electronics and Communication Engineering (IOSR-JECE). e-ISSN: 2278-2834. p-ISSN: 2278-8735. pp. 38-46.
14. Polychroniou, O. and Ross, K.A. A Comprehensive Study of Main-memory Partitioning and its Application to Large-scale Comparison- and Radix-sort. Print SIGMOD. 2014. pp. 755-766.

15. Schuhknecht, F.M., Khanchandani, P., and Dittrich, J. On the Surprising Difficulty of Simple Things: the case of radix partitioning. *VLDB*. 8(9): 2015. pp. 934–937.
16. Wu, L., Barker, R.J., Kim, M.A., and Ross, K.A. Navigating Big Data with High-throughput, Energy-efficient Data Partitioning. *Print SIGARCH*. volume 41, 2013. pp. 249–260.
17. Wang, Z., He, B., and Zhang, W. A Study of Data Partitioning on OpenCL-based FPGAs. In *FPL*. 2015. pp. 1–8.
18. Ke, Q., Prabhakaran, V., Xie, Y., Yu, Y., Wu, J., Yang, J. Optimizing Data Partitioning for Data-Parallel Computing. *Hot Topics in Operating Systems (HotOS XIII) | May 2011*. Published by USENIX.
19. Cieslewicz, J. and Ross, K. Data Partitioning on Chip Multiprocessors. *DaMoN '08: Proceedings of the 4th international workshop on Data management on new hardware June 2008*. pp. 25–34. DOI: 10.1145/1457150.1457156.
20. Zhong, Z. et al. Data Partitioning on Heterogeneous Multicore and Multi-gpu Systems Using Functional Performance Models of Data-parallel Applications in Cluster. 2012. pp. 191–199.
21. Kara, K., Giceva, J., and Alonso, G. FPGA-based Data Partitioning. *Proceedings of the 2017 ACM International Conference on Management of Data*. May 2017. pp. 433–445. DOI: 10.1145/3035918.3035946.
22. Zhong, Z., Rychkov, V., and Lastovetsky, V. Data Partitioning on Multicore and Multi-GPU Platforms Using Functional Performance Models. *IEEE Transactions on Computers*. Volume 64. Issue 9. Sept. 1 2015. Doi: 10.1109/TC.2014.2375202.
23. Alanazi, H.O., Zaidan, B.B., Zaidan, A.A., Jalab, H.A., Shabbir, M., and Al-Nabhani, Y. New Comparative Study Between DES, 3DES and AES. *J. of computing*. volume 2. issue 3. March 2010. ISSN 2151-9617.
24. Farooq, U. and Faisal Aslam, M. Comparative Analysis of Different AES Implementation Techniques for Efficient Resource Usage and better Performance of an FPGA. *Journal of King Saud University - Computer and Information Sciences*. Volume 29. Issue 3. July 2017. pp. 295-302.
25. Sen, K. and Viswanathan, M. Model Checking Multithreaded Programs with Asynchronous Atomic Methods. In *CAV*. 2006. pp. 300–314.
26. Kidd, N., Jagannathan, S., and Vitek, J. One Stack to Run Them all: Reducing Concurrent Analysis to Sequential Analysis under Priority Scheduling. In *SPIN '10: Proc. of the 17th International Workshop on Model Checking Software*. volume 6349 of *LNCS*, Springer. 2010. pp. 245–261.
27. Lian, X., Huang, Y., Li, Y., and Liu, J. Asynchronous Parallel Stochastic Gradient for Nonconvex Optimization. *NIPS 2015*: pp. 2737-2745.
28. Alba, E. and Troya, J.M. Analyzing Synchronous and Asynchronous Parallel Distributed Genetic Algorithms, *Future Generation Computer Systems*. Vol. 17. Issue 4. January 2001. pp. 451–465.

Bui Phuc — Ph.D., Postgraduate student, Vietnam National University. Research interests: cyber security, embedded software optimization, software technology. The number of publications — 3. phucbh.hvan@gmail.com; 144, Xuan Tu St., 11311, Hanoi, Viet Nam; office phone: +84 24 3858 4615.

Le Minh — Ph.D., Dr.Sci., Head of department, Department of information system security, Information Technology Institute – Vietnam National University. Research interests: cyber security, reliability of information system. The number of publications — 40. quangminh@vnu.edu.vn; 144, Xuan Tu St., 11311, Hanoi, Viet Nam; office phone: +84 24 3858 4615.

Hoang Binh — Software engineer, Technological Application and Production One Member Limited Liability company. Research interests: mobile and embedded software engineering, AIoT, IoT optimization, mobile apps. binhht@teca.vn; 18A, Republic St., 700901, Ho Chi Minh City, Russia; office phone: +84-28 3811 0181.

Ngoc Nguyen — Ph.D., Dr.Sci., Professor, Vice-president, Kyoto College of Graduate Studies for Informatics (KCGI). Research interests: computer science, software engineering, embedded systems & software, machine learning, data mining, information security. The number of publications — 114. nn_binh@kcg.edu; 7, Tanaka Monzencho St., 606-8225, Kyoto, Japan; office phone: +81 75-711-0161.

Pham Huong — Ph.D., Dr.Sci., Vice dean of the faculty, Faculty of information technology, Academy of Cryptography. Research interests: machine learning in information security, cloud computing, AIoT and IoT optimization. The number of publications — 30. huongpv@actvn.edu.vn; 141, Victory St., 100915, Hanoi, Viet Nam; office phone: +84 24 3854 4244.

Ф.Х. Буй, М.К. Ле, Б.Т. Хоанг, Н.Б. Нгок, Х.В. Фам
**РАЗДЕЛЕНИЕ ДАННЫХ И АСИНХРОННАЯ ОБРАБОТКА ДЛЯ
ПОВЫШЕНИЯ ПРОИЗВОДИТЕЛЬНОСТИ ВСТРОЕННОГО
ПРОГРАММНОГО ОБЕСПЕЧЕНИЯ НА МНОГОКОДНЫХ
ПРОЦЕССОРАХ**

Буй Ф.Х., Ле М.К., Хоанг Б.Т., Нгок Н.Б., Фам Х.В. Разделение данных и асинхронная обработка для повышения производительности встроенного программного обеспечения на многокодных процессорах.

Аннотация. Сегодня обеспечение информационной безопасности крайне неизбежно и актуально. Мы также наблюдаем активное развитие встраиваемых IoT-систем. В результате основное внимание уделяется исследованиям по обеспечению информационной безопасности встроенного программного обеспечения, особенно в задаче повышения скорости процесса шифрования. Однако исследованиям по оптимизации встроенного программного обеспечения на многоядерных процессорах для обеспечения информационной безопасности и повышения производительности встроенного программного обеспечения не уделялось особого внимания. В статье предлагается и развивается метод повышения производительности встроенного программного обеспечения на многоядерных процессорах на основе разделения данных и асинхронной обработки в задаче шифрования данных. Данные используются глобально для извлечения любыми потоками. Данные разбиты на разные разделы, также программа устанавливается по многопоточной модели. Каждый поток обрабатывает раздел разделенных данных. Размер каждой части данных пропорционален скорости обработки и размеру кэша ядра многоядерного процессора. Потоки работают параллельно и не нуждаются в синхронизации, но необходимо совместно использовать глобальную общую переменную для проверки состояния выполнения системы. Наше исследование встроенного программного обеспечения основано на безопасности данных, поэтому мы протестировали и оценили метод с несколькими блочными шифрованиями, такими как AES, DES и т. д. На Raspberry Pi3. В нашем результате средний показатель повышения производительности составил около 59,09%. В частности, наши экспериментальные результаты с алгоритмами шифрования показали: AES - 51,78%, DES - 57,59%, Triple DES - 66,55%.

Ключевые слова: повышение производительности встроенного программного обеспечения, многоядерные процессы, многопоточность, разделение данных, асинхронная обработка.

Буй Фук Хуу — Ph.D., аспирант, Вьетнамский национальный университет. Область научных интересов: кибербезопасность, оптимизация встроенного программного обеспечения, программные технологии. Число научных публикаций — 3. phucbh.hvan@gmail.com; Суан Туи, 144, 11311, Ханой, Вьетнам; р.т.: +84 24 3858 4615.

Ле Минь Куанг — д-р техн. наук, заведующий кафедрой, кафедра безопасности информационных систем, Институт информационных технологий Вьетнамского национального университета. Область научных интересов: кибербезопасность, надежность информационных систем. Число научных публикаций — 40. quangminh@vnu.edu.vn; Суан Туи, 144, 11311, Ханой, Вьетнам; р.т.: +84 24 3858 4615.

Хоанг Бинь Тхань — инженер-программист, Общество с ограниченной прикладной и технической продукцией - TECAPRO. Область научных интересов: мобильное и встроенное программное обеспечение, AIoT, оптимизацию IoT, мобильные приложения. binhht@teca.vn; Республика, 18А, 700901, Хошимин, Россия; р.т.: +84-28 3811 0181.

Нгок Нгуен Бинь — д-р техн. наук, профессор, вице-президент, Киотский колледж аспирантуры по информатике. Область научных интересов: компьютерные науки, разработка программного обеспечения, встроенные системы и программное обеспечение, машинное обучение, интеллектуальный анализ данных, информационная безопасность. Число научных публикаций — 114. nn_binh@kcg.edu; Танака Монзенчо, 7, 606-8225, Киото, Япония; р.т.: +81 75-711-0161.

Фам Хуонг Ван — д-р техн. наук, заместитель декана факультета, факультет информационных технологий, Академия криптографии. Область научных интересов: машинное обучение в области информационной безопасности, облачных вычислений, AIoT и оптимизации IoT. Число научных публикаций — 30. huongpv@actvn.edu.vn; Победы, 141, 100915, Ханой, Вьетнам; р.т.: +84 24 3854 4244.

Литература

1. Yao, Y. Power and Performance Optimization for Network-on-Chip based Many-Core Processors. PhD thesis. KTH. School of Electrical Engineering and Computer Science (EECS). 2019.
2. Lim, G. and Suh, S.-B. User-Level Memory Scheduler for Optimizing Application Performance in NUMA-Based Multicore Systems. IEEE 5th International Conference on Software Engineering and Service Science. 2014. 10.1109/ICSESS.2014.6933553.
3. Wei, X., Ma, L., Zhang, H. & Liu, Y. Multi-core, Multi-thread based Optimization Algorithm for Large-scale Traveling Salesman Problem. Alexandria Engineering Journal 60, 2021, pp. 189-197.
4. Khalib, Z.I.A., Ng, H.Q. Elshaikh, M., and Othman, M.N., Optimizing Speedup on Multicore Platform with OpenMP Schedule Clause and Chunk Size, IOP Conference Series. 2020. Materials Science and Engineering 767, 012037.
5. Lingampalli, S., Mirza, F., Raman, T. and Agonafer, D. Performance Optimization of Multi-core Processors using Core Hopping - Thermal and Structural. Proc. of the 28th Annual IEEE Semiconductor Thermal Measurement and Management Symposium (SEMI-THERM). 2012. pp. 112-117.
6. Gunther, N.J., Subramanyam, S., and Parvu, S. A Methodology for Optimizing Multithreaded System Scalability on Multi-cores. Programming Multi-core and Many-core Computing Systems. 2011. CoRR abs/1105.4301.
7. Rengasamy, V., Fu, T.-Y., Lee, W.-C., and Madduri, K. Optimizing Word2Vec Performance on Multicore Systems. Proceedings of the Seventh Workshop on Irregular Applications. 2017. Architectures and Algorithms. Association for Computing Machinery. New York. NY. USA.
8. Wipe, E., Miller, J.E., Choi, I., Yeung, D. Amarasinghe, S.P., and Agarwal, A. Multicore Performance Optimization Using Partner Cores. 2011. in Michael McCool & Mendel Rosenblum. 'HotPar'. USENIX Association.
9. Zhou, Y., He, F., Hou, N., and Qiu, Y. Parallel Ant Colony Optimization on Multi-core SIMD CPUs. Future Generation Computer Systems 79. 2018. pp. 473-487.
10. Emmi, M., Lal, A., and Qadeer, S. Asynchronous Programs with Prioritized Task-buffers. SIGSOFT FSE. 2012. 48.

11. Emmi, M., Ganty, P., Majumdar, R., Rosa-Velardo, F. Analysis of Asynchronous Programs with Event-Based Synchronization ESOP 2015: Programming Languages and Systems. 2015. pp. 535-559.
12. Kornaros, G. Multi-Core Embedded Systems. CRC Press. 2010. Inc., USA.
13. Bodake, V. and Gawande, R.M. A Review on An Encryption Engines For Multi Core Processor Systems. IOSR Journal of Electronics and Communication Engineering (IOSR-JECE). e-ISSN: 2278-2834. p-ISSN: 2278-8735. pp. 38-46.
14. Polychroniou, O. and Ross, K.A. A Comprehensive Study of Main-memory Partitioning and its Application to Large-scale Comparison- and Radix-sort. Print SIGMOD. 2014. pp. 755-766.
15. Schuhknecht, F.M., Khanchandani, P., and Dittrich, J. On the Surprising Difficulty of Simple Things: the case of radix partitioning. VLDB. 8(9): 2015. pp. 934-937.
16. Wu, L., Barker, R.J., Kim, M.A., and Ross, K.A. Navigating Big Data with High-throughput, Energy-efficient Data Partitioning. Print SIGARCH. volume 41. 2013. pp. 249-260.
17. Wang, Z., He, B., and Zhang, W. A Study of Data Partitioning on OpenCL-based FPGAs. In FPL. 2015. pp. 1-8.
18. Ke, Q., Prabhakaran, V., Xie, Y., Yu, Y., Wu, J., Yang, J. Optimizing Data Partitioning for Data-Parallel Computing. Hot Topics in Operating Systems (HotOS XIII) | May 2011. Published by USENIX.
19. Cieslewicz, J. and Ross, K. Data Partitioning on Chip Multiprocessors. DaMoN '08: Proceedings of the 4th international workshop on Data management on new hardware June 2008. pp. 25-34. DOI: 10.1145/1457150.1457156.
20. Zhong, Z. et al. Data Partitioning on Heterogeneous Multicore and Multi-gpu Systems Using Functional Performance Models of Data-parallel Applications in Cluster. 2012. pp. 191-199.
21. Kara, K., Giceva, J., and Alonso, G. FPGA-based Data Partitioning. Proceedings of the 2017 ACM International Conference on Management of Data. May 2017. pp. 433-445. DOI: 10.1145/3035918.3035946.
22. Zhong, Z., Rychkov, V., and Lastovetsky, V. Data Partitioning on Multicore and Multi-GPU Platforms Using Functional Performance Models. IEEE Transactions on Computers. Volume 64. Issue 9. Sept. 1 2015. Doi: 10.1109/TC.2014.2375202.
23. Alanazi, H.O., Zaidan, B.B., Zaidan, A.A., Jalab, H.A., Shabbir, M., and Al-Nabhani, Y. New Comparative Study Between DES, 3DES and AES. J. of computing. volume 2. issue 3. March 2010. ISSN 2151-9617.
24. Farooq, U. and Faisal Aslam, M. Comparative Analysis of Different AES Implementation Techniques for Efficient Resource Usage and better Performance of an FPGA. Journal of King Saud University - Computer and Information Sciences. Volume 29. Issue 3. July 2017. pp. 295-302.
25. Sen, K. and Viswanathan, M. Model Checking Multithreaded Programs with Asynchronous Atomic Methods. In CAV. 2006. pp. 300-314.
26. Kidd, N., Jagannathan, S., and Vitek, J. One Stack to Run Them all: Reducing Concurrent Analysis to Sequential Analysis under Priority Scheduling. In SPIN '10: Proc. of the 17th International Workshop on Model Checking Software. volume 6349 of LNCS, Springer. 2010. pp. 245-261.
27. Lian, X., Huang, Y., Li, Y., and Liu, J. Asynchronous Parallel Stochastic Gradient for Nonconvex Optimization. NIPS 2015: pp. 2737-2745.
28. Alba, E. and Troya, J.M. Analyzing Synchronous and Asynchronous Parallel Distributed Genetic Algorithms, Future Generation Computer Systems. Volume 17. Issue 4. January 2001. pp. 451-465.

И.А. ЛУБКИН, В.В. ЗОЛОТАРЕВ
**КОМПЛЕКСНАЯ СИСТЕМА ЗАЩИТЫ ОТ УЯЗВИМОСТЕЙ,
ОСНОВАННЫХ НА ВОЗВРАТНО-ОРИЕНТИРОВАННОМ
ПРОГРАММИРОВАНИИ**

Лубкин И.А., Золотарев В.В. Комплексная система защиты от уязвимостей, основанных на возвратно-ориентированном программировании.

Аннотация. Затруднительно или невозможно создать программное обеспечение, не содержащее ошибок. Ошибки могут приводить к тому, что переданные в программу данные вызывают нештатный порядок выполнения ее машинного кода. Разбиение на подпрограммы приводит к тому, что инструкции возврата из подпрограмм могут использоваться для проведения атаки. Существующие средства защиты в основном требуют наличия исходных текстов для предотвращения таких атак. Предлагаемая методика защиты направлена на комплексное решение проблемы. Во-первых, затрудняется получение атакующим контроля над исполнением программы, а во-вторых, снижается количество участков программ, которые могут быть использованы в ходе атаки. Для затруднения получения контроля над исполнением применяется вставка защитного кода в начало и конец подпрограмм. При вызове защищенной подпрограммы производится защита адреса возврата, а при завершении – восстановление – при условии отсутствия повреждения его атакующим. Для снижения количества пригодных для атак участков применяются синонимичные замены инструкций, содержащие опасные значения. Предложенные меры не изменяют алгоритм работы исходного приложения. Для проверки описанных решений была выполнена программная реализация и проведено ее тестирование с использованием синтетических тестов, тестов производительности и реальных программ. Тестирование показало правильность принятых решений, что обеспечивает устранение пригодных для атак участков и невозможность использования штатных инструкций возврата для проведения атак. Тестирование производительности показало 14 % падения скорости работы, что находится на уровне ближайших аналогов. Сравнение с аналогами показало, что количество реализуемых сценариев атаки для предложенного решения меньше, а применимость выше.

Ключевые слова: уязвимость, удаленное исполнение кода, защита кода, RoP, вставка кода.

1. Введение. История развития информационных технологий и их современное состояние позволяют сделать вывод о том, что разработка программного обеспечения (ПО) в общем случае не может исключать появление в нем ошибок. Отражением такого состояния дел являются работы по оценке нормального количества ошибок на определенное количество строк кода [1].

Ошибки ПО являются основой для реализации угроз. Объективно не может быть создана методика разработки ПО, не содержащего ошибок. Устранение ошибок, как разработчиками, так и специалистами по информационной безопасности также объективно не может быть выполнено полностью. Поэтому должны применяться меры, препят-

ствующие использованию ошибок в качестве уязвимостей. Для этого выполняются мероприятия по устранению условий, специфичных для определенных классов уязвимостей. В данной работе рассматривается класс *RoP*-уязвимостей. Уязвимости такого типа основаны на передаче в программу вызывающих сбой данных. Некорректная их обработка нарушает граф потока управления (далее – ГПУ) и приводит к реализации уязвимостей удаленного исполнения кода.

Предлагаемый подход к защите ставит своей целью устранение фрагментов программ – так называемых «гаджетов» (их виды и критерии описаны в [2]), используемых для реализации уязвимостей *RoP*-класса. Защита актуальна для программ, с которыми может взаимодействовать атакующий (например, сетевые сервисы, доступные через Интернет). Устранение реализуется за счет удаления гаджетов или их перевода в непригодное для атакующего состояние. При этом логика работы программы не нарушается.

Предлагаемый метод защиты не требует наличия исходных текстов программ. Это обеспечивает расширение области применимости по сравнению с теми существующими методами, которые могут быть использованы исключительно на этапе компиляции.

Структурно данная работа изложена следующим образом: во втором разделе приведён обзор наиболее близких аналогов предлагаемого подхода по выявлению и устранению уязвимостей. Из их анализа следует актуальность данной работы. В третьем разделе представлены основные сведения о рассматриваемом аппаратном и программном окружении защищаемых программ. Его необходимо учитывать для реализуемости системы защиты. Тут же описываются ограничения на применимость предлагаемого метода. Далее формулируются модель атаки и следующие из нее критерии определения гаджетов. В четвертом разделе предлагается алгоритм предотвращения использования штатных инструкций в составе гаджетов. В пятом разделе перечисляются виды структурных элементов ПО, которые могут быть интерпретированы атакующим как гаджеты. В подразделах приводятся конкретные меры защиты, и обосновывается отсутствие их влияния на целевой алгоритм ПО. В шестом разделе описывается программная реализация метода, и в седьмом – способы ее проверки на корректность и эффективность. В заключительном разделе будут сделаны выводы и сформулированы дальнейшие перспективы работы.

2. Обзор существующих подходов. Проанализируем существующие решения по защите от *RoP*-атак для формирования требований к разрабатываемой системе защиты. Они направлены на устранение условий проведения атаки. Согласно [2], таковыми являются:

- перехват управления атакующим в результате повреждения адреса возврата в стеке;
- передача управления на цепочку гаджетов, адрес размещения которых известен атакующему;
- наличие дополнительной уязвимости, которая позволяет читать фрагменты адресного пространства атакуемой программы (примитив чтения), в частности для определения адресов гаджетов.

В защищаемой системе могут присутствовать как программы с доступным эксплуатанту исходным кодом, так и без такового. Вследствие этого требование наличия исходного кода и проведения перекомпиляции должно быть включено в критерии сравнения, так как оно влияет на применимость средства защиты.

В таблице 1 приведены ближайшие обнаруженные аналоги, направленные на противодействие *RoP*-уязвимостям.

Таблица 1. Анализ средств противодействия *RoP*-уязвимостям

Наименование решения	Необходимость перекомпиляции	Защита от перехвата управления	Снижение числа гаджетов	Уязвимо при раскрытии информации	Примечание
<i>Control-flow Enforcement Technology (CET)</i> [3, 4]	Да	За счет теневого стека	Не требуется за счет контроля над ГПУ	Нет	Поддерживается с 2020 г. [5], но не всеми процессорами не всех изготовителей
<i>PaXgrsecurity RAP</i> [6] и схожий по принципу <i>StackGuard</i>	Да	За счет контроля целостности адреса возврата	Нет	Да [7]	—
ИСП Обфускатор [8, 9]	Да	За счет косвенной адресации	Нет	Да [10]	—

Продолжение Таблицы 1

Наименование решения	Необходимость перекомпиляции	Защита от перехвата управления	Снижение числа гаджетов	Уязвимо при раскрытии информации	Примечание
<i>Selfrando</i> (и схожие решения [11, 12])	Да	Нет	За счет случайного размещения при запуске	Да [13]	—
<i>Shuffler</i> [14]	Нет	За счет косвенной адресации	За счет периодического случайного перемещения	Да, до момента следующего перемещения	Уменьшение периода негативно влияет на производительность
<i>RuntimeASLR</i> [15]	Нет	Нет	Да, в дочерних процессах	Да	—
<i>G-free</i> [16] и работа [17]	Да	Нет	Да, кроме технологических участков	Нет	—
<i>Scylla</i> [18]	Да	Нет	Да, за счет шифрования участков кода и случайного перемещения	Да, в рамках расшифрованного участка	—

Из анализа приведенной подборки существующих решений противодействия *RoP*-атакам можно заключить, что первым проработанным направлением являются аппаратные системы защиты, но они требуют новейшего или специфичного оборудования и представлены не на всех платформах. Они неприменимы для эксплуатируемых систем, аппаратное обеспечение которых не подвергается модернизации (в том числе по экономическим причинам). Вторым направлением являются встраиваемые на этапе компиляции системы защиты, но они не могут быть применены при отсутствии исходных текстов ПО. Рассмотренные решения, которые применимы в описанной ситуации

(*Shuffler* и *RuntimeASLR*), уязвимы к получению атакующим информации о содержимом памяти, так как не ставят своей целью устранение гаджетов, а лишь пытаются сделать место их размещения неизвестным. Методика защиты, решающая перечисленные проблемы, не должна требовать наличия исходных текстов и должна обеспечивать отсутствие в памяти программы пригодных для использования гаджетов. Это послужило основанием для создания новой предлагаемой методики защиты. Она должна быть применима для архитектуры процессоров AMD64 (и ее аналога *Intel 64*), так как именно для них актуальны *RoP*-уязвимости, что подтверждается рассмотренными источниками.

3. Постановка задачи. Введем ограничения на область применимости предлагаемой методики. Введение ограничений осмысленно, так как абсолютные защиты нереализуемы. Ограничимся атаками типа *RoP*. Не рассматриваются по причине отсутствия или крайней редкости в рамках штатной эксплуатации программных средств:

- исполнимые участки, доступные для записи – и наоборот, у доступных для записи участков отсутствует атрибут исполнимости;
- программы, содержащие вручную написанные ассемблерные вставки (вследствие возможной их нестандартной структуры, что делает затраты на реализацию системы защиты неприемлемыми);
- ситуации отладки, когда возможна модификация контекста исполнения и данных в адресном пространстве программ внешней программой (в основном отладчиком);
- приложения с содержащимися в них программными закладками (в силу возможности внедрения такого вредоносного кода, который позволит обойти произвольную систему защиты из-за отсутствия механизмов разграничения доступа к регионам памяти в рамках одного адресного пространства).

Рассматриваются прикладные программы, написанные на компилируемых языках программирования (например, Си и Си++), так как они удовлетворяют перечисленным условиям. В общем виде они состоят из набора подпрограмм, которые в ходе выполнения вызывают другие подпрограммы. Структура подпрограмм и порядок обмена данными между ними (передача аргументов и получение результата выполнения) определяются бинарным интерфейсом приложений (далее – БИП). Типовая структура подпрограмм включает пролог, основную часть и эпилог. В прологе обычно сохраняются регистры, которые не должны быть повреждены при вызове (в типовом случае – *RBP*), устанавливается база фрейма стека для данной подпрограммы (в регистр *RBP* сохраняется значение регистра *RSP*) и выделяется память под локальные переменные (любые из перечисленных операций могут отсут-

ствовать). В основной части выполняется целевое содержимое подпрограммы (при этом баланс работы со стеком должен быть нулевым – сколько было помещено, столько должно быть извлечено). В эпилоге выполняются в противоположном порядке обратные операции, содержащиеся в прологе. В завершение вызывается инструкция возврата из подпрограммы.

При рассмотрении содержимого стека программы в произвольный момент работы он будет состоять из последовательно записанных фреймов. Пронумеруем их, начиная с нулевого, который соответствует подпрограмме, с которой началось исполнение. Введем обозначения для подпрограмм P_i и адресов возврата R_i , где i – некое неотрицательное число. Нижний индекс показывает положение фрейма стека относительно фрейма начальной подпрограммы. Например, если у текущей рассматриваемой подпрограммы фрейм i , то у вызывающей подпрограммы будет фрейм $i-1$, а у вызываемой – $i+1$. Описываемые структурные элементы приведены на рисунке 1.

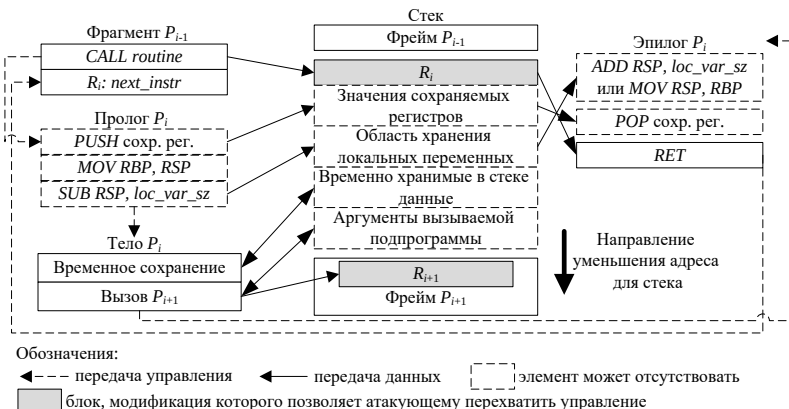


Рис. 1. Типовая структура подпрограмм

Для упрощения описания введем понятие опасной инструкции (далее – ОИ). Под ОИ будем понимать инструкции передачи управления по аргументу, который считывается из ОЗУ или регистров процессора. Примерами ОИ являются инструкции возврата *RET* или перехода по адресу, содержащегося в регистре *JMP <регистр>*. Инструкции перехода по фиксированному смещению не относятся к ОИ. Дополнительно введем понятие опасного значения (далее – ОЗ): это последовательность байтов в составе инструкции, которая не является ОИ, при

передаче управления на которые они будут интерпретированы процессором как ОИ.

На основании результатов анализа данных о структуре подпрограмм [3] и классификации гаджетов из [2] выделим критерии применимости участков программ в качестве гаджетов, которые могут использоваться злоумышленником и подлежат защите:

- последняя инструкция гаджета является ОИ или ОЗ. Она необходима для передачи управления на следующий гаджет. Адрес передачи управления должен быть предсказуем для атакующего;
- перед последней инструкцией гаджет должен содержать инструкции, модифицирующие содержимое регистров или памяти, эффект выполнения которых необходим для достижения целей атакующего;
- гаджет не должен содержать данных, которые интерпретируются процессором как некорректные инструкции или повреждают данные, необходимые атакующему.

Невыполнение любого из перечисленных критериев делает гаджет непригодным для использования в рамках атаки. Сформулируем модель атаки для определения целей атакующего. Атака осуществляется путем обмена данными с уязвимым приложением локально или удаленно. Модель включает в себя следующие действия атакующего:

- запись при выполнении тела P_i в стек массива данных такого размера, чтобы как минимум произошло повреждение R_i . В результате этого он будет заменен на необходимый атакующему адрес R_{gadget} , который передаст управление на первый гаджет в цепочке. Повреждение становится возможным при условии некорректной обработки поступающей информации атакуемым приложением. Дополнительно в стек по положительному смещению от места хранения адреса возврата записываются данные для дальнейших этапов эксплуатации уязвимости (адреса следующих гаджетов и их аргументы). Примем, что у атакующего в ходе выполнения P_i нет возможности модификации $R_{i-1} \dots R_0$ без повреждения R_i вследствие последовательной записи в стек. Из невыполнения данного условия следует наличие у атакующего примитива записи по произвольному адресу, что не позволяет противодействовать такой атаке из-за возможности произвольным образом задавать ГПУ (например, внося коррекции в таблицы виртуальных функций);
- по завершении выполнения эпилога уязвимой P_i выполняется инструкция RET , которая передает управление по адресу R_{gadget} , записанному атакующим вместо штатного адреса возврата R_i . Гаджеты должны заканчиваться инструкцией RET , которая считывает адрес

следующего гаджета и передает на него управление (саму цепочку формирует атакующий перед атакой);

- выполнение цепочки гаджетов, которое имеет конечной целью вызов функций ОС, которые необходимы атакующему для нарушения безопасности. Для этого атакующему необходимо записать аргументы функции в регистры, используемые для передачи аргументов (определяются БИП), и передать управление на её точку входа из последнего гаджета;

- делается предположение об отсутствии у атакующего одновременно и возможности контролировать значения произвольных регистров, и передавать управление по произвольному адресу. Невыполнение данного предположения делает *RoP*-атаку ненужной и позволяет, например, сделать область стека исполнимой и передать на нее управление посредством классической атаки на переполнение буфера;

- делается предположение о возможности у атакующего сформировать в стеке массив данных, содержащий адреса необходимых гаджетов, и передать управление на первый гаджет не только через повреждение R_i . Например, путем записи в единичный регистр, значение которого контролирует атакующий, адреса гаджета и последующего выполнения инструкции вида *JMP* <регистр> (*JoP*-атака). Другим примером является повреждение данных в куче (например, адреса таблицы виртуальных функций с последующим вызовом метода, для которого была выполнена подмена адреса);

- делается предположение о наличии в программе и доступности для атакующего уязвимости, реализующей примитив чтения. Подавая на вход некорректные данные, атакующий сможет получить фрагменты адресного пространства атакуемой программы. Предположим, что за раз считывается фрагмент, соизмеримый с размером фрейма подпрограммы. Примитив чтения позволяет до эксплуатации *RoP*-уязвимости сформировать дамп содержимого исполнимой памяти программы и определить адреса содержащихся в ней гаджетов;

- принимается (с учетом редкости уязвимостей ПО), что в рамках одного вызова подпрограммы не могут быть последовательно реализованы и уязвимость примитива чтения, и уязвимость повреждения R_i . Ситуация, при которой в рамках одной подпрограммы атакующий либо читает участок памяти ограниченного размера, либо в рамках нее же повреждает стек, принимается как актуальная.

Использование такой модели позволяют априорно сделать предлагаемые меры устойчивыми к большему числу сценариев атак. Безопасность решения не основывается на неизвестности атакующему сведений о программе или невозможности влиять на управление ина-

че, кроме как путем повреждения стека. Построенная на таких исходных посылках защита не станет уязвима из-за наличия у атакующего любого примитива чтения.

Можно сформулировать требования к предлагаемой системе защиты от *RoP*-уязвимостей:

- исходные коды защищаемой программы не требуются;
- преобразования защищаемого приложения не должны влиять на логику его работы (т.е. на результаты работы приложения);
- должна быть обеспечена невозможность использования штатных инструкций возврата из подпрограмм в качестве гаджетов. То есть после применения мер защиты при замене R_i на R_{gadget} в рамках выполнения тела P_i недопустим переход по адресу R_{gadget} при выполнении штатной инструкции возврата;
- должны быть устранены ОЗ. То есть инструкция, в составе данных которой присутствует ОЗ, должна быть заменена на одну или более инструкций, байтовое представление которых не содержит ОЗ.

Для реализации предлагаемой методики нужны средства резервирования в программе участков для вставки кода системы защиты без нарушения ее алгоритма и без использования исходных текстов. Пример такого средства и принцип его работы описан в [19, 20]. Предлагаемые меры защиты приведены далее.

4. Защита штатных инструкций возврата из подпрограмм.

Из рассмотрения существующих подходов к решению данной проблемы можно выделить два направления: замена ОИ с косвенным указанием R_i и контроль целостности R_i . Недостатком первого подхода является утрата бинарной совместимости кода со стандартными библиотеками ОС, невозможность использования в качестве функций обратного вызова, невозможность реализации исключений. Второй подход не имеет указанных недостатков, но уязвим к специфичным для метода атакам. Для контроля целостности адреса возврата используется некое значение (условно ключ системы защиты), на основе которого и R_i в прологе P_i вычисляется значение, которое будет использовано для контроля целостности в эпилоге. Если атакующий компрометирует ключ (общий для всех подпрограмм), то любая из подпрограмм становится уязвима. Для компрометации достаточно прочитать любой фрейм стека посредством примитива чтения. Для устранения данного недостатка предлагается метод защиты, который является развитием второго подхода к защите.

Если для защиты ОИ из состава P_i при каждом запуске формируется непредсказуемое или трудно предсказуемое (для простоты далее называется случайным) значение K_i , то для возможности подмены

R_i необходимо получение содержимого фрейма стека, содержащего это значение. Техническим ограничением является скорость генерации K_i . Сравнение возможных источников приведено в таблице 2.

Таблица 2. Сравнение возможных источников K_i

Источник K_i	Характеристика источника	Количество случайных бит	Возможность компроматации	Примечание
Инструкция <i>RDRAND</i>	Декларируется как криптографический стойкий ГПСЧ	32	Аппаратные закладки или уязвимости	Доступна в процессорах <i>Intel</i> с 2013 года, а производства <i>AMD</i> – с 2015
Инструкция <i>RDTSO</i>	Недетерминированность из-за прерываний, промахов кэша, затрат на синхронизацию, возврата значений для разных ядер и т.п.	Минимум 4 (по результатам экспериментов авторов)	Нет из-за невозможности построения физической модели процессора	Доступна на всех процессорах с архитектурой <i>AMD64</i>
ГПСЧ, предоставляемые ОС или в составе эпилогов	Криптографически стойкий	Не менее 32	Нет из-за изоляции процессов в ОС	Имеет неприемлемую скорость генерации
ГПСЧ, функционирующий на выделенном ядре	Криптографически стойкий	Не менее 32	Локальные атаки	Ресурсы одного ядра становятся недоступны
Прочие источники случайности	Отсутствуют гарантии непредсказуемости в общем случае	Не определено	Нет оценки	Например, «мусор» в стеке

Помимо качества случайного значения с точки зрения эксплуатации, важна скорость работы. По результатам сравнения были оценены как перспективные варианты *RDRAND*, *RDTSO* и ГПСЧ на отдельном ядре. Дальнейшее их сравнение проводилось по результатам апробации. Для снижения накладных расходов защитное преобразова-

ние R_i также должно занимать минимальные ресурсы. При этом криптографическая стойкость такого преобразования не требуется. Вследствие этого может быть выбрана любая обратимая операция. В рамках рассмотренных существующих решений в подавляющем большинстве случаев используется XOR . Данный вариант был также выбран для предлагаемой системы защиты.

Согласно модели атаки, повреждение R_i может быть осуществлено в ходе выполнения тела P_i . Вследствие этого защитное преобразование должно быть выполнено до тела, в прологе, а контроль целостности адреса – после тела, в эпилоге. Вследствие того, что при выполнении атаки посредством повреждения R_i штатное продолжение функционирования программы будет невозможно в любом случае, то допустимо в защищенной программе аварийно завершать ее функционирование при выявлении попытки получения злоумышленником контроля над ГПУ. Вследствие этого, если при повреждении R_i управление будет передано кодом системы защиты на непредсказуемый адрес, то это будет приемлемым решением. Таким образом, для контроля целостности в эпилоге R_i заменяется на $R_i XOR K_i$, а в прологе на основании того же ключа за счет повторной операции преобразования восстанавливается корректный адрес возврата. Если атакующий заменит хранимое в стеке значение $R_i XOR K_i$ на R_{gadget} , то в результате работы кода системы защиты управление будет передано на адрес $R_{gadget} XOR K_i$, что не позволит провести успешную RoP -атаку, так как значение K_i неизвестно атакующему.

Для корректной работы такой системы защиты сгенерированное в прологе значение K_i должно быть сохранено до эпилога для восстановления адреса. Место хранения должно выбираться так, чтобы оно не было доступно атакующему. Для этого могут быть использованы регистры общего назначения (далее – РОН) или регистры расширения. При этом программа должна быть проанализирована на наличие гаджетов и штатных инструкций, позволяющих атакующему влиять на место хранения. Подверженные влиянию регистры не должны использоваться для защиты. Отметим, что хранение в стеке неприемлемо, так как при возможности переполнения буфера атакующий может заменить хранящееся значение K_i на нулевое, а R_i на R_{gadget} . Тогда управление будет передано на нужный ему адрес.

Проблемой описанного подхода является то, что вследствие ограниченности числа регистров процессора сгенерированное в прологе P_{i-1} значение K_{i-1} должно где-то храниться на время вызова P_i (и для этого не могут использоваться те же регистры). Решением является сохранение K_{i-1} в стеке в рамках пролога P_i и возвращение в регистр-

хранилище в рамках эпилога. При этом хранение защищаемого значения в непреобразованном виде создает изъян безопасности (описан ранее), вследствие чего стоит хранить значение в виде « $K_{i-1} \text{ XOR } K_i$ » (решение проблемы первого в цепочке значения K_0 приведено далее). Для реализации данного варианта в ходе встраивания системы защиты должен проводиться анализ неиспользуемых регистров. При вызове внешних подпрограмм текущей подпрограммой ее K_i либо должно быть сохранено в стеке, либо для вызываемой внешней подпрограммы должно быть известно, что она не модифицирует хранилище K_i , либо для вызываемой библиотеки должна быть применена описываемая система защиты.

Анализ защищенности приведенного выше решения показал, что при наличии у атакующего примитива чтения в P_j может быть прочитан фрейм стека P_i , где $j > i$. Таким образом, атакующий получит значение $R_i \text{ XOR } K_i$. Исходя из предсказуемости выполнения трасс программ, атакующий может определить значение R_i и вследствие этого получить значение K_i . В том же фрейме хранится и K_i в открытом виде. Если в P_i содержится уязвимость переполнения буфера, то атакующий, заменив защищенный адрес возврата на $R_{\text{gadget}} \text{ XOR } K_i$, перехватит управление над выполнением программы.

Для защиты от реализации данного сценария необходимо выполнить в рамках пролога P_i дополнительное защитное преобразование над местом хранения защищенного R_{i-1} с использованием K_i . Также K_{i-1} в рамках пролога должен помещаться в стек в защищенном виде как $K_{i-1} \text{ XOR } K_i$. Запишем в формальном виде защитные преобразования в прологе P_i , выполняемые в i и $i-1$ фрейме стека:

$$\begin{aligned} R_i &\rightarrow R_i \text{ XOR } K_i \\ R_{i-1} \text{ XOR } K_{i-1} &\rightarrow R_{i-1} \text{ XOR } K_{i-1} \text{ XOR } K_i \\ K_{i-1} &\rightarrow K_{i-1} \text{ XOR } K_i \end{aligned}$$

При выполнении пролога P_{i+1} в фрейме i будет выполнена модификация:

$$R_i \rightarrow R_i \text{ XOR } K_i \text{ XOR } K_{i+1}.$$

В эпилогах будут выполнены обратные преобразования. В частности, обратная операция размещается непосредственно перед ОИ, что делает невозможным использование штатных инструкций подпрограмм в качестве гаджетов. Попытка повреждения памяти, хранящей R_i

в защищенном виде, приведет к переходу на непредсказуемый для атакующего адрес и аварийное завершение программы.

При отсутствии повреждения защищенных адресов возвратов аргумент инструкции *RET* будет аналогичен оригинальной программе, и управление будет передано вызывающей подпрограмме. Таким образом, предлагаемая мера защиты не искажает ГПУ.

При выполнении P_0 предшествующего фрейма не существует. Для решения данной проблемы в рамках кода подготовки программы к запуску генерируется значение K_{init} , которое сохраняется в стеке и указывается как адрес возврата вызывающей подпрограммы для P_0 . В ходе выполнения P_0 хранимое значение K_{init} будет заменено на $K_{init} \text{ XOR } K_0$. Таким образом, значение K_{init} не хранится в незащищенном виде и может быть восстановлено атакующим только при знании K_0 . Допустимо принять, что ГПУ стабилен и определен для атакующего, поэтому R_i известно атакующему (так как вызывающая подпрограмма определяется ГПУ), а, следовательно, он может вычислить « $K_i \text{ XOR } K_{i-1}$ » для любого i .

Отметим, что переходы по адресам, содержащимся в регистрах, также не могут быть устранены или заменены. Используемые не в качестве начального элемента *JoP*-атаки, они не являются самодостаточными без предшествующего помещения атакующим значений в регистры и передачи управления на соответствующие инструкции в рамках *RoP*-атаки. Следовательно, повышение устойчивости программы к *RoP*-атакам автоматически повышает устойчивость к *JoP*-атакам. При этом первичное получение контроля на ГПУ в рамках *JoP*-атаки возможно согласно модели атаки.

Рассмотрим возможные атаки, направленные на преодоление предложенных мер по защите возвратных инструкций.

1. Повреждение стека с перезаписью адреса. Для успешной реализации требует чтения значения $R_i \text{ XOR } K_i$ и переполнения стека в рамках одной подпрограммы с записью в стек значения $R_{gadget} \text{ XOR } K_i$, где R_{gadget} – адрес первого гаджета в *RoP*-цепочке. При этом между утечкой значения защищенного адреса и применением эксплойта должно выполняться формирование полезной нагрузки эксплойта (то есть массива данных, который будет некорректно обработан и приведет к повреждению адреса возврата). Согласно модели атак, данный вариант принят как неактуальный.

2. Угадывание значения K_i . При использовании *RDRAND* или *RDTS* при неизвестном предшествующем значении вероятность 2^{-32} (вследствие использования 32-битного случайного значения). При из-

вестном значении $RDTSC$ для вызывающей подпрограммы можно пессимистично оценить как 2^{-4} .

3. Повреждение стека вызывающей подпрограммы с заменой $R_i \text{ XOR } K_i \text{ XOR } K_{i+1}$ на $R_{gadget} \text{ XOR } K_i \text{ XOR } K_{i+1}$. Данный вид атаки эффективен, но требует, согласно модели атаки, знания K_i для корректного возврата из текущей подпрограммы для последующего перехвата управления в вызывающей подпрограмме.

4. Атака с чтением начальной области стека. Атакующий использует примитив чтения в n -й подпрограмме для получения значений $K_{init} \text{ XOR } K_0$ и пар значений: $\langle K_i \text{ XOR } K_{i-1}, K_i \text{ XOR } K_{i+1} \rangle$, где $0 < i < m$, $m < n$. Полученная последовательность не позволяет сформировать эксплойт для i -й подпрограммы, так как атакующему требуется значение K_i , которое невычислимо из полученных пар при использовании качественного источника случайных чисел.

5. Атака с чтением фрагментов полного стека. Использование примитива чтения в n -й подпрограмме для получения значений $K_{init} \text{ XOR } K_0$ и пар значений: $\langle K_i \text{ XOR } K_{i-1}, K_i \text{ XOR } K_{i+1} \rangle$, где $0 < i < n$, и значения K_n атака будет возможна, если все перечисленные значения могут быть получены без повторного вызова подпрограмм с i -й по n -ю, так как в противном случае будет выполнено пересоздание ключей на указанном участке. При пересоздании ключей в распоряжении атакующего будет фактическая последовательность: $K_{init} \text{ XOR } K_0$ и пары значений $\langle K_i \text{ XOR } K_{i-1}, K_i \text{ XOR } K_{i+1} \rangle$ для $0 < i < k$, пары значений $\langle K'_i \text{ XOR } K'_{j-1}, K'_j \text{ XOR } K'_{j+1} \rangle$ для $k < j < n$. Индекс k соответствует границе прочитанных фрагментов стека. Чтение стека производится атакующими фрагментами согласно модели атаки. Полученная последовательность не позволяет сформировать эксплойт для i -й подпрограммы, так как требует значения K_i , для вычисления которого необходимы значения $K_{i+1}, \dots, K_n$.

6. При наличии уязвимой к переполнению буфера подпрограммы и примитива чтения в подпрограммах не в единой последовательности вызываемых подпрограмм атака перестает быть возможна (то есть нахождение в разных ветках дерева вызовов ГПУ).

5. Меры по устранению опасных значений. Проанализируем возможные места размещения ОЗ и эффективные меры по их устранению. В соответствии с разделом 3, такими местами являются:

1. Модификация последовательностей инструкций перед RET должна осуществляться таким образом, чтобы исключить чтение со сдвигом. Такое чтение в ходе атаки позволяет интерпретировать фрагменты существующих инструкций как необходимые атакующему другие инструкции. Защита обеспечивается за счет формирования перед

RET такой последовательности инструкций, которая не может быть интерпретирована атакующим как содержащая полезные для атаки операции.

2. Фрагменты инструкций, не входящие в множество ОИ, подлежат синонимической замене с сохранением алгоритма оригинальной программы.

3. Области размещения данных, по технологическим причинам являющиеся исполнимыми, должны быть лишены признака исполнимости.

Далее приводятся детали реализации для каждой из предложенных мер.

5.1. Противодействие исполнению со сдвигом. Первой задачей противодействия является оценка возможности использования атакующим кода системы защиты, находящегося в эпилоге, при исполнении его со сдвигом.

Второй задачей является поиск в коде оригинальной программы гаджетов, не входящих в состав эпилогов и получаемых за счет исполнения со сдвигом. Включающие гаджеты инструкции подлежат замене на синонимичные, которые исключают их использование атакующим.

Для решения обеих задач используется алгоритм *GALILEO*, описанный в [21] и позволяющий определить пригодность использования ОЗ и расположенных перед ним инструкций в качестве гаджета. Данный алгоритм обеспечивает исключение из рассмотрения таких ОЗ, перед которыми содержатся последовательности байт, которые не могут быть интерпретированы как инструкции или приводят к аварийному завершению программы. Код системы защиты подлежит анализу для определения того, увеличивает ли он число гаджетов относительно оригинального приложения.

5.2. Синонимическая замена опасных значений в составе инструкций. Для защиты от неверной интерпретации код программы должен быть проанализирован в соответствии с разделом 5.1 для определения инструкций, подлежащих защите. ОЗ включают инструкции возврата в пределах сегмента (0xC3, 0xC2) и инструкции возврата за пределы сегмента (двухбайтовые последовательности «0x48 0xCA» и «0x48 0xCB»). В результате анализа оригинальной программы формируется перечень инструкций, подлежащих защите путем синонимической замены.

Для каждой из них необходимо определить, в каком качестве в составе инструкции содержится ОЗ, для того чтобы сформировать эквивалентные замены. Из проверки исключаются ОИ из состава эпилогов, так как они защищаются вставкой кода применения к ним K_i .

Для контроля корректности работы алгоритма поиска могут использоваться существующие средства поиска гаджетов (например, [22]). Критерием корректности является то, что состав найденных по предлагаемой методике ОЗ должен быть не меньше, чем у существующих средств.

Под синонимической заменой инструкции, содержащей ОЗ, будем понимать такую замену, чтобы состояние используемых регистров и ячеек ОЗУ после исполнения оригинального и модифицированного участка не отличалось для любого исходного состояния используемых регистров и ячеек ОЗУ на начало рассматриваемого участка. Критерий использования сформулирован далее.

Рассмотрим виды составных частей инструкций, которые могут содержать ОЗ. Учет их специфики необходим для формирования порядка генерации синонимичных замен. Непосредственно байт, который интерпретируется как ОЗ, не может входить в состав префиксов (так как префикс не может совпадать с существующими операциями) и кодов операций первого уровня. Согласно [4, 19], ОЗ могут встречаться в следующих участках инструкций:

- аргумент данных операции, являющийся адресом (относительным или абсолютным);
- параметры операции, задающие режим ее выполнения, или операнды (*ModRM*, *SIB* компоненты инструкции);
- аргумент данных операции, не являющийся адресом;
- второй и третий байт многобайтного префикса (значения *0xC2* и *0xC3* являются допустимыми значениями, а *0xCA* и *0xCB* неприменимы, так как необходимый им префикс *0x48* не входит в область допустимых второго байта трехбайтного префикса и не может являться первым байтом ни двухбайтового, ни трехбайтового префикса);
- код операции второго или большего уровня, интерпретируемый как опасная операция.

Вариант с выявлением в полях, содержащих адрес необходимого атакующему значения, решается следующими способами:

1. При вхождении ОЗ в состав относительного адреса до (если относительный адрес является отрицательным числом) или после (если он положительный) в содержащую ОЗ подпрограмму выполняется вставка последовательности байт «0x90» такого размера, чтобы относительный адрес после коррекции с учётом вставки не содержал ОЗ;

2. При нахождении ОЗ в младшем байте абсолютного адреса выполняется перемещение целевого участка в конец содержащей его

секции (исходное место не используется). Новое размещение выбирается так, чтобы его адрес не содержал ОЗ;

3. При нахождении ОЗ не в младшем байте абсолютного адреса мер по его защите не требуется, так как вследствие работы механизма *ASLR* значение при начале работы программы будет изменено на случайное.

Процесс устранения опасных относительных адресов проводится итерационно, с контролем сформированных заменяющих адресов на возможное порождение новых уязвимых последовательностей. Возможна вставка региона, размер которого является диапазоном. Это может привести к порождению опасного участка. Тогда размер вставки изменяется и проверяется следующий размер. Сходимость процесса обеспечивается тем, что количество подлежащих перемещению участков монотонно убывает из-за того, что перед перемещением проверяется порождение опасных участков.

При нахождении опасного значения в префиксе или расширенном коде операции для его использования атакующим необходимо передать управление до начала опасной инструкции. Поэтому для устранения данного гаджета достаточно перед инструкцией вставить цепочку однобайтовых инструкций, которые делают некорректное исполнение невозможным (например, *NOP*). Для определения размера необходимой вставки используется подход, описанный в подразделе 5.1 для оценки возможности использования байта в инструкции, описанной ранее. Перебираются размеры от 1 до 14.

Для ситуаций, когда опасный байт входит в состав аргумента данных инструкции, или указания её операндов опасная инструкция должна быть заменена на две и более инструкции, совокупное выполнение которых дает эквивалентное преобразование состояний программы тому, которое осуществляла защищаемая инструкция, и не содержит опасных последовательностей. В ходе такой замены необходимо контролировать, чтобы выполнение результатов синонимической замены не оказывало влияние на логику целевой программы. Это требует использования только таких регистров или ячеек ОЗУ, которые не используются целевой программой. Если таких нет, то требуется сохранение с последующим восстановлением ресурсов, для которых присутствует конфликт использования. Критерием отсутствия использования является то, что значение регистра в дальнейшем в рамках анализируемой подпрограммы будет утрачено вследствие затирания другим значением без предшествующего этому чтению. Вследствие этого оно никак не может повлиять на работу программы. Место под

замещающую последовательность выделяется тем же алгоритмом, что используется для резервирования места в ходе защиты эпилогов.

В рамках опасного значения в составе операндов эквивалентная замена достигается путем замены операндов инструкции. Опасное значение для байта *ModRM* состоит из: режима, равного 3 (операндами являются регистры), поля, определяющего регистровый аргумент «*reg*», равного 2 или 3 (что соответствует регистрам *RDX/R10* и *RBX/R11*), и поля «*r/m*», равного 0 или 1 (что соответствует регистрам *RAX/R8* и *RCX/R9*). Для байта «*SIB*» аналогичные опасные значения должны содержаться в полях «*base*» и «*index*» с теми же значениями.

Для того чтобы сделать непригодным к использованию данный гаджет, достаточно заменить любой из операндов на регистр, который не входит в состав опасных. Если инструкция принимает только входной регистр или один из регистров является входным, а другой выходным, то эквивалентной заменой будет вставка перед защищаемой инструкцией команды *MOV*, которая копирует данные из опасного регистра-операнда в свободный регистр, который не входит в состав опасных. В самой инструкции опасный регистр заменяется на свободный, который к моменту выполнения инструкции будет содержать то же значение, что и в оригинальной программе.

Если единственным операндом инструкции является выходной регистр или регистр, который одновременно является и входным, и выходным, то выполняется вставка двух дополнительных инструкций: перед – запись в свободный регистр значения из опасного, после – запись в опасный регистр значения из свободного. Отметим, что инструкции перемещения значений между регистрами не влияют на регистр флагов, что не нарушает логики работы программы. В самой инструкции производится аналогичная замена опасного регистра на один из свободных и не входящих в перечень опасных. Алгоритм определения свободных регистров приведен далее.

При устранении ОЗ в составе непосредственного операнда данных принципиальным является то, имеет ли использующая опасное значение в качестве операнда инструкция альтернативную форму, в которой вместо непосредственного значения используется операнд в регистре. Необходимо отметить, что подавляющее большинство инструкций, которые могут использовать непосредственный операнд, имеют альтернативную форму с регистровым операндом.

Если альтернативная форма с регистровым операндом существует (например, инструкция *ADD*), то оригинальная инструкция заменяется на последовательность:

- если свободных регистров нет, то выполняется временное освобождение регистра путем записи его значения в стек (в конце последовательности значение должно быть восстановлено из стека);

- загрузка в свободный регистр преобразованного операнда. Для устранения опасного значения операнд хранится в преобразованном виде (например, инвертированном или циклически сдвинутом) так, чтобы преобразованный вид не содержал опасных значений;

- восстановление значения операнда путем выполнения обратной операции – той, которая использовалась для его преобразования;

- выполнение регистровой формы оригинальной инструкции с оригинальным операндом.

Если не существует альтернативной формы оригинальной инструкции (например, инструкция *ENTER*), то производится дополнение со стороны меньших адресов защитными инструкциями, которые исключают неправильную трактовку защищаемой инструкции. Дополнительно необходимо учитывать особенности:

- не для всех инструкций опасное значение входит в область допустимых значений. Например, приведенная инструкция *ENTER* принимает в качестве аргументов два непосредственных операнда. Опасное значение не может содержаться в младшем байте первого операнда, так как это нарушает выравнивание. Кроме того, опасное значение не может содержаться во втором операнде, так как его максимальное значение составляет 32. Маловероятно, что опасное значение будет содержаться в старшем байте первого операнда, так как фреймы стека размером в 32 кБайта представляют редкость. Вероятность появления снижает то, что опасные значения составляют 2 из 255 возможных значений старшего байта первого операнда, и то, что компиляторы в принципе не используют такие инструкции (что не сужает область применимости относительно указанной в третьем разделе);

- вторая и третья найденные инструкции, не имеющие альтернативной регистровой формы (*LWPINS* и *LWPVAL*), применимы только на этапе профилирования приложения и в генерируемом коде компилятора не встречаются (что не сужает область применимости относительно указанной в третьем разделе);

- иных инструкций, помимо перечисленных только с формой, использующей операнд, согласно документации [4, 19], не существует.

Поиск существующих решений, позволяющих определить используемые программой ресурсы процессора, показал, что единствен-

ной публично доступной реализацией является [21], принцип описан в [22]. Анализ данного решения показал, что оно не учитывает БИП, вследствие чего каждый вызов подпрограммы приводит к пометке всех регистров как используемых, что ведет к неверной оценке состава используемых регистров. Было принято решение разработать новый алгоритм определения свободных ресурсов, учитывающий БИП. Перед формулированием алгоритма рассмотрим теоретические сведения, необходимые для этого.

Для определения свободных ресурсов процессора необходим анализ полной последовательности инструкций подпрограммы, в которой находится подлежащая защите инструкция. Это связано с тем, что ряд регистров используется достаточно редко и при анализе неполного содержимого подпрограммы может быть сделан неверный вывод об отсутствии их использования. Анализуются входные и выходные аргументы инструкций для определения регистров, которые могут быть задействованы в модифицированном участке и не повлияют на эквивалентность вычислений. Регистр может считаться свободным, если содержащееся в нем значение не может быть использовано в дальнейших вычислениях. Так как описанные выше меры оперируют регистрами целиком, то для того, чтобы регистр считался свободным, должны быть свободны все разряды регистра.

Пример: значение, содержащееся в регистре после считывания, в дальнейшем не считывается повторно, а вместо этого перезаписывается другим значением. Приведенный пример показывает, что необходимо рассматривать свободные ресурсы строго с точки зрения шагов выполнения программы (то есть свобода ресурса может быть выявлена только для диапазона конкретных инструкций). Учитывая БИП, максимальным участком анализа разумно выбрать подпрограмму. Подпрограмма представляется в виде набора линейно исполняющихся участков кода – базовых блоков (далее – ББ), переходы между которыми осуществляются согласно ГПУ.

Определение свободы ресурсов в рамках ББ (то есть когда последовательность выполнения детерминирована) является тривиальной задачей. При необходимости определения в границах больших, чем ББ (например, при необходимости вставки перед первой инструкцией ББ), необходимо учитывать все возможные переходы на все рассматриваемые ББ (если достоверно не определено, что возможен только один вариант перехода).

Если в ходе анализа используемых ресурсов встречаются инструкции, создающие сторонние эффекты, то возможны следующие варианты действий:

– при анализе инструкций вызова подпрограмм из состава защищаемого приложения должен быть учтен БИП как максимально ограничивающий состав свободных ресурсов. БИП определяет, из каких регистров считываются данные, в какие регистры записывается результат исполнения, какие сохраняют свое значение и какие будут иметь неопределенное значение в результате вызова подпрограмм. При недостатке ресурсов может быть проведен дополнительный анализ использования регистров в начале и в конце вызываемой подпрограммы. Если соответствующие ресурсы находятся в свободном состоянии, то должен быть сделан вывод о том, что они свободны, до и после вызова анализируемой вызываемой подпрограммы;

– если для системных вызовов и вызовов библиотечных функций не определен БИП, должно быть принято утверждение, что модифицируются все регистровые ресурсы. Если значения каких-либо ресурсов должны быть сохранены в интересах системы защиты, то до исполнения они должны быть сохранены (не обязательно непосредственно перед вызовом) и восстановлены на следующей инструкции, после создающей сторонние эффекты. Вследствие того, что в работе учитывается БИП, это дает априорную информацию об используемых ресурсах. Это позволяет сохранять только необходимые ресурсы, что снижает накладные расходы. Для ряда ситуаций сохранение значений непосредственно перед исполнением является недопустимым, так как нарушает эквивалентность состояний. Например, рассмотрим сохранение единичного восьмибайтового значения непосредственно перед инструкцией *CALL*, которая передает управление на подпрограмму со стандартным прологом, которая принимает аргументы через стек. Тогда адрес *RBP-18₁₆* будет указывать вместо первого аргумента на сохраненное значение, что нарушит эквивалентность программ. Такая проблема ассоциирована со стеком, так как код программ содержит массово работу по относительным адресам при передаче аргументов и хранении локальных переменных. Безопасны для помещения значений в стек только участки подпрограмм двух видов: а) после выделения всех локальных переменных и до выделения первого блока аргументов для вызова вложенной подпрограммы, а также б) после восстановления стека, следующего за вызовом подпрограммы и помещением в стек аргументов для следующей подпрограммы.

Если модифицированный участок требует специфического набора ресурсов, и они на участке вставки не являются свободными, то они должны быть сформированы путем сохранения текущих целевых значений в прологе модифицированного участка и восстановления в эпилоге. Примером является инструкция *RDTSC*, которая записывает

результат своей работы в фиксированные регистры, что приводит к безвозвратной потере содержащихся в них целевых значений. Теоретически, возможным вариантом являлось бы сохранение значений всех регистров (например, последовательностью операций *PUSH* или инструкцией *PUSHAD* для 32-разрядного режима), но это существенно повышает накладные расходы. А с учетом массовой вставки кода для защиты инструкция возврата замедляет программу гарантированно неприемлемым образом.

Необходимо отметить, что в анализе ресурса памяти в предлагаемой системе защиты нет необходимости, вследствие чего данный вопрос не рассматривался.

Непосредственное определение свободных ресурсов, во-первых, сводится к определению регистров, запись в которые не приведет к нарушению эквивалентности состояний программы, а во-вторых, записанные в интересах системы защиты значения не будут повреждены в ходе вычислений целевой программы. На основе данных утверждений можно сформулировать критерий отнесения ресурсов к свободным:

- регистр считается свободным до инструкции, которая произведет запись в него;
- если после записи значение из регистра не считывается до выполнения повторной записи, то регистр должен рассматриваться как свободный с момента первой записи. При этом сама первая операция должна считаться избыточной в части модификации регистра;
- после помещения в регистр неопределенного значения (например, после возврата из подпрограммы в несохраняемых регистрах, согласно БИП) он считается свободным до записи в него значения.

Модификация свободных регистров не влияет на выполнение программы. Вследствие этого их использование в рамках кода системы защиты гарантированно не модифицирует алгоритм оригинальной программы.

Состояние одного регистра не влияет на состояние другого. То есть используемые регистры определяются машинным кодом и не связаны с результатами вычислений, так как архитектура команд не предусматривает прямое указание регистров-операндов. Вследствие этого поиск свободных ресурсов допустимо проводить для каждого регистра независимо. Дальнейшее описание приведено для одиночного регистра, для получения полной информации процедура повторяется для всех необходимых регистров.

Пусть анализируемая подпрограмма состоит из набора инструкций f_j , $1 \leq j \leq n$, где n – количество инструкций в подпрограмме. Перед

определением состояния ресурсов выполним подготовительные операции. Согласно документации на процессор и БИП, определим для каждой f_j , является ли анализируемый регистр входным аргументом (читается), и запишем результат в массив RI размером n , индексы элементов которого соответствуют индексам инструкций. Выполним аналогичное определение с точки зрения того, является ли регистр выходным, и запишем результат в массив RO . Элементы массивов могут принимать значения: используется (Y), не используется (N), не определен (U). Результат определения записывается в массив RF , имеющий аналогичный размер, элементы которого принимают значения: свободен (F), занят (U).

БИП определяет порядок передачи аргументов, но не все используемые для этого регистры могут быть задействованы. Для их определения может быть либо проведен анализ вызываемой подпрограммы, либо проанализирована семантика инструкций подпрограммы. При втором подходе все инструкции между записывающей в регистр неопределенное значение и потенциально считывающей его в качестве аргумента вызываемой подпрограммы должны быть помечены как не использующие регистр. Для этого вводится массив RR размером n , значения которого принимают значения: инструкция предшествует чтению (Y) или нет (N).

Для проведения анализа необходимы сведения о ГПУ подпрограммы. Множество инструкций разбивается на ББ $b_k = (j_{bk}, j_{ek})$, где j_b и j_e – индексы первой и последней инструкции ББ, $1 \leq k \leq n_b$, n_b – число ББ. Для ББ определяются связи между ними (для каждого определяется перечень, ссылающихся на него ББ). Также определяется перечень эпилогов, которыми заканчивается выполнение анализируемой подпрограммы, $E = \{e_1, \dots, e_m\}$, где e – номера ББ, включающие эпилоги, а m – число эпилогов в подпрограмме. Для учета порядка обработанных блоков используется массив PBV размером n_b , элементы которого принимают значения: обработан (P), запланирован к обработке (S), не обработан (U , значение по умолчанию).

Разработанный алгоритм определения свободных ресурсов с учетом БИП приведен на рисунке 2.

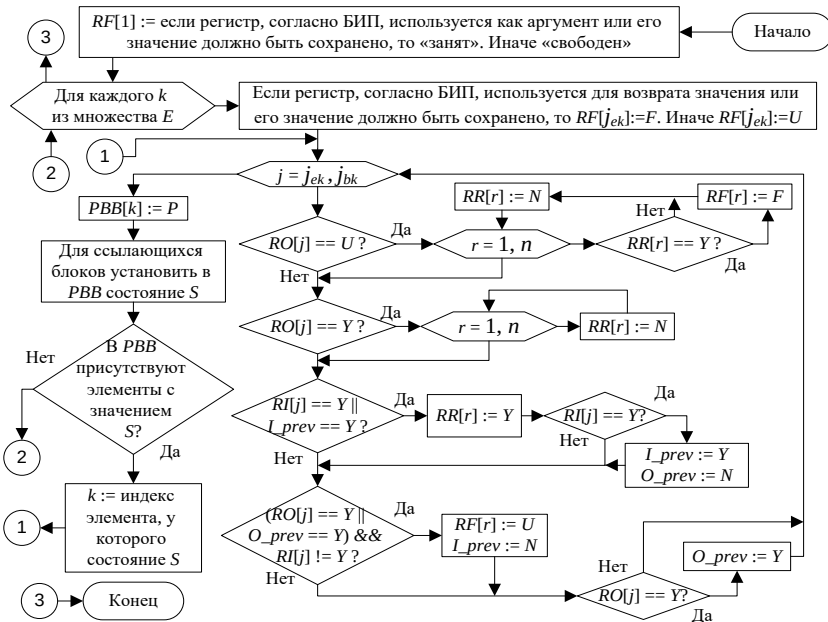


Рис. 2. Алгоритм определения свободных ресурсов

В результате выполнения для каждого регистра указанного алгоритма каждой инструкции анализируемой подпрограммы будет составлен перечень свободных ресурсов. Для опасной инструкции он используется для формирования эквивалентной замены.

При описанном выше порядке вызываемые подпрограммы рассматриваются как «черный ящик»; выдвигаются предположения о том, что для передачи аргументов и возврата значений используется максимально возможный состав регистров в рамках БИП. Состав входных и выходных регистров для вызова подпрограммы может быть уточнен при необходимости (например, при недостатке свободных ресурсов) путем анализа содержимого вызываемой подпрограммы. Порядок такого анализа совпадает с описанным ранее, а сведения о считываемых и модифицируемых ресурсах вызываемой подпрограммы интерпретируются как эффект от выполнения инструкции вызова подпрограммы.

Отсутствие же использования ресурса в рамках всех инструкций программы свидетельствует о том, что любые операции с ним не влияют на эквивалентность программ. Если имеется априорная информация об отсутствии использования ресурса в рамках кода программы и не входящих в состав программы библиотек (например,

вследствие ограничений на состав используемых компилятором инструкций) или гарантий отсутствия модификации, используемых системой защиты ресурсов системными библиотеками и системными вызовами, то она должны быть учтена для снижения времени анализа. Априорно не используемые ресурсы помечаются свободными.

5.3. Исключение данных из состава исполнимой памяти. Для реализации этой меры защиты неисполняемые данные должны быть исключены из перечня исполнимых областей. Тогда содержащиеся в них байты с опасными значениями не могут быть использованы атакующими как гаджеты. Для определения атрибута исполнимости применяется технология «*NX-bit*». Она позволяет разрешать исполнение инструкций только для конкретных участков программ (как минимум участка, содержащего машинный код). Перечень исполнимых участков содержится в заголовке образа программы.

Единицей указания атрибута исполнимости является страница памяти (в типовом случае размером 4096 байт). Для того чтобы в исполнимую память не были включены данные, применяется выравнивание. Перед участком кода и после него помещается массив байтов со значениями, которые не могут быть применены в составе гаджетов. Размер массива выбирается так, чтобы адреса начала и конца исполнимого участка памяти были кратны размеру страницы. В качестве значений байтов используется «0xCC», так как передача на них управления вызовет отладочное прерывание. Исходная секция, включающая данные до исполнимого кода, сам исполнимый код и данные после него разбиваются на три соответствующие секции, причем исполнимой (бит *NX* сброшен) является только вторая.

6. Программная реализация. Предлагаемые решения были апробированы на ОС *Debian 9*, относящейся к семейству *GNU/Linux*. Данная ОС использует БИП «*ABI Linux-AMD64*» [24]. Для реализации программной реализации использовалась среда *QtCreator*, компилятор *gcc 6.3.0*.

Общая архитектура программного средства (далее – ПС) приведена на рисунке 3. Заимствуемые в рамках данной работы компоненты обеспечивают анализ структуры и связей программы, а также восстановление корректной ссылочной структуры с учетом добавляемого кода системы защиты. Анализатор подпрограмм обеспечивает определение границ подпрограмм и их параметров, таких как: границы прологов и эпилогов подпрограмм, выделяемое в стеке место под сохраняемые в ходе работы подпрограммы регистров и локальные переменные. Алгоритм определения ресурсов приведен ранее. Логика функционирования средства поиска опасных участков описана ранее.

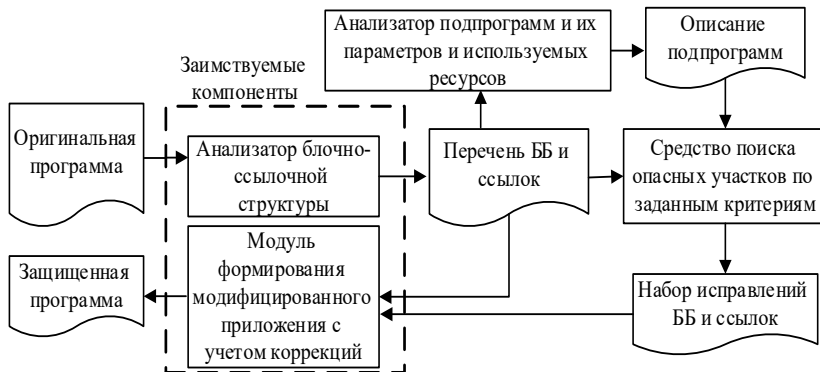


Рис. 3. Архитектура ПС

Для обеспечения автоматического тестирования ПС для управления был выбран интерфейс командной строки, а для вывода данных – текстовое представление. Посредством интерфейса командной строки задаются основные параметры, такие как имя оригинальной программы, имя выходного файла для записи защищенной программы, вид используемого источника K_i , детальность выводимой диагностической информации. Для отражения успешности применения системы защиты используется код завершения процесса (0 в случае успеха, иначе значение, указывающее на этап, вызвавший сбой). В консоль выводится диагностическая информация с настраиваемым уровнем детализации. Например, выводится информация о сегментах и секциях оригинальной программы, ББ, ссылках, найденных ОЗ, подпрограммах, а также об операциях по формированию защищенной программы. При сбое работы выводится диагностическая информация для устранения проблемы.

Для реализации защиты эпилогов используются вставки, представленные в таблице 3. Приведен пример при использовании инструкции *RDRAND*. Защите подлежат только подпрограммы, хотя бы один эпилог которых заканчивается инструкцией *RET* (не все подпрограммы заканчиваются такими инструкциями). Код приведенных вставок не содержит ОЗ.

Таблица 3. Вставки кода для защиты штатных инструкций возврата

Место вставки	Мнемоническое представление	Байтовое представление, шестнадц.	Примечание
В месте завершения пролога	rdrand r10 movd r11 , xmm14 xor [r11], r10 push r11 lea r11 , [rsp+X] movd xmm14, r11 xor [r11], r10 movd r11 , xmm15 xor r11 , r10 push r11 movd xmm15, r10	49 0F C7 F2 66 4D 0F 7E F3 4D 31 13 41 53 4C 8D 9C 24 XX XX XX XX 66 4D 0F 6E F3 4D 31 13 66 4D 0F 7E FB 4D 31 D3 41 53 66 4D 0F 6E FA	X – смещение до адреса возврата; свободные РОН – R10 и R11; для хранения текущего ключа и адреса возврата выбраны не используемые в программе регистры расширения XMM15 и XMM14
Перед эпилогом	pop r9 pop r8 movd xmm14, r8	41 59 41 58 66 4D 0F 6E F0	Регистры R8, R9 свободны перед началом эпилога, согласно БИП
Перед инструкцией <i>RET</i>	movd r11, xmm15 xor [r8], r11 xor r9 , r11 movd xmm15, r9 xor [rsp], r11	66 4D 0F 7E FB 4D 31 18 4D 31 D9 66 4D 0F 6E F9 4C 31 1C 24	Регистр R11 также свободен к завершению подпрограммы
В начале подпрограммы <i>_start</i>	rdrand eax push rax movd xmm14, rsp movd xmm15, rax	0F C7 F0 50 66 4C 0F 6E F4 66 4C 0F 6E F8	Подпрограмма <i>_start</i> не требует защиты, так как заканчивается инструкцией <i>HLT</i>

Анализ вставляемого перед инструкцией *RET*-кода показывает следующие возможности интерпретации при передаче управления за указанное количество байтов до самой инструкции *RET*:

– 1 байт – последовательность «AND AL, C3» (шестн. 24 C3), которая не может быть использована как гаджет (то есть байт C3 интерпретируется как часть инструкции «AND» и не может быть исполнен);

– 2 байта – «*SBB AL, 24₁₆; RET*» (шестн. 1C 24 C3) – потенциально может быть использована как гаджет, выполняет вычитание из младшего байта регистра *RAX* константы 24₁₆;

– 3 байта – «*XOR [RSP], EBX, RET*» – такая инструкция не может быть использована как гаджет, так как не модифицирует регистров, необходимых для вычислений или управления программой, но при известном атакующему значении в регистре *EBX* может быть использована для передачи управления на другой гаджет (путем записи на вершину стека значения $R_{\text{gadget}} \text{ XOR } EBX$ при предшествующем повреждении стека);

– более 3 байт – перед *RET* неизменно будет инструкция «*XOR [RSP], R11*», которая преобразует значение адреса возврата на неизвестное атакующему значение (регистр *RSP* указывает на адрес возврата, так как защитный код вставлен перед инструкцией *RET*).

При применении описанного комплекса мер даже если атакующий получит возможность влиять на ГПУ (например, посредством атак на таблицы виртуальных функций), то в его распоряжении будет всего один гаджет, влияющий на РОН, не используемый для передачи аргументов в подпрограммы, и возможность передавать управление на другие такие же гаджеты, чего недостаточно для эксплуатации уязвимости, согласно модели атаки.

7. Экспериментальная проверка корректности и эффективности программного средства. Предлагаемый метод защиты влияет на стек приложения, а для *RoP*-атак даже минимальное изменение расположения данных в стеке приводит к их неработоспособности. Вследствие этого проверка со сравнением работоспособности эксплойта для оригинального и защищенного приложения не является показательной. Проверка работоспособности эксплойтов требует разработки отдельной методики и выходит за рамки данной статьи. Проверка корректности реализации проводилась по трем направлениям: неизменность логики работы программ, тестирование производительности и устранение пригодных для проведения атак гаджетов.

Для контроля неизменности логики работы применялись приложения (синтетические тесты, программы тестирования производительности и программы, снабженные модульными тестами), для которых исполняются все участки, для которых вносились изменения (это проверялось средствами контроля покрытия кода тестами *gcov*). Вследствие того, что выше приведено обоснование эквивалентности вносимых изменений (то есть отсутствия влияния на логику работы), то для оценки корректности не важна вариативность входных данных. Неизменность логики работы оценивалась по идентичности выводи-

мых программой данных при одинаковых входных (для оригинальной и защищенной программ) и отсутствию сбоев в работе. Результаты показали корректность и неизменность работы алгоритмов. Примеры эквивалентных замен для устранения гаджетов приведены в таблице 4. Для контроля устранения пригодных для использования гаджетов использовалось средство их поиска «ROPgadget» [25]. Анализ показал, что во всех обработанных программах отсутствуют иные гаджеты, кроме приведенных в разделе «Программная реализация». Примеры оригинальных и защищенных программ приведены в публичном репозитории github.com/LubkinIvan/zpk/tree/main/examples

Таблица 4. Примеры эквивалентных замен для устранения гаджетов

Содержимое памяти оригинальной программы	Опасный участок оригинальной программы	Опасный участок после эквивалентной замены	Содержимое памяти защищенной программы	Примечание
01 C2	ADD EDX EAX	MOV RCX, RAX ADD EDX, ECX	48 89 C1 01 CA	Опасное значение в <i>ModRM</i>
48 83 C3 01	ADD RBX, 1	MOV RCX, RBX ADD RCX, 1 MOV RBX, RCX	48 89 D9 48 83 C1 01 48 89 CB	Опасное значение в <i>ModRM</i>
E8 C3 FF FF FF	CALL -3D	NOP NOP CALL -3F	90 90 E8 C1 FF FF FF	Опасное значение в операнде <i>Imm32</i>

Сравнение производительности выполнялось программным средством «coremark» [26]. Оно было выбрано, так как дополнительно контролирует целостность результатов расчетов. Замеры производительности показали падение относительно оригинальной программы: при использовании *RDRAND* – 91 %, *RDTS* – 53 %, внешний источник ГПСЧ – 14 % (число запусков 100 для каждого). Характеристики тестовой системы: процессор AMD A6-6310 1,8 ГГц, 4 ядра, 4 ГБ ОЗУ, НЖМД 7200 об/мин.

Прямое сравнение предлагаемого средства с аналогами не является в полной мере корректным, так как они рассчитаны на несовпадающую область применения. Дополнительной проблемой является то, что для перечисленных в разделе 2 ближайших программных ана-

логов не опубликованы результаты их применения, вследствие чего возможно сравнение только на качественном уровне. Отметим также, что сравнение производительности носит справочный характер, так как если аналог имеет меньшие накладные расходы, но не обеспечивает защиту, то это достоинство не играет роли. Результаты сравнения приведены в таблице 5.

Таблица 5. Сравнение решений для противодействия RoP-уязвимостям

Средство защиты	Не анализируемые области	Устранение гаджетов	Уязвимость при наличии примитива чтения	Треб. исх. текст
<i>PaXgrsecurity RAP</i>	Исполнимые данные	Нет	Да	Да
Обфускатор машинного кода	Технологические участки и исполнимые данные	В анализируемых областях	Да	Да
<i>Selfrando</i>	Отсутствуют	Нет	Да	Да
<i>Shuffler</i>	Технологические участки и исполнимые данные	В части эпилогов	Между интервалами перерасшифрования	Нет
<i>RuntimeASLR</i>	Отсутствуют	Нет	Да	Нет
<i>G-free</i>	Технологические участки и исполнимые данные	Да	Нет	Да
<i>Scylla</i>	Технологические участки и исполнимые данные	Нет	В рамках расшифрованной области	Да
Разработанное решение	Отсутствуют	Да	Использование примитива чтения и повреждение стека в рамках вызова одной подпрограммы	Нет

8. Заключение. В данной статье были предложены оригинальная методика защиты от RoP-уязвимостей и её алгоритмическое обеспечение.

Предложенная методика поиска уязвимых участков и меры по их устранению обеспечивают выявление и эквивалентные замены участков, которые могут использоваться для атак. Научная новизна

заключается в обеспечении защиты в условиях наличия у злоумышленника почти полной информации об атакуемой программе (кроме ключей защиты K_i).

Разработанный алгоритм определения свободных ресурсов позволяет встраивать код защиты, не нарушая логику работы программ. Данный алгоритм учитывает БИП, что позволяет, в отличие от аналогов, получать более точные сведения о свободных регистрах. Методика защиты эпилогов от использования в составе гаджетов, с одной стороны, затрудняет получение атакующим контроля над ГПУ, а с другой стороны, не позволяет использовать их как гаджеты, если контроль над ГПУ был получен иным способом. Отличием от существующих средств является отсутствие общих для подпрограмм ключей защиты, компрометация которых делает уязвимой всю программу. Алгоритмическое обеспечение методики позволило реализовать ее в виде ПС.

Проведенная экспериментальная оценка программного средства показала, что данное средство обеспечивает устранение гаджетов, которые не основаны на инструкциях *RET* из состава эпилогов. Для эпилогов обеспечивается блокирование работы программы при попытке их использования атакующим. Для случаев, когда атакующий проводит атаки по сети, падение производительности находится на уровне менее защищенных аналогов.

Предлагаемое решение в ходе обеспечения защиты создает накладные расходы. Вследствие этого не имеет смысла его массовое применение. Наиболее целесообразным является применение для защиты сетевых сервисов, с которыми может взаимодействовать атакующий, и для приложений, обрабатывающих поступающие извне защищенной информационной системы данные. В описанных ситуациях успешная эксплуатация уязвимостей удаленного исполнения кода создает опасность нанесения неприемлемого ущерба, что перевешивает создаваемое системой защиты замедление работы.

Дальнейшее развитие видится в апробации для ОС *Windows* и разработке методики для оценки применимости реальных эксплойтов для защищенных приложений с учетом вносимых системой защиты в стек изменений.

Литература

1. Гласс, Р. Факты и заблуждения профессионального программирования // СПб.: Символ-Плюс. 2007. 240 с.
2. Вишняков А.В. Классификация ROP-гаджетов // Труды ИСП РАН. 2016. Т. 28. Вып. 6, с. 27–36. DOI: 10.15514/ISPRAS-2016-28(6)-2
3. Vedvyas Shanbhogue, Deepak Gupta, and Ravi Sahita. Security Analysis of Processor Instruction Set Architecture for Enforcing Control-Flow Integrity // Proceedings of the 8th International Workshop on Hardware and Architectural Support for Security and

- Privacy (HASP '19). Association for Computing Machinery, New York, NY, USA. 2019. Article 8, 1–11. DOI: <https://doi.org/10.1145/3337167.3337175>
4. Intel 64 and IA-32 Architectures Software Developer's Manual Combined Volumes: 1, 2A, 2B, 2C, 2D, 3A, 3B, 3C, 3D, and 4 // <https://software.intel.com/content/dam/develop/external/us/en/documents-tps/325462-sdm-vol-1-2abcd-3abcd.pdf>
5. Intel Launches World's Best Processor for Thin-and-Light Laptops: 11th Gen Intel Core // <https://www.intel.com/content/www/us/en/newsroom/news/11th-gen-tiger-lake-evo.html>
6. RAP: RIP ROP 2015 // <https://pax.grsecurity.net/docs/PaXTeam-H2HC15-RAP-RIP-ROP.pdf>
7. Koo, Z.Z., Ayop, Zakiah, Abidin, Z.Z. Analysis of ROP attack on grsecurity / PaX linux kernel security variables // International Journal of Applied Engineering Research. 2017. no. 12. pp. 13179–13185.
8. Иванников В., Курмангалеев Ш., Белеванцев А., Нурмухаметов А., Савченко В., Матевосян Р., Аветисян А. Реализация запутывающих преобразований в компиляторной инфраструктуре LLVM // Труды ИСП РАН. 2014. Т. 26. Вып. 1. С. 327–342.
9. Нурмухаметов А.Р., Курмангалеев Ш.Ф., Каушан В.В., Гайсарян С.С. Применение компиляторных преобразований для противодействия эксплуатации уязвимостей программного обеспечения // Труды ИСП РАН. 2014. Т. 26. Вып. 3. С. 113–126. DOI: 10.15514/ISPRAS-2014-26(3)-6.
10. ИСП Обфускатор. Технология запутывания кода для защиты от эксплуатации уязвимостей // https://www.ispras.ru/technologies/isp_obfuscator/
11. Нурмухаметов А.Р., Жаботинский Е.А., Курмангалеев Ш.Ф., Гайсарян С.С., Вишняков А.В. Мелкогранулярная рандомизация адресного пространства программы при запуске // Труды ИСП РАН. 2017. Т. 29. Вып. 6. С. 163–182. DOI: 10.15514/ISPRAS-2017-29(6)-9.
12. S. Crane, A. Homescu, P. Larsen. Code randomization: Haven't we solved this problem yet? Cybersecurity Development (SecDev), IEEE. 2016.
13. M. Conti, S. Crane, T. Frassetto et al. Selfrando: Securing the tor browser against de-anonymization exploits // PoPETs. 2016. no. 4. pp. 454–469.
14. D. Williams-King, G. Gobieski, K. Williams-King et al. Shuffler: Fast and deployable continuous code re-randomization // Proceedings of the 12th USENIX Conference on Operating Systems Design and Implementation. 2016. pp. 367–382.
15. Kangjie Lu, Stefan Nürnberger, Michael Backes, Wenke Lee. How to Make ASLR Win the Clone Wars: Runtime Re-Randomization // Proceedings of the 23rd Annual Network and Distributed System Security Symposium. 2016.
16. Onarlioglu K., Bilge L., Lanzi A., Balzarotti D., Kidra E. G-Free: Defeating return-oriented programming through gadget-less binaries // Proceedings of ACSAC: M. Franz and J. McDermott, Eds. ACM Press. 2010. pp. 49–58.
17. Jinku Li, Zhi Wang, Xuxian Jiang, Mike Grace, Sina Bahram. Defeating return-oriented rootkits with «return-less» kernels. // Proceedings of EuroSys. 2010, edited by G. Muller. ACM Press. pp. 195–208.
18. Dean Sullivan, Orlando Arias, David Gens, Lucas Davi, Ahmad-Reza Sadeghi, Yier Jin. 2017. Execution Integrity with In-Place Encryption. arXiv preprint arXiv:1703.02698 (2017).
19. Lubkin I.A., Subbotin N.A. Technique of verified program module modification with algorithm preservation // IEEE Xplore Digital Library. 2017. 11th International IEEE scientific and technical conference "Dynamics of systems, mechanisms and machines" (Dynamics), 2017. pp. 1–5.

20. Lubkin I.A., Bazhenov I. O. Methodology of software code decomposition analysis // Dynamics of systems, mechanisms and machines. Omsk. 2018. pp. 1–5.
21. Hovav Shacham. The Geometry of Innocent Flash on the Bone: Return-into-libc without Function Calls (on the x86). 2007. ACM Conference on Computer and Communications Security (CCS), Proceedings of CCS, 2007. pp. 552–561.
22. Статья Permutation conditions. URL: <https://z0mbie.dreamhosters.com/pcond.txt> (дата обращения 01.09.2021).
23. Репозиторий с исходным кодом библиотеки eXtended Disassembler Engine (version 1.02). URL: <https://github.com/nimrood/xde> (дата обращения 01.09.2021).
24. AMD64 Architecture Processor Supplement Draft Version 0.99.7 // https://www.uclibc.org/docs/psABI-x86_64.pdf
25. Инструмент ROPgadget. Репозиторий с исходным кодом. URL: <https://github.com/JonathanSalwan/ROPgadget> (дата обращения 01.09.2021).
26. Coremark. Программа оценки производительности. URL: <https://github.com/eembc/coremark> (дата обращения 01.09.2021).

Лубкин Иван Александрович — старший преподаватель, кафедры безопасности информационных технологий, СибГУ им. М.Ф. Решетнёва. Область научных интересов: безопасность операционных систем, защита программного кода от несанкционированного использования, противодействие уязвимостям. Число научных публикаций — 27. lubkin@rambler.ru; проспект им. газеты Красноярский рабочий, 48Б, 660037, Красноярск, Россия; р.т.: +7 (391) 222-76-39; факс: +7(391)222-76-39.

Золотарев Вячеслав Владимирович — канд. техн. наук, доцент, заведующий кафедрой, кафедра безопасности информационных технологий, СибГУ им. М.Ф. Решетнёва. Область научных интересов: управление информационной безопасностью, противодействие уязвимостям. Число научных публикаций — 98. zolotarev@sibsau.ru; проспект им. газеты Красноярский рабочий, 48Б, 660037, Красноярск, Россия; р.т.: +7(391)222-76-39.

Поддержка исследований. Исследование выполнено при финансовой поддержке Минобрнауки России (грант ИБ). Соглашение № 21/2020.

I. LUBKIN, V. ZOLOTAREV
**COMPREHENSIVE DEFENSE SYSTEM AGAINST
VULNERABILITIES BASED ON RETURN-ORIENTED
PROGRAMMING**

Lubkin I., Zolotarev V. Comprehensive Defense System against Vulnerabilities Based on Return-Oriented Programming.

Abstract. It is difficult or impossible to develop software without included errors. Errors can lead to an abnormal order of machine code execution during data transmission to a program. Program splitting into routines causes possible attacks by using return instructions from these routines. Most of existing security tools need to apply program source codes to protect against such attacks. The proposed defensive method is intended to a comprehensive solution to the problem. Firstly, it makes it difficult for an attacker to gain control over program execution, and secondly, the number of program routines, which can be used during the attack, decreases. Specific security code insertion is used at the beginning and end of the routines to make it complicated to gain control over the program execution. The return address is kept secure during a call of the protected routine, and the protected routine is restored after its execution if it was damaged by the attacker. To reduce the number of suitable routines for attacks, it was suggested to use synonymous substitutions of instructions that contain dangerous values. It should be mentioned that proposed defensive measures do not affect the original application's algorithm. To confirm the effectiveness of the described defensive method, software implementation and its testing were accomplished. Acknowledging controls were conducted using synthetic tests, performance tests and real programs. Results of testing have demonstrated the reliability of the proposed measures. It ensures the elimination of program routines suitable for attacks and ensures the impossibility of using standard return instructions for conducting attacks. Performance tests have shown a 14 % drop in the operating speed, which approximately matches the level of the nearest analogues. The application of the proposed solution declines the number of possible attack scenarios, and its applicability level is higher in comparison with analogues.

Keywords: vulnerability, remote code execution, code protection, RoP, code insertion.

Lubkin Ivan — Senior lecturer, Department of information technologies security, Reshetnev Siberian State University of Science and Technology. Research interests: security of operating systems, protection of program code from unauthorized use, countering vulnerabilities. The number of publications — 27. lubkin@rambler.ru; 48Б, newspaper Krasnoyarskiy rabochiy Ave., 660037, Krasnoyarsk, Russia; office phone: +7 (391) 222-76-39; fax: +7(391)222-76-39.

Zolotarev Vyacheslav — Ph.D., Associate Professor, Head of department, Department of information technologies security, Reshetnev Siberian State University of Science and Technology. Research interests: information security management, vulnerabilities protection. The number of publications — 98. zolotarev@sibsau.ru; 48Б, newspaper Krasnoyarskiy rabochiy Ave., 660037, Krasnoyarsk, Russia; office phone: +7(391)222-76-39.

Acknowledgements. The reported study was funded by Russian Ministry of Science (information security). Contract № 21/2020.

References

1. Glass, R. [Facts and misconceptions of professional programming] Fakty i zabluzhdenija professional'nogo programmirovaniya // Saint Petersburg: Symbol-Plus. 2007. 240 p. (InRuss.)
2. Vishnjakov, A.V. [Classification of ROP-gadgets] Klassifikacija ROP-gadzhetov // Proceedings of the ISP RAS. 2016. Vol. 28. Issue 6, pp. 27–36. DOI: 10.15514/ISPRAS-2016-28(6)-2 (InRuss.)
3. Vedvyas Shanhogue, Deepak Gupta, and Ravi Sahita. Security Analysis of Processor Instruction Set Architecture for Enforcing Control-Flow Integrity // Proceedings of the 8th International Workshop on Hardware and Architectural Support for Security and Privacy (HASP '19). Association for Computing Machinery, New York, NY, USA. 2019. Article 8, 1–11. DOI:<https://doi.org/10.1145/3337167.3337175>
4. Intel 64 and IA-32 Architectures Software Developer's Manual Combined Volumes: 1, 2A, 2B, 2C, 2D, 3A, 3B, 3C, 3D, and 4 // <https://software.intel.com/content/dam/develop/external/us/en/documents-tps/325462-sdm-vol-1-2abcd-3abcd.pdf>
5. Intel Launches World's Best Processor for Thin-and-Light Laptops: 11th Gen Intel Core // <https://www.intel.com/content/www/us/en/newsroom/news/11th-gen-tiger-lake-evo.html>
6. RAP:RIP ROP 2015 // <https://pax.grsecurity.net/docs/PaXTeam-H2HC15-RAP-RIP-ROP.pdf>
7. Koo, Z.Z. & Ayop, Zakiah & Abidin, Z.Z. Analysis of ROP attack on grsecurity / PaX linux kernel security variables // International Journal of Applied Engineering Research. 2017. no. 12. pp. 13179–13185.
8. Ivannikov V., Kurmangaleev Sh., Belevancev A., Nurmuhametov A., Savchenko V., Matevosjan R., Avetisjan A. [Implementation of confusing transformations in the computer infrastructure LLVM] Realizacija zaputyvajushih preobrazovanij v kompiljatoroj infrastrukture LLVM // Proceedings of the ISP RAS. 2014. Vol. 26. Issue 1. pp. 327–342. (InRuss.)
9. Nurmuhametov A.R., Kurmangaleev Sh.F., Kaushan V.V., Gajsarjan S.S. [] Primenenie kompiljatornyh preobrazovanij dlja protivodejstviya jekspluatacii ujazvimostej programmnoho obespechenija // Proceedings of the ISP RAS. 2014. Vol. 26. Issue 3. pp. 113–126. DOI: 10.15514/ISPRAS-2014-26(3)-6. (InRuss.)
10. [ISP Obfuscator. Code obfuscation technology to protect against exploiting vulnerabilities] ISP Obfuscator. Tehnologija zaputyvanija koda dlja zashhity ot jekspluatacii ujazvimostej // https://www.ispras.ru /technologies/isp_obfuscator/ (InRuss.)
11. Nurmuhametov A.R., Zhabotinskij E.A., Kurmangaleev Sh.F., Gajsarjan S.S., Vishnjakov A.V. [Fine-grained randomization of the program's address space at startup] Melkogranuljarnaja randomizacija adresnogo prostranstva programmy pri zapuske // Proceedings of the ISP RAS. 2017. Vol. 29. Issue. 6. pp. 163–182. DOI: 10.15514/ISPRAS-2017-29(6)-9. (InRuss.)
12. S. Crane, A. Homescu, P. Larsen. Code randomization: Haven't we solved this problem yet? Cybersecurity Development (SecDev), IEEE. 2016.
13. M. Conti, S. Crane, T. Frassetto et al. Selfrando: Securing the tor browser against de-anonymization exploits // PoPETs. no. 4. 2016. pp. 454–469.
14. D. Williams-King, G. Gobieski, K. Williams-King et al. Shuffler: Fast and deployable continuous code re-randomization // Proceedings of the 12th USENIX Conference on Operating Systems Design and Implementation. 2016. pp. 367–382.
15. Kangjie Lu, Stefan Nürnberger, Michael Backes, and Wenke Lee. How to Make ASLR Win the Clone Wars: Runtime Re-Randomization // Proceedings of the 23rd Annual Network and Distributed System Security Symposium. 2016.

16. Onarlioglu K., Bilge L., Lanzi A., Balzarotti D., Kidra E. G-Free: Defeating return-oriented programming through gadget-less binaries // Proceedings of ACSAC: M. Franz and J. McDermott, Eds. ACM Press, 2010. pp. 49–58.
17. Jinku LI, Zhi WANG, Xuxian JIANG, Mike GRACE, and Sina BAHARAM. Defeating return-oriented rootkits with «return-less» kernels // Proceedings of EuroSys: Edited by G. Muller. ACM Press, 2010. pp. 195–208.
18. Dean Sullivan, Orlando Arias, David Gens, Lucas Davi, Ahmad-Reza Sadeghi, Yier Jin. Execution Integrity with In-Place Encryption // arXiv preprint arXiv:1703.02698. 2017.
19. Lubkin I.A., Subbotin N.A. Technique of verified program module modification with algorithm preservation // 11th International IEEE scientific and technical conference "Dynamics of systems, mechanisms and machines" (Dynamics): IEEE Xplore Digital Library, 2017. pp. 1–5.
20. Lubkin I.A., Bazhenov I.O. Methodology of software code decomposition analysis // Dynamics of systems, mechanisms and machines. Omsk, 2018. pp. 1–5.
21. Hovav Shacham. The Geometry of Innocent Flash on the Bone: Return-into-libc without Function Calls (on the x86). 2007 ACM Conference on Computer and Communications Security (CCS) // Proceedings of CCS. 2007. pp. 552–561.
22. Article Permutation conditions. URL: <https://z0mbie.dreamhosters.com/pcond.txt> (accessed 01.09.2021)
23. Sources of library eXtended Disassembler Engine (version 1.02). URL: <https://github.com/nimrood/xde> (accessed 01.09.2021)
24. AMD64 Architecture Processor Supplement Draft Version 0.99.7 // https://www.uclibc.org/docs/psABI-x86_64.pdf
25. Instrument ROPgadget. Source code repository. URL: <https://github.com/JonathanSalwan/ROPgadget> (accessed 01.09.2021).
26. Coremark. Benchmark software. URL: <https://github.com/eembc/coremark> (accessed 01.09.2021)

M. ABRAMOV, E. TSUKANOVA, A. TULUPYEV, A. KOREPANOVA,
S. ALEKSANIN
**IDENTIFICATION OF DETERIORATION CAUSED BY AHF,
MADS OR CE BY RR AND QT DATA CLASSIFICATION**

Abramov M., Tsukanova E., Tulupyevev A., Korepanova A., Aleksanin S. **Identification of Deterioration caused by AHF, MADS or CE by RR and QT Data Classification.**

Abstract. A sharp deterioration of the patient's condition against the backdrop of the development of life-threatening arrhythmias with symptoms of acute heart failure (AHF), multiple organ dysfunction syndrome (MODS) or cerebral edema (CE) can lead to the death of the patient. Since the known methods of automated diagnostics currently cannot accurately and promptly determine that the patient is in a life-threatening condition leading to the fatal outcome caused by AHF, MODS or CE, there is a need to develop appropriate methods. One of the ways to identify predictors of such a state is to apply machine learning methods to the collected datasets. In this article, we consider using data analysis methods to test the hypothesis that there is a predictor of death risk assessment, which can be derived from the previously obtained values of the ECG intervals, which gives a statistically significant difference for the ECG of the two groups of patients: those who suffered deterioration leading to the fatal outcome caused by MODS, AHF or CE, and those with favorable outcome. A method for unifying ECG data was proposed, which allow, based on the sequence of RR and QT intervals, to the construct of a number that is a characteristic of the patient's heart condition. Based on this characteristic, the patients are classified into groups: the main (patients with fatal outcome) and control (patients with favorable outcome). The resulting classification method lays the potential for the development of methods for identifying the patient's health condition, which will automate the detection of its deterioration. The novelty of the result lies in the confirmation of the hypothesis stated above, as well as the proposed classification criteria that allow solving the urgent problem of an automatic detection of the deterioration of the patient's condition.

Keywords: ECG-based patient classification, deterioration identification, medical prediction, ECG analysis, machine learning, artificial intelligence, data science, logistic regression, mortality risk stratification, RR interval, QT interval.

1. Introduction. Currently, in medical practice, the generally accepted method for analyzing an electrocardiogram (ECG) is a visual assessment of changes in ECG curves. Despite the great diagnostic efficiency of such an analysis, it obviously has a subjective component, and in the area of borderline states of health, its diagnostic ability is extremely unsatisfactory [1, 2]. Usually, 6 waves can be distinguished on the ECG record: P, Q, R, S, T, U. The intervals between the waves, the duration, shape, amplitude of the waves carry information that allows doctors to draw conclusions whether the patient suffers from diseases associated with the cardiovascular system. The RR interval is an indicator that characterizes the duration of the cardiac cycle and is measured in any ECG lead. The availability of research on this indicator is undeniable and has long been used in assessing heart rate variability [3, 4, 5]. For example, a decrease in

SDNN (standard deviation of RR intervals) of less than 50 milliseconds is a highly specific (88%) sign in predicting fatal outcomes in patients with myocardial infarction [6].

The QT interval, which reflects the time between the beginning of the depolarization process (the beginning of the Q wave) and the completion of the process of repolarization of the ventricular myocardium (the end of the T wave), is one of the most significant ECG parameters in cardiology. Upward or downward deviation of its value from the norm is associated with a high risk of malignant ventricular arrhythmias. Any genesis of changes in the QT interval (congenital or acquired) is equally dangerous as a risk factor for the development of ventricular tachyarrhythmias [7]. It should be noted that cardiac fluctuations are not strictly periodic and are characterized by rhythm variability [33].

There are good results in research on the spectral characteristics of heart activity [34–36], which uses mathematical modeling of the human cardiovascular system and shows its connection with different patient states. There are also results in the field of data analysis and artificial intelligence, allowing a posteriori in an automatic mode to form a class of patients with an increased risk of stroke [8] or to predict the state of the pilot, using the results of the ECG as an indicator of the pilot's physical condition [9]. Thus, the study of the ECG has shown itself to be promising in predicting the various conditions of patients; however, until now, the question of the relationship between the state of deterioration of health in patients, manifested by symptoms of multiple organ dysfunction syndrome (MODS), acute heart failure (AHF) or cerebral edema (CE), and their ECG data, has not yet been studied. Research in this area would potentially make it possible to predict changes in this condition on the basis of ECG signal. Timely identification of such conditions and subsequent prediction of their changes are urgent tasks. Moreover, with the development of such conditions, deterioration occurs quickly, without leaving sufficient time to prepare for treatment, and standard visual analysis does not allow a sufficiently accurate forecast of the patient's condition.

Thus, on the one hand, the problem of automating the assessment of the patient's condition leading to a fatal outcome from the development of MODS, AHF or CE is urgent, its solutions are in demand; on the other hand, there are no existing studies aimed at the automation of this assessment. The aim of this study is to automate the assessment of the patient's condition based on ECG data using machine learning methods. This article presents the result of the first stage of the study, namely, the verification using data analysis methods of the hypothesis that there is a predictor of death risk assessment, which can be derived from the

previously obtained values of the ECG intervals, which gives a statistically significant difference for the ECG of the two groups of patients: those who suffered deterioration leading to the fatal outcome caused by MODS, AHF or CE, and those with a favorable outcome.

The practical significance of the study lies in the formation of the potential for solving the problem of automated assessment of the patient's condition based on ECG data and predicting outcomes in case of deterioration in health using machine learning models. Such a diagnosis could allow earlier recognition and faster response to potentially life-threatening conditions of patients. Doctors could take the necessary measures before the moment of critical changes. The theoretical significance of the study lies in the proposition of the estimation method for classifying patients into the main (with MODS, AHF or CE, which led to a fatal outcome) and control (recovered) groups based on the ECG records, as well as the formation of a basis for the development of a classification methodology that allows achieving higher accuracy metrics, and predicting deterioration. The novelty of the study is attributable to the lack of methods, models and algorithms for identifying the state of deterioration and predicting its changes as a result of the development of the syndrome of acute cardiovascular, multiple organ failure or cerebral edema. The new feature based on the sequence of ECG records allowing patient's classification into main and control group is proposed, its level of significance in a logistic regression model is 3.64×10^{-6} , a new method is proposed for solving the urgent problem of identifying the deterioration of the patient's condition on the basis of ECG analyzes, which differs from the existing ones in the formulation of the problem being solved, and in the accuracy of the results obtained.

Machine learning and data analysis methods are increasingly used today in applications to problems from the field of medicine and healthcare [10–24]. These are generally the tasks of diagnostics, predicting the course of the disease development, assistance in prescribing treatment, etc. [10–24]. An important role among these tasks is those related to monitoring and diagnostics of the most probable scenarios of the course of the disease. For example, in [10], on the basis of data collected using a smartphone (voice data on calls, data on the use of smartphones, reflecting social and physical activity), authors draw the conclusions about the user's mood, and the prospects for bipolar disorder. In [11], machine learning methods are used in diabetes monitoring to determine the risk of diabetes based on the personal medical history of patients. In addition, machine learning methods are also used to study the brain, as, for example, in [21] — for the motor imagery EEG classification.

There are other examples of the successful use of artificial intelligence and data analysis methods in tasks from the field of medicine, in particular, a wide range of works using image analysis methods to search for anomalies on X-rays and other images [16].

One of the applications of machine learning in medicine is also the automation of the recognition of critical cases for various diseases of patients. [19–21, 22, 23]. For example, [19] proposed a model for automating the analysis of breast cancer from images. In [20], the authors applied unsupervised machine learning methods to cluster ICU patients to identify patients at increased risk of mortality based on laboratory data. The authors of [21] proposed a new method for analyzing the relationship between the results of electroencephalography and human cognitive functions. The work [23] proposed methods for identifying patients with osteoporosis who have an increased risk of hip fracture.

Another common area of application of machine learning and data analysis methods in medicine is the processing of information related to the human cardiovascular system [12, 13, 17, 18, 22, 24]. Authors of [24] proposed a model for heart disease prediction; they trained three classifiers on the Cleveland dataset of 303 data instances that contains the following patient features: gender, type of chest pain, cholesterol level, maximum heart rate, etc. The authors used the following ECG features: the slope of the peak exercise ST segment V1, ST depression induced by exercise relative to rest, maximum heart rate achieved, the characteristics of resting ECG. The results showed that LR and SVM have high effectiveness (87% and 85% accuracy, respectively). The work [12] is focused on the automation of several waveforms detection on the ECG signal with the use of machine learning methods. In [13], these methods are used to develop a classification of heart diseases based on patient data from the Cleveland dataset. The author used models different from those in [24]. 6 machine learning models were used to classify patients into healthy and suffering from cardiovascular diseases subjects; CNN showed the best performance with the highest accuracy in 85.86%. [17] is focused on the ECG anomaly leading to different types of heart diseases automatic detection with MATLAB. The following features were considered: interval and amplitude values of P wave, Q wave, T wave, QRS complex, and ST segment, FIR filters showed good results. Results of [18] showed a great performance of machine learning techniques in automatic coronary artery disease detection based on ECG time series. The authors used 50 ECG features, including average and standard deviation of RR interval values, corrected QT interval values, etc. Authors of [22] proposed a method for heart disease identification with missing data handling: they replaced missing values with

the mean values during pre-processing. Three ML models were trained on the data from the Cleveland dataset.

There are also works on heart disease detection which don't use ECG information. In [14], authors tried to predict cardiovascular and cerebrovascular events in patients with hypertension based on data on the status of use of medical resources, detailed information about diseases, their types, treatment methods, prescriptions, and the status of the clinic in which the treatment was carried out. In [15], machine learning methods are used to assess the development of coronary heart disease in breast cancer survivors.

Thus, the application of machine learning to ECG analysis has shown good results in identifying conditions associated with the cardiovascular system. However, the task of the analysis of ECG measurements to identify the subsequent death of the patient as a result of the development of MODS, AHF or CE has not yet been addressed.

The purpose of the general study is to automate the assessment of the patient's condition based on ECG data using machine learning methods. As the first stage of the study, this work solves the problem of testing the hypothesis using data analysis methods that there is a predictor of death risk assessment, which can be derived from the basis of previously obtained values of the ECG intervals, which gives a statistically significant difference for the ECG of the two groups of patients: those who suffered deterioration leading to the fatal outcome caused by MODS, AHF or CE, and those with a favorable outcome. The dataset provided by the Federal State Public Enterprise Nikiforov's All-Russian Center for Emergency and Radiation Medicine of the Emergencies Ministry of Russia (the Nikiforov's ARCERM) is characterized by a wide variation in the number of ECG tests taken (from 6 to 407) in different patients. This variation is explained by the different number of hospitalization days of patients, different reasons for admission to the hospital, and an increase in the frequency of testing when the condition worsens. Since the analysis of such data by statistical methods is difficult, it is necessary to unify them, that is, from the sets of values of the RR and QT intervals for each patient, obtain one value. In other words, in this work, the following problem was solved: on the basis of all the obtained values of the intervals, it is necessary to form one, which will be a characteristic of the patient's ECG records, and the differences in which for two groups (patients who died as a result of deterioration with symptoms of AHF, MODS and CE, and patients with a favorable outcome) will be statistically significant.

2. Materials and methods. The dataset was collected, the criterion derived from ECG data was proposed. The formulated hypothesis was tested using logistic regression.

2.1. Dataset. The first stage of the study was carried out on the basis of anonymized data on 120 patients of both sexes who were consecutively hospitalized in the Federal State Public Enterprise Nikiforov's All-Russian Center for Emergency and Radiation Medicine of the Emergencies Ministry of Russia (the Nikiforov's ARCERM) in the period from 01.01.2015 to 31.12.2017. The main group (MG) consisted of 60 patients with a lethal outcome as a result of deterioration; the control group (CG) consisted of 60 patients with a favorable outcome. The dataset provided by ARCERM at the first stage consisted of tables containing the following information: patient diagnoses, data on all ECG records, general data about the patient, data on events significant for the patient's condition that occurred during his stay in the hospital (surgery, life-threatening change in the rhythm of heart contractions, etc.). Below we will take a closer look at each of the tables:

Patient diagnosis table was formed on the basis of epicrisises; each patient in the dataset had from 1 to 37 diagnoses.

ECG diaries contained records of all available amenable to analysis of ECG records in the form of diaries. All patients underwent a standard 12-lead ECG: 4 electrodes were placed on the arms and legs distal from the shoulder and thigh; 6 electrodes were placed on the chest surface at the following points (V1 – fourth intercostal space along the right sternal line, V2 – fourth intercostal space along the left sternal line, V3 – in the middle of the distance between V2 and V4, V4 – fifth intercostal space along the midclavicular line, V5 – in horizontal plane V4 along the anterior axillary line, V6 – in the horizontal plane V4 along the mid-axillary line). The duration of the QT interval on the resting ECG using the interval editing module was measured by the classical method of E. Lepeshkin and B. Surawicz [25], which is based on drawing a straight line tangent along the line of maximum slope of the descending part of the T-wave until it intersects with the isoline (slope method). Measurements were performed in standard lead II and, in the case of a pronounced U-wave, in chest lead V5. For ECG measurements, MAC 1600 ECG Analysis System and MUSE NX were used. The table contains the following information:

- Date and time of measurement;
- Data on the rhythm of heart contractions. In the presented data set, 4 main types of rhythm were distinguished: 1) sinus; 2) atrial fibrillation and atrial flutter; 3) ectopic; and 4) with a permanent pacemaker;
- Whether the rhythm has changed since the last diary. The exact time of the change is unknown; the rhythm can change from one of the first three types described above to another of the first three types;
- Whether the patient was in the intensive care unit at the time of the ECG recording;

- Characteristics of the general condition of the patient (satisfactory, moderate, severe, etc.);
- Data on the patient's hemodynamics (stable, unstable), on the patient's breathing and consciousness, which together determine whether the patient is in deterioration. This information was extracted from complete diaries kept by doctors in a free-form text while observing patients;
- Whether the patient experienced deterioration during the ECG recording;
- Was the operation performed on the patient close to the time of the ECG recording;
- Whether sympathomimetic drugs have been administered to the patient.

General patient data contained information on the following patient characteristics: age, gender, date of admission, date of a discharge / fatal outcome, number of days the patient was hospitalized, cause of fatal outcome, and hospitalization department.

RR and QT intervals contained data on the RR and QT intervals corresponding to the measurement diaries. Each diary corresponds to a different number of pairs of RR and QT intervals; for each pair of intervals, their lengths are known, as well as the presence of supraventricular and ventricular extrasystoles.

Next, let's consider the primary statistics of the dataset. Patients were treated in the departments of cardiology, cardiovascular surgery, intensive care, neurology, neurosurgery, emergency surgery, therapy, traumatology, hematology, and rehabilitation. In the main group, 26 (43.3%) patients had multiple organ dysfunction syndrome as the cause of fatal outcome, 28 (46.7%) patients had an acute cardiovascular failure, and 6 (10%) had cerebral edema. The gender distribution in both groups was the same and amounted to 50% of men and 50% of women. The duration of hospitalization ranged from 1 to 172 bed-days. Among the patients of the main group, diseases of the circulatory system prevailed as the main pathology and amounted to about 50% of cases, taking into account the concomitant pathology, the number of those suffering from diseases of the circulatory system is greater – 83%. In the control group, diseases of the circulatory system also prevailed, amounting to more than 60% (Table 1), also taking into account the concomitant pathology, the number of those suffering from diseases of the circulatory system was greater – 82%.

During hospitalization, all patients underwent electrocardiogram testing. Thus, each patient received from 1 to 38 ECGs. Due to the variability of heart rate between ECG tests, and also due to the fact that the quality of the records did not always allow using the entire ECG record,

each processed record contained a different number of pairs of RR and QT intervals: from 3 to 28.

Table 1. Clinical characteristics of patients

Pathology (ICD code)	MG n (%)	CG n (%)
Main pathology		
Diseases of the circulatory system (I00–I99)	29 (48,3)	41 (68,3)
Neoplasms (C00–D48)	16 (26,7)	5 (8,3)
Endocrine system diseases (E00–E90)	2 (3,3)	5 (8,3)
Infectious and parasitic diseases (A00–B99)	5 (8,3)	0,0
Trauma (S00–T98)	6 (10,0)	0,0
Diseases of the digestive system (K00–K93)	2 (3,3)	9 (15,0)
Concomitant pathology		
Diseases of the circulatory system (I00–I99)	23 (38,3)	9 (15,0)
Neoplasms (C00–D48)	3(5,0)	2 (3,3)
Endocrine system diseases (E00–E90)	18 (30,0)	15 (25,0)
Diseases of the Genitourinary System (N00–N99)	26 (43,3)	12 (20,0)
Diseases of the nervous system (G00–G99)	7 (11,7)	7 (11,7)
Diseases of the digestive system (K00–K93)	22 (36,7)	18 (30,0)
Diseases of the respiratory system (J00–J99)	23 (38,3)	7 (11,7)

The distribution by the type of the rhythm for all ECG records was comparable in both groups and amounted to about 60% for sinus rhythm, about 25-30% for atrial fibrillation (Table 2). The original dataset contained 3 patients who received a permanent pacemaker. Since the pacemaker affects the heart rhythm, the changes in the ECG of these patients may be of a different nature, so these patients were excluded from the study at the moment. The final dataset contained 117 patients: 58 from the main group and 59 from the control group.

Table 2. Electrocardiographic characteristics of patients

Rhythm	Main group	Control group
	n (%)	n (%)
Sinus rhythm	166 (61,7)	123 (65,0)
Atrial fibrillation	63 (23,3)	63 (33,3)
Permanent pacemaker	9 (3,3)	3 (1,7)
Ectopic rhythm	31 (11,7)	0 (0,0)

At the second stage of the study, the model was tested on new data of 38 patients of both sexes who were consecutively admitted to the Federal

State Public Enterprise Nikiforov's All-Russian Center for Emergency and Radiation Medicine of the Emergencies Ministry of Russia (the Nikiforov's ARCERM) during the period from 01.01.2018 to 31.12.2019. The main group (MG) consisted of 19 patients with a lethal outcome, the control group (CG) consisted of 19 patients with a favorable outcome. Sex distribution in both groups (main and control): 9 men and 10 women. The duration of the hospitalization ranged from 1 to 105 bed-days.

2.2. Data preprocessing. The dataset is characterized by a high degree of heterogeneity, for example, the table with diagnoses contains 1986 diagnoses, 1621 of them are unique within the dataset. Also, each patient has a variable number of diagnoses, days of hospitalization, ECG records, RR and QT intervals, etc. An additional difficulty when working with this data is the small size of the dataset. Thus, most of the data is either inapplicable for analysis, or requires a significant degree of preprocessing, so a subset of features was extracted from the general patient data set.

Based on the information about patients collected at the first stage, a database was formed, which contained information about the gender and age of patients, the number of days of hospitalization, pairs of RR and QT intervals from all ECG records.

Based on the available data, two tables were generated. In the first table, each row contains the patient identification number; the group to which the patient belongs (MG or CG); the patient's gender; age; the number of days of hospitalization, and information about whether the type of the heart rate has changed or not. Table 3 shows the general view of the row in the table.

Table 3. General view of a table row containing information about patients

Patient number	Patient group	Patient sex	Patient age	Number of days of hospitalization	Was there a change in rhythm?
$n : n \in 1..117$	$g : g \in \{0;1\}$ (0 for MG; 1 for CG)	$p : p \in \{0,1\}$ (0 for female; 1 for male)	$v : v \in [34, 94]$	$k : k \in 0..172$	$r : r \in \{0;1\}$ (0 if not; 1 if it was)

The second dataset table contains information about all measured intervals (RR and QT interval sets). Since each patient underwent several ECGs, and each ECG contains several RR and QT intervals, there are several rows in the table for each patient, each of which records the patient's number, as well as the values of the RR and QT intervals. For each patient,

there are as many rows as there are RR and QT intervals (from 6 to 407). Table 4 shows a general format of the row of the analyzed dataset.

Table 4. General format of the table row containing information about the RR and QT intervals

Patient ID	RR interval value	QT interval value
$n : n \in 1..117$	$rr : rr \in 244..3042$	$qt:qt \in 141..692$

2.3. The proposed method for the unification of ECG data. To simplify the further presentation of the research results, we introduce the following notation:

Notation. Each patient number i has n_i pairs of RR and QT intervals.

Notation. For the patient number i the value of the j -th RR interval is denoted by RR_{ij} .

Notation. For the patient number i the value of the j -th QT interval is denoted by QT_{ij} .

Notation. For the patient number i the set of all the values of the patient's RR intervals is denoted by RR_i .

Notation. For the patient number i the set of all the values of the patient's QT is denoted by QT_i .

In addition, we will use the notation for some numerical characteristics, used for unifying the data on the recorded ECG.

Notation. For patient number i the mean value of cubes of the values of RR intervals is denoted by $\overline{RR_i}^{(3)}$. To put it another way,

$$\sum_{j=1}^{n_i} \frac{RR_{ij}^3}{n_i} = \overline{RR_i}^{(3)}.$$

Notation. For patient number i the mean value of cubes of the values of QT intervals is denoted by $\overline{QT_i}^{(3)}$. To put it another way,

$$\sum_{j=1}^{n_i} \frac{QT_{ij}^3}{n_i} = \overline{QT_i}^{(3)}.$$

Notation. For patient number i as the **criterion value** will be referred the special dimensionless constant K_i , which is denoted by the following: $K_i = \log_{\overline{QT_i}^{(3)}} \overline{RR_i}^{(3)}$.

This criterion was chosen because it reflects the relationship between the RR and QT intervals in all of the patient's ECGs. This criterion was

chosen among a set of functions with similar properties, as having the highest classifying ability.

It is proposed to consider the values of the K_i criterion or the analyzed dataset as a marker for classifying patients according to the values of the RR and QT intervals into two classes: the main group and the control group. Figures 1, 4 show the distribution of test values in both groups. Figures 2, 3 show the distribution of the criterion value in the MG and CG groups, respectively. According to Shapiro-Wilk normality test, criterion value isn't distributed normally, p-value is equal 0.2833 for the distribution of the criterion value in both groups, 0.4829 for MG, and 0.7372 for CG.

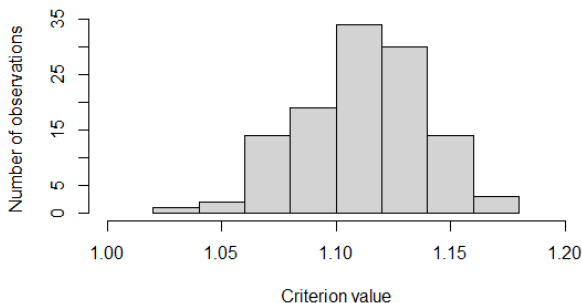


Fig. 1. Distribution of criterion values in both group

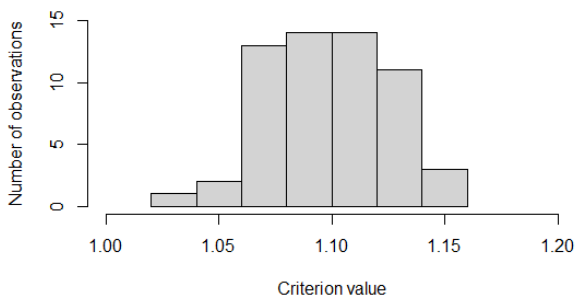


Fig. 2. Distribution of criterion values in the MG group

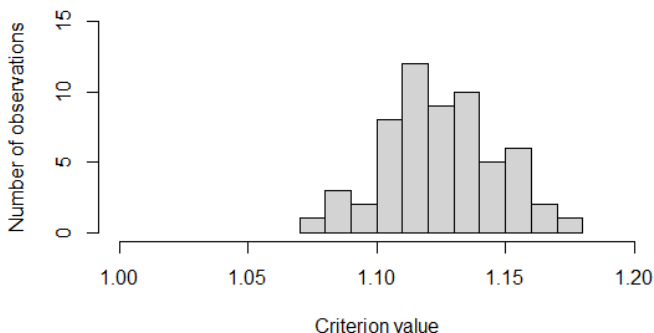


Fig. 3. Distribution of criterion values in the CG group

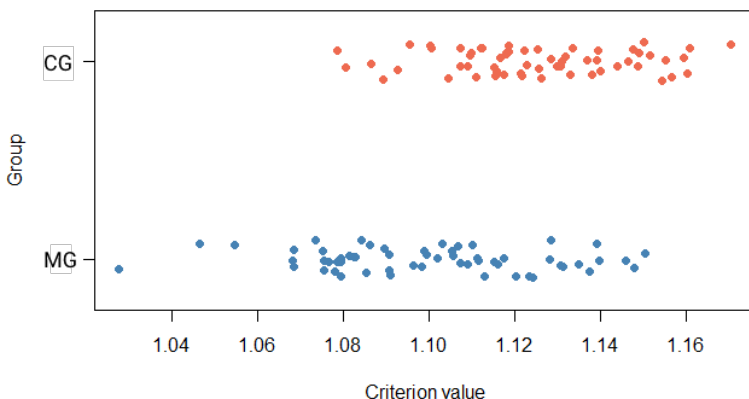


Fig. 4. Scatterplot of criterion value distribution in two groups (1 – MG, 0 – CG)

Hypothesis. There is a predictor of death risk assessment, which can be derived from the previously obtained values of the ECG intervals, which gives a statistically significant difference for the ECG of the two groups of patients: those who suffered deterioration leading to the fatal outcome caused by MODS, AHF or CE, and those with a favorable outcome. The criterion value can be used for patient classification into one of two groups.

To test the hypothesis, a logistic regression model was applied, this model doesn't require specific distribution of predictors. Logistic regression is one of the machine learning methods [25], a statistical model used to predict the likelihood of an object belonging to one of two classes using a logistic function. The following assumptions are made about the probability of an object belonging to classes 1 and 0:

$$P(y = 1 | x) = h_w(x),$$

$$P(y = 0 | x) = 1 - h_w(x),$$

$$h_w(x) = \frac{1}{1 + e^{-(w_0 + w_1x_1 + \dots + w_nx_n)}},$$

where $x_j, j \in 1..n$ denotes attributes of the object x , $w_j : j \in 0..n$ — regression coefficients. In our case, the patients are the objects, the criterion value is the attribute of the object, the patient groups are the classes (1 for MG, and 0 for CG).

3. Results. This section describes the results of testing the proposed criterion during the first and the second stages of the study.

3.1. Testing the criterion. Table 5 shows the characteristics of the logistic regression model between the patient group and the criterion value. Figure 5 shows a plot of sensitivity and specificity (ROC-curve) for the initial data, figure 6 – graph of the logistic function.

Since the *p-value* of the test of the association between criterion value and the patient group is 3.64×10^{-6} , which is less than 0.05, then we can talk about the presence of a statistically significant association. As the size of the dataset was bigger than 30 subjects, and the maximum likelihood estimates are asymptotically normal, for the p-value calculation z-score was applied that uses normal distribution.

Table 5. Results of logistic regression between the group and the value of the criterion for the initial data

	Estimate	Std. Error	z value	p-value
(Intercept)	-49.791	10.762	-4.626	3.72×10^{-6}
slope	44.753	9.664	4.631	3.64×10^{-6}

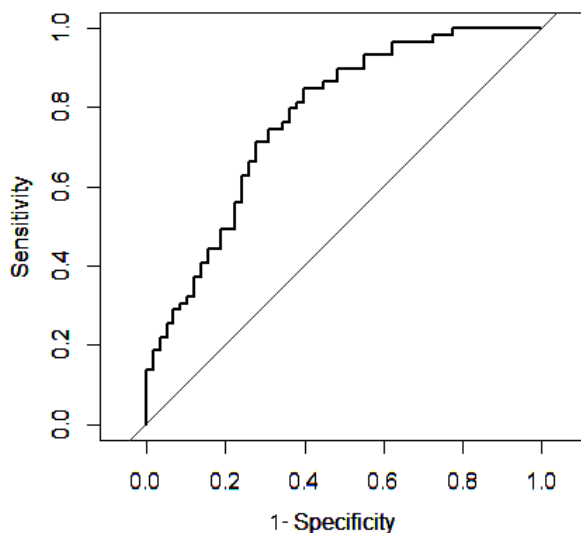


Fig. 5. A plot of sensitivity and specificity for the initial data (ROC-curve)

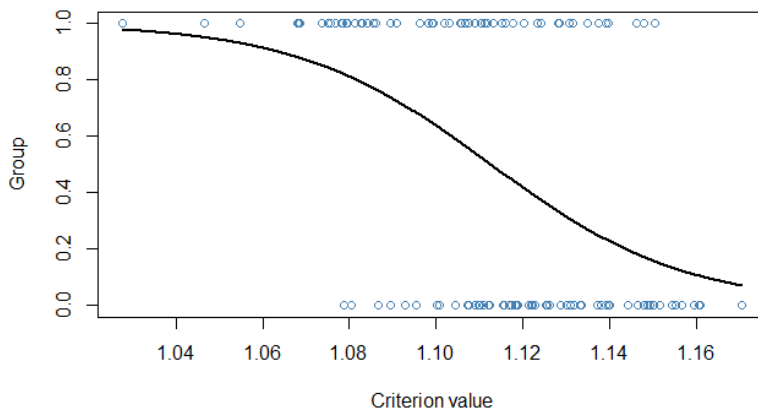


Fig. 6. Logistic function graph

The area under the ROC-curve was 0.772, which testifies to the good quality of the classification. Thus, it is possible to accept the hypothesis

formed for the initial data. An increase in the criterion value by 1 entails an increase in the logarithm of the odds ratio of dying by 44.7.

Now, knowing that the value of the patient's criterion is a statistically significant indicator in determining the patient's group, it is possible to test how well this model is able to classify patients into the groups indicated above (MG and CG). In the role of the metric for assessing the quality of classification, we used sensitivity, that is, the proportion of all correctly classified patients relative to all patients in this group in the sample (Table 6).

Table 6. Results of classification using logistic regression.

Patient group	Number of correctly classified	Number of all patients	Classification sensitivity in percentage
All patients	82	117	75%
MG patients	40	58	78%
CG patients	42	59	73%

3.2. Testing the criterion on new data. The proposed model was tested on new data obtained at the second stage of data collection. The new dataset contained information on 38 patients, 19 of whom belonged to the MG group and 19 to the CG group. Table 7 shows the results of applying the resulting logistic regression model to new data.

Table 7. Testing the model on new data

Patient group	Number of correctly classified	Number of all patients	Classification sensitivity in percentage
All patients	25	38	65%
MG patients	14	19	73%
CG patients	11	19	53%

3.3. Testing the criterion for patients with different causes of a fatal outcome. As a result of the work done, a model was built to classify patients into two groups: MG (with a fatal outcome as a result of deterioration) and CG (with a favorable outcome), based on data from the sequence of ECG tests. In the resulting dataset, deterioration manifested itself as symptoms of the following three conditions: acute heart failure, multiple organ dysfunction syndrome, and cerebral edema. Knowing the patient's risk of fatal outcome for whatever reason could help in treating the patient in a life-threatening condition. Also, deterioration leading to fatal outcome for various reasons can have a different clinical picture. From these points of view, it is of interest to evaluate the ability of criterion to

classify patients who had deterioration leading to the different causes of fatal outcome. Figure 7 shows a scatter diagram of the criterion values for patients with different causes of fatal outcome, Table 8 shows the results of the classification of patients with different causes of fatal outcome on a combined dataset consisting of data obtained at the first and second stages. For better visualization, groups of patients with different causes of death are clearly separated along the OX axis in the diagram.

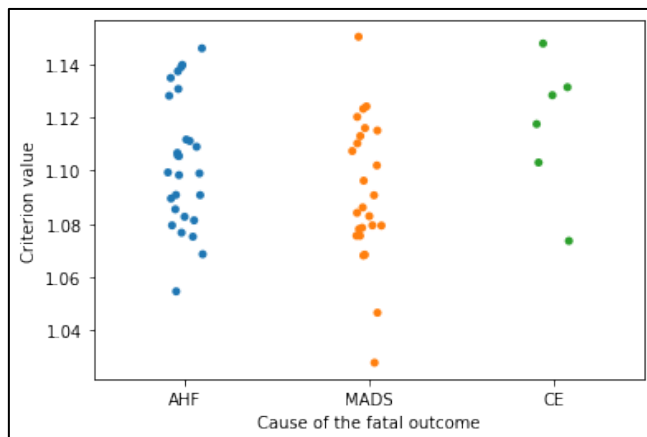


Fig. 7. Scatter diagram of the criterion values for patients with different causes of fatal outcome

Table 8. Testing the model for patients with different causes of fatal outcome

Patient group	Number of correctly classified	Number of all patients	Classification sensitivity in percentage
All patients	107	155	69%
Patients with MODS	22	30	73%
Patients with AHF	30	41	73%
Patients with CE	2	6	33%
CG patients	53	78	68%

Thus, according to the results of testing the model for patients with different causes of fatal outcome, this model best categorizes patients with a life-threatening state of deterioration that subsequently led to multiple organ dysfunction syndrome, and the worst categorizes those whose cause of fatal outcome was cerebral edema.

3.4. Software implementation. The developed model was applied in the implementation of a prototype web service that allows using ECG data to determine whether a patient suffered from a life-threatening deterioration. The web service was developed using PHP for the server-side, HTML, CSS, JavaScript for the client-side and is hosted on the page at <https://ecg.dscs.pro>. To use it, you need to upload a file in excel format with data on the values of the pairs of RR and QT intervals of the patient in the following format, presented in Table 9:

Table 9. File format with data on the values of the patient's RR and QT intervals for uploading to the web service form

RR value	QT value
rr : rr $\in$ 244..3042	qt:qt $\in$ 141..692

Figure 8 shows the main page and interface of the web service prototype, Figure 9 shows the screen with the output of the result.



Fig. 8. Web service prototype interface. Translation: «Analysis of the patient's condition according to ECG data», «Upload a document», «Calculate»

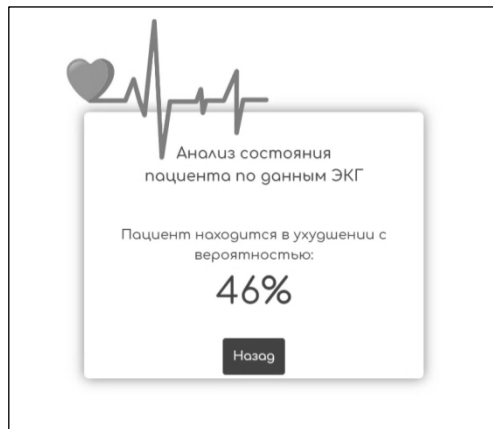


Fig. 9. A screen displaying the result. Translation: «Analysis of the patient's condition according to ECG data», «The patient is in deterioration with a probability of 46%»

4. Discussion. The difference in the distribution of the criterion value in two groups (MG, CG) is statistically significant, so the proposed method for unifying the RR and QT interval values can be used as an independent criterion for classifying patients into MG (with a fatal outcome as a result of worsening with symptoms of MODS, AHF and CE) and CG (with a favorable outcome). This follows from the fact that for the initial data, the *p-value* in the logistic regression was close to 0, and the area under the ROC curve was 0.772. Thus, the hypothesis that there is a predictor of death risk assessment, which can be derived from the previously obtained values of the ECG intervals, which gives a statistically significant difference for the ECG of the two groups of patients: those who suffered deterioration leading to the fatal outcome caused by MODS, AHF or CE, and those with favorable outcome was confirmed. Each ECG analysis results in a set of RR and QT interval values. Thus, based on a sequence of ECG records, the value of the K_i criterion can be calculated to identify whether the patient had a state of deterioration leading to fatal outcome using a logistic regression model. The results of classifying patients into risk groups on the new dataset are worse than the results on the original dataset; nevertheless, the likelihood of correct identification of deterioration in patients is quite high. The model was tested on patients with fatal outcomes for various causes. The model showed itself in the best way when identifying deterioration in patients whose cause of fatal outcome was the

syndrome of multiple organ failure: the sensitivity of determining such patients was 78%.

On the one hand, the approach proposed in the article is limited by the fact that it does not take into account the individual characteristics of the patient, on the other hand, taking into account the high-quality characteristics of the proposed criterion, it is very likely that the proposed indicator sufficiently aggregates information about the state of the patient's cardiovascular system. The main distinctive feature of this study is that it aims to find common features of patients with different causes of fatal outcome and derive the predictor of deterioration from their ECG data. The results obtained show that the proposed estimation method succeeds in finding similar features in ECGs of patients with MODS and AHF, but works poorly for patients with CE. Given the small number of the patients with this condition, we probably will further exclude them from the study.

The methodological choices were constrained by the size of the dataset as well as the characteristics of data received from ARCERM. In the future, the dataset will be expanded, and the number of considered features will also be increased, which will allow the use of a wider toolkit, which may positively affect the results of the study.

5. Conclusion. The logistic regression model used in the article demonstrated the possibility of accepting the hypothesis that the proposed method for unifying the values of the RR and QT intervals for one patient makes a statistically significant feature for determining the patient group (the main one – with a fatal outcome and the control one – with a favorable outcome). The constructed estimation method showed good results on new data as well. Thus, the hypothesis is accepted that there is a predictor of death risk assessment, which can be derived from the previously obtained values of the ECG intervals, which gives a statistically significant difference for the ECG of the two groups of patients: those who suffered deterioration leading to the fatal outcome caused by MODS, AHF or CE, and those with a favorable outcome.

The obtained criterion allows, with a sensitivity of at least 65%, to identify patients who had a state of deterioration as a result of the development of acute heart failure syndrome, multiple organ failure or cerebral edema, based on the results of ECG analysis. The model showed the best sensitivity (83%) when identifying the condition of patients whose cause of death was multiple organ dysfunction syndrome. As a result, the constructed model was used to develop a prototype web service for assessing the patient's condition and automatically determining life-threatening deterioration of condition based on ECG data. The practical significance of the result is based on the fact that such a diagnosis can allow

taking the necessary measures before the moment of critical changes. Patients may show signs of deterioration long before they become unstable. The ability to recognize these signs can lead to a decrease in mortality. The theoretical significance of the result lies in the construction of a mathematical model for classifying patients into those who had a favorable outcome, and those who have a fatal outcome as a result of the development of MODS, AHF or CE, as well as the basic formation for the development of a methodology for predicting the development of life-threatening conditions. Several new questions emerge in light of the discoveries presented here. One of the possible further research directions lies in the inclusion of the obtained results into a more complex model that takes into account other information about the patient, such as gender, age, the presence of a change in rhythms, etc. Bayesian trust networks or algebraic Bayesian networks can serve as such a model [26–29]. The advantage of the latter is that they can work with incomplete and inaccurate data, which can be an important advantage in the further work with medical data [28, 30, 31]. This will allow more information about the patient to be used, which can increase the chance of correct classification of the patient's risk group. The second direction can be the study of the dynamics of changes in ECG, in other words, the study of changes in RR and QT interval values over time, and building a model based on time series, which describes the dynamics of changes in the studied interval values. It will also provide new information about the patient that can help predict disease outcomes. The third possible direction is the construction of a model for classifying deterioration in terms of probable causes of fatal outcome, that is, a model that allows predicting the risk of the development of acute heart failure, multiple organ dysfunction syndrome, or cerebral edema.

References

1. [Preventive cardiology. Under. ed. G.M. Kositsky.]. Meditsina [Medicine]; 1987. 512 p. (InRuss.).
2. Mjasnikov A.L. Jeksperimental'nye nekrozy miokarda [Experimental myocardial necrosis]. M. Medicina. 1963. 204 p. (InRuss.).
3. Brodsky M., Wu D., Penes P., Kanakis Ch., Rosen K. Arrhythmias documented by 24 hour continuous electrocardiographic monitoring in 50 male medical students without apparent heart diseases. The American journal of cardiology. 1977. vol. 39. No. 3. pp. 390–395. doi: 10.1097/00132586-197710000-00002
4. Umetini K., Singer D., McCarty R., Atkinson M. 24 Hour time domain heart rate variability and heart rate: relations to age and gender over nine decades. Journal of the American College of Cardiology. 1998. vol. 31. No. 3. pp. 593–601. doi: 10.1016/s0735-1097(97)00554-8
5. Makarov L.M. Holterovskoe monitorirovanie. [Holter monitoring.] M: ID "Medpraktika – M". 2008. 504 p. (InRuss.).
6. Kleiger R.E., Miller J.P., Bigger Jr J.T., Moss A.J. Decreased heart rate variability and its association with increased mortality after acute myocardial infarction. The

- American journal of cardiology. 1987. vol. 59. №. 4. pp. 256–262. doi: 10.1016/0002-9149(87)90795-8
7. Anischenko V.S. Znakomstvo s nelinejnoj dinamikoj [Acquaintance with nonlinear dynamics.]. Izdatelstvo LKI. 2018. 224 p. (InRuss.).
8. Rathakrishnan K., Min S.N., Park S.J. Evaluation of ECG Features for the Classification of Post-Stroke Survivors with a Diagnostic Approach. *Applied Sciences*. 2021. vol. 11. No. 1: 192.
9. Wang X., Gong G., Li N., Ding L., Ma Y. Decoding pilot behavior consciousness of EEG, ECG, eye movements via an SVM machine learning model. *International Journal of Modeling, Simulation, and Scientific Computing*. 2020. vol. 11. №. 04: 2050028.
10. Antosik-Wojcinska A.Z., Dominiak M., Chojnacka M. Kaczmarek-Majer K., Opara K.R., Radziszewska. W., Olwert A., Swiecicki L. Smartphone as a monitoring tool for bipolar disorder: a systematic review including data analysis, machine learning algorithms and predictive modelling. *International Journal of Medical Informatics*. 2020. vol. 138: 104131.
11. Srivastava A.K., Kumar Y., Singh P.K. A Rule-Based Monitoring System for Accurate Prediction of Diabetes: Monitoring System for Diabetes. *International Journal of E-Health and Medical Communications (JEHMC)*. 2020. vol. 11. №. 3. pp. 32–53.
12. Aziz, S., Ahmed, S. & Alouini, M.S. ECG-based machine-learning algorithms for heartbeat classification. *Sci Rep* 11. 2021. vol. 18738. <https://doi.org/10.1038/s41598-021-97118-5>.
13. Tougui I., Jilbab A., El Mhamdi J. Heart disease classification using data mining tools and machine learning techniques. *Health and Technology*. 2020. vol. 10. pp. 1137–1144.
14. Park J., Kim J.W., Ryu B., Heo E., Jung S.Y., Yoo S. Patient-Level Prediction of Cardio-Cerebrovascular Events in Hypertension Using Nationwide Claims Data. *Journal of medical Internet research*. 2019. vol. 21. №. 2: e11757.
15. Guo A., Zhang K.W., Reynolds K., Foraker R.E. Coronary heart disease and mortality following a breast cancer diagnosis. *BMC medical informatics and decision making*. 2020. vol. 20: 88.
16. Kovalev M.S., Utkin L.V., Kasimov E.M. SurvLIME: A method for explaining machine learning survival models. *Knowledge-Based Systems*. 2020. vol. 203. pp. 106–164.
17. Sohal H., Jain S. Interpretation of cardio vascular diseases using electrocardiogram: A study. 2018 Fifth International Conference on Parallel, Distributed and Grid Computing (PDGC). IEEE, 2018. pp. 159–164.
18. Yao L.K., Liu C.C., Li P., Wang J.K., Liu Y.Y., Li W., Wang X.P., Li H., Zhang H. Enhanced Automated Diagnosis of Coronary Artery Disease Using Features Extracted From QT Interval Time Series and ST–T Waveform. *IEEE Access*. 2020. vol. 8. pp. 129510–129524.
19. George K., Sankaran P., Joseph K.P. Computer assisted recognition of breast cancer in biopsy images via fusion of nucleus-guided deep convolutional features. *Computer Methods and Programs in Biomedicine*. 2020. vol. 194: 105531.
20. Hyun S., Kaewprag P., Cooper C., Hixon B., Moffatt-Bruce S. Exploration of Critical Care Data by Using Unsupervised Machine Learning. *Computer Methods and Programs in Biomedicine*. 2020. vol. 194: 105507.
21. Hang W.L., Feng W., Liang S., Wang Q., Liu X.J., Choi K.S. Deep Stacked Support Matrix Machine Based Representation Learning for Motor Imagery EEG Classification. *Computer Methods and Programs in Biomedicine*. 2020. vol. 193: 105466.
22. Louridi N., Amar M., El Ouahidi B. Identification of Cardiovascular Diseases Using Machine Learning. *Proceedings of the 7th Mediterranean Congress of*

- Telecommunications 2019, CMT 2019, Fez, Morocco, 24–25 October 2019. 2019. pp. 1–6
23. Villamor E., Monserrat C., Del Rio L., Romero-Martin J.A., Ruperez M.J. Prediction Of Osteoporotic Hip Fracture In Postmenopausal Women Through Patient-Specific FE Analyses And Machine Learning. Computer Methods and Programs in Biomedicine. 2020. vol. 193: 105484.
 24. Perumal R., Kaladevi A.C. Early Prediction of Coronary Heart Disease from Cleveland Dataset using Machine Learning Techniques. Int. J. Adv. Sci. Technol. 2020. vol. 29. pp. 4225–4234.
 25. Lepeshkin E., Surawicz B. The measurement of the QT interval of the electrocardiogram. Circulation, 1952;6:378-388 <https://doi.org/10.1161/01.cir.6.3.378>
 26. Rymarczyk T., Kozłowski E., Kłosowski G., Niderla K. Logistic Regression for Machine Learning in Process Tomography. Sensors. 2019. vol. 19. No. 15: 3400. doi: 10.3390/s19153400
 27. Bishop C. Pattern Recognition and Machine Learning. Springer-Verlag New York, 2006. 738 p.
 28. Delen D., Topuz K., Eryarsoy E. Development of a Bayesian Belief Network-based DSS for predicting and understanding freshmen student attrition. Operations research and management sciences. 2020. vol. 281. pp. 575–587.
 29. Kharitonov N.A., Maximov A.G., Tulupyev A.L. Algebraic Bayesian networks: the use of parallel computing while maintaining various degrees of consistency. International Conference on Information Technologies. Springer, Cham. 2019. pp. 446–456.
 30. Nielsen T.D., Jensen F.V. Bayesian networks and decision graphs. Springer Science Business Media. 2009. 464 p.
 31. Tulupyev A.L., Stolyarov D.M., Mentyukov M.V. Representation of the local and global structure of an algebraic Bayesian network in Java applications. Proceedings of SPIIRAS. 2017. vol. 5. No. 0. pp. 71–99.
 32. Kharitonov N., Malchevskaja E., Zolotin A., Abramov M. External consistency maintenance algorithm for chain and stellate structures of algebraic bayesian networks: statistical experiments for running time analysis. International Conference on Intelligent Information Technologies for Industry. Springer, Cham. 2018. pp. 23–30.
 33. Baevskij R.M., Ivanov G.G., Gavrilushkin A.P., Dovgalevskij P.Ja., Kukushkin Ju.A., Mironova T.F., Priluckij D.A., Semenov A.V., Fedorov V.F., Flejshman A.N., Medvedev M.M., Chirejkin L.V. [Analysis of Heart Rate Variability When Using Different Electrocardiographic Systems (Part 1)]. Vestnik aritmologii [Bulletin of Arithmology]. 2002. No. 24. pp 65–86.
 34. Ponomarenko V.I., Karavaev A.S., Borovkova E.I., Hramkov A.N., Kiselev, A.R., Prokhorov M.D., Penzel T. Decrease of coherence between the respiration and parasympathetic control of the heart rate with aging. Chaos: An Interdisciplinary Journal of Nonlinear Science. 2021. vol. 31. no. 7: 073105.
 35. Karavaev A.S., Ishbulatov Yu Y.M., Prokhorov M.D., Ponomarenko V.I., Kiselev A.R., Runnova A.E., Hramkov A.N., Semyachkina-Glushkovskaya O.V., Kurths J., Penzel T. Simulating dynamics of circulation in the awake state and different stages of sleep using non-autonomous mathematical model with time delay. Frontiers in Physiology. 2020. vol. 11: 1656.
 36. Karavaev A.S., Borovkova E.I., Runnova A.E., Kiselev A.R., Zhuravlev M.O., Ponomarenko V.I., Prokhorov M.D., Koronovskii A.A., Hramov A.E. Experimental Observation of the Self-Oscillatory Dynamics of the Regulation Contours of the Cardiovascular System. Radiophysics and Quantum Electronics. 2019. vol. 61. no. 8. pp. 681-688.

Abramov Maxim — Ph.D., Head of laboratory, senior researcher, Laboratory of theoretical and interdisciplinary computer science, St Petersburg Federal Research Center of the Russian Academy of Sciences; Assistant professor, Computer science department, faculty of mathematics and mechanics, SPbSU. Research interests: information security, social engineering attacks, analysis of the security of users of information systems from social engineering attacks of intruders; analysis of the dissemination of information in social networks based on the models used in the analysis of the security of users of information systems from social engineering attacks, analysis and modeling of social networks, client-server technologies, researching the relationship between content posted by users on social networks and offline behavior, business analytics, social computing, business intelligence. The number of publications — 100. mva@dscs.pro; 39, 14-th Line V.O., 199178, St. Petersburg, Russia; office phone: +7(812)328-3337.

Tsukanova Ekaterina — Chief specialist, The Federal State Budgetary Institute "The Nikiforov Russian Center of Emergency and Radiation Medicine" The Ministry of Russian Federation for Civil Defense, Emergencies and Elimination of Consequences of Natural Disasters (EMERCOM of Russia). Research interests: statistical research methods in medicine, integrative approach in functional diagnostics, decision support systems for early detection, prevention, selection and assessment of the effectiveness of therapy for life-threatening arrhythmias, rational organization of the diagnostic service in a modern hospital. The number of publications — 8. Eka_77@bk.ru; 4/2, Academician Lebedev St., 194044, St. Petersburg, Russia; office phone: +7(812)607-5900.

Tulupyev Alexander — Ph.D., Dr.Sci., Professor, Chief researcher, Laboratory of theoretical and interdisciplinary computer science, St Petersburg Federal Research Center of the Russian Academy of Sciences; Professor, Computer science department, faculty of mathematics and mechanics, SPbSU. Research interests: uncertain knowledge and data representation and processing, application of mathematics and computer science in sociocultural and epidemiological studies, software technologies and development of information systems with databases. The number of publications — 400. alt@dscs.pro; 39, 14-th Line V.O., 199178, St. Petersburg, Russia; office phone: +7(812)328-3337.

Korepanova Anastasia — Junior researcher, laboratory of theoretical and interdisciplinary computer science, St Petersburg Federal Research Center of the Russian Academy of Sciences. Research interests: information security, social engineering attacks, analysis of the security of users of information systems from social engineering attacks of intruders, analysis and modeling of social networks, machine learning, data analysis. The number of publications — 13. aak@dscs.pro; 39, 14-th Line V.O., 199178, St. Petersburg, Russia; office phone: +7(812)328-3337.

Aleksanin Sergei — Ph.D., Dr.Sci., Professor, Corresponding member, Director, The Federal State Budgetary Institute "The Nikiforov Russian Center of Emergency and Radiation Medicine" The Ministry of Russian Federation for Civil Defense, Emergencies and Elimination of Consequences of Natural Disasters (EMERCOM of Russia). Research interests: disaster medicine, problems of a person's functional state in extreme and unfavorable environmental conditions, occupational pathology, radiation medicine, preventive and restorative medicine. The number of publications — 356. medicine@nrcrm.ru; 4/2, Academician Lebedev St., 194044, St. Petersburg, Russia; office phone: +7(812)607-5900.

Acknowledgements. The work was carried out within the framework of the project under the state order of the St. Petersburg Federal Research Center of the Russian Academy of Sciences SPIIRAN No. FFZF-2022-0003.

М.В. АБРАМОВ, Е.И. ЦУКАНОВА, А.Л. ТУЛУПЬЕВ, А.А. КОРЕПАНОВА,
С.С. АЛЕКСАНИН

ИДЕНТИФИКАЦИЯ КЛИНИЧЕСКОГО УХУДШЕНИЯ В РЕЗУЛЬТАТЕ РАЗВИТИЯ ОСН, СПОН ИЛИ ОГМ ПОСРЕДСТВОМ КЛАССИФИКАЦИИ НА ОСНОВЕ ДАННЫХ ОБ ИНТЕРВАЛАХ RR И QT

Абрамов М.В., Цуканова Е.И., Тулупьев А.Л., Корепанова А.А., Алексанин С.С.
Идентификация клинического ухудшения в результате развития ОСН, СПОН или ОГМ посредством классификации на основе данных об интервалах RR и QT.

Аннотация. Резкое ухудшение состояния на фоне развития жизнеугрожающих аритмий с симптомами острой сердечной недостаточности (ОСН), синдрома полиорганной недостаточности (СПОН) или отёка головного мозга (ОГМ) может привести к гибели пациента. Поскольку известные методы автоматизированной диагностики в настоящий момент не могут достаточно точно и своевременно определить, что пациент находится в жизнеугрожающем состоянии, ведущем к летальному исходу от ОСН, СПОН или ОГМ, существует необходимость в разработке соответствующих методов. Одним из способов выявить предикторы такого состояния является применение методов машинного обучения к накопленным наборам данных. В данной статье решалась задача проверки с помощью методов анализа данных гипотезы о наличии зависимости между результатами измерения ЭКГ и последующим летальным исходом пациента в результате развития СПОН, ОСН или ОГМ. Был предложен метод комбинирования данных, сводящейся к тому, чтобы на основе характеристик ЭКГ для каждого пациента предложить алгоритм, на вход которого подаются пары интервалов RR и QT, а на выходе получается число, которое является характеристикой состояния пациента. На основе полученной характеристики производится классификация пациентов на группы: основную (пациенты с летальным исходом) и контрольную (выжившие пациенты). Полученная модель классификации закладывает потенциал для разработки методов идентификации клинического состояния пациента, что позволит автоматизировать получение сигнала о его ухудшении. Новизна результата заключается в подтверждении гипотезы о наличии зависимости между результатами измерения ЭКГ и последующим летальным исходом пациента в результате развития СПОН, ОСН или ОГМ, а также предложенном критерии и модели классификации, которые позволяют решать актуальную задачу автоматической фиксации ухудшения состояния пациентов.

Ключевые слова: классификация пациентов на основе ЭКГ, идентификация ухудшения клинического состояния, прогнозирование по медицинским данным, анализ ЭКГ, машинное обучение, искусственный интеллект, наука о данных, логистическая регрессия, стратификация риска смертности, интервал RR, интервал QT.

Абрамов Максим Викторович — канд. техн. наук, руководитель лаборатории, старший научный сотрудник, лаборатория теоретических и междисциплинарных проблем информатики, Федеральное государственное бюджетное учреждение науки "Санкт-Петербургский Федеральный исследовательский центр Российской академии наук"; доцент, кафедра компьютерных наук, математико-механический факультет, СПбГУ. Область научных интересов: информационная безопасность, соционженерные атаки, анализ защищенности пользователей информационных систем от соционженерных атак злоумышленников, анализ распространения информации в социальных сетях на основе моделей, используемых при анализе защищенности

пользователей информационных систем от атак социальной инженерии, анализ и моделирование социальных сетей, клиент-серверные технологии, исследование взаимосвязи между контентом, размещаемым пользователями в социальных сетях, и поведением в оффлайне; бизнес-аналитика, социальные вычисления, бизнес-аналитика. Число научных публикаций — 100. mva@dscs.pro; 14-я линия В.О., 39, 199178, Санкт-Петербург, Россия; р.т.: +7(812)328-3337.

Цуканова Екатерина Игоревна — главный специалист, ФГБУ Всероссийский центр экстренной и радиационной медицины имени А.М. Никифорова МЧС России. Область научных интересов: статистические методы исследования в медицине, интегративный подход в функциональной диагностике, системы поддержки принятия решений для раннего выявления, профилактики, выбора и оценки эффективности терапии жизнеугрожающих аритмий, рациональная организация диагностической службы в современном стационаре. Число научных публикаций — 8. Eka_77@bk.ru; Академика Лебедева, 4/2, 194044, Санкт-Петербург, Россия; р.т.: +7(812)607-5900.

Тулупьев Александр Львович — д-р физ.-мат. наук, профессор, главный научный сотрудник, лаборатория теоретических и междисциплинарных проблем информатики, Федеральное государственное бюджетное учреждение науки "Санкт-Петербургский Федеральный исследовательский центр Российской академии наук"; профессор, кафедра компьютерных наук, математико-механический факультет, СПбГУ. Область научных интересов: представление и обработка неопределенных знаний и данных, применение математики и информатики в социокультурных и эпидемиологических исследованиях, программные технологии и разработка информационных систем с базами данных. Число научных публикаций — 400. alt@dscs.pro; 14-я линия В.О., 39, 199178, Санкт-Петербург, Россия; р.т.: +7(812)328-3337.

Корепанова Анастасия Андреевна — младший научный сотрудник, лаборатория теоретических и междисциплинарных проблем информатики, Федеральное государственное бюджетное учреждение науки "Санкт-Петербургский Федеральный исследовательский центр Российской академии наук". Область научных интересов: информационная безопасность, социоинженерные атаки, анализ защищенности пользователей информационных систем от социоинженерных атак злоумышленников, анализ и моделирование социальных сетей, машинное обучение, анализ данных. Число научных публикаций — 13. aak@dscs.pro; 14-я линия В.О., 39, 199178, Санкт-Петербург, Россия; р.т.: +7(812)328-3337.

Алексанин Сергей Сергеевич — д-р мед. наук, профессор, член-корреспондент, директор, ФГБУ Всероссийский центр экстренной и радиационной медицины имени А.М. Никифорова МЧС России. Область научных интересов: медицина катастроф, проблемы функционального состояния человека в экстремальных и неблагоприятных условиях окружающей среды, профессиональная патология, радиационная медицина, профилактическая и восстановительная медицина. Число научных публикаций — 356. medicine@nrccrm.ru; Академика Лебедева, 4/2, 194044, Санкт-Петербург, Россия; р.т.: +7(812)607-5900.

Поддержка исследований. Работа выполнена в рамках проекта по государственному заказу Санкт-Петербургского Федерального исследовательского центра РАН (СПб ФИЦ РАН) - СПИИРАН № FFZF-2022-0003.

Литература

1. [Preventive cardiology. Under. ed. G.M. Kositsky.]. Meditsina [Medicine]; 1987. 512 p. (InRuss.).
2. Mjasnikov A.L. Jeksperimental'nye nekrozy miokarda [Experimental myocardial necrosis]. M. Medicina. 1963. 204 p. (InRuss.).
3. Brodsky M., Wu D., Penes P., Kanakis Ch., Rosen K. Arrhythmias documented by 24 hour continius electrocardiographic monitoring in 50 male medical students without apparent heart diseases. The American journal of cardiology. 1977. vol. 39. No. 3. pp. 390–395. doi: 10.1097/00132586-197710000-00002
4. Umetini K., Singer D., McCarty R., Atkinson M. 24 Hour time domain heart rate variability and heart rate: relations to age and gender over nine decades. Journal of the American College of Cardiology. 1998. vol. 31. No. 3. pp. 593–601. doi: 10.1016/s0735-1097(97)00554-8
5. Makarov L.M. Holterovskoe monitorirovanie. [Holter monitoring.] M: ID "Medpraktika – M". 2008. 504 p. (InRuss.).
6. Kleiger R.E., Miller J.P., Bigger Jr J.T., Moss A.J. Decreased heart rate variability and its association with increased mortality after acute myocardial infarction. The American journal of cardiology. 1987. vol. 59. №. 4. pp. 256–262. doi: 10.1016/0002-9149(87)90795-8
7. Anischenko V.S. Znakomstvo s nelinejnoj dinamikoj [Acquaintance with nonlinear dynamics.]. Izdatelstvo LKI. 2018. 224 p. (InRuss.).
8. Rathakrishnan K., Min S.N., Park S.J. Evaluation of ECG Features for the Classification of Post-Stroke Survivors with a Diagnostic Approach. Applied Sciences. 2021. vol. 11. No. 1: 192.
9. Wang X., Gong G., Li N., Ding L., Ma Y. Decoding pilot behavior consciousness of EEG, ECG, eye movements via an SVM machine learning model. International Journal of Modeling, Simulation, and Scientific Computing. 2020. vol. 11. №. 04: 2050028.
10. Antosik-Wojcinska A.Z., Dominiak M., Chojnacka M. Kaczmarek-Majer K., Opara K.R., Radziszewska. W., Olwert A., Swiecicki L. Smartphone as a monitoring tool for bipolar disorder: a systematic review including data analysis, machine learning algorithms and predictive modelling. International Journal of Medical Informatics. 2020. vol. 138: 104131.
11. Srivastava A.K., Kumar Y., Singh P.K. A Rule-Based Monitoring System for Accurate Prediction of Diabetes: Monitoring System for Diabetes. International Journal of E-Health and Medical Communications (IJEHMC). 2020. vol. 11. №. 3. pp. 32–53.
12. Aziz, S., Ahmed, S. & Alouini, M.S. ECG-based machine-learning algorithms for heartbeat classification. Sci Rep 11. 2021. vol. 18738. <https://doi.org/10.1038/s41598-021-97118-5>.
13. Tougui I., Jilbab A., El Mhamdi J. Heart disease classification using data mining tools and machine learning techniques. Health and Technology. 2020. vol. 10. pp. 1137–1144.
14. Park J., Kim J.W., Ryu B., Heo E., Jung S.Y., Yoo S. Patient-Level Prediction of Cardio-Cerebrovascular Events in Hypertension Using Nationwide Claims Data. Journal of medical Internet research. 2019. vol. 21. №. 2: e11757.
15. Guo A., Zhang K.W., Reynolds K., Foraker R.E. Coronary heart disease and mortality following a breast cancer diagnosis. BMC medical informatics and decision making. 2020. vol. 20: 88.
16. Kovalev M.S., Utkin L.V., Kasimov E.M. SurvLIME: A method for explaining machine learning survival models. Knowledge-Based Systems. 2020. vol. 203. pp. 106–164.

17. Sohal H., Jain S. Interpretation of cardio vascular diseases using electrocardiogram: A study. 2018 Fifth International Conference on Parallel, Distributed and Grid Computing (PDGC). IEEE, 2018. pp. 159–164.
18. Yao L.K., Liu C.C., Li P., Wang J.K., Liu Y.Y., Li W., Wang X.P., Li H., Zhang H. Enhanced Automated Diagnosis of Coronary Artery Disease Using Features Extracted From QT Interval Time Series and ST–T Waveform. IEEE Access. 2020. vol. 8. pp. 129510–129524.
19. George K., Sankaran P., Joseph K.P. Computer assisted recognition of breast cancer in biopsy images via fusion of nucleus-guided deep convolutional features. Computer Methods and Programs in Biomedicine. 2020. vol. 194: 105531.
20. Hyun S., Kaewprag P., Cooper C., Hixon B., Moffatt-Bruce S. Exploration of Critical Care Data by Using Unsupervised Machine Learning. Computer Methods and Programs in Biomedicine. 2020. vol. 194: 105507.
21. Hang W.L., Feng W., Liang S., Wang Q., Liu X.J., Choi K.S. Deep Stacked Support Matrix Machine Based Representation Learning for Motor Imagery EEG Classification. Computer Methods and Programs in Biomedicine. 2020. vol. 193: 105466.
22. Louridi N., Amar M., El Ouahidi B. Identification of Cardiovascular Diseases Using Machine Learning. Proceedings of the 7th Mediterranean Congress of Telecommunications 2019, CMT 2019, Fez, Morocco, 24–25 October 2019. 2019. pp. 1–6
23. Villamor E., Monserrat C., Del Rio L., Romero-Martin J.A., Ruperez M.J. Prediction Of Osteoporotic Hip Fracture In Postmenopausal Women Through Patient-Specific FE Analyses And Machine Learning. Computer Methods and Programs in Biomedicine. 2020. vol. 193: 105484.
24. Perumal R., Kaladevi A.C. Early Prediction of Coronary Heart Disease from Cleveland Dataset using Machine Learning Techniques. Int. J. Adv. Sci. Technol. 2020. vol. 29. pp. 4225–4234.
25. Lepeshkin E., Surawicz B. The measurement of the QT interval of the electrocardiogram. Circulation, 1952;6:378-388 <https://doi.org/10.1161/01.cir.6.3.378>
26. Rymarczyk T., Kozłowski E., Kłosowski G., Niderla K. Logistic Regression for Machine Learning in Process Tomography. Sensors. 2019. vol. 19. No. 15: 3400. doi: 10.3390/s19153400
27. Bishop C. Pattern Recognition and Machine Learning. Springer-Verlag New York, 2006. 738 p.
28. Delen D., Topuz K., Eryarsoy E. Development of a Bayesian Belief Network-based DSS for predicting and understanding freshmen student attrition. Operations research and management sciences. 2020. vol. 281. pp. 575–587.
29. Kharitonov N.A., Maximov A.G., Tulupyev A.L. Algebraic Bayesian networks: the use of parallel computing while maintaining various degrees of consistency. International Conference on Information Technologies. Springer, Cham. 2019. pp. 446–456.
30. Nielsen T.D., Jensen F.V. Bayesian networks and decision graphs. Springer Science Business Media. 2009. 464 p.
31. Tulupyev A.L., Stolyarov D.M., Mentyukov M.V. Representation of the local and global structure of an algebraic Bayesian network in Java applications. Proceedings of SPIRAS. 2017. vol. 5. No. 0. pp. 71–99.
32. Kharitonov N., Malchevskaya E., Zolotin A., Abramov M. External consistency maintenance algorithm for chain and stellate structures of algebraic bayesian networks: statistical experiments for running time analysis. International Conference on Intelligent Information Technologies for Industry. Springer, Cham. 2018. pp. 23–30.

33. Baevskij R.M., Ivanov G.G., Gavrilushkin A.P., Dovgalevskij P.Ja., Kukushkin Ju.A., Mironova T.F., Priluckij D.A., Semenov A.V., Fedorov V.F., Flejshman A.N., Medvedev M.M., Chirejkin L.V. [Analysis of Heart Rate Variability When Using Different Electrocardiographic Systems (Part 1)]. *Vestnik aritmologii* [Bulletin of Arithmology]. 2002. No. 24. pp 65–86.
34. Ponomarenko V.I., Karavaev A.S., Borovkova E.I., Hramkov A.N., Kiselev, A.R., Prokhorov M.D., Penzel T. Decrease of coherence between the respiration and parasympathetic control of the heart rate with aging. *Chaos: An Interdisciplinary Journal of Nonlinear Science*. 2021. vol. 31. no. 7: 073105.
35. Karavaev A.S., Ishbulatov Yu Y.M., Prokhorov M.D., Ponomarenko V.I., Kiselev A.R., Runnova A.E., Hramkov A.N., Semyachkina-Glushkovskaya O.V., Kurths J., Penzel T. Simulating dynamics of circulation in the awake state and different stages of sleep using non-autonomous mathematical model with time delay. *Frontiers in Physiology*. 2020. vol. 11: 1656.
36. Karavaev A.S., Borovkova E.I., Runnova A.E., Kiselev A.R., Zhuravlev M.O., Ponomarenko V.I., Prokhorov M.D., Koronovskii A.A., Hramov A.E. Experimental Observation of the Self-Oscillatory Dynamics of the Regulation Contours of the Cardiovascular System. *Radiophysics and Quantum Electronics*. 2019. vol. 61. no. 8. pp. 681-688.

Г.И. АЛГАЗИН, Д.Г. АЛГАЗИНА
**МОДЕЛИРОВАНИЕ ДИНАМИКИ КОЛЛЕКТИВНОГО
ПОВЕДЕНИЯ В РЕФЛЕКСИВНОЙ ИГРЕ С ПРОИЗВОЛЬНЫМ
ЧИСЛОМ ЛИДЕРОВ**

Алгазин Г.И., Алгазина Д.Г. Моделирование динамики коллективного поведения в рефлексивной игре с произвольным числом лидеров.

Аннотация. Рассматривается олигополия с произвольным числом лидеров по Штакельбергу в условиях неполной, асимметричной информированности агентов и неадекватности предсказаний ими действий конкурентов. Исследуются модели процессов принятия агентами индивидуальных решений. Теоретической основой для построения и аналитического исследования моделей процессов являются теория рефлексивных игр и теория коллективного поведения. Они дополняют друг друга тем, что рефлексивные игры позволяют использовать процедуры коллективного поведения и результаты размышлений агентов, приводящие к равновесию Нэша. Динамический процесс принятия решений рассматривается как повторяемые статические игры на диапазоне допустимых ответов агентов на ожидаемые действия окружения с учетом в каждой игре реальных экономических ограничений и конкурентоспособности. Каждый рефлексиирующий агент в каждой игре рассчитывает свое текущее положение цели и изменяет свое состояние, делая шаги в направлении текущего положения цели так, чтобы получить положительную собственную прибыль или минимизировать потери. Основным результатом работы являются достаточные условия сходимости процессов в дискретном времени для случая линейных издержек агентов и линейного спроса. Получены новые аналитические выражения для диапазонов величин текущих шагов агентов, при которых гарантируется сходимость моделей коллективного поведения к статичному равновесию Нэша. Что позволяет каждому агенту максимизировать собственную прибыль, предполагая полное (совершенное) знание среди агентов. Анализируются также процессы, когда агент выбирает свой наилучший ответ. Последние могут не давать сходящиеся траектории. Подробно обсуждается случай дуополии в сравнении с современными результатами. Приведены необходимые математические леммы, утверждения и их доказательства.

Ключевые слова: рефлексивные игры, олигополия, лидер по Штакельбергу, неполная информированность, коллективное поведение, равновесие, условия сходимости.

1. Введение. Представляет повышенный интерес поведение неоднородных групп в типичных повторяющихся ситуациях, когда каждый из агентов, независимо от других в рамках собственной информированности выбирает свою стратегию, а функции выигрыша и равновесие характеризуют успех стратегий [1–7].

Конкурентный олигопольный рынок дает широкий спектр таких ситуаций. Чтобы успешно конкурировать на рынке, фирмы-агенты должны уметь прогнозировать действия своих конкурентов. Поэтому математические исследования, направленные на повышение адекватности таких прогнозов, являются актуальными для современных рын-

ков. В соответствующих теоретико-игровых моделях используются различные предположения о взаимной информированности и лидерстве агентов в принятии решений.

Впервые аналитический подход к исследованию взаимодействия фирм (агентов) на конкурентном рынке предложен Курно [8]. Он полагал, что на рынке дуополии с целью максимизации собственной прибыли каждой фирме следует устанавливать объем выпуска, считая неизменными объемы выпуска конкурентов, т. е. другим фирмам не выгодно отклониться от равновесия для получения «мгновенной прибыли». Фирмы рациональны в том смысле, что используют полную информацию о конъюнктуре рынка и не только знают все условия окружения, но знают, что все фирмы об этом знают и используют это для максимизации собственной прибыли. Идеи, заложенные в основание подхода Курно, определили направления дальнейших исследований. Так в «классическом» определении равновесия Нэша [9] используется концепция общего (а также полного или совершенного) знания, полагающая, что вся существенная информация и принципы принятия решений агентами всем им известны, всем известно, что всем это известно и т. д. до бесконечности [3, 10]. Так, если теоретико-игровая модель формализована в виде игры в нормальной форме, то агенты всегда выберут равновесные по Нэшу стратегии. В современных терминах равновесие Курно, это – статическое равновесие при полной информации или некооперативное (некоалиционное) равновесие Курно-Нэша в статической игре.

Вместе с тем, многочисленные исследования рынков свидетельствуют, что условие о наличии общего знания, как правило, невыполнимо. В конкурентной среде агенты часто не заинтересованы раскрывать другим агентам свою информацию, которая является существенной для принятия ими адекватных решений. Так же, как показано в экспериментах [11–14], достижению равновесия, предсказанного теорией, могут препятствовать такие факторы: ограниченность когнитивных возможностей агентов, необходимость уверенности каждого агента в том, что все остальные могут вычислить равновесие Нэша и делают это, неполная информированность, наличие нескольких равновесий.

Отказ от предположения о наличии среди агентов общего знания приводит к тому, что каждый агент в рамках своей информированности следует некоторой повторяемой процедуре принятия индивидуальных решений. Рациональность поведения агента заключается в желании максимизировать свою целевую функцию. Однако его наилучшее действие (решение) зависит в общем случае от того, какое

действие выберет любой другой агент, что трудно однозначно знать априори, и поэтому он вынужден предсказывать поведение конкурентов и выбирать свои действия уже с учетом своего прогноза. При этом равновесие Нэша «превращается в более общее информационное равновесие Нэша, в рамках которого каждый агент осуществляет информационную рефлексию – при принятии решений использует не только свою информацию о существенных параметрах, но и свои представления о представлениях других агентов об этих параметрах, представления о представлениях о представлениях и т. д. [3]».

Имеется значительное число работ, в которых в дополнение к фирмам, действующим по Курно, вводится фирма, действующая по особым правилам. Особенность заключается в том, что она устанавливает свой уровень производства, максимизируя собственную прибыль при явном учете реакции остальных фирм на изменения ее объема выпуска. Остальные же фирмы, как и раньше, максимизируют свою прибыль, используя предположение Курно о неизменности производства других фирм. Эту фирму называют лидером или фирмой, действующей по Штакельбергу, так как он первым рассмотрел такую стратегию поведения [15]. Потенциально, лидер имеет возможность получить большую прибыль и поэтому на конкурентном олигопольном рынке каждый из рациональных агентов, ради увеличения собственной прибыли, стремится стать лидером. Реализуя свои лидерские амбиции, они вносят на рынок неоднородность.

В своем подавляющем большинстве проводимые исследования предполагают только одного лидера на рынке, но интерес представляют также рефлексивной игры, когда реализованы лидерские амбиции нескольких агентов [7, 16–23]. Изучается лидерство как на одном [13, 17, 18, 24, 25], так и на разных уровнях [7, 10, 20–23]. Сетевые модели рынка с неоднородным по составу и нефиксированными ролями участников при полной информации рассматриваются в монографии [26]. В этих и многих других современных исследованиях рынка растет понимание того, что в условиях асимметричной информированности агентов и неопределенности выбора конкурентов равновесие достигается не в результате однократного выбора агентами своих действий, а как исход итерационного процесса рефлексивного принятия решений.

В настоящей статье рассматривается проблема достижения равновесия на рынке олигополии на основе математического моделирования процессов принятия рациональных решений в условиях неполной, асимметричной информированности агентов и неадекватности предсказаний ими действий конкурентов. Работы в этом направлении яв-

ляются актуальными ввиду значимости понимания процессов принятия решений, происходящих на реальных современных рынках, и сближения с ними теоретических моделей. Несмотря на множество различных рефлексивных моделей и определения игр на них, на сегодняшний день отсутствует более-менее завершённые исследования и пока не существует универсального аппарата аналитических решений для широких классов задач рефлексивного поведения, в комплексе учитывающих меняющиеся ситуации по экономическим ограничениям, конкуренцию, несовпадение экономических интересов и неполную, асимметричную информированность хозяйствующих субъектов при принятии решений. Пока основные успехи ограничены набором частных и достаточны простых моделей.

Особенностью исследований в статье является то, что динамический процесс принятия решений осуществляется не путем оптимальных ответов агентов на их ожидаемые действия, а как повторяемые статические игры на диапазоне допустимых ответов с учетом в каждой игре реальных экономических ограничений по их прибыли и конкурентоспособности. Традиционное для теории игр оптимизационное поведение, когда каждый агент на каждом шаге процесса принятия решений выбирает свой наилучший ответ, зачастую не позволяет получить сходящиеся к равновесию траектории. Здесь же каждый из рефлексизирующих агентов рассчитывает свое текущее положение цели и изменяет свое состояние в направлении текущего положения цели, рассчитывая при выборе ответов на ожидаемые действия конкурентов на положительную собственную ожидаемую прибыль или минимизацию потерь.

Теоретической основой для построения и исследования процессов являются теория рефлексивных игр [3, 10] и теория коллективного поведения [2, 27]. Они дополняют друг друга тем, что рефлексивные игры позволяют использовать процедуры коллективного поведения и результаты размышлений агентов, приводящие к равновесию Нэша.

Основным новым результатом проведенного аналитического исследования являются достаточные условия сходимости процедур рефлексивного коллективного поведения, к равновесию для олигополии с произвольным числом лидеров по Штакельбергу в классе линейных издержек агентов и линейного спроса.

Результаты, представленного в статье исследования, могут иметь прикладное значение для понимания группового поведения агентов на современных конкурентных рынках. Что является особенно важным в условиях цифровизации экономики, нуждающейся в новых моделях взаимодействия хозяйствующих субъектов, осуществляющих

совместную деятельность, позволяющие оперативное реагирование и регулирование рынков.

2. Постановка задач исследования. Рассматривается динамическая система, состоящая из n взаимосвязанных целенаправленных агентов и функционирующая в дискретном времени. Пусть состояние системы в момент времени t дается n -мерным вектором $q^t = (q_1^t, \dots, q_i^t, \dots, q_n^t)$, $t = 0, 1, 2, \dots$, и текущее положение цели i -го агента $x_i(q_{-i}^t)$ ($i \in N = \{1, \dots, n\}$) зависит от действий остальных агентов, где $q_{-i}^t = (q_1^t, \dots, q_{i-1}^t, q_{i+1}^t, \dots, q_n^t)$ – состояние окружения для i -го агента, вектор q^t без i -й компоненты. Текущее положение цели агента – такое его действие, которое максимизировало бы его целевую функцию при условии, что в текущий момент времени остальные агенты выбрали бы те же действия, что и в предыдущий [2, 9, 27, 28].

Пусть смена состояния системы при переходе от предыдущего момента времени t к последующему $(t+1)$ -му моменту, т.е. преобразование вектора q^t в q^{t+1} , представляется в виде:

$$q_i^{t+1} = q_i^t + \gamma_i^{t+1} (x_i(q_{-i}^t) - q_i^t), i \in N, t = 0, 1, 2, \dots \quad (1)$$

Здесь $\gamma_i^{t+1} \in [0; 1]$ – параметры, выбираемые агентами.

Определим такое итерационное преобразование действий агентов как *процесс 1*.

Модель (1) является наиболее распространенной моделью динамики коллективного поведения. На сегодняшний день имеется немало прикладных моделей, иллюстрирующих эффекты коллективного поведения по модели (1) [2, 3, 10, 24, 25, 27, 29–32].

Модель (1) может являться основой для определения более простых частных процессов. Так при $\gamma_i^{t+1} \equiv 1$ получаем преобразование:

$$q_i^{t+1} = x_i(q_{-i}^t), i \in N, t = 0, 1, 2, \dots \quad (2)$$

Определим итерационное преобразование компонент вектора q^t в компоненты вектора q^{t+1} по формуле (2) как *процесс 2*.

Такой процесс можно условно назвать «оптимизационным», так как каждый агент в каждый момент времени выбирает свой наилуч-

ший ответ на действия окружения. Однако такая динамика часто не является сходящейся [2, 3, 25, 27, 29].

Требование неотрицательности действий агентов, возникающее, например, с точки зрения экономических ограничений, может быть реализовано преобразованием вида:

$$q_i^{t+1} = \begin{cases} x_i(q_{-i}^t), & x_i(q_{-i}^t) > 0; \\ 0, & x_i(q_{-i}^t) \leq 0. \end{cases} \quad (3)$$

Определим преобразование (3) действий агентов как *процесс 3*.

Исследования таких процессов можно посмотреть, например, в работах [25, 29].

Модель (1) может являться основой для определения более сложных процессов. Так в следующем процессе учтено условие неотрицательности действий агентов:

$$q_i^{t+1} = \begin{cases} q_i^t + \gamma_i^{t+1}(x_i(q_{-i}^t) - q_i^t), & x_i(q_{-i}^t) > 0; \\ 0, & x_i(q_{-i}^t) \leq 0. \end{cases} \quad (4)$$

Здесь также $\gamma_i^{t+1} \in [0; 1]$.

Определим итерационное преобразование действий агентов по формуле (4) как *процесс 4*.

Ряд прикладных моделей коллективного поведения, когда агенты обнуляют свои действия, если их текущее положение цели равно нулю или отрицательно, можно найти в работах [25, 29]. Другие модификации модели (1) встречаются в монографиях [2, 3, 27].

В основу итерационного процесса вычисления новых действий агентов положена следующая единая для моделей коллективного поведения процедура [2, 3, 27]:

1. Каждый агент i ($i \in N$) в текущий $(t+1)$ -й момент времени наблюдает действия других агентов $\{q_i^t\}_{i \in N \setminus \{i\}}$, выбранные ими в предыдущий t -й момент времени (начальные действия агентов $\{q_i^0\}_{i \in N}$ считаются заданными);
2. Каждый агент рассчитывает свое текущее положение цели $x_i(q_{-i}^t)$.

3. Каждый агент в момент времени $(t+1)$ уточняет свое предыдущее действие, делая от него «шаг» $\gamma_i^{t+1} \in [0; 1]$ к текущему положению цели.

4. Процесс повторяется с п. 1 для следующего момента времени.

В (2) и (3) агент всегда делает полный шаг, полагая $\gamma_i^{t+1} \equiv 1$, тем самым, выбирает свой наилучший ответ. В (1) и (4) агент выбором параметра $\gamma_i^{t+1} \in [0; 1]$ может делать «неполный шаг» от своего предыдущего состояния к текущему положению цели. Естественно, чтобы при $\gamma_i^{t+1} \equiv 1$ (1) переходило в (2), а (4) – в (3).

Характерно, что теория коллективного поведения исследует динамику поведения однородных групп в типичных повторяющихся конфликтных ситуациях, когда каждый рациональный агент принимает решения при достаточно слабых предположениях относительно его информированности.

Приведенная здесь для коллективного (группового) поведения процедура основана на идее метода градиентов для нахождения экстремума аналитической функции. Как и метод градиента, она при одних значениях диапазона величины шага движения к цели может сходиться, при других – расходиться.

Текущее положение цели $x_i(q_{-i}^t)$ и, соответственно, текущий ответ каждого агента, рассчитываемый по формулам (1)–(4), зависят от того, какие действия выберут другие агенты в тот же самый $(t+1)$ -й момент времени, относительно которых трудно однозначно сказать априори. Что вынуждает агентов рефлексировать, т. е. предсказывать поведение других агентов и выбирать свои действия уже с учетом этого прогноза.

Определяющим эффектом рефлексии является достижение равновесия, под которым понимается устойчивый в том или ином (оговариваемом в каждом конкретном случае) смысле исход взаимодействия агентов – вектор их равновесных действий.

Целью настоящей статьи является применение моделей рефлексивного коллективного поведения для описания и прогнозирования группового поведения агентов на рынке олигополии и выявление условий достижения равновесия на их основе.

В нашем исследовании каждый агент может реагировать на действия конкурентов одним из двух способов, которые являются tradi-

ционными для моделей олигополистического рынка: рефлексировать как ведомый или рефлексировать как лидер.

Согласно приведенной в [3] классификации ведомый агент имеет нулевой ранг рефлексии, а лидер – первый ранг рефлексии. Агент с нулевым рангом рефлексии выбирает свои действия, считая, что действия остальных агентов будут такими же, что и в предыдущий момент времени. Агент с первым рангом рефлексии считает всех остальных агентов обладающими нулевым рангом рефлексии и что он точно предсказывает их выбор.

В условиях неопределенности выбора конкурентов рефлексивная модель олигополии с лидером имеет свою особенность по сравнению с классической иерархической игрой Штакельберга. В игре Штакельберга лидер делает первым ход, который становится известен другим агентам. В рефлексивной модели выбор реальных действий всеми агентами осуществляется синхронно (одновременно), другие агенты не знают ход лидера, синхронный своему ходу. Подобный прием упрощает реальный процесс последовательных реакций [24, 25, 33], он оправдан и адекватен в случае, когда достигнутое равновесие стабильно [33].

В рефлексивной игре ведомый агент даже не знает, что у него есть лидер, полагая, что он, как и другие агенты, оставит свой объем выпуска неизменным (например, считая остальных агентов менее «интеллектуальными», чем они сами, либо что оппоненты достигли равновесия и им не выгодно от него отклониться). Ведомый агент не знает, что другие такие агенты действуют подобным образом.

В статье рассматривается задача применения рефлексивных повторяющихся игр и моделей динамики коллективного поведения для описания и прогнозирования группового поведения агентов на рынке олигополии в классе линейных функций спроса и издержек агентов. Допускается, что может быть произвольное число ведомых агентов и лидеров. Для лидеров уровни (ранги) лидерства в статье не рассматриваются.

С вычислительной точки зрения динамики для ведомых и лидеров различаются расчетом текущего положения цели. Соответствующие формулы для изучаемой в статье прикладной модели рынка получены в следующем разделе.

Условия сходимости процессов к положению равновесия относятся к начальным приближениям $q^0 = (q_1^0, \dots, q_n^0)$, параметрам γ_i^{t+1} , общему числу агентов на рынке и к числу агентов, действующих как ведомые и как лидеры.

Другая основная задача статьи для изучаемой прикладной модели олигопольного рынка – получение аналитических выражений для диапазонов величин шагов агентов, а также условий на начальные приближения q^0 , при которых гарантируется сходимость моделей рефлексивного коллективного поведения к равновесию.

Будем полагать, если для динамической системы равновесие $q^* = (q_1^*, \dots, q_i^*, \dots, q_n^*)$ существует, то $q_i^* > 0, \forall i \in N$. Последнее условие означает, что все агенты конкурентоспособны в равновесии. Для модели олигополии в случае линейных издержек агентов и линейного спроса равновесие существует и единственно. Под равновесием понимается равновесие Нэша. Полагаем также, что $q^0 > 0$.

Под сходимостью к равновесию понимается сходимость по евклидовой норме.

3. Прикладная модель для олигопольного рынка. В качестве прикладной модели рефлексивного коллективного поведения рассмотрим классическую модель олигополии, состоящей из n , конкурирующих объемами выпуска однородной продукции, агентов с целевыми функциями:

$$\Pi_i(p(Q), q_i) = p(Q)q_i - \varphi_i(q_i) \rightarrow \max_{q_i}, \quad i \in N, \quad (5)$$

линейными функциями затрат:

$$\varphi_i(q_i) = c_i q_i + d_i, \quad i \in N, \quad (6)$$

и линейной обратной функции спроса вида:

$$p(Q) = a - bQ. \quad (7)$$

Здесь: q_i – выпуск (действие) i -го агента, $Q = \sum_{i \in N} q_i$ – суммарный объем выпуска, c_i, d_i – предельные и постоянные издержки агентов, $p(Q)$ – единая рыночная цена, a, b – параметры спроса. Полагается, что весь выпуск реализуется, ограничения мощности и коалиции отсутствуют.

Определим расчетные формулы для положения цели $x_i(q_{-i}^t)$ i -го агента в текущий момент времени ($t = 0, 1, 2, \dots$).

Оптимальный объем активности агента можно определить из условия $\frac{\partial \Pi_i^t}{\partial q_i^t} = \frac{\partial p^t}{\partial q_i^t} \cdot q_i^t + p^t - \frac{\partial \varphi_i}{\partial q_i^t} = 0$. Отсюда $\frac{\partial p^t}{\partial q_i^t} = \frac{1}{q_i^t} \cdot \left(\frac{\partial \varphi_i}{\partial q_i^t} - p^t \right)$. Из этого равенства с учетом (7) имеем $\frac{\partial p^t}{\partial q_i^t} = -b \cdot \frac{\partial Q^t}{\partial q_i^t} = \frac{1}{q_i^t} \cdot (c_i - a + b \cdot Q^t)$ или $1 + \frac{\partial Q_{-i}^t}{\partial q_i^t} = \frac{1}{q_i^t} \cdot h_i - 1 - \frac{1}{q_i^t} \cdot Q_{-i}^t$, где использованы следующие обозначения:

$$h_i = \frac{a - c_i}{b}, \quad (8)$$

$$Q_{-i}^t = \sum_{j \neq i} q_j^t, \quad (9)$$

т.е. Q_{-i}^t – суммарный объем выпуска без i -го агента.

Получаем $q_i^t \cdot \left(2 + \frac{\partial Q_{-i}^t}{\partial q_i^t} \right) = h_i - Q_{-i}^t$, и, окончательно:

$$q_i^t = \frac{h_i - Q_{-i}^t}{2 + \frac{\partial Q_{-i}^t}{\partial q_i^t}} \quad (i \in N; t = 0, 1, 2, \dots). \quad (10)$$

Вначале рассмотрим применение формулы (10) к ведомому агенту.

Ведомый агент, поведение которого известно как поведение по Курно [8], устанавливает объем выпуска, полагая, что другие агенты оставят свои объемы выпуска неизменными. Тогда, согласно определению дифференциала, имеем $dQ_{-i}^t = \sum_{j \in N} \frac{\partial Q_{-i}^t}{\partial q_j^t} dq_j^t$, и из того чтобы $dQ_{-i}^t = 0$ следует, что для i -го агента должны быть равны нулю не

только dq_j^t ($j \neq i$), но и $\frac{\partial Q_{-i}^t}{\partial q_i^t}$, так как в противном случае при

$dq_i^t \neq 0$ будет $dQ_{-i}^t \neq 0$. Итак, $\frac{\partial Q_{-i}^t}{\partial q_i^t} = 0$ ($i \in N$), а по (10) имеем

$$q_i^t = \frac{h_i - Q_{-i}^t}{2}.$$

Обозначим через N_c – множество ведомых агентов. Получаем выражение для оптимального ответа ведомого агента или его текущего положения цели (см. также [25, 29]):

$$x_i(q_{-i}^t) = \frac{h_i - Q_{-i}^t}{2} \quad (i \in N_c; t = 0, 1, 2, \dots). \quad (11)$$

Теперь рассмотрим применение формулы (10) к агенту-лидеру.

Лидер устанавливает свой объем выпуска, считая всех остальных агентов ведомыми и полагая, что точно знает их реакцию (объемы выпуска) на его действие.

Допустим, что k -й агент является лидером. Из предположения, что остальные действуют как ведомые, следует, что $dq_i^t = 0$ и $dQ_{-i}^t = 0$ при $i \in N \setminus \{k\}$. Имеем $\frac{\partial Q_{-i}^t}{\partial q_i^t} = 0$, а также $q_i^t = \frac{h_i - Q_{-i}^t}{2}$. Из последнего

равенства получаем $\frac{\partial q_i^t}{\partial q_k^t} = -\frac{1}{2} \cdot \frac{\partial Q_{-i}^t}{\partial q_k^t} = -\frac{1}{2} \cdot \frac{\partial Q^t}{\partial q_k^t} + \frac{1}{2} \cdot \frac{\partial q_i^t}{\partial q_k^t} = -\frac{1}{2} -$
 $-\frac{1}{2} \cdot \frac{\partial Q_{-k}^t}{\partial q_k^t} + \frac{1}{2} \cdot \frac{\partial q_i^t}{\partial q_k^t}$, или $\frac{\partial q_i^t}{\partial q_k^t} = -1 - \frac{\partial Q_{-k}^t}{\partial q_k^t}$, при $i \in N \setminus \{k\}$. Суммируя

левые и правые части последних равенств по $i \in N \setminus \{k\}$, получаем, что

$$\frac{\partial Q_{-k}^t}{\partial q_k^t} = -(n-1) - (n-1) \frac{\partial Q_{-k}^t}{\partial q_k^t}, \quad \text{т.е.} \quad \frac{\partial Q_{-k}^t}{\partial q_k^t} = -\frac{n-1}{n}. \quad \text{Тогда по (10)}$$

$$q_k^t = \frac{h_k - Q_{-k}^t}{2 - \frac{n-1}{n}}, \quad \text{или} \quad q_k^t = \frac{n(h_k - Q_{-k}^t)}{n+1}.$$

Обозначим через N_s – множество агентов-лидеров. Получаем выражение для оптимального ответа лидера или его текущего положения цели [25, 29]:

$$x_i(q_{-i}^t) = \frac{n(h_i - Q_{-i}^t)}{1+n} \quad (i \in N_s; t = 0, 1, 2, \dots). \quad (12)$$

Преобразуем (5) с учетом (6)–(9) к виду $\Pi_i^{t+1} = b(h_i - Q_{-i}^{t+1} - q_i^{t+1})q_i^{t+1} - d_i$.

Полагая неизменным выпуск остальных агентов, ведомый агент $i \in N_c$ при $h_i - Q_{-i}^{t+1} = h_i - Q_{-i}^t > 0$, используя параметр $\gamma_i^{t+1} \in [0; 1]$, по (1) выбирает выпуск q_i^{t+1} (по (4) выбирает положительный выпуск) в направлении текущего положения цели, которое определяется по формуле (11). Если $\gamma_i^{t+1} = 0$, то он не ожидает изменения своей прибыли. Если $\gamma_i^{t+1} = 1$, то агент выбирает оптимальный отклик (11) на ожидаемые действия конкурентов, максимизируя ожидаемую прибыль, так как $\frac{\partial \Pi_i^{t+1}}{\partial q_i^{t+1}} = b(h_i - Q_{-i}^t - 2q_i^{t+1}) = 0$. Агент также выбором параметра $\gamma_i^{t+1} \in [0; 1]$ может делать «неполный шаг». В модели (4) при $h_i - Q_{-i}^t \leq 0$ положительный выпуск дает отрицательную валовую прибыль (т.е. прибыль без учета постоянных издержек d_i) и, чтобы минимизировать ожидаемые потери, агент выбирает нулевой выпуск.

Агент-лидер $i \in N_s$ при $h_i - Q_{-i}^t > 0$, используя параметр $\gamma_i^{t+1} \in [0; 1]$, по (1) выбирает выпуск q_i^{t+1} (или по (4) выбирает положительный выпуск) в направлении текущего положения цели, которое определяется по (12). Если $\gamma_i^{t+1} = 0$, то он не ожидает изменения своей прибыли. Если $\gamma_i^{t+1} = 1$, то агент выбирает оптимальный отклик (12) на ожидаемые действия окружения, максимизируя ожидаемую прибыль, так как $\frac{\partial \Pi_i^{t+1}}{\partial q_i^{t+1}} = b\left(-\frac{\partial Q_{-i}^{t+1}}{\partial q_i^{t+1}} - 1\right)q_i^{t+1} + b(h_i - Q_{-i}^{t+1} - q_i^{t+1}) = b\left(\frac{n-1}{n} - 1\right)q_i^{t+1} +$

$$+b\left(h_i - Q_{-i}^{t+1} - q_i^{t+1}\right) = b\left(h_i - Q_{-i}^{t+1} - \frac{1+n}{n}q_i^{t+1}\right) = 0. \quad \text{В модели (4) при}$$

$h_i - Q_{-i}^t \leq 0$ положительный выпуск дает отрицательную валовую прибыль (т.е. прибыль без учета постоянных издержек d_i) и, чтобы минимизировать ожидаемые потери, агент выбирает нулевой выпуск.

Полагаем, что все агенты точно знают собственные затраты и целевую функцию, собственную функцию реакции, включающую параметры спроса a и b , ранее произведенный выпуск другими агентами, но не располагают достоверной априорной информацией относительно ожидаемых объёмов их выпуска, множеств допустимых действий, функций затрат и целевых функций конкурентов. Лидеры также знают общее число агентов на рынке.

Пример 1. Рассмотрим численный пример для процессов (1) – (4) на рынке с 4 агентами $N = \{1, 2, 3, 4\}$. Пусть $\{1, 2\} \in N_s$ и $\{3, 4\} \in N_c$. Пусть также $q^0 = (q_1^0, q_2^0, q_3^0, q_4^0) = (250, 250, 250, 200)$, $c = (c_1, c_2, c_3, c_4) = (20, 25, 20, 30)$, $a = 100$, $b = 0,1$. По (8) находим, что

$$h = (h_1, h_2, h_3, h_4) = (800, 750, 800, 700). \quad \text{Отношение } \frac{a}{b} = 1000 \text{ определя-$$

ет «ёмкость рынка» (считается, если суммарный объём выпуска Q превышает ёмкость рынка, то фирмы несут потери в объёме полных издержек), а h_i – объём совершенно конкурентного рынка при ценообразовании по предельным издержкам $p(Q) = c_i$, называемый «совершенно конкурентный объём фирмы i ».

Величины $q_i^t, h_i, \frac{a}{b}, x_i(q_{-i}^t)$ имеют натуральные единицы измерения (тонны, штуки и пр.), c_i, a – стоимостные.

В таблице 1 приведены первые двадцать итераций для процесса 4 с четырьмя агентами, из которых первые два действуют по Штакельбергу, другие два агента – по Курно. Объёмы выпуска q_i^t рассчитываются по (4), текущие цели $x_i(q_{-i}^t)$ по (11) и (12). Шаги γ_i^{t+1} принимают случайным образом одно из двух значений «0,5» или «1».

Таблица 1. Начальный фрагмент численного примера для процесса 4

Итерации	Выпуски				Текущие цели				Параметры шагов			
t	q_1	q_2	q_3	q_4	x_1	x_2	x_3	x_4	γ_1	γ_2	γ_3	γ_4
0	250,0	250,0	250,0	200,0	80,0	40,0	50,0	-25,0				
1	80,0	145,0	150,0	0,0	404,0	416,0	287,5	162,5	1,0	0,5	0,5	1,0
2	242,0	416,0	287,5	81,3	12,2	111,4	30,4	-122,8	0,5	1,0	1,0	0,5
3	127,1	263,7	30,4	0,0	404,7	474,0	204,6	139,4	0,5	0,5	1,0	0,5
4	404,7	368,9	117,5	69,7	195,2	126,5	-21,7	-95,5	1,0	0,5	0,5	0,5
5	299,9	247,7	0,0	0,0	441,9	360,0	126,2	76,2	0,5	0,5	0,5	1,0
6	370,9	303,8	126,2	38,1	265,5	171,8	43,6	-50,5	0,5	0,5	1,0	0,5
7	265,5	237,8	84,9	0,0	381,8	319,7	148,3	55,9	1,0	0,5	0,5	0,5
8	323,7	278,8	116,6	55,9	279,0	203,1	70,8	-9,5	0,5	0,5	0,5	1,0
9	279,0	240,9	70,8	0,0	390,6	320,1	140,1	54,6	1,0	0,5	1,0	0,5
10	334,8	280,5	105,4	27,3	309,4	226,0	78,7	-10,4	0,5	0,5	0,5	0,5
11	309,4	253,3	92,1	0,0	363,7	278,9	118,7	22,7	1,0	0,5	0,5	0,5
12	363,7	278,9	105,4	22,7	314,5	206,6	67,4	-24,0	1,0	1,0	0,5	1,0
13	339,1	206,6	67,4	0,0	420,9	274,8	127,2	43,5	0,5	1,0	1,0	1,0
14	380,0	240,7	127,2	21,7	328,3	176,9	78,8	-23,9	0,5	0,5	1,0	0,5
15	354,2	208,8	78,8	0,0	409,9	253,6	118,5	29,1	0,5	0,5	1,0	0,5
16	382,0	253,6	98,7	14,6	346,5	203,8	74,9	-17,2	0,5	1,0	0,5	0,5
17	346,5	228,7	86,8	0,0	387,6	253,4	112,4	19,0	1,0	0,5	0,5	1,0
18	367,1	253,4	99,6	9,5	350,0	219,1	85,0	-10,0	0,5	1,0	0,5	0,5
19	358,5	219,1	92,3	0,0	390,9	239,3	111,2	15,0	0,5	1,0	0,5	0,5
20	374,7	229,2	101,7	7,5	369,2	212,8	94,3	-2,8	0,5	0,5	0,5	0,5

При полной информированности агентов статичное равновесие Нэша q^* может быть найдено как решение следующей системы линейных алгебраических уравнений:

$$\begin{cases} q_i + nQ = nh_i & (i \in N_s), \\ q_i + Q = h_i & (i \in N_c). \end{cases}$$

Для нашего примера такая система уравнений имеет вид:

$$\begin{cases} q_1 + 4Q = 3200, \\ q_2 + 4Q = 3000, \\ q_3 + Q = 800, \\ q_4 + Q = 700. \end{cases}$$

Отсюда находим решение: $q^* = (400, 200, 100, 0)$.

По таблице 1 видно, что для процесса 4 имеет место тенденция сходимости к статичному равновесию Нэша. Особенностью процесса является обнуление выпуска агентами, если их текущее положение цели равно нулю или отрицательно. Так, например, на первой и третьей итерациях обнуляется выпуск четвертого агента, а на пятой – выпуски третьего и четвертого агентов.

Приведем также основные выводы по данному примеру, относящиеся к сходимости других процессов при таком же начальном выпуске q^0 .

Для процесса 1 при таких же шагах γ_i^{t+1} , как и для процесса 4, также прослеживается тенденция сходимости к равновесию. Последние две итерации для процесса 1 имеют вид $q^{19} = (342, 3; 224, 7; 87, 8; -20, 0)$, $q^{20} = (374, 2; 248, 3; 107, 2; 1, 3)$.

Процесс 2 расходится. Уже на четвертой и пятой итерациях имеем следующие выпуски: $q^4 = (640, 0; 600, 0; 400, 0; 350, 0)$, $q^5 = (-1805, 0; -1939, 4; -1406, 6; -1505, 0)$. Далее колебания процесса только увеличиваются.

Процесс 3 не сходится ввиду заикливания. Начиная с шестой итерации, каждые две соседние итерации такой же вид, как на четвертой и пятой: $q^4 = (640, 0; 600, 0; 400, 0; 350, 0)$, $q^5 = (0; 0; 0; 0)$.

Напомним, что процессы 2 и 3 можно рассматривать как частный случай процесса 1 и процесса 4, соответственно, при $\gamma_i^{t+1} \equiv 1$. Ниже будут получены условия, когда процессы 2 и 3 сходятся.

4. Вспомогательные леммы и их доказательства. Введем функции-индикаторы [2] $\alpha_i^t = 2(x_i(q_{-i}^t) - q_i^t)$, характеризующие отклонения текущих выпусков от текущих оптимумов агентов, для агентов с реакцией по Курно ($i \in N_c$) и $\alpha_i^t = \frac{1+n}{n}(x_i(q_{-i}^t) - q_i^t)$ для лидеров ($i \in N_s$). Коэффициенты «2» и « $\frac{1+n}{n}$ » введены для последующих удобств. Известно [23, 25], что $h_i = Q^* + q_i^*$ ($i \in N_c$) и $h_i = Q^* + \frac{q_i^*}{n}$

($i \in N_s$). Эти формулы можно также получить из предельных выражений для (11) и (12) при $t \rightarrow \infty$. Тогда имеем:

$$\alpha_i^t = Q^* + q_i^* - Q^t - q_i^t, \quad i \in N_c, \quad (13)$$

$$\alpha_i^t = Q^* + \frac{1}{n} q_i^* - Q^t - \frac{1}{n} q_i^t, \quad i \in N_s. \quad (14)$$

Лемма 1. Вектор действий агентов $q^t = (q_1^t, \dots, q_i^t, \dots, q_n^t)$ является статичным равновесием $q^ = (q_1^*, \dots, q_i^*, \dots, q_n^*)$ модели олигополии (5)–(7) тогда и только тогда, когда $\alpha_i^t = 0 \quad \forall i \in N$.*

Доказательство леммы 1. Обозначим через n_c – число ведомых агентов на рынке, а через n_s – число агентов-лидеров, $n_c + n_s = n$. Используя формулы (13) и (14), суммированием α_i^t по индексу $i \in N_c$ и $n\alpha_i^t$ по индексу $i \in N_s$, получаем $\sum_{i \in N_c} \alpha_i^t + n \sum_{i \in N_s} \alpha_i^t = (n_c + nn_s + 1)(Q^* - Q^t)$. Пусть $\exists t$, что $\alpha_i^t = 0 \quad \forall i \in N$.

Тогда $Q^* = Q^t$ и по (13) и (14) имеем $q_i^t = q_i^* \quad (i \in N)$. Решением однородной системы уравнений (13)–(14), когда $\alpha_i^t = 0$, является равновесный выпуск.

Если $q_i^t = q_i^* \quad \forall i \in N$, то опять по (13) и (14) получаем, что $\alpha_i^t = 0 \quad \forall i \in N$.

Лемма 1 доказана.

Лемма показывает, что функции-индикаторы α_i^t можно рассматривать в качестве оценки «удаленности» агентов от положения равновесия.

Введем более удобную для последующих преобразований замену параметров:

$$\lambda_i^{t+1} = \frac{\gamma_i^{t+1}}{2}, \quad i \in N_c; \quad \lambda_i^{t+1} = \frac{\gamma_i^{t+1}n}{1+n}, \quad i \in N_s. \quad (15)$$

С учетом введенных обозначений (13)–(15) перепишем (1) для процесса 1 в виде:

$$q_i^{t+1} = q_i^t + \lambda_i^{t+1} \alpha_i^t, \quad i \in N_c, \quad \lambda_i^{t+1} \in \left(0; \frac{1}{2}\right]; \quad (16)$$

$$q_i^{t+1} = q_i^t + \lambda_i^{t+1} \alpha_i^t, \quad i \in N_s, \quad \lambda_i^{t+1} \in \left(0; \frac{n}{1+n}\right]. \quad (17)$$

Далее потребуются следующие соотношения для процесса 1:

$$\alpha_i^{t+1} = \left(1 - \lambda_i^{t+1}\right) \alpha_i^t - \sum_{j \in N} \lambda_j^{t+1} \alpha_j^t, \quad i \in N_c, \quad (18)$$

$$\alpha_i^{t+1} = \left(1 - \frac{\lambda_i^{t+1}}{n}\right) \alpha_i^t - \sum_{j \in N} \lambda_j^{t+1} \alpha_j^t, \quad i \in N_s. \quad (19)$$

Приведем вывод (19) для агентов с реакцией по Штакельбергу, для агентов с реакцией по Курно эта формула выводится аналогично.

Используя формулу (14) для α_i^{t+1} и α_i^t , а также (17), имеем

$$\alpha_i^{t+1} - \alpha_i^t = -\frac{1}{n} \left(q_i^{t+1} - q_i^t \right) - Q^{t+1} + Q^t = -\frac{1}{n} \lambda_i^{t+1} \alpha_i^t - Q^{t+1} + Q^t = \left(1 - \frac{\lambda_i^{t+1}}{n} \right) \alpha_i^t - Q^{t+1} + Q^t. \quad \text{Суммированием (16), (17) по индексу } i \text{ получаем}$$

$$Q^{t+1} = Q^t + \sum_{j \in N} \lambda_j^{t+1} \alpha_j^t \text{ и, окончательно, (19).}$$

Приведем аналогичные формулы для процесса 4.

Введем новые обозначения: $N_1^t = \left\{ i \mid x_i^t > 0, i \in N \right\}$,

$N_2^t = \left\{ i \mid x_i^t \leq 0, i \in N \right\}$. Тогда $N_1^t \cap N_2^t = \emptyset$ и $N_1^t \cup N_2^t = N$.

С учетом этих и ранее введенных обозначений (13)–(15), а также (11), (12) запишем (4) как:

$$q_i^{t+1} = \begin{cases} q_i^{t+1} = q_i^t + \lambda_i^{t+1} \alpha_i^t, & i \in N_c \cap N_1^t, \quad \lambda_i^{t+1} \in \left(0; \frac{1}{2}\right]; \\ 0, & i \in N_c \cap N_2^t, \end{cases} \quad (20)$$

$$q_i^{t+1} = \begin{cases} q_i^t + \lambda_i^{t+1} \alpha_i^t, & i \in N_s \cap N_1^t, \quad \lambda_i^{t+1} \in \left(0; \frac{n}{1+n}\right]; \\ 0, & i \in N_s \cap N_2^t. \end{cases} \quad (21)$$

Из формул (13), (14) и (20), (21) имеем:

$$Q^* + q_i^* - Q^{t+1} - q_i^{t+1} = Q^* + q_i^* - Q^t - q_i^t - \lambda_i^{t+1} \alpha_i^t - Q^{t+1} + Q^t, \quad i \in N_c \cap N_1^t;$$

$$Q^* + \frac{1}{n} q_i^* - Q^{t+1} - \frac{1}{n} q_i^{t+1} = Q^* + \frac{1}{n} q_i^* - Q^t - \frac{1}{n} q_i^t - \frac{1}{n} \lambda_i^{t+1} \alpha_i^t - Q^{t+1} + Q^t, \quad i \in N_s \cap N_1^t.$$

Тогда для процесса 4 с учетом того, что по (20), (21) $Q^{t+1} = Q^t + \sum_{j \in N_1^t} \lambda_j^{t+1} \alpha_j^t - \sum_{j \in N_2^t} q_j^t$, имеем:

$$\alpha_i^{t+1} = (1 - \lambda_i^{t+1}) \alpha_i^t - \sum_{j \in N_1^t} \lambda_j^{t+1} \alpha_j^t + \sum_{j \in N_2^t} q_j^t, \quad i \in N_c \cap N_1^t; \quad (22)$$

$$\alpha_i^{t+1} = \left(1 - \frac{\lambda_i^{t+1}}{n}\right) \alpha_i^t - \sum_{j \in N_1^t} \lambda_j^{t+1} \alpha_j^t + \sum_{j \in N_2^t} q_j^t, \quad i \in N_s \cap N_1^t. \quad (23)$$

Также по формулам (13), (14) и (20), (21) получаем:

$$Q^* + q_i^* - Q^{t+1} - q_i^{t+1} = Q^* + q_i^* - Q^t - q_i^t + q_i^t - Q^{t+1} + Q^t, \quad i \in N_c \cap N_2^t;$$

$$Q^* + \frac{1}{n} q_i^* - Q^{t+1} - \frac{1}{n} q_i^{t+1} = Q^* + \frac{1}{n} q_i^* - Q^t - \frac{1}{n} q_i^t + \frac{1}{n} q_i^t - Q^{t+1} + Q^t, \quad i \in N_s \cap N_2^t.$$

Тогда для процесса 4 имеем:

$$\alpha_i^{t+1} = \alpha_i^t + q_i^t - \sum_{j \in N_1^t} \lambda_j^{t+1} \alpha_j^t + \sum_{j \in N_2^t} q_j^t, \quad i \in N_c \cap N_2^t; \quad (24)$$

$$\alpha_i^{t+1} = \alpha_i^t + \frac{1}{n} q_i^t - \sum_{j \in N_1^t} \lambda_j^{t+1} \alpha_j^t + \sum_{j \in N_2^t} q_j^t, \quad i \in N_s \cap N_2^t. \quad (25)$$

Лемма 2. Если в последовательности $\{\alpha_i^t, i \in N\}$ есть не только положительные члены, то для процессов 1 и 4 имеет место неравенство $\max_{i,j \in N} \{\alpha_i^{t+1} - \alpha_j^{t+1}\} \leq \eta^t \max_{i,j \in N} \{\alpha_i^t - \alpha_j^t\}$. Здесь $\eta^t \in (0; 1)$.

Доказательство леммы 2 можно найти в [25, 29].

Лемма 3. Пусть для процесса 1 $\max_{i \in N} \sum_{j \neq i} \lambda_j^{t+1} < 1$ и $\forall i \in N \alpha_i^t \geq 0$ или $\forall i \in N \alpha_i^t \leq 0$. Тогда $\exists \mu^t \in (0; 1)$, что $\mu^t \max_{i \in N} |\alpha_i^t| \geq \max_{i \in N} |\alpha_i^{t+1}|$.

Доказательство леммы 3.

Пусть вначале $\alpha_i^t \geq 0$ ($\forall i \in N$) и не все члены равны нулю. Когда в $\{\alpha_i^t, i \in N\}$ только нулевые члены, по лемме 1 процесс находится в состоянии равновесия. Если для некоторого i -го агента $\alpha_i^{t+1} \geq 0$, то по (18), (19) имеем, что $(1 - \lambda_i^{t+1})\alpha_i^t > \alpha_i^{t+1}$ ($i \in N_c$) и $\left(1 - \frac{\lambda_i^{t+1}}{n}\right)\alpha_i^t > \alpha_i^{t+1}$ ($i \in N_s$). Если $\alpha_i^{t+1} \leq 0$, то $\sum_{j \in N, j \neq i} \lambda_j^{t+1} \cdot \max_{j \in N, j \neq i} \alpha_j^t > -\alpha_i^{t+1}$ ($i \in N$).

Пусть $\alpha_i^t \leq 0$ ($\forall i \in N$). Если для некоторого i -го агента $\alpha_i^{t+1} \geq 0$, то по (18), (19) $\sum_{j \in N, j \neq i} \lambda_j^{t+1} \cdot \max_{j \in N, j \neq i} (-\alpha_j^t) > \alpha_i^{t+1}$. Если $\alpha_i^{t+1} \leq 0$, то $(1 - \lambda_i^{t+1})(-\alpha_i^t) > -\alpha_i^{t+1}$ ($i \in N_c$) и $\left(1 - \frac{\lambda_i^{t+1}}{n}\right)(-\alpha_i^t) > -\alpha_i^{t+1}$ ($i \in N_s$).

Полученные неравенства могут быть записаны в требуемом виде $\mu^t \max_{i \in N} |\alpha_i^t| \geq \max_{i \in N} |\alpha_i^{t+1}|$, где $\mu^t \in (0; 1)$.

Лемма 3 доказана.

Лемма 4. Пусть для процесса 4 $\forall i \in N \alpha_i^t \geq 0$ и $\max_{i \in N_1^t} \sum_{j \in N_1^t, j \neq i} \lambda_j^{t+1} < 1$, или $\forall i \in N \alpha_i^t, \alpha_i^{t+1} \leq 0$. Тогда $\exists \mu^t \in (0; 1)$, что $\mu^t \max_{i \in N} |\alpha_i^t| \geq \max_{i \in N} |\alpha_i^{t+1}|$.

Доказательство леммы 4.

Пусть вначале $\alpha_i^t \geq 0$ ($\forall i \in N$). Если для некоторого i -го агента $\alpha_i^t > 0$, то $i \in N_1^t$. Если $\alpha_i^t = 0$, то при $i \in N_1^t$ будет $q_i^t > 0$, а при $i \in N_2^t$ будет $x_i(q_{-i}^t) = 0$ и $q_i^t = 0$.

Тогда, если для некоторого i -го агента $\alpha_i^{t+1} \geq 0$, то по (22), (23) можем записать $(1 - \lambda_i^{t+1})\alpha_i^t > \alpha_i^{t+1}$ ($i \in N_c$) и $\left(1 - \frac{\lambda_i^{t+1}}{n}\right)\alpha_i^t > \alpha_i^{t+1}$ ($i \in N_s$). А по (24), (25) следует, что $\alpha_i^{t+1} < 0$ ($i \in N_2^t$).

Пусть $\alpha_i^t \geq 0$ ($\forall i \in N$) и для некоторого i -го агента $\alpha_i^{t+1} \leq 0$. Тогда по (22)–(25) имеем $\sum_{j \in N_1^t, j \neq i} \lambda_j^{t+1} \cdot \max_{j \in N_1^t, j \neq i} \alpha_j^t > -\alpha_i^{t+1}$.

Пусть теперь $\alpha_i^t \leq 0$ ($\forall i \in N$) и для некоторого i -го агента $\alpha_i^{t+1} \leq 0$. По (22), (23) получаем, что $(1 - \lambda_i^{t+1})(-\alpha_i^t) > -\alpha_i^{t+1}$ ($i \in N_c \cap N_1^t$) и $\left(1 - \frac{\lambda_i^{t+1}}{n}\right)(-\alpha_i^t) > -\alpha_i^{t+1}$ ($i \in N_s \cap N_1^t$), а по (24), (25) $-\alpha_i^t - q_i^t > -\alpha_i^{t+1}$ ($i \in N_c \cap N_2^t$), $-\alpha_i^t - \frac{q_i^t}{n} > -\alpha_i^{t+1}$ ($i \in N_s \cap N_2^t$).

Во всех рассмотренных ситуациях полученные неравенства также могут быть записаны в обобщенном виде $\mu^t \max_{i \in N} |\alpha_i^t| \geq \max_{i \in N} |\alpha_i^{t+1}|$, где $\mu^t \in (0; 1)$.

Лемма 4 доказана.

5. Результаты и их обсуждение на рефлексивной модели олигополии. Основным результатом данной статьи являются доказанные в этом разделе утверждения для олигополии с произвольным числом лидеров по Штакельбергу.

Утверждение 1. Пусть в олигополии (5)–(7) с произвольным числом лидеров по Штакельбергу для процесса 1 в последовательности $\{\alpha_i^t\}$ есть члены с разными знаками. Тогда в последовательности $\{\alpha_i^{t+1}\}$ также будут члены с разными знаками, если:

$$\gamma_i^{t+1} \in \left(0; \frac{2}{1+n_+^t}\right] \left(i \in N_c \cap N_+^t\right), \gamma_i^{t+1} \in \left(0; \frac{2}{1+n_-^t}\right] \left(i \in N_c \cap N_-^t\right), \quad (26)$$

$$\gamma_i^{t+1} \in \left(0; \frac{1+n}{1+nn_+^t}\right] \left(i \in N_s \cap N_+^t\right), \gamma_i^{t+1} \in \left(0; \frac{1+n}{1+nn_-^t}\right] \left(i \in N_s \cap N_-^t\right). \quad (27)$$

Здесь N_c – множество агентов с реакцией по Курно, N_s – множество агентов с реакцией по Штакельбергу, $N_+^t(N_-^t)$ – множество положительных (отрицательных) членов в последовательности $\{\alpha_i^t\}$, $n_+^t(n_-^t)$ – число положительных (отрицательных) членов в последовательности $\{\alpha_i^t\}$.

Доказательство утверждения 1. Из (18) и (19) для процесса 1 имеем
$$\sum_{j \in N_+^t} \alpha_j^{t+1} = \sum_{j \in N_+^t \cap N_c} (1 - \lambda_j^{t+1}) \alpha_j^t + \sum_{j \in N_+^t \cap N_s} \left(1 - \frac{\lambda_j^{t+1}}{n}\right) \alpha_j^t - n_+^t \sum_{j \in N} \lambda_j^{t+1} \alpha_j^t = \sum_{j \in N_+^t \cap N_c} (1 - \lambda_j^{t+1} (1 + n_+^t)) \alpha_j^t + \sum_{j \in N_+^t \cap N_s} \left(1 - \lambda_j^{t+1} \left(\frac{1}{n} + n_+^t\right)\right) \alpha_j^t - n_+^t \sum_{j \in N_-^t} \lambda_j^{t+1} \alpha_j^t.$$
 Поэтому, если $1 - \lambda_j^{t+1} (1 + n_+^t) \geq 0$ ($j \in N_+^t \cap N_c$), $1 - \lambda_j^{t+1} \frac{1 + nn_+^t}{n} \geq 0$ ($j \in N_+^t \cap N_s$) и в $\{\alpha_i^t\}$ есть хотя бы один отрицательный член, то $\sum_{j \in N_+^t} \alpha_j^{t+1} > 0$ и среди α_j^{t+1} ($j \in N_+^t$) всегда будет хотя бы один положительный член, т. е. его знак сохранится. Также

$$\begin{aligned} \sum_{j \in N_-^t} \alpha_j^{t+1} &= \sum_{j \in N_-^t \cap N_c} (1 - \lambda_j^{t+1}) \alpha_j^t + \sum_{j \in N_-^t \cap N_s} \left(1 - \frac{\lambda_j^{t+1}}{n}\right) \alpha_j^t - n_-^t \sum_{j \in N} \lambda_j^{t+1} \alpha_j^t = \\ &= \sum_{j \in N_-^t \cap N_c} (1 - \lambda_j^{t+1} (1 + n_-^t)) \alpha_j^t + \sum_{j \in N_-^t \cap N_s} \left(1 - \lambda_j^{t+1} \left(\frac{1}{n} + n_-^t\right)\right) \alpha_j^t - n_-^t \sum_{j \in N_+^t} \lambda_j^{t+1} \alpha_j^t. \end{aligned}$$

Если $1 - \lambda_j^{t+1} (1 + n_-^t) \geq 0$ ($j \in N_-^t \cap N_c$), $1 - \lambda_j^{t+1} \frac{1 + nn_-^t}{n} \geq 0$ ($j \in N_-^t \cap N_s$) и в $\{\alpha_i^t\}$ есть хотя бы один положительный член, то $\sum_{j \in N_-^t} \alpha_j^{t+1} < 0$ и среди α_j^{t+1} ($j \in N_-^t$) всегда будет отрицательный член.

Используя замену переменных λ на γ по формуле (15), приходим к условиям (26) и (27) утверждения 1.

Утверждение 1 доказано.

Утверждение 2. В олигополии (5)–(7) с произвольным числом лидеров по Штакельбергу процесс 1 сходится к равновесию при любых начальных выпусках агентов $\{q_i^0, i \in N = \{1, \dots, n\}, n \geq 2\}$, если, начиная с некоторого момента времени t_0 , при $t > t_0$, выполнены условия (26) и (27).

Доказательство утверждения 2. Пусть в текущий t -й момент времени в $\{\alpha_i^t\}$ есть члены с разными знаками и

$$\lambda_i^{t+k} \in \left(0; \frac{1}{1 + n_+^{t+k-1}}\right] \left(i \in N_c \cap N_+^{t+k-1}\right), \lambda_i^{t+k} \in \left(0; \frac{1}{1 + n_-^{t+k-1}}\right] \left(i \in N_c \cap N_-^{t+k-1}\right),$$

$$\lambda_i^{t+k} \in \left(0; \frac{n}{1 + nn_+^{t+k-1}}\right] \left(i \in N_s \cap N_+^{t+k-1}\right), \lambda_i^{t+k} \in \left(0; \frac{n}{1 + nn_-^{t+k-1}}\right] \left(i \in N_s \cap N_-^{t+k-1}\right)$$

при $k = 1, 2, \dots$. Тогда в $\{\alpha_i^{t+k}\}$ также будут члены с разными знаками. По лемме 2 будет $0 \leq \dots \leq \eta^{t+k} \max_{i, j \in N} \{\alpha_i^{t+k} - \alpha_j^{t+k}\} \leq$

$$\leq \eta^{t+k} \eta^{t+k-1} \max_{i, j \in N} \{\alpha_i^{t+k-1} - \alpha_j^{t+k-1}\} \leq \dots \leq \eta^{t+k} \eta^{t+k-1} \dots \eta^{t+1} \max_{i, j \in N} \{\alpha_i^{t+1} - \alpha_j^{t+1}\} \leq$$

$$\leq \eta^{t+k} \eta^{t+k-1} \dots \eta^{t+1} \eta^t \max_{i, j \in N} \{\alpha_i^t - \alpha_j^t\}. \text{ Таким образом, если } \eta^t \in (0; 1) \text{ и } \eta^t$$

не стремится к нулю, то $\max_{i, j \in N} \{\alpha_i^t - \alpha_j^t\} \rightarrow 0$ при $t \rightarrow \infty$. Поскольку

знаки минимального и максимального членов в $\{\alpha_i^t\}$ не совпадают, то

$\forall i \in N \alpha_i^t \rightarrow 0$ при $t \rightarrow \infty$. По лемме 1 процесс будет сходиться к равновесному выпуску q^* .

Пусть $\alpha_i^t \geq 0 \forall i \in N$ или $\alpha_i^t \leq 0 \forall i \in N$ и есть отличные от нуля члены (по лемме 1 равенство всех членов последовательности нулю означает равновесие). Если в $\{\alpha_i^{t+1}\}$ будут члены с разными знаками, то, как выше показано, в условиях утверждения 1 такой процесс сходится. Если в $\{\alpha_i^{t+1}\}$ будут только неотрицательные или только неположительные члены, то по лемме 3 $\exists \mu^t \in (0; 1)$, что $\mu^t \max_{i \in N} |\alpha_i^t| \geq \max_{i \in N} |\alpha_i^{t+1}|$. Если такая ситуация будет повторяться, т. е. при $k = 0, 1, 2, \dots$ $\alpha_i^{t+k} \geq 0 \forall i \in N$ или $\alpha_i^{t+k} \leq 0 \forall i \in N$, то при $\max_{i \in N} \sum_{j \neq i} \lambda_j^{t+k+1} < 1$ получим $0 \leq \dots \leq \mu_i^{t+k} \max_{i \in N} |\alpha_i^{t+k}| \leq \mu_i^{t+k} \mu_i^{t+k-1} \max_{i \in N} |\alpha_i^{t+k-1}| \leq \dots \leq \mu_i^{t+k} \mu_i^{t+k-1} \dots \mu_i^{t+1} \max_{i \in N} |\alpha_i^{t+1}| \leq \mu_i^{t+k} \mu_i^{t+k-1} \dots \mu_i^{t+1} \mu_i^t \max_{i \in N} |\alpha_i^t|$. В условиях утверждения 1 неравенства $\max_{i \in N} \sum_{j \neq i} \lambda_j^{t+k+1} < 1$ будут выполнены. Из полученной цепочки неравенств следует, что α_i^{t+k} сходятся к нулю, поскольку $\mu_i^{t+k} \in (0; 1) \forall i \in N, \forall k \geq 0$ и μ_i^{t+k} не стремятся к нулю. По лемме 1 процесс будет сходиться к равновесному выпуску q^* .

В доказательстве утверждения 2 фигурировал произвольный момент времени, при $t = 0$ имеем начальные условия процесса.

Таким образом, показана сходимость процесса 1 при любых начальных выпусках агентов $\{q_i^0, i \in N\}$, если выполнены условия (26), (27).

Утверждение 2 доказано.

Пример 2. Рассмотрим процесс 1 на рынке с 4 агентами $N = \{1, 2, 3, 4\}$. Пусть $\{1, 3\} \in N_c, \{2, 4\} \in N_s$, и в текущий момент времени $\alpha_1^t > 0, \alpha_2^t > 0, \alpha_3^t < 0, \alpha_4^t < 0$. Уравнения такого процесса:

$$\alpha_1^{t+1} = (1 - 2\lambda_1^{t+1})\alpha_1^t - \lambda_2^{t+1}\alpha_2^t - \lambda_3^{t+1}\alpha_3^t - \lambda_4^{t+1}\alpha_4^t,$$

$$\alpha_2^{t+1} = \left(1 - \frac{3}{2}\lambda_2^{t+1}\right)\alpha_2^t - \lambda_1^{t+1}\alpha_1^t - \lambda_3^{t+1}\alpha_3^t - \lambda_4^{t+1}\alpha_4^t,$$

$$\alpha_3^{t+1} = (1 - 2\lambda_3^{t+1})\alpha_3^t - \lambda_1^{t+1}\alpha_1^t - \lambda_2^{t+1}\alpha_2^t - \lambda_4^{t+1}\alpha_4^t,$$

$$\alpha_4^{t+1} = \left(1 - \frac{3}{2}\lambda_4^{t+1}\right)\alpha_4^t - \lambda_1^{t+1}\alpha_1^t - \lambda_2^{t+1}\alpha_2^t - \lambda_3^{t+1}\alpha_3^t.$$

При $\lambda_1^{t+1}, \lambda_3^{t+1} \leq \frac{1}{3}$ и $\lambda_2^{t+1}, \lambda_4^{t+1} \leq \frac{2}{5}$ имеем:

$$\alpha_1^{t+1} + \alpha_2^{t+1} = (1 - 3\lambda_1^{t+1})\alpha_1^t + \left(1 - \frac{5}{2}\lambda_2^{t+1}\right)\alpha_2^t - 2\lambda_3^{t+1}\alpha_3^t - 2\lambda_4^{t+1}\alpha_4^t > 0,$$

$$\alpha_3^{t+1} + \alpha_4^{t+1} = (1 - 3\lambda_3^{t+1})\alpha_3^t + \left(1 - \frac{5}{2}\lambda_4^{t+1}\right)\alpha_4^t - 2\lambda_1^{t+1}\alpha_1^t - 2\lambda_2^{t+1}\alpha_2^t < 0.$$

Таким образом, в последовательности $\{\alpha_i^{t+1}\}$ есть отрицательные и положительные члены.

В таблице 2 приведен фрагмент числового примера процесса 1 с четырьмя агентами, в которых два агента (второй и четвертый) являются лидерами по Штакельбергу.

В расчетах значения параметров λ полагаются равными верхним границам диапазонов сходимости процесса, которые приведены в утверждении 1. Перерасчет λ обусловлен изменением знаков членов в последовательностях $\{\alpha_i^t\}$. Так на пятой итерации знак α_3^5 изменился с «-» на «+», а на седьмой, наоборот, с «+» на «-», что вызвало перерасчет λ_i^6 и λ_i^8 , соответственно. Также смена знаков некоторых функций-индикаторов произошла на девятнадцатой, двадцать первой, двадцать восьмой и двадцать девятой итерациях. О сходимости процесса указывает то, что $\max_{i,j \in N} \{\alpha_i^t - \alpha_j^t\} \rightarrow 0$ при возрастании t .

Следствие 1. В дуополии (5)–(7) процесс 1 сходится к положению равновесия при $\gamma_i^{t+1} \in (0; 1]$ ($i \in N, t = 0, 1, 2, \dots$) и любых начальных условиях $q^0 > 0$, независимо от того, агент рефлексивирует по Курно или по Штакельбергу.

Доказательство. Если в текущий t -й момент времени у α_1^t и α_2^t разные знаки, то $n_+^t = n_-^t = 1$ и по (26)–(27) имеем, что

$\gamma_i^{t+1} \in (0; 1]$. Если один член нулевой или α_1^t и α_2^t имеют одинаковые знаки, то, как было показано выше для произвольного числа агентов на рынке, сходимость процесса доказывается с применением лемм 3 и 1.

Таблица 2. Фрагмент сходящегося процесса 1 для четырех агентов

Итерации	Значения функций-индикаторов				Параметры шагов								$\max\{\alpha_i - \alpha_j\}$
	α_1	α_2	α_3	α_4	λ_1	λ_2	λ_3	λ_4	γ_1	γ_2	γ_3	γ_4	
0	30,00	20,00	-18,00	-50,00									80,00
1	28,00	24,00	-4,00	-32,00	0,33	0,40	0,33	0,40	0,67	0,50	0,67	0,50	60,00
2	13,87	14,40	-7,47	-30,40	0,33	0,40	0,33	0,40	0,67	0,50	0,67	0,50	44,27
3	13,51	15,79	-0,71	-20,05	0,33	0,40	0,33	0,40	0,67	0,50	0,67	0,50	33,56
4	6,45	10,07	-3,03	-18,60	0,33	0,40	0,33	0,40	0,67	0,50	0,67	0,50	25,05
5	6,57	10,33	0,25	-12,61	0,33	0,40	0,33	0,40	0,67	0,50	0,67	0,50	22,94
6	8,66	12,58	3,89	-4,68	0,25	0,29	0,33	0,67	0,50	0,36	0,67	0,83	17,26
7	2,56	6,85	-1,34	-7,06	0,25	0,29	0,33	0,67	0,50	0,36	0,67	0,83	13,90
8	1,38	5,16	-1,21	-5,97	0,33	0,40	0,33	0,40	0,67	0,50	0,67	0,50	11,12
9	1,19	4,39	-0,54	-4,50	0,33	0,40	0,33	0,40	0,67	0,50	0,67	0,50	8,90
10	0,62	3,34	-0,53	-3,77	0,33	0,40	0,33	0,40	0,67	0,50	0,67	0,50	7,12
11	0,56	2,82	-0,21	-2,88	0,33	0,40	0,33	0,40	0,67	0,50	0,67	0,50	5,69
12	0,28	2,16	-0,23	-2,39	0,33	0,40	0,33	0,40	0,67	0,50	0,67	0,50	4,56
13	0,26	1,81	-0,08	-1,84	0,33	0,40	0,33	0,40	0,67	0,50	0,67	0,50	3,64
14	0,13	1,40	-0,10	-1,52	0,33	0,40	0,33	0,40	0,67	0,50	0,67	0,50	2,92
15	0,13	1,16	-0,03	-1,17	0,33	0,40	0,33	0,40	0,67	0,50	0,67	0,50	2,33
16	0,06	0,90	-0,04	-0,97	0,33	0,40	0,33	0,40	0,67	0,50	0,67	0,50	1,87
17	0,06	0,74	-0,01	-0,75	0,33	0,40	0,33	0,40	0,67	0,50	0,67	0,50	1,49
18	0,03	0,58	-0,02	-0,61	0,33	0,40	0,33	0,40	0,67	0,50	0,67	0,50	1,19
19	0,03	0,48	0,00	-0,48	0,33	0,40	0,33	0,40	0,67	0,50	0,67	0,50	0,96
20	0,20	0,58	0,18	-0,14	0,25	0,29	0,33	0,67	0,50	0,36	0,67	0,83	0,73
21	-0,03	0,32	-0,06	-0,28	0,25	0,29	0,33	0,67	0,50	0,36	0,67	0,83	0,60
22	-0,03	0,26	-0,05	-0,24	0,25	0,33	0,25	0,29	0,50	0,42	0,50	0,36	0,50
23	-0,02	0,22	-0,04	-0,21	0,25	0,33	0,25	0,29	0,50	0,42	0,50	0,36	0,43
24	-0,02	0,18	-0,03	-0,18	0,25	0,33	0,25	0,29	0,50	0,42	0,50	0,36	0,36
25	-0,01	0,15	-0,02	-0,15	0,25	0,33	0,25	0,29	0,50	0,42	0,50	0,36	0,30
26	-0,01	0,13	-0,02	-0,13	0,25	0,33	0,25	0,29	0,50	0,42	0,50	0,36	0,26
27	-0,01	0,11	-0,01	-0,11	0,25	0,33	0,25	0,29	0,50	0,42	0,50	0,36	0,22
28	0,00	0,09	-0,01	-0,09	0,25	0,33	0,25	0,29	0,50	0,42	0,50	0,36	0,18
29	0,00	0,08	0,00	-0,07	0,33	0,40	0,33	0,40	0,67	0,50	0,67	0,50	0,15
30	0,03	0,09	0,02	-0,02	0,25	0,29	0,33	0,67	0,50	0,36	0,67	0,83	0,11

Очевидно следующее следствие.

Следствие 2. В дуополии (5)–(7) процесс 2 сходится к положению равновесия при любых начальных условиях $q^0 > 0$, независимо от того, агент рефлексировал по Курно или по Штакельбергу.

Для процесса 4 в дуополии докажем следующее утверждение.

Утверждение 3. В дуополии (5)–(7) процесс 4 сходится к положению равновесия при $\gamma_i^{t+1} \in (0; 1]$ и любых начальных условиях $q^0 > 0$, независимо от того, агент рефлексировал по Курно или по Штакельбергу, за исключением случаев а) если $q_1^t = 0$, то $\gamma_2^{t+1} = 1$, б) если $q_2^t = 0$, то $\gamma_1^{t+1} = 1$.

Доказательство утверждения 3. Вначале покажем, что в дуополии для агентов с разными знаками, имеет место утверждение, аналогичное утверждению 1 для процесса 1.

Поскольку каждый агент может рефлексировать по Курно или по Штакельбергу, то вначале рассмотрим случай, когда оба агента рефлексировали по Курно, и в текущий t -й момент времени α_1^t и α_2^t имеют разные знаки. Покажем, если $\alpha_i^t > 0$, то $\alpha_i^{t+1} > 0$, и, если $\alpha_i^t < 0$, то $\alpha_i^{t+1} < 0$. Пусть, для определенности, $\alpha_1^t > 0$, $\alpha_2^t < 0$. Так как $\alpha_1^t > 0$, то из формулы $\alpha_i^t = 2(x_i(q_{-i}^t) - q_i^t)$ имеем $x_1(q_{-1}^t) > 0$, и поэтому $\{1\} \in N_1^t$. В свою очередь, для второго агента есть две возможности $\{2\} \in N_1^t$ или $\{2\} \in N_2^t$. Если $\{2\} \in N_1^t$, то уравнения процесса 4 совпадают с соответствующими уравнениями процесса 1, которые были рассмотрены в примере 1, и $\alpha_1^{t+1} > 0$, $\alpha_2^{t+1} < 0$.

Если $\{2\} \in N_2^t$, то поскольку $x_2(q_{-2}^t) \leq 0$ и $0 < \lambda_1^{t+1} \leq \frac{1}{2}$, используя соотношение $\alpha_2^t = 2(x_2(q_{-2}^t) - q_2^t)$, по формулам (22) и (24) имеем при $q_2^t \neq 0$ или $\lambda_1^{t+1} \neq \frac{1}{2}$.

$$\begin{aligned}\alpha_1^{t+1} &= (1 - \lambda_1^{t+1})\alpha_1^t - \lambda_1^{t+1}\alpha_1^t + q_2^t = (1 - 2\lambda_1^{t+1})\alpha_1^t + q_2^t > 0, \\ \alpha_2^{t+1} &= \alpha_2^t + q_2^t - \lambda_1^{t+1}\alpha_1^t + q_2^t = \alpha_2^t + 2x_2(q_{-2}^t) - \alpha_2^t - \lambda_1^{t+1}\alpha_1^t = 2x_2(q_{-2}^t) - \\ &- \lambda_1^{t+1}\alpha_1^t < 0.\end{aligned}$$

Таким образом, в дуополии Курно для процесса 4, если α_1^t и α_2^t имеют разные знаки, то такие же знаки будут у α_1^{t+1} и α_2^{t+1} . Если $0 < \lambda_i^{t+k} \leq \frac{1}{2}$ ($k = 1, 2, \dots$) и $q_2^{t+k-1} \neq 0$, то знаки α_1^{t+k} и α_2^{t+k} будут разными и совпадать со знаками α_1^t и α_2^t , соответственно. По леммам 3 и 1 процесс сходится.

Пусть $q_2^t = 0$ и $\lambda_1^{t+1} = \frac{1}{2}$. По предположению $\alpha_2^t = 2(x_2(q_{-2}^t) - q_2^t) = 2x_2(q_{-2}^t) < 0$. Поэтому по (4) $q_2^{t+1} = 0$. Но $\alpha_2^{t+1} < 0$, следовательно, $2(x_2(q_{-2}^{t+1}) - q_2^{t+1}) = 2x_2(q_{-2}^{t+1}) < 0$ и поэтому $q_2^{t+2} = 0$. Поскольку в данной ситуации $\alpha_1^{t+1} = 0$, то $2(x_1(q_{-1}^{t+1}) - q_1^{t+1}) = 0$. По предположению $x_1(q_{-1}^t) > 0$, поэтому $q_1^{t+1} > 0$ и $x_1(q_{-1}^{t+1}) > 0$. Тогда $\alpha_1^{t+2} = (1 - 2\lambda_1^{t+2})\alpha_1^{t+1} + q_2^{t+1} = (1 - 2\lambda_1^{t+2})\alpha_1^{t+1} = 0$ и $\alpha_2^{t+2} = \alpha_2^{t+1} + 2q_2^{t+1} - \lambda_1^{t+2}\alpha_1^{t+1} = 2x_2(q_{-2}^{t+1}) - \lambda_1^{t+2}\alpha_1^{t+1} = 2x_2(q_{-2}^{t+1}) < 0$. Таким образом, имея исходные предположения $\{1\} \in N_1^t$, $\{2\} \in N_2^t$, $\alpha_1^t > 0$, $\alpha_2^t < 0$, $q_2^t = 0$ и $\lambda_1^{t+1} = \frac{1}{2}$, далее получаем $\alpha_1^{t+1} = 0$, $\alpha_1^{t+2} = 0$ и $q_2^{t+1} = 0$, $q_2^{t+2} = 0$. Очевидно, если процесс будет продолжаться таким образом, то сходится не может.

Рассмотрим случай разнорефлексирующих агентов. Пусть, для определенности, первый агент рефлексируют по Курно, второй – по Штакельбергу и $\alpha_1^t > 0$, $\alpha_2^t < 0$. Покажем, что $\alpha_1^{t+1} > 0$ и $\alpha_2^{t+1} < 0$.

Для $\{2\} \in N_1^t$ это следует из системы равенств (18)–(19).

Для $\{2\} \in N_2^t$ при $0 < \lambda_1^{t+1} \leq \frac{1}{2}$, используя соотношение $\alpha_2^t = \frac{3}{2}(x_2(q_2^t) - q_2^t)$, по формулам (22) и (25) имеем при $q_2^t \neq 0$ или $\lambda_1^{t+1} \neq \frac{1}{2}$.

$$\alpha_1^{t+1} = (1 - \lambda_1^{t+1})\alpha_1^t - \lambda_1^{t+1}\alpha_1^t + q_2^t = (1 - 2\lambda_1^{t+1})\alpha_1^t + q_2^t > 0,$$

$$\alpha_2^{t+1} = \alpha_2^t + \frac{1}{2}q_2^t - \lambda_1^{t+1}\alpha_1^t + q_2^t = \frac{3}{2}x_2(q_2^t) - \lambda_1^{t+1}\alpha_1^t < 0.$$

Таким образом, в такой дуополии для процесса 4, если α_1^t и α_2^t имеют разные знаки, то такие же знаки будут у α_1^{t+1} и α_2^{t+1} . Если $0 < \lambda_1^{t+k} \leq \frac{1}{2}$, $0 < \lambda_2^{t+k} \leq \frac{2}{3}$, ($k = 1, 2, \dots$) (что означает $0 < \gamma_i^{t+k} \leq 1$), то знаки α_1^{t+k} и α_2^{t+k} будут разными и совпадать со знаками α_1^t и α_2^t , соответственно. По леммам 2 и 1 процессы сходятся.

Аналогично показывается, что при $q_2^t = 0$ и $\lambda_1^{t+1} = \frac{1}{2}$ процесс не сходится.

Пусть оба агента рефлексируют по Штакельбергу и $\alpha_1^t > 0, \alpha_2^t < 0$. Покажем, что $\alpha_1^{t+1} > 0$ и $\alpha_2^{t+1} < 0$.

Для $\{2\} \in N_1^t$ это следует из системы двух равенств вида (19).

Для $\{2\} \in N_2^t$ при $0 < \lambda_1^{t+1} \leq \frac{2}{3}$ по формуле (25) имеем при $q_2^t \neq 0$ или $\lambda_1^{t+1} \neq \frac{2}{3}$:

$$\alpha_1^{t+1} = \left(1 - \frac{1}{2}\lambda_1^{t+1}\right)\alpha_1^t - \lambda_1^{t+1}\alpha_1^t + q_2^t = \left(1 - \frac{3}{2}\lambda_1^{t+1}\right)\alpha_1^t + q_2^t > 0,$$

$$\alpha_2^{t+1} = \alpha_2^t + \frac{1}{2}q_2^t - \lambda_1^{t+1}\alpha_1^t + q_2^t = \frac{3}{2}x_2(q_2^t) - \lambda_1^{t+1}\alpha_1^t < 0.$$

Поэтому для процесса 4, если α_1^t и α_2^t имеют разные знаки, то такие же знаки будут у α_1^{t+1} и α_2^{t+1} . Если $0 < \lambda_i^{t+k} \leq \frac{2}{3}$, что означает $0 < \gamma_i^{t+k} \leq 1$ ($k = 1, 2, \dots$), то знаки α_1^{t+k} и α_2^{t+k} будут разными и совпадать со знаками α_1^t и α_2^t , соответственно. По леммам 2 и 1 процесс сходится.

Аналогично показывается, что при $q_2^t = 0$ и $\lambda_1^{t+1} = \frac{2}{3}$ процесс не сходится.

Итак, мы рассмотрели случаи, когда в текущий t -й момент времени у α_1^t и α_2^t разные знаки. Если в дуополии $\alpha_1^t \geq 0, \alpha_2^t \geq 0$, или $\alpha_1^t, \alpha_2^t \leq 0$, $\alpha_1^{t+1}, \alpha_2^{t+1} \leq 0$, то сходимость процесса 4 показывается с применением лемм 4 и 1. Поэтому остается рассмотреть случай $\alpha_1^t, \alpha_2^t \leq 0$, $\alpha_1^{t+1}, \alpha_2^{t+1} \geq 0$.

Пусть для данного случая один агент действует по Штакельбергу (для определенности первый), другой – по Курно. Другие возможности, когда оба агента ведущие или оба ведомые рассматриваются аналогично.

Пусть $\alpha_1^t, \alpha_2^t \leq 0$, $\alpha_1^{t+1}, \alpha_2^{t+1} \geq 0$.

Допустим $\{1, 2\} \in N_1^t$. По (22), (23) имеем:

$$\begin{aligned}\alpha_1^{t+1} &= \left(1 - \frac{3}{2}\lambda_1^{t+1}\right)\alpha_1^t - \lambda_2^{t+1}\alpha_2^t, \\ \alpha_2^{t+1} &= \left(1 - 2\lambda_2^{t+1}\right)\alpha_2^t - \lambda_1^{t+1}\alpha_1^t.\end{aligned}$$

Отсюда следует, что $\lambda_2^{t+1}(-\alpha_2^t) \geq \alpha_1^{t+1}$ и $\lambda_1^{t+1}(-\alpha_1^t) \geq \alpha_2^{t+1}$. Поэтому $\max\{\lambda_1^{t+1}, \lambda_2^{t+1}\} \max_{i \in \{1, 2\}} |\alpha_i^t| \geq \max_{i \in \{1, 2\}} |\alpha_i^{t+1}|$.

Пусть теперь $\{1, 2\} \in N_2^t$. По (24), (25) имеем:

$$\alpha_1^{t+1} = \alpha_1^t + \frac{3}{2}q_1^t + q_2^t,$$

$$\alpha_2^{t+1} = \alpha_2^t + 2q_2^t + q_1^t.$$

Так как $\{1, 2\} \in N_2^t$, то $\alpha_1^t + \frac{3}{2}q_1^t \leq 0$ или $q_1^t \leq -\frac{2}{3}\alpha_1^t$ и $\alpha_2^t + 2q_2^t \leq 0$ или $q_2^t \leq -\frac{1}{2}\alpha_2^t$. Поэтому $-\frac{1}{2}\alpha_2^t \geq \alpha_1^{t+1}$, $-\frac{2}{3}\alpha_1^t \geq \alpha_2^{t+1}$ и $\frac{2}{3} \max_{i \in \{1, 2\}} |\alpha_i^t| \geq \max_{i \in \{1, 2\}} |\alpha_i^{t+1}|$.

Пусть теперь $\{1\} \in N_1^t$, $\{2\} \in N_2^t$. Случай $\{1\} \in N_2^t$, $\{2\} \in N_1^t$ рассматривается аналогично. По (23), (24) имеем:

$$\alpha_1^{t+1} = \left(1 - \frac{3}{2}\lambda_1^{t+1}\right)\alpha_1^t + q_2^t,$$

$$\alpha_2^{t+1} = \alpha_2^t + 2q_2^t - \lambda_1^{t+1}\alpha_1^t.$$

Так как $\alpha_2^t + 2q_2^t \leq 0$, то $-\frac{1}{2}\alpha_2^t \geq \alpha_1^{t+1}$ и $-\lambda_1^{t+1}\alpha_1^t \geq \alpha_2^{t+1}$. Поэтому $\max \left\{ \lambda_1^{t+1}, \frac{1}{2} \right\} \max_{i \in \{1, 2\}} |\alpha_i^t| \geq \max_{i \in \{1, 2\}} |\alpha_i^{t+1}|$.

Итак, для случая $\alpha_1^t, \alpha_2^t \leq 0$, $\alpha_1^{t+1}, \alpha_2^{t+1} \geq 0$ в дуополии также имеет место неравенство $\mu^t \max_{i \in N} |\alpha_i^t| \geq \max_{i \in N} |\alpha_i^{t+1}|$, где $\mu^t \in (0; 1)$. Что указывает на сходимость процесса 4.

Утверждение 3 доказано.

Следствие 3. В дуополии (5)–(7) процесс 3 сходится к положению равновесия при любых начальных условиях $q^0 > 0$, независимо от того агент рефлексирует по Курно или по Штакельбергу.

Результат, представленный в утверждении 1, обобщает известные аналоги. Так в работах [24, 25, 29] частными его случаями являются достаточные условия сходимости вида $\gamma_i^t \in \left(0; \frac{2}{1+n}\right)$ для всех агентов в модели Курно, такие же условия для агентов с реакцией по Курно и $\gamma^t \in \left(0; \frac{1}{n}\right)$ для лидера в модели с одним лидером по Шта-

кельбергу, условия $\gamma_i^t \in \left(0; \frac{1}{n}\right)$ для всех агентов в модели олигополии

с одними лидерами. Для олигополии с произвольным числом лидеров аналогичные модельные исследования для сравнения нам не известны.

Результаты, сформулированные для дуополии в утверждении 2 и утверждении 3, полностью согласуются с известными аналогами и модельными экспериментами [2, 25, 29, 34–38]. Что может свидетельствовать о корректности применяемых методов и для олигополии при $n > 2$. Новыми здесь являются метод доказательства этих утверждений, а также выявленные особенности, когда в дуополии процесс 4 сходиться не может.

6. Заключение. В статье рассмотрена проблема достижения равновесия Нэша на рынке олигополии при отсутствии общего знания среди агентов с применением рефлексивных повторяющихся игр и моделей динамики коллективного поведения. Агенты, основываясь на наблюдении за текущими объемами выпуска конкурентов и размышлениях о наилучшем собственном действии с учетом наилучших откликов конкурентов, при повторении игры уточняют по модели коллективного поведения свои объемы выпуска, делая шаги в направлении текущего оптимального выпуска. Развитие динамики направляется выбором агентами величины шагов.

Научная новизна результатов проведенного исследования заключается в получении новых аналитических выражений для диапазонов величин текущих шагов агентов, при которых гарантируется сходимость моделей коллективного поведения к статичному равновесию Нэша. В нашем случае условия сходимости получены для олигополии с произвольным числом лидеров по Штакельбергу в классе линейных функций спроса и издержек агентов. Агенты при выборе величины шагов принимают во внимание текущие экономические ограничения по прибыльности и конкурентоспособности, поэтому длины шагов могут меняться во времени. Такая политика выбора шага позволяет каждому агенту максимизировать собственную прибыль на основе доступной ему информации, предполагая полное (совершенное) знание среди агентов. Также показано, что динамики сходятся при любых начальных объемах выпуска. Подробно обсуждается случай дуополии в сравнении с современными результатами. Приведены необходимые математические леммы, утверждения и их доказательства.

Полученные результаты свидетельствуют о том, что в условиях неполной информированности и отсутствии общего знания при определенном диапазоне значений действий агентов можно реализовать равновесные объемы выпуска при любом числе лидеров за счет введе-

ния рефлексующих агентов. С точки зрения теории принятия решений, введение рефлексующих агентов расширяет множество векторов действий, которые могут выбрать агенты, осуществляющие совместную деятельность, и позволяет существенным образом изменять состояние рынка. Полученные результаты могут иметь прикладное значение для понимания рефлексивного группового поведения агентов на современных конкурентных рынках.

Перспективным направлением дальнейших исследований представляется ослабление условий (26), (27) на индивидуальный выбор каждого агента и выработка единого для всей совокупности агентов условия на коллективный (групповой) выбор. Также актуален и перспективен поиск аналитических решений для нелинейных моделей рынка и известных типовых рефлексивных моделей принятия коллективных решений с лидерами на разных уровнях. Представляют несомненный интерес модели, которые бы хорошо интерпретировали содержательный смысл и психологию выбора агентом того или иного шага движения к своей текущей цели.

Литература

1. Shapiro C. Theories Oligopoly Behavior / Handbook of Industrial Organization // Elsevier. 1989. vol. 1. chapter 6. pp. 329–414.
2. Малишевский А.В. Качественные модели в теории сложных систем // М.: Наука. 1998. 528 с.
3. Novikov D.A., Chkhartishvili A.G. Reflexion and Control: Mathematical Models // Leiden: CRC Press. 2014. 298 p.
4. Новиков Д.А. Модели динамики психических и поведенческих компонент деятельности в коллективном принятии решений // Управление большими системами: М. ИПУ РАН. 2020. вып. 85. С.206–237.
5. Kalashnikov V.V., Bulavsky V.A., Kalashnykova N.I. Existence of the Nash-Optimal Strategies in the Meta-Game // Studies in Systems, Decision and Control. 2018. vol. 100. pp. 95–100.
6. Berger U., De Silva H., Fellner-Rohling G. Cognitive Hierarchies in the Minimizer Game // Journal of Economic Behavior and Organization. 2016. vol. 130. pp. 337–348.
7. Mallozzi L., Messalli R. Multi-Leader Multi-Follower Model with Aggregative Uncertainty // Games. 2017. vol. 8(3). pp. 1–14.
8. Cournot A. Researches into the Mathematical Principles of the Theory of Wealth // London: Hafner. 1960. (Original 1838). 235 p.
9. Nash J. Non-Cooperative Games // Annals of Mathematics. 1951. no. 54. pp. 286–295.
10. Novikov D.A., Chkhartishvili A.G. Mathematical Models of Informational and Strategic Reflexion: a Survey // Advances in Systems Science and Applications. 2014. no. 3. pp. 254–277.
11. The Handbook of Experimental Economics / Ed. by Kagel J. and Roth A. // Princeton: Princeton University Press. 1995. 744 p.

12. Wright J., Leyton-Brown K. Beyond Equilibrium: Predicting Human Behavior in Normal Form Games // Proceedings of Conference on Associations for the Advancement of Artificial Intelligence (AAAI-10). 2010. pp. 461–473.
13. Askar S., Simos T. Tripoly Stackelberg Game Model: One Leader Versus Two Followers // Applied Mathematics and Computation. 2018. vol. 328. pp. 301–311.
14. Askar S. On Complex Dynamics of Cournot-Bertrand Game with Asymmetric Market Information // Applied Mathematics and Computation. 2021. vol. 393(3) // <https://doi.org/10.1016/j.amc.2020.125823>.
15. Stackelberg H. Market Structure and Equilibrium / Transl. into English by Basin D., Urch L. & Hill. R. // New York. Springer. 2011. (Original 1934). 134 p.
16. Sherali H.D. Multiple Leader Stackelberg Model and Analysis // Operations Research. 1984. vol. 32. issue 2. pp. 390–404.
17. Julien L.A. On Noncooperative Oligopoly Equilibrium in the Multiple Leader – Follower Game // European Journal of Operational Research. 2017. vol. 256. issue 2. pp. 650–662.
18. Solis C.U., Clempner J.B., Poznyak A.S. Modeling Multileader – Follower Noncooperative Stackelberg Games // Cybernetics and Systems. 2017. vol. 47. no. 8. pp. 650–673.
19. Geras'kin M.I. Approximate Calculation of Equilibria in the Nonlinear Stackelberg Oligopoly Model: A Linearization Based Approach // Automation and Remote Control. 2020. vol. 81 no. 9. pp. 1659–1678.
20. Castiglioni M., Marchesi A., Gatti N. Committing to Correlated Strategies with Multiple Leaders // Artificial Intelligence. 2021. vol. 300 // <https://doi.org/10.1016/j.artint.2021.103549>.
21. Zewde A.B., Kassa S.M. Multilevel Multi-Leader Multiple-Follower Games with Nonseparable Objectives and Shared Constraints // Computational Management Science. 2021. vol. 18(4). pp. 455–475.
22. Zewde A.B., Kassa S.M. Hierarchical Multilevel Optimization with Multiple-Leaders Multiple-Followers Setting and Nonseparable Objectives // RAIRO –Operations Research. 2021. vol. 55(5). pp. 2915–2939.
23. Wu R., Van Gorder R.A. Nonlinear Dynamics of Discrete Time Multi-Level Leader-Follower Games // Applied Mathematics and Computation. 2018. vol. 320. pp. 240–250.
24. Algazin G.I., Algazina D.G. Collective Behavior in the Stackelberg Model under Incomplete Information // Automation and Remote Control. 2017. vol. 78. no. 9. pp. 1619–1630.
25. Algazin G.I., Algazina D.G. Reflexive Processes and Equilibrium in an Oligopoly Model with a Leader // Automation and Remote Control. 2020. vol. 81. no. 7. pp. 1258–1270.
26. Алгазина Д.Г., Алгазин Г.И. Модельные исследования сетевого взаимодействия на конкурентных рынках с нефиксированными ролями участников // Барнаул: Изд-во Алтайского ун-та. 2015. 146 с.
27. Опойцев В.И. Равновесие и устойчивость в моделях коллективного поведения // М.: Наука. 1977. 248 с.
28. Yoo T.-H., Ko W., Rhee C.-H., Park J.-K. The Incentive Announcement Effect of Demand Response on Market Power Mitigation in the Electricity Market // Renewable and Sustainable Energy Reviews. 2017. vol. 76. pp. 545–554.
29. Algazin G.I., Algazina Yu.G. Reflexive Dynamics in the Cournot Oligopoly under Uncertainty // Automation and Remote Control. 2020. vol. 81. no. 2. pp. 345–359.
30. Alcantara-Jiménez G., Clempner J.B. Repeated Stackelberg Security Games: Learning with Incomplete State Information // Reliability Engineering and System Safety. 2020. vol. 195 // <https://doi.org/10.1016/j.ress.2019.106695>.

31. Wei L., Wang H., Wang J., Hou J. Dynamics and Stability Analysis of a Stackelberg Mixed Duopoly Game with Price Competition in Insurance Market // *Discrete Dynamics in Nature and Society*. 2021. vol. 2021. pp. 1–18.
32. Fedyanin D.N. Monotonicity of Equilibriums in Cournot Competition with Mixed Interactions of Agents and Epistemic Models of Uncertain Market // *Procedia Computer Science*. 2021. vol. 186(3). pp. 411–417.
33. Geras'kin M.I., Chkhartishvili A.G. Analysis of Game-Theoretic Models of an Oligopoly Market under Constrains on the Capacity and Competitiveness of Agents // *Automation and Remote Control*. 2017. vol. 78. no. 11. pp. 2025–2038.
34. Askar S.S., Elettreybc M.F. The Impact of Cost Uncertainty on Cournot Oligopoly Games // *Applied Mathematics and Computation*. 2017. vol. 312. pp. 169–176.
35. Алгазин Г.И., Алгазина Д.Г. Динамика рефлексивного коллективного поведения в модели олигополии с лидерами // *Известия Алтайского государственного университета*. 2018. № 1(99). С. 64–68.
36. Ueda M. Effect of Information Asymmetry in Cournot Duopoly Game with Bounded Rationality // *Applied Mathematics and Computation*. 2019. vol. 362 // <https://doi.org/10.1016/j.amc.2019.06.049.124535>
37. Long J., Huang H. A Dynamic Stackelberg-Cournot Duopoly Model with Heterogeneous Strategies through One-Way Spillovers // *Discrete Dynamics in Nature and Society*. 2020. vol. 2. pp. 1–11.
38. Elsadany A.A. Dynamics of a Cournot Duopoly Game with Bounded Rationality Based on Relative Profit Maximization // *Applied Mathematics and Computation*. 2017. vol. 294. pp. 253–263.

Алгазин Геннадий Иванович — д-р физ.-мат. наук, профессор, кафедра теоретической кибернетики и прикладной математики, ФГБОУ ВО "Алтайский государственный университет" (АлтГУ). Область научных интересов: математическое моделирование социально-экономических процессов, теория игр, исследование операций, математическая концепция системного компромисса, информационное управление. Число научных публикаций — 180. algaz46@yandex.ru; Ленина, 61, 656049, Барнаул, Россия; р.т.: 7(385)229-81-51.

Алгазина Дарья Геннадьевна — канд. техн. наук, доцент, кафедра цифровых технологий и бизнес-аналитики, ФГБОУ ВО "Алтайский государственный университет" (АлтГУ). Область научных интересов: моделирование конкурентных рынков, управление в организационных системах, цифровые технологии франчайзинга. Число научных публикаций — 40. darya.algazina@mail.ru; Ленина, 61, 656049, Барнаул, Россия; р.т.: +7(385)229-65-46.

G. ALGAZIN, D. ALGAZINA
**MODELING THE DYNAMICS OF COLLECTIVE BEHAVIOR IN A
REFLEXIVE GAME WITH AN ARBITRARY NUMBER OF
LEADERS**

Algazin G., Algazina D. Modeling the Dynamics of Collective Behavior in a Reflexive Game with an Arbitrary Number of Leaders.

Abstract. An oligopoly with an arbitrary number of Stackelberg leaders under incomplete, asymmetrical agents' awareness and inadequacy of their predictions of competitors' actions is considered. Models of individual decision-making processes by agents are studied. The reflexive games theory and collective behavior theory are the theoretical basis for construction and analytical study process models. They complement each other in that reflexive games allow using the collective behavior procedures and the results of agents' reflections, leading to a Nash equilibrium. The dynamic decision-making process considered repeated static games on a range of agents' feasible responses to the expected actions of the environment, considering current economic restrictions and competitiveness in each game. Each reflexive agent in each game calculates its current goal position and changes its state, taking steps towards the current position of the goal to obtain positive profit or minimize losses. Sufficient conditions for the convergence of processes in discrete time for the case of linear costs of agents and linear demand is the main result of this work. New analytical expressions for the agents' current steps' ranges guarantee the convergence of the collective behavior models to static Nash equilibrium is obtained. That allows each agent to maximize their profit, assuming common knowledge among the agents. The processes when the agent chooses their best response are also analyzed. The latter may not give converging trajectories. The case of the duopoly in comparison with modern results is discussed in detail. Necessary mathematical lemmas, statements, and their proofs are presented.

Keywords: reflexive games, oligopoly, Stackelberg leader, incomplete awareness, collective behavior, equilibrium, convergence conditions.

Algazin Gennady — Ph.D., Dr.Sci., Professor, Department of theoretical cybernetics and applied mathematics, Altai State University. Research interests: mathematical modeling of socio-economic processes, game theory, operations research, the mathematical concept of systemic compromise, informational control. The number of publications — 180. algaz46@yandex.ru; 61, Lenin St., 656049, Barnaul, Russia; office phone: 7(385)229-81-51.

Algazina Daria — Ph.D., Associate professor, Department of digital technologies and business analytics, Altai State University. Research interests: modeling of competitive markets, organization control, digital technologies of franchising. The number of publications — 40. darya.algazina@mail.ru; 61, Lenin St., 656049, Barnaul, Russia; office phone: +7(385)229-65-46.

References

1. Shapiro C. Theories Oligopoly Behavior. Handbook of Industrial Organization. Elsevier. 1989. vol. 1. chapter 6. pp. 329–414.
2. Malishevski A.V. Kachestvennyye modeli v teorii slozhnykh sistem [Qualitative Models in the Theory of Complex Systems]. Moscow: Nauka, 1998. 528 p. (In Russ.).
3. Novikov D.A., Chkhartishvili A.G. Reflexion and Control: Mathematical Models. Leiden: CRC Press. 2014. 298 p.

4. Novikov D.A. [Dynamics Models of Mental and Behavioral Components of Activity in Collective Decision-Making]. *Upravleniye bol'shimi sistemami: Sb. nauch. tr. [Large-Scale Systems Control: Collected papers]*. Moscow: Institute of Control Sciences of Russian Academy of Sciences, 2020. no. 85. pp. 206–237. (In Russ.).
5. Kalashnikov V.V., Bulavsky V.A., Kalashnykova N.I. Existence of the Nash-Optimal Strategies in the Meta-Game. *Studies in Systems, Decision and Control*. 2018. vol. 100. pp. 95–100.
6. Berger U., De Silva H., Fellner-Rohling G. Cognitive Hierarchies in the Minimizer Game. *Journal of Economic Behavior and Organization*. 2016. vol. 130. pp. 337–348.
7. Mallozzi L., Messalli R. Multi-Leader Multi-Follower Model with Aggregative Uncertainty. *Games*. 2017. vol. 8(3). pp. 1–14.
8. Cournot A. *Researches into the Mathematical Principles of the Theory of Wealth*. London: Hafner. 1960. (Original 1838). 235 p.
9. Nash J. Non-Cooperative Games. *Annals of Mathematics*. 1951. no. 54. pp. 286–295.
10. Novikov D.A., Chkhartishvili A.G. Mathematical Models of Informational and Strategic Reflexion: a Survey. *Advances in Systems Science and Applications*. 2014. no. 3. pp. 254–277.
11. *The Handbook of Experimental Economics*. Ed. by Kagel J. and Roth A. Princeton: Princeton University Press. 1995. 744 p.
12. Wright J., Leyton-Brown K. Beyond Equilibrium: Predicting Human Behavior in Normal Form Games. *Proceedings of Conference on Associations for the Advancement of Artificial Intelligence (AAAI-10)*. 2010. pp. 461–473.
13. Askar S., Simos T. Tripoly Stackelberg Game Model: One Leader Versus Two Followers. *Applied Mathematics and Computation*. 2018. vol. 328. pp. 301–311.
14. Askar S. On Complex Dynamics of Cournot-Bertrand Game with Asymmetric Market Information. *Applied Mathematics and Computation*. 2021. vol. 393(3). 125823.
15. Stackelberg H. *Market Structure and Equilibrium*. Transl. into English by Basin D., Urch L. & Hill. R. New York: Springer. 2011. (Original 1934). 134 p.
16. Sherali H.D. Multiple Leader Stackelberg Model and Analysis. *Operations Research*. 1984. vol. 32. issue 2. pp. 390–404.
17. Julien L.A. On Noncooperative Oligopoly Equilibrium in the Multiple Leader – Follower Game. *European Journal of Operational Research*. 2017. vol. 256. issue 2. pp. 650–662.
18. Solis C.U., Clempner J.B., Poznyak A.S. Modeling Multi Leader – Follower Noncooperative Stackelberg Games. *Cybernetics and Systems*. 2017. vol. 47. no. 8. pp. 650–673.
19. Geras'kin M.I. Approximate Calculation of Equilibria in the Nonlinear Stackelberg Oligopoly Model: A Linearization Based Approach. *Automation and Remote Control*. 2020. vol. 81 no. 9. pp. 1659–1678.
20. Castiglioni M., Marchesi A., Gatti N. Committing to Correlated Strategies with Multiple Leaders. *Artificial Intelligence*. 2021. vol. 300. 103549.
21. Zewde A.B., Kassa S.M. Multilevel Multi-Leader Multiple-Follower Games with Nonseparable Objectives and Shared Constraints. *Computational Management Science*. 2021. vol. 18(4). pp. 455–475.
22. Zewde A.B., Kassa S.M. Hierarchical Multilevel Optimization with Multiple-Leaders Multiple-Followers Setting and Nonseparable Objectives. *RAIRO – Operations Research*. 2021. vol. 55(5). pp. 2915–2939.
23. Wu R., Van Gorder R.A. Nonlinear Dynamics of Discrete Time Multi-Level Leader-Follower Games. *Applied Mathematics and Computation*. 2018. vol. 320. pp. 240–250.

24. Algazin G.I., Algazina D.G. Collective Behavior in the Stackelberg Model under Incomplete Information. Automation and Remote Control. 2017. vol. 78. no. 9. pp. 1619–1630.
25. Algazin G.I., Algazina D.G. Reflexive Processes and Equilibrium in an Oligopoly Model with a Leader. Automation and Remote Control. 2020. vol. 81. no. 7. pp. 1258–1270.
26. Algazina D.G., Algazin G.I. Model'nyye issledovaniya setevogo vzaimodeystviya na konkurentnykh rynkakh s nefiksirovannymi rolyami uchastnikov [Model Studies of Network Interaction in Competitive Markets with Unfixed Roles of Participants]. Barnaul: Publishing House of Altay State University, 2015. 146 p. (In Russ.).
27. Opoitsev V.I. Ravnovesiye i ustoychivost' v modelyakh kolektivnogo povedeniya [Equilibrium and Stability in Models of Collective Behavior]. Moscow: Nauka, 1977. 248 p. (In Russ.).
28. Yoo T.-H., Ko W., Rhee C.-H., Park J.-K. The Incentive Announcement Effect of Demand Response on Market Power Mitigation in the Electricity Market. Renewable and Sustainable Energy Reviews. 2017. vol. 76. pp. 545–554.
29. Algazin G.I., Algazina Yu.G. Reflexive Dynamics in the Cournot Oligopoly under Uncertainty. Automation and Remote Control. 2020. vol. 81. no. 2. pp. 345–359.
30. Alcantara-Jiménez G., Clempner J.B. Repeated Stackelberg Security Games: Learning with Incomplete State Information. Reliability Engineering and System Safety. 2020. vol. 195. 106695.
31. Wei L., Wang H., Wang J., Hou J. Dynamics and Stability Analysis of a Stackelberg Mixed Duopoly Game with Price Competition in Insurance Market. Discrete Dynamics in Nature and Society. 2021. vol. 2021. pp. 1–18.
32. Fedyanin D.N. Monotonicity of Equilibriums in Cournot Competition with Mixed Interactions of Agents and Epistemic Models of Uncertain Market. Procedia Computer Science. 2021. vol. 186(3). pp. 411–417.
33. Geras'kin M.I., Chkhartishvili A.G. Analysis of Game-Theoretic Models of an Oligopoly Market under Constrains on the Capacity and Competitiveness of Agents. Automation and Remote Control. 2017. vol. 78. no. 11. pp. 2025–2038.
34. Askar S.S., Elettreby M.F. The Impact of Cost Uncertainty on Cournot Oligopoly Games. Applied Mathematics and Computation. 2017. vol. 312. pp. 169–176.
35. Algazin G.I., Algazina D.G. [Dynamics of Reflexive Collective Behavior in the Oligopoly Model with Leaders]. Izvestiya Altayskogo gosudarstvennogo universiteta – Izvestiya of Altay State University. 2018. vol. 99. no. 1. pp. 64–68. (In Russ.).
36. Ueda M. Effect of Information Asymmetry in Cournot Duopoly Game with Bounded Rationality. Applied Mathematics and Computation. 2019. vol. 362. 124535.
37. Long J., Huang H. A Dynamic Stackelberg-Cournot Duopoly Model with Heterogeneous Strategies through One-Way Spillovers. Discrete Dynamics in Nature and Society. 2020. vol. 2. pp. 1–11.
38. Elsadany A.A. Dynamics of a Cournot Duopoly Game with Bounded Rationality Based on Relative Profit Maximization. Applied Mathematics and Computation. 2017. vol. 294. pp. 253–263.

М.В. БОБЫРЬ, А.Е. АРХИПОВ, С.В. ГОРБАЧЕВ, Ц. ЦАО, С.Б. БХАТТАЧАРЬЯ
**НЕЧЕТКО-ЛОГИЧЕСКИЕ МЕТОДЫ В ЗАДАЧЕ
ДЕТЕКТИРОВАНИЯ ГРАНИЦ ОБЪЕКТОВ**

Бобырь М.В., Архипов А.Е., Горбачев С.В., Цао Ц., Бхаттачарья С.Б. **Нечетко-логические подходы в задаче детектирования границ объектов.**

Аннотация. Рассматривается задача уменьшения вычислительной сложности методов выделения контуров на изображениях. Решение поставленной задачи достигается модификацией детектора Канни двумя нечетко-логическими методами, позволяющими сократить число проходов по исходному изображению: в первом случае, путем исключения двух проходов, связанных с определением наличия соседства претендующего на границу пикселя со смежными в рамке размером 3×3 , а во втором случае, исключением операции определения угла направления градиента путем формирования данной величины комбинацией нечетких правил. Целью работы является повышение производительности вычислительных операций в процессе детектирования границ объектов путем уменьшения числа проходов по исходному изображению. Интеллектуализация процесса детектирования границ осуществляется частичным повтором вычислительных операций, используемых в детекторе Канни, с дальнейшей заменой наиболее сложных вычислительных процедур. В предлагаемых методах после определения величины градиента и угла его направления осуществляется фазификация восьми входных переменных, в качестве которых используется разность градиентов между центральной и смежными ячейками в рамке размером 3×3 . Затем строится база нечетких правил. В первом методе в зависимости от угла направления градиента используются четыре нечетких правила и исключается один проход. Во втором методе шестнадцать нечетких правил сами задают угол направления градиента, при этом исключается два прохода вдоль изображения. Разность градиентов между центральной ячейкой и смежными ячейками позволяет учитывать форму распределения градиента. Затем на основе метода центра тяжести осуществляется дефазификация результирующей переменной. Дальнейшее использование нечетких α -срезов позволяет осуществить бинаризацию результирующего изображения с выделением на нем границ объектов. Для оценивания вычислительной скорости работы предложенных нечетких методов детектирования границ в среде Microsoft Visual Studio было разработано программное обеспечение. Представленные экспериментальные результаты показали, что уровень шума зависит от величины α -среза и параметров меток трапециевидных функций принадлежности. Ограничением двух методов является использование кусочно-линейных функций принадлежности. Экспериментальные исследования работоспособности предложенных методов детектирования контуров показали, что время первого нечеткого метода на 18% быстрее по сравнению с детектором Канни и на 2% по отношению ко второму нечеткому методу. Однако при визуальной оценке установлено, что второй нечеткий метод лучше определяет границы объектов.

Ключевые слова: нечёткая логика, детектор Канни, выделение границ, оператор Собеля, центр тяжести.

1. Введение. Границы объектов на цифровых изображениях предоставляют важную информацию, необходимую для анализа и дальнейших интерпретаций, таких как процедур как сегментация, распознавание лиц, распознавание объектов, отслеживание, поиск изоб-

ражений или стереозрение и т.д. [1, 2]. Чтобы извлечь контур объекта на изображении необходимо получить полную информацию о границах, поэтому обнаружение границ является неотъемлемой частью обработки изображения [3]. При этом интенсивность изменения градиента на изображении позволяет оценивать внезапные изменения яркости на изображении и тем самым детектировать границы на нем. Обнаружение границ включает в себя ряд вычислительных процедур, например, вычисление производной функции интенсивности изображения в заданном местоположении пикселя. Если величина производной функции интенсивности изображения относительно высока, то пиксель классифицируется как краевой пиксель. Важным свойством такого подхода обнаружения границ является способность выделять точную грань границы с хорошей ориентацией на рассматриваемом изображении.

За последние два десятилетия было опубликовано много статей по обнаружению границ. Существует достаточное число подходов и способов обнаружения границ. Большинство из них базируются на процедуре вычисления величины градиента пикселя в качестве меры обнаружения интенсивности границ [4]. Другие методы обнаружения границ, такие как, фильтрация Прюитта, лапласианская фильтрация Гаусса, операторы на основе моментов, оператор Шена и Кастанана, детектор Канни и Де-Риша, [5], фильтрация Собеля [6] используют различные оценки дискретного приближения производной функции. Но некоторые общие проблемы этих методов заключаются в большом объеме вычислений чувствительности к шуму. Один из подходов для обнаружения границ представлен в статье [7]. Некоторые детекторы, основанные на оптимизации, представлены в научных работах [8, 9]. Использование статистических методов для детектирования границ проиллюстрировано в статье [10]. Другие интеллектуальные подходы такие как, генетические алгоритмы [11], радиально-базисные нейронные сети [12], байесовский подход [13], алгоритмы, основанные на остаточном анализе [14] и ортогональных проекциях [15], также широко применяются в задачах выделения контуров на изображениях. Некоторые авторы пытались изучить влияние шума на изображениях на производительность краевых детекторов [16], предлагая детектор границ, в котором пороговое значение выполняется с использованием статистических принципов.

В исследовательской статье [17] рассматривается метод построения обнаружения границ на основе клеточного автомата с использованием нечеткой эвристической функции, который приводит к лучшей

производительности для некоторых линейных правил и подходящих параметров, на основе метода оптимизации роя частиц.

В научной статье [18] описана стратегия выявления границ опухоли головного мозга на МРТ-изображениях пациента. Метод включает в себя некоторые функции шумоподавления, улучшения баланса контраста (BCET) и кластеризации Fuzzy c-Means (FCM). Для обнаружения тонких границ применялся детектор Канни, что обеспечивает некоторую устойчивость к шуму.

Путем экспериментального сравнения различных методов обнаружения границ цифрового изображения, можно сделать вывод, что точность сегментации описанных методов в ряде случаев остается недостаточно высокой относительно экспертных оценок.

Алгоритм Канни хорошо известен как оптимальный метод обнаружения границ. Он работает по трем основным принципам: низкий уровень ошибок, хорошая локализация граничных точек и четкий ответ на наличие границы. Чтобы улучшить старые методы обнаружения границ, Канни предложил в своем алгоритме два новых подхода: немаксимальное подавление и двойное определение порога для выбора граничных точек. Из-за плохого освещения границы областей на изображении могут стать расплывчатыми, создавая неточности в градиентном изображении. Однако эти два порога, используемые для сегментации градиентного изображения, устанавливаются экспериментально для каждого изображения.

Метод адаптивного обнаружения границ, основанный на операторе Канни, был представлен в работе [19]. Для определения пороговых значений в нем использовался метод пороговой обработки Оцу. В статье [20] авторы предложили метод обнаружения границ, основанный на нечетких рассуждениях, учитывающих зрительные характеристики человека. Сяо и др. в научном исследовании [21] предложили улучшенную версию детектора границ Канни, специально разработанную для изображений, искаженных Гауссовским шумом. В научном исследовании [22] детектор краев Канни использовался для сегментации медицинских изображений.

Многие авторы использовали различные модификации этого метода. В [23] предложен алгоритм, основанный на концепции нечетких множеств типа 2 с целью обработки неопределенностей, который автоматически выбирает пороговые значения, необходимые для сегментации градиентного изображения.

Гибридизация методов детектирования границ на МРТ-изображениях представлена Shah в его научной статье [24]. Еще одним

примером модификации метода является преобразование Шерлета с цветовым кодированием [25].

Современные алгоритмы, основанные на обучении, обнаруживают границы с помощью контролируемых моделей и созданных вручную функций. Например, Д. Ну и др. [26] используют иерархическую модель для поиска многомасштабных объектов, объединенных закрытым условным случайным полем. Ж. Не и др. [27, 28] предлагают структуру двунаправленной каскадной сети (BDCN) для обнаружения ребер в разных масштабах. Они обучают сеть, используя другие помеченные ребра для каждого масштаба.

Новейшие алгоритмы обнаружения границ, основанные на глубоком обучении [29-37], фокусируются на точном обнаружении границ объектов, которые могут предоставлять семантические подсказки для дальнейшей обработки, такие как обнаружение объектов, сегментация и отслеживание [38-46].

Из-за плохого освещения, низкого качества изображения или других возможных факторов граница между различными областями изображения может быть нечеткой. Это делает большинство граничных точек неопределенными, что приводит к неправильному определению границ существующими методами. Таким образом, разработка методов обнаружения границ изображений, особенно с разрывами [47], является актуальной задачей.

Согласно [48], степень нечеткости изображения является одним из ключевых факторов, влияющих на производительность определения пороговых значений сегментации изображения. Очевидно, неопределенность, присутствующая в изображении, делает неопределенным и его градиентное изображение. В таких случаях процесс выбора пороговых значений из гистограммы градиентного изображения для обнаружения границ становится затруднительным, что требует разработки новой модели детектора выделения границ, основанной на нечеткой логике [49-54].

2. Цель исследования. Описанные выше методы базируются на определении изменения градиентов пикселей и выполняются за довольно продолжительное время, что вызвано выполнением большого количества последовательных процедур, необходимых для поиска границ. В то же время методы, базирующиеся на вейвлет-преобразованиях и нейронных сетях с обучением, хоть и дают лучшие результаты, требуют для своей реализации специальные вычислительные устройства, в которых при увеличении количества входных переменных возможно возникновение ошибки «проклятие размерности».

Таким образом, в представленной работе решается задача поиска границ, ограничивающих области одного цвета, за два этапа. На первом этапе частично используются операции детектора Канни. На втором этапе в первой модификации используется набор из четырех нечетких правил для каждого направления угла градиента, которые позволяют с использованием метода центра тяжести осуществить детектирование контуров на изображениях. Во второй модификации расчет угла градиента не требуется, так как разработанная структура базы знаний, состоящая из шестнадцати нечетких правил, самостоятельно формирует углы распределения градиента. И также с помощью центра тяжести осуществляется выделение границ. При этом в первой модификации по сравнению с детектором Канни исключаются два прохода по изображению, во второй модификации три прохода.

Целью исследования является уменьшение времени детектирования границ объектов на фото - видеоизображениях, за счет уменьшения вычислительной сложности применяемых методов, при сохранении качества детектированных границ. При этом экспериментальная оценка включает среднюю оценку времени выделения контуров по ста изображениям. Для качества выделения границ использовалась метрика Прэтта.

3. Метод детектирования границ Канни. Детектор Канни состоит из следующих вычислительных операций.

Шаг 1. Преобразование пикселей в градации серого. В ходе этой операции значение интенсивности каждого пикселя определяется следующей формулой [55]

$$I_{x,y} = 0,299R_{x,y} + 0,587G_{x,y} + 0,114B_{x,y}, \quad (1)$$

где $I_{x,y}$ – интенсивность яркости градации серого в пикселе ($I_{x,y} \in [0, 255]$) с координатами вдоль оси абсцисс ($x=1..w$) и вдоль оси ординат ($y=1..h$); w – ширина изображения в пикселях; h – высота изображения в пикселях; R – значение интенсивности красного цвета в пикселе; G – значение интенсивности зелёного цвета в пикселе; B – значение интенсивности синего цвета в пикселе.

Шаг 2. Сглаживание. Размытие изображения фильтром Гаусса, преобразованного в градации серого, с целью удаления шума вычисляется путем поэлементного умножения элементов $I_{x,y}$ матрицы I , и элементов $H_{x,y}$ матрицы Гаусса H с размером окна 5×5 :

$$K_{x,y} = \sum_{x=-2}^2 \sum_{y=-2}^2 \frac{1}{b} H_{x,y} I_{x,y}, \quad (2)$$

где $K_{x,y}$ – элементы значений интенсивности для каждого пикселя после сглаживания; b – коэффициент нормировки, равен сумме элементов матрицы H ($b=159$); H – матрица Гаусса

$$H = \begin{bmatrix} 2 & 4 & 5 & 4 & 2 \\ 4 & 9 & 12 & 9 & 4 \\ 5 & 12 & 15 & 12 & 5 \\ 4 & 9 & 12 & 9 & 4 \\ 2 & 4 & 5 & 4 & 2 \end{bmatrix}.$$

Шаг 3. Поиск градиента. На данной операции вычисляются приближенные значения градиента яркости изображения по всему изображению путем свертки изображения целочисленными матрицами в вертикальном и горизонтальном направлениях по формулам [56]:

$$GX_{x,y} = \sum_{y=-1}^1 \sum_{x=-1}^1 Ver_{x,y} K_{x,y}, \quad (3)$$

$$GY_{x,y} = \sum_{y=-1}^1 \sum_{x=-1}^1 Gor_{x,y} K_{x,y}, \quad (4)$$

где $Ver_{x,y}$, $Gor_{x,y}$ элементы вертикальной Ver и горизонтальной Gor матриц Собеля, соответственно:

$$Ver = \begin{bmatrix} -1 & 0 & 1 \\ -2 & 0 & 2 \\ -1 & 0 & 1 \end{bmatrix}, \quad Gor = \begin{bmatrix} -1 & -2 & -1 \\ 0 & 0 & 0 \\ 1 & 2 & 1 \end{bmatrix}.$$

После операции свертки изображения осуществляется расчет градиента для каждого пикселя по формуле [57]:

$$G_{x,y} = \sqrt{GX_{x,y}^2 + GY_{x,y}^2}, \quad (5)$$

Шаг 4. Определение угла направления градиента, который после округления принимает только четыре значения 0° , 45° , 90° и 135° градусов.

$$\Theta_{x,y} = \text{round} \left(\text{atan} \left(\frac{GY_{x,y}}{GX_{x,y}} \right) \right), \quad (6)$$

где round – функция округления результата $\Theta_{x,y}$ до целого; atan – функция определения арктангенса отношения градиентов.

Шаг 5. Подавление немаксимумов. В зависимости от значения угла направления градиента осуществляется проверка значений градиентов в смежных ячейках рамки 3×3 относительно её центральной точки:

Правило 1. Если $\Theta = 0^\circ$ (направление $\uparrow$). Центральная ячейка является границей, если её градиент больше, чем у верхней и нижней ячеек:

$$\text{edge}_{x,y} = \begin{cases} 1, & \text{если } (G_{x,y} > G_{x,y-1}) \& (G_{x,y} > G_{x,y+1}) \\ 0, & \text{иначе} \end{cases}. \quad (7)$$

Правило 2. Если $\Theta = 45^\circ$ (направление $\nearrow$). Центральная ячейка является границей, если её градиент больше, чем у верхней левой и нижней правой ячеек:

$$\text{edge}_{x,y} = \begin{cases} 1, & \text{если } (G_{x,y} > G_{x-1,y-1}) \& (G_{x,y} > G_{x+1,y+1}) \\ 0, & \text{иначе} \end{cases}. \quad (8)$$

Правило 3. Если $\Theta = 90^\circ$ (направление $\rightarrow$). Центральная ячейка является границей, если её градиент больше, чем у левой и правой ячеек:

$$\text{edge}_{x,y} = \begin{cases} 1, & \text{если } (G_{x,y} > G_{x-1,y}) \& (G_{x,y} > G_{x+1,y}) \\ 0, & \text{иначе} \end{cases}. \quad (9)$$

Правило 4. Если $\Theta = 135^\circ$ (направление $\searrow$). Центральная ячейка является границей, если её градиент больше, чем у нижней левой и верхней правой ячеек:

$$edge_{x,y} = \begin{cases} 1, & \text{если } (G_{x,y} > G_{x-1,y+1}) \& (G_{x,y} > G_{x+1,y-1}) \\ 0, & \text{иначе} \end{cases} \quad (10)$$

Поясним на примере операцию подавления немаксимумов. На рисунках 1 и 2 представлена рамка размером 5×5 со значениями градиента в каждой ячейке и графическим обозначением его угла. Почти все пиксели имеют угол наклона градиента равный 0° . Следовательно, наличие границы в этой точке определяется в зависимости от величины градиента в ячейках, расположенных сверху ($G_{x,y-1}=5$) и снизу ($G_{x,y+1}=2$) относительно центральной ($G_{x,y}=6$). Значение градиента в центральной ячейке больше смежных. Следовательно, данный пиксель является границей. На рисунке 2 красным шрифтом выделены пиксели, являющиеся границей.

3 ↑	2 ↖	2 ↑	2 ↖	2 ↖
4 ↑	3 ↑	5 ↑	7 ↑	7 ↑
6 ↑	7 ↑	6 ↑	4 ↑	5 ↑
4 ↑	5 ↑	2 ↑	3 ↑	4 ↑
2 ↗	3 ↗	2 ↗	1 ↗	1 ↗

Рис. 1. До операции подавления немаксимумов

0	0	0	0	0
0	0	0	1	1
1	1	1	0	0
0	0	0	0	0
0	0	0	0	0

Рис. 2. После операции подавления немаксимумов

Шаг 6. Пороговая фильтрация. На этой операции в зависимости от пороговых значений для каждого пикселя уточняется, является ли он границей на изображении или нет. В зависимости от условия формулы (11) принимается решение о достоверности границы:

$$edge'_{x,y} = \begin{cases} 1, & \text{если } G_{x,y} > G_{\max} \cdot T_h \\ 0.5, & \text{если } (G_{x,y} < G_{\max} \cdot T_h) \& (G_{x,y} > G_{\max} \cdot T_l) \\ 0, & \text{если } G_{x,y} < G_{\max} \cdot T_l \end{cases} \quad (11)$$

где T_h , T_l – пороговые значения; $G_{\max} = \max_{x=1..w, y=1..h} (G_{x,y})$, – максимальное значение градиента в изображении размером $w \times h$.

Второе условие формулы (11), при котором значение градиента находится в диапазоне между двумя порогами, уточняется на седьмом шаге, для которого требуется дополнительный проход вдоль изображения.

Шаг 7. Уточнение промежуточной границы. На данной операции уточняется значение переменной, вычисленной по формуле (11) в случае, когда $edge'_{x,y} = 0,5$.

$$edge''_{x,y} = \begin{cases} 1, & \text{если } (G_{x-1,y-1}=1) \parallel (G_{x,y-1}=1) \parallel (G_{x+1,y-1}=1) \parallel (G_{x-1,y}=1) \parallel \\ & \parallel (G_{x+1,y}=1) \parallel (G_{x-1,y+1}=1) \parallel (G_{x,y+1}=1) \parallel (G_{x+1,y+1}=1) , \\ 0 & \text{иначе} \end{cases}, \quad (12)$$

где $\parallel$ – знак, обозначающий операцию логического ИЛИ.

То есть анализируемый пиксель является границей, если он соприкасается с границей по одному из возможных восьми направлений.

Недостатки детектора Канни. Во-первых, определение пороговых значений (шаг 6) является одним из недостатков фильтра Канни, так как при высоком пороге T_h могут игнорироваться слабые края, у которых значения градиентов в смежных точках не особенно отличаются от величины градиента центральной ячейки. При низком пороге T_l может выделяться больше границ, что приводит к появлению дополнительного шума на изображении. Во-вторых, фильтр Канни выделяет границы только в случае, если градиент имеет выпуклую форму (рисунок 3а). Изменение остальных видов границ фильтр Канни не детектирует. Например, на рисунке 3а представлено распределение интенсивностей градиаций серого, вычисленных по формуле (1). На рисунке 3б распределение градиентов, определенных по формуле (5), при этом угол направления градиента был вычислен по формуле (6) и составил $\Theta=0^\circ$. Поэтому по формуле (7) осуществляется сравнение значений градиентов центральной ячейки ($G_{x,y}=209$) с верхней ($G_{x,y-1}=200$) и нижней ($G_{x,y+1}=36$) относительно неё. Так как значение градиента в центральной ячейке, больше чем в смежных, то центральная ячейка детектируется как граница (рисунок 3в). Однако, при незначительных изменениях значений градиаций серого (рисунок 3г) в четвертой строке происходит изменение величин градиентов (рисунок 3д). Далее с учетом уравнения (7) делается вывод, что центральная ячейка не является границей (рисунок 3е). Хотя в действительности в данной точке граница существует. То есть фильтр Канни слабо детектирует распределение выпуклой формы реальных границ объектов на

изображении. Вогнутая форма также плохо детектируется фильтром Канни.

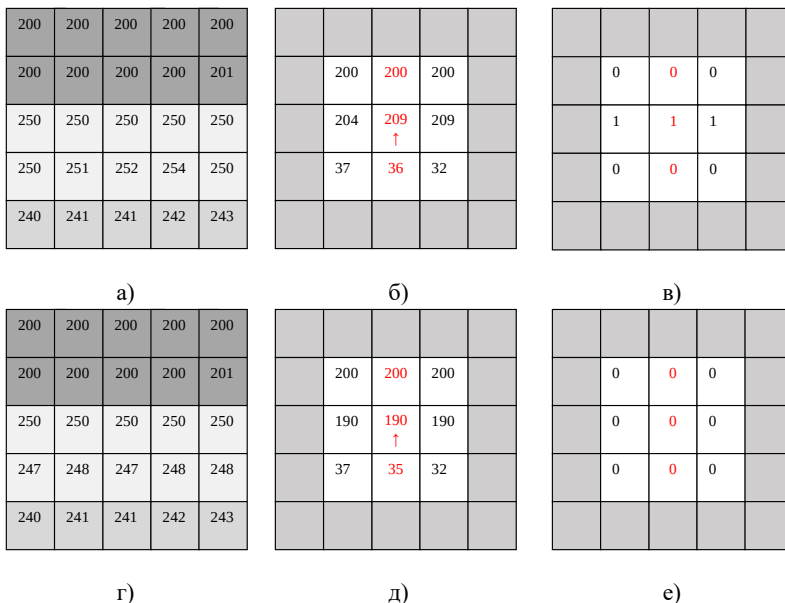


Рис. 3. Детектирование форм границ: а) матрица градацией серого цвета; б) матрица градиентов; в) детектирование границ; г) матрица градацией серого цвета после изменения; д) матрица градиентов после изменения; е) детектирование границ после изменения

В-третьих, для реализации алгоритма Канни для шестого и седьмого шагов требуется два дополнительных прохода по изображению, что снижает его производительность.

4. Нечетко-логические методы детектирования границ. Рассмотрим нечетко-логические методы детектирования границ, основанные на модификации фильтра Канни.

4.1. Метод 1. Нечетко-логический подход детектирования границ с учетом угла направления градиента.

Шаг 1. Повторение вычислений, аналогичных детектору Канни с 1 по 4 шаги.

Шаг 2. Фаззификация входных и выходной переменных.

Для формирования входных переменных используется разность градиентов смежных ячеек относительно центральной (рисунок 4).

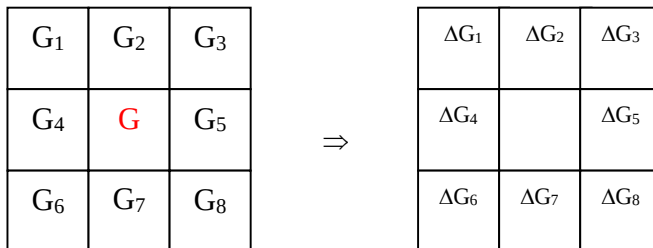


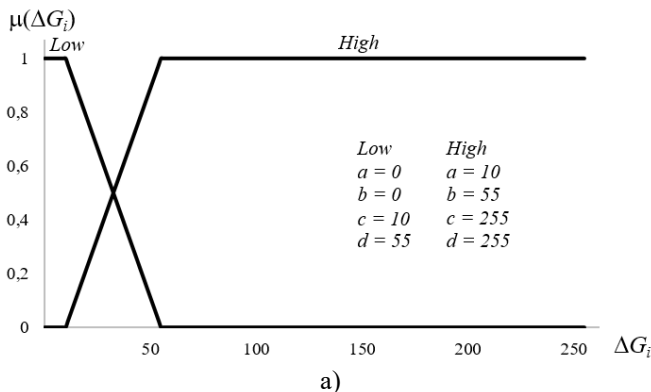
Рис. 4. Разность градиентов относительно центральной ячейки и их обозначение

Тогда входными параметрами являются восемь переменных, определяющих разность между градиентом центральной ячейки с её смежными восьмью ячейками, определяемых по формуле:

$$\Delta G_i = G - G_i, \quad (13)$$

где $i = 1 \dots 8$ – номер ячейки относительно центральной (рисунок 4).

Пусть функции принадлежности имеют трапециевидный вид и задаются четырьмя параметрами. Графики входных и выходной функций принадлежности и псевдокод вычисления степеней функций принадлежности представлен на рисунке 5 и Листинге 1, соответственно.



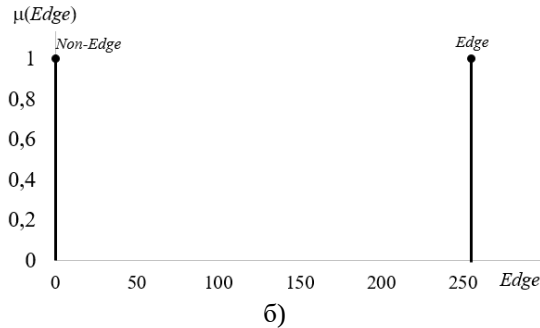


Рис. 5. Фаззификация: а) график входных функций принадлежности; б) график выходной функции принадлежности

Procedure Degrees_of_Membership_Function

Input: ΔG_i , a , b , c , d – параметры трапецевидной функции принадлежности (рисунок 5а)

Output: $\mu(\Delta G_i)$ – степень функции принадлежности ($i=1..8$)

Begin

If ($\Delta G_i \geq a \ \&\& \ \Delta G_i \leq b$)

Return $(\Delta G_i - a) / (b - a)$;

Else If ($\Delta G_i \geq b \ \&\& \ \Delta G_i \leq c$)

Return 1;

Else If ($\Delta G_i \geq c \ \&\& \ \Delta G_i \leq d$)

Return $(d - \Delta G_i) / (d - c)$;

Else

Return 0;

End

Листинг 1. Расчет степеней функций принадлежности

В ходе выполнения данного шага рассчитываются восемь значений степеней функций принадлежности.

Выходная функция принадлежности задается двумя синглтонными (одноэлементными) функциями принадлежности с термами $\mu(\text{Out}) = \{\text{Edge}(255), \text{Non-edge}(0)\}$ (рисунок 5б).

Шаг 3. Формирование базы нечетких правил.

Правила нечеткой базы формируются в зависимости от направления угла градиента, рассчитанного по формуле (6), которая для упрощения расчета была трансформирована в формулу (14), позволяющую четко интерпретировать угол направления градиента в зависимости от его возможных выходных значений:

$$\Theta_{x,y} = \text{round} \left(\frac{\text{atan} \left(\frac{GY_{x,y}}{GX_{x,y}} \right)}{\frac{\pi}{4}} \right) \Rightarrow \begin{cases} 0^\circ, & \text{если } \Theta_{x,y} = 4, \Theta_{x,y} = -4, \Theta_{x,y} = 0 \\ 45^\circ, & \text{если } \Theta_{x,y} = 1, \Theta_{x,y} = -3 \\ 90^\circ, & \text{если } \Theta_{x,y} = 2, \Theta_{x,y} = -2 \\ 135^\circ, & \text{если } \Theta_{x,y} = -1, \Theta_{x,y} = 3 \end{cases} \quad (14)$$

Нечеткие правила сведены в таблицу 1.

Таблица 1. База нечетких правил для метода 1

Rule	$\Theta_{x,y}$	$\mu(\Delta G_1)$	$\mu(\Delta G_2)$	$\mu(\Delta G_3)$	$\mu(\Delta G_4)$	$\mu(\Delta G_5)$	$\mu(\Delta G_6)$	$\mu(\Delta G_7)$	$\mu(\Delta G_8)$	Out
R_1	0°		Low					High		Edge
R_2			High					Low		Edge
R_3			High					High		Edge
R_1	45°	Low							High	Edge
R_2		High							Low	Edge
R_3		High							High	Edge
R_1	90°				Low	High				Edge
R_2					High	Low				Edge
R_3					High	High				Edge
R_1	135°			Low			High			Edge
R_2				High			Low			Edge
R_3				High			High			Edge

Для нахождения значений степеней нечетких правил с учетом данных, представленных в таблице 1, используется операция нечеткого минимума:

$$R_i = \min [\mu(\Delta G_i), \mu(\Delta G_i)]. \quad (15)$$

Шаг 4. Дефаззификация четкого значения. На этом шаге осуществляется дефаззификация и бинаризация выходного изображения с помощью нечеткого α -среза и определения на нем границ объектов. Псевдокод данной операции представлен в Листинге 2.

В отличие от детектора Канни данная база правил и использование α -среза позволяет реагировать на различные изменения градиента центральной ячейки относительно смежных ячеек, и тем самым предложенный метод становится более чувствительным к распознаванию

форм (вогнутой и выпуклой) распределения градиентов на исходном изображении.

Procedure Defuzzification_I

Input: R_1, R_2, R_3 – степени нечетких правил

Threshold – пороговое значение (α -срез)

Output: Out – бинаризованный код: 1 – есть граница; 0 – нет.

Begin

$DeFuzzy = \max(R_1, R_2, R_3);$

If ($DeFuzzy > \text{Threshold}$)

Return 1;

Else

Return 0;

End

Листинг 2. Процедура дефаззификации и бинаризации выходного изображения по методу 1

Например, на рисунках 3г, 3д, 3е представлена ситуация, когда граница не детектируется. Рассмотрим вычисление проблемной ситуации свойственной детектору Канны на основе предложенного подхода.

1. С учетом процедуры (листинг 1, $\text{Threshold}=0,4$), вычисляется угол направления градиента $\Theta_{x,y} = 0^\circ$ (рисунок 3д). Далее с учетом базы правил таблицы 1 определяются степени принадлежности для второй (ΔG_2) и седьмой (ΔG_7) входных переменных:

$$Low\{\mu(\Delta G_2)\} = 1$$

$$High\{\mu(\Delta G_2)\} = 0$$

$$Low\{\mu(\Delta G_7)\} = 0$$

$$High\{\mu(\Delta G_7)\} = 1$$

2. Далее осуществляется расчет степеней нечетких правил (таблица 1):

$$R_1 = \min[Low\{\mu(\Delta G_2)\}, High\{\mu(\Delta G_7)\}] = \min[1, 1] = 1,$$

$$R_2 = \min[High\{\mu(\Delta G_2)\}, Low\{\mu(\Delta G_7)\}] = \min[0, 0] = 0,$$

$$R_3 = \min[High\{\mu(\Delta G_2)\}, High\{\mu(\Delta G_7)\}] = \min[0, 1] = 0.$$

3. На шаге дефаззификации, используя процедуру (листинг 2), детектируется граница для центральной ячейки:

$$DeFuzzy = \max(R_1, R_2, R_3) = \max(1, 0, 0) = 1,$$

Out = 1, так как $DeFuzzy > \text{Threshold}$.

Таким образом, становится возможно определение различных форм (вогнутая или выпуклая) распределения градиента относительно центральный ячейки.

4.2. Метод 2. Нечетко-логический подход детектирования границ без учета угла направления градиента.

Шаг 1. Повторение вычислений, аналогичных детектору Канни, с 1 по 3 шаги.

Шаг 2. Фаззификация входных и выходной переменных. Данная операция осуществляется аналогично Шагу 2 Метода 1.

Шаг 3. Формирование базы нечетких правил.

На данном шаге задается база нечетких правил, состоящая из шестнадцати комбинаций величин разности градиента относительно центральной ячейки. В процессе обработки данных осуществляется сопоставление противоположных величин разности градиента относительно центральной ячейки, что позволяет не использовать расчет угла направления градиента, так как он фактически задается структурой нечетких правил. При этом данная структура правил так же, как и в методе 1 позволяет реагировать на любое изменение градиента относительно центральной ячейки. Поэтому данный подход также чувствителен к различению форм распределения градиента исходного изображения.

Таблица 2. База нечетких правил для метода 2

Rule	$\mu(\Delta G_1)$	$\mu(\Delta G_2)$	$\mu(\Delta G_3)$	$\mu(\Delta G_4)$	$\mu(\Delta G_5)$	$\mu(\Delta G_6)$	$\mu(\Delta G_7)$	$\mu(\Delta G_8)$	Out
R_1	High	High						Low	Edge
R_2	High			High				Low	Edge
R_3		High	High					Low	Edge
R_4				High		High		Low	Edge
R_5	High	High					Low		Edge
R_6	High			High			Low		Edge
R_7		High	High				Low		Edge
R_8				High		High	Low		Edge
R_9	High	High			Low				Edge
R_{10}	High			High	Low				Edge
R_{11}		High	High		Low				Edge
R_{12}				High	Low	High			Edge
R_{13}			Low			High	High		Edge
R_{14}			Low				High	High	Edge
R_{15}	High		Low	High					Edge
R_{16}			Low	High		High			Edge

Для нахождения значений степеней нечетких правил с учетом данных, представленных в таблице 2, используется операция нечеткого минимума:

$$R_{i=1...16} = \min[\mu(\Delta G_i), \mu(\Delta G_i), \mu(\Delta G_i)]. \quad (16)$$

Шаг 4. Дефаззификация четкого значения. На этом шаге с учетом комбинации модели дефаззификации центра тяжести и модели нечеткого α -среза осуществляется бинаризация выходного изображения и определение на нем границ объектов. Операция производится по процедуре, представленной в Листинге 3.

Procedure Defuzzification_II

Input: $R_1, R_2, R_3, R_4, R_5, R_6, R_7, R_8, R_9, R_{10}, R_{11}, R_{12}, R_{13}, R_{14}, R_{15}, R_{16}$ - степени нечетких правил

Threshold – пороговое значение (α -срез)

Output: Out – бинаризованный код: 1 – есть граница; 0 – нет.

Begin

$DeFuzzy = (\sum_{i=1...16} R_i \cdot \mu(Edge)) / \sum_{i=1...16} R_i$;

If ($DeFuzzy \geq Threshold$)

Return 1;

Else

Return 0;

End

Листинг 3. Процедура дефаззификации и бинаризации выходного изображения по методу 2

Как и в предыдущем случае, на основе предложенного подхода рассмотрим вычисление проблемной ситуации, свойственной детектору Канни (рисунок 3д).

1. С учетом процедуры (листинг 1) вычисляются степени принадлежности для восьми входных переменных (таблица 3):

Таблица 3. Степени принадлежности входных переменных

$Low\{\mu(\Delta G_1)\}=1$	$Low\{\mu(\Delta G_5)\}=1$	$High\{\mu(\Delta G_1)\}=0$	$High\{\mu(\Delta G_5)\}=0$
$Low\{\mu(\Delta G_2)\}=1$	$Low\{\mu(\Delta G_6)\}=0$	$High\{\mu(\Delta G_2)\}=0$	$High\{\mu(\Delta G_6)\}=0$
$Low\{\mu(\Delta G_3)\}=1$	$Low\{\mu(\Delta G_7)\}=0$	$High\{\mu(\Delta G_3)\}=0$	$High\{\mu(\Delta G_7)\}=0$
$Low\{\mu(\Delta G_4)\}=1$	$Low\{\mu(\Delta G_8)\}=0$	$High\{\mu(\Delta G_4)\}=0$	$High\{\mu(\Delta G_8)\}=0$

2. После расчета степеней функций принадлежности вычисляются значения степеней нечетких правил (таблица 2):

$$\begin{aligned} R_1 &= \min[High\{\mu(\Delta G_1)\}, High\{\mu(\Delta G_2)\}, Low\{\mu(\Delta G_3)\}] = \min[0, 0, 0] = 0, \\ R_2 &= \min[0, 0, 0] = 0, \quad R_3 = \min[0, 0, 0] = 0, \quad R_4 = \min[0, 1, 0] = 0, \quad R_5 = \min[0, 0, 0] = 0, \\ R_6 &= \min[0, 0, 0] = 0, \quad R_7 = \min[0, 0, 0] = 0, \quad R_8 = \min[0, 1, 0] = 0, \quad R_9 = \min[0, 0, 1] = 0, \\ R_{10} &= \min[0, 0, 1] = 0, \quad R_{11} = \min[0, 0, 1] = 0, \quad R_{12} = \min[0, 1, 1] = 0, \quad R_{13} = \min[1, 1, 1] = 1, \\ R_{14} &= \min[1, 1, 1] = 1, \quad R_{15} = \min[0, 1, 0] = 0, \quad R_{16} = \min[1, 0, 1] = 0. \end{aligned}$$

3. На шаге дефаззификации, используя процедуру (листинг 3), детектируется граница для центральной ячейки:

$$DeFuzzy = 255,$$

Out = 1, так как $DeFuzzy > Threshold$.

Таким образом, данный метод также позволяет определять различные формы распределения градиента относительно центральной ячейки.

5. Экспериментальные результаты. Практическое применение разработанных двух методов детектирования границ на изображениях выполнено в виде комплекса программного обеспечения, реализованного в среде Microsoft Visual Studio 2019 на языке программирования C#. Для проведения эксперимента использовался персональный компьютер Intel(R) Core(TM) i5-8600K CPU 3.60GHz, ОЗУ 16 ГБ, операционная система Win10. В ходе эксперимента определялось время, необходимое для выделения границ на изображениях. Эксперименты повторялись 100 раз для каждой из картинок. Выделения контуров на трех картинках (трактор, машина, процессор) сведены в таблицу 4. В таблице 5 указано время по 100 экспериментам, необходимое для выделения контуров по каждой из картинок. Средневзвешенный показатель IMP рассчитан по трем картинкам для каждого из операторов. Расчет сведен в таблицу 5.

Формула для расчета IMP имеет вид [56]:

$$IMP = \frac{1}{\max(I_I, I_A)} \sum_{i=1}^{I_A} \frac{1}{1 + \alpha \cdot d_i^2}, \quad (17)$$

где I_I, I_A – количество граничных точек в эталонном и полученном в процессе детектирования контуре объекта, соответственно; α – коэффициент, величины штрафа за смещение граничной точки (по умолчанию $\alpha=1/9$);

d_i – расстояние от граничной точки эталонного контура до граничной точки полученного в результате детектирования.

Таблица 4. Детектирование границ на изображении

Исходное изображение	Размер	Результат детектирования границ			
		Оператор Собеля	Детектор Канни	Метод 1	Метод 2
	1920 × 1080				
	728 × 410				
	1830 × 1029				

Таблица 5. Анализ производительности вычислительных алгоритмов

Изображение	Оператор Собеля, сек	Детектор Канни, сек	Нечеткий метод 1, сек	Нечеткий метод 2, сек
Трактор	66	84	69	70
Машина	9	17	10	10
Процессор	60	78	63	64
IMR	0,73	0,76	0,79	0,80

В ходе эксперимента, как и предполагалось, установлено, что наилучшее быстродействие детектирования границ имеет оператор Собеля. Но изображения, полученные с помощью данного оператора, представлены в градациях серого, и для дальнейшей бинаризации выделенных контуров необходимы дополнительные вычисления. Наилучшие результаты с точки зрения выделения контуров и снижения времени обработки изображений получены на основе первого нечетко-логического метода. Он в среднем на 18% производительней относительно детектора Канни и на 2% производительней второго нечетко-логического метода выделения контуров. Также в ходе эксперимента было установлено, что при обработке картинок с меньшим разрешением (картинка – машина) первый метод не полностью выделяет границы на объектах. При увеличении разрешения до формата full hd первый метод достаточно хорошо выделяет границы. Так, например, на картинке *процессор* детектор Канни и второй нечеткий метод в верхнем левом углу потеряли часть контуров, в то время как первый нечет-

кий метод смог выделить контуры. Наилучшее выделение контуров было получено с помощью второго нечеткого метода, хотя по времени он проигрывает первому нечетко-логическому подходу, но выигрывает у детектора Канни по быстродействию на 16%. Для повышения точности выделения контуров в ходе эксперимента было установлено, что первый и второй нечеткие методы детектирования границ зависят от параметров трапецевидной функции принадлежности. Для терма *Low* (рисунок 5а) наибольшее влияние оказывает расположение меток *c* и *d*. Для терма *High* (рисунок 5а) наибольшее влияние оказывает расположение меток *a* и *b*. Также количество выделяемых контуров зависит от переменной *Threshold* (листинги 2 и 3) и от структуры нечетких правил (таблицы 1 и 2), которые позволяют детектировать границу в центральной ячейке обрабатываемой рамки изображения 3×3. Установлено, что реакцией на детектирование границ можно управлять изменением структуры нечетких правил. Например, в таблице 1 правило R_2 (при $\Theta_{x,y}=0^\circ$) реагирует только на перепад двух соседних ячеек. Если предположить, что граница имеет большую длину и зависит не от двух, а от трех ячеек, то точность выделения контуров увеличится. В соответствии с этим предположением были модифицированы нечеткие правила, представленные в таблице 1, и сведены в таблицу 6. При этом число нечетких правил было увеличено до шестнадцати.

Таблица 6. База модифицированных нечетких правил по методу 1

Rul e	$\Theta_{x,y}$	$\mu(\Delta G_1)$	$\mu(\Delta G_2)$	$\mu(\Delta G_3)$	$\mu(\Delta G_4)$	$\mu(\Delta G_5)$	$\mu(\Delta G_6)$	$\mu(\Delta G_7)$	$\mu(\Delta G_8)$	Out
R_1	0°	High	High					Low		Edge
R_2		High			High			Low		Edge
R_3			High	High				Low		Edge
R_4				High		High		Low		Edge
R_1	45°	High	High				Low			Edge
R_2			High	High			Low			Edge
R_3				High		High	Low			Edge
R_4						High	Low		High	Edge
R_1	90°	High	High			Low				Edge
R_2		High			High	Low				Edge
R_3					High	Low	High			Edge
R_4						Low	High	High		Edge
R_1	135°	High	High						Low	Edge
R_2		High			High				Low	Edge
R_3			High	High					Low	Edge
R_4					High		High		Low	Edge

Для нахождения степеней нечетких правил использовалось уравнение (16), а для дефаззификации и бинаризации процедура, представленная в Листинге 3. При имитационном моделировании были изменены значения меток входных трапециевидных функций принадлежности: для терм $Low = \{0, 0, 20, 75\}$; для терма $High = \{20, 75, 255, 255\}$. Переменная $Threshold = 0.01$. Результаты детектирования представлены в таблице 7.

Таблица 7. Детектирование границ после модификации нечетко-логического метода I

Изображение до модификации по методу I	Изображение после модификации по методу I	Изображение по методу II
		

Время детектирования границ составило 10 секунд для обработки 100 картинок, так же, как и без модификации, но качество выделения контуров существенно улучшилось. Таким образом, можно сделать вывод о том, что поставленная цель повышения производительности при детектировании границ достигнута.

6. Заключение. Описаны два нечетко-логических метода детектирования границ на изображениях. Разработка данных методов была обусловлена необходимостью распознавания большего количества границ, так как наиболее оптимальный с точки зрения выделения контуров фильтр Канни имеет ряд систематических ошибок. К одной из таких ошибок относится отсутствие реакции на разные формы изменения градиента. В статье было установлено, что к изменениям выпуклой и вогнутой формы распределения градиента детектор Канни мало чувствителен. Поэтому возникла необходимость компенсации данной ошибки. Для реакции детектора на скорость изменения градиента и формы распределения градиента наиболее лучше подходит нечеткая логика, так как она позволяет с помощью трапециевидных функций принадлежности учитывать изменение скорости градиента, а с помощью нечетких правил задавать угол его распределения. В первом нечетком методе так же, как и в детекторе Канни, используется расчет градиента и угла его направления, в зависимости от значения которых осуществляется выбор четырех нечетких правил, дефаззификация осуществляется на основе метода центра тяжести, бинаризация с ис-

пользование нечеткого α -среза. Во-втором нечетком подходе расчет угла направления градиента не рассчитывается, его направление фактически задают шестнадцать нечетких правил. Высокая производительность представленных методов обуславливается уменьшением количества проходов по исходному изображению, так в детекторе Канни для реализации правил используется три прохода, в предлагаемых методах только один. Математическое обоснование устойчивости и сходимости предложенного метода является дальнейшим направлением нашего научного исследования.

Литература

1. Yuksel M.E. Edge detection in noisy images by neuro-fuzzy processing. AEU – Int J Electron Commun 2007; 61(2, no. 1): 8289. <http://dx.doi.org/10.1016/j.aeue.2006.02.006>.
2. Kang C.C., Wang W.J. A novel edge detection method based on the maximizing objective function. Pattern Recogn 2007; 40(2): 609–18. <http://dx.doi.org/10.1016/j.patcog.2006.03.016>.
3. Lopez-Molina C., De Baets B., Bustince H. Generating fuzzy edge images from gradient magnitudes. Comput Vision Image Understand 2011; 115(11): 1571–80. <http://dx.doi.org/10.1016/j.cviu.2011.07.003>.
4. Bovik A. Handbook of image and video processing. New York: Academic; 2000.
5. Canny J. A computational approach to edge detection. IEEE Trans Pattern Anal Mach Intell 1986; 8(6): 679–98.
6. Sobel E. Camera Models and Machine Perception. Ph.D thesis. Stanford University, Stanford, California; 1970.
7. Chen G., Yang Y.H.H. Edge detection by regularized cubic B-spline fitting. IEEE Trans Syst, Man Cybern 1995; 25: 636–43.
8. Canny J. A computational approach to edge detection. IEEE Trans Pattern Anal Mach Intell 1986; PAMI-8: 679–97.
9. Shen J., Castan S. An optimal linear operator for step edge detection. Graph Models Image Process 1992; 54(1): 112–33.
10. deSouza P. Edge detection using sliding statistical tests. Comput Vis, Graph Image Process 1983; 23(1).
11. Bhandarkar S.M., Zhang Y., Potter W.D. An edge detection technique using genetic algorithm based optimization. Pattern Recognit 1994; 27(9): 1159–80.
12. Srinivasan V., Bhatia P., Ong S.H. Edge detection using neural network. Pattern Recognit 1995; 27(12): 1653–62.
13. Chen M.H., Lee D., Pavlidis T. Residual analysis for feature detection. IEEE Trans Pattern Anal Mach Intell 1991; 13: 30 – 40.
14. Hebert T.J., Malagre D. Edge detection using a priori model. Int Conf Image Process 1994; 94: 303–7.
15. Mejias A., Romero S., Moreno F. A new algorithm to extract the lines and edges through orthogonal projections. Digital Signal Process. - 2012; 22(1): 147–52.
16. Rakesh R.R., Chaudhuri P., Murthy C.A. Thresholding in edge detection: a statistical approach. IEEE Trans Image Process 2004; 13(7): 927–36. <http://dx.doi.org/10.1109/TIP.2004.828404>.
17. S. Uguz, U. Sahin, F. Sahin. Edge detection with fuzzy cellular automata transition function optimized by PSO. Computers and Electrical Engineering. - 2015; 43. – pp.180–192

18. Alexander Zotin, Konstantin Simonov, Mikhail Kurako, Yousif Hamad, Svetlana Kirillova Edge detection in MRI brain tumor images based on fuzzy C-means clustering. *Procedia Computer Science*. – 2018; 126. – pp. 1261–1270
19. Er-sen L., Shu-long Z., Bao-shan Z., Yong Z., Chao-gui X., Li-hua S. An Adaptive Edge Detection Method Based on The Canny Operator. *IEEE Int. Conf. Environmental Sci. and Inform. Applicat. Technology* 2009; p. 265–269.
20. Cho S.M., Cho J.H. Thresholding for Edge Detection using Fuzzy Reasoning Technique. *IEEE Int. Conf. Computational Sci. Proc.* 1994; p. 1121–1124.
21. Xiao W., Hui X. An Improved Canny Edge Detection Algorithm Based on Predisposal Method for Image Corrupted by Gaussian Noise. *IEEE World Automation Congr.* 2010; p. 113–116.
22. Wang H.R., Yang J.L., Sun H.J., Chen D., Liu X.L. An improved Region Growing Method for Medical Image Selection and Evaluation Based on Canny Edge Detection. *IEEE Int. Conf. Manage. and Service Sci.* 2011; p. 1–4, DOI: 10.1109/ICMSS.2011.5999180.
23. Ranita Biswas, Jaya Sil An Improved Canny Edge Detection Algorithm Based on Type-2 Fuzzy Sets. *Procedia Technology*. 2012; 4. – pp. 820 – 824
24. Shah, Hemang J. Detection of Tumor in MRI Images using Image Segmentation. *International Journal of Advance Research in Computer Science and Management Studies*. – 2014; 2(6).- pp. 53-56.
25. Zotin Alexander, Konstantin Simonov, Fedor Kapsargin, Tatyana Cherepanova, Alexey Kruglyakov, and Luis Cadena. “Techniques for Medical Images Processing Using Shearlet Transform and Color Coding”, in Favorskaya M. and Jain L. (eds) *Computer Vision in Control Systems-4. Intelligent Systems*. – 2018: 136, Springer, Cham.
26. D. Xu, W. Ouyang, X. Alameda-Pineda, E. Ricci, X. Wang, and N. Sebe, “Learning deep structured multi-scale features using attention gated crfs for contour prediction,” in *Conference on Neural Information Processing Systems*, 2017, pp. 3964–3973.
27. J. He, S. Zhang, M. Yang, Y. Shan, and T. Huang, “Bi-directional cascade network for perceptual edge detection,” in *IEEE Conference on Computer Vision and Pattern Recognition*, 2019, pp. 3828–3837.
28. Jianzhong He et al. “Bi-directional cascade network for perceptual edge detection”. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. 2019, pp. 3828–3837.
29. Yang Liu, Zongwu Xie, and Hong Liu. “An Adaptive and Robust Edge Detection Method Based on Edge Proportion Statistics”. In: *IEEE Transactions on Image Processing* 29 (2020), pp. 5206–5215.
30. S. Yun, J. Choi, Y. Yoo, K. Yun, and J. Young Choi, “Action-decision networks for visual tracking with deep reinforcement learning,” in *IEEE Conference on Computer Vision and Pattern Recognition*, 2017, pp. 2711–2720.
31. K.H. Choi and J.E. Ha, “Edge detection based-on U-Net using edge classification CNN,” *Journal of Institute of Control, Robotics and Systems*, vol. 25, no. 8, pp. 684–689, 2019 (in Korean).
32. Yun Liu et al. “Richer convolutional features for edge detection”. In: *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*. 2017, pp. 3000–3009.
33. Xavier Soria Poma, Edgar Riba, and Angel Sappa. “Dense extreme inception network: Towards a robust cnn model for edge detection”. In: *Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision*. 2020, pp. 1923–1932.
34. Raihan F., Ce W. “PCB defect detection USING OPENCV with image subtraction method”. In *International Conference on Information Management and Technology*, 2017, pp. 204–209.

35. Lee D.H., Chen P.Y., Yang F.J., et al. "High-Efficient Low-Cost VLSI Implementation for Canny Edge Detection". *Journal of Information Science & Engineering*, vol. 36, no. 3, 2020, pp. 34-57.
36. François Chollet. "Xception: Deep learning with depthwise separable convolutions". In: *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*. 2017, pp. 1251–1258.
37. Debotosh Bhattacharjee and Hiranmoy Roy. "Pattern of local gravitational force (PLGF): A novel local image descriptor". In: *IEEE Transactions on Pattern Analysis and Machine Intelligence* 43.2 (2019), pp. 595–607.
38. Mengtian Li et al. "Photo-sketching: Inferring contour drawings from images". In: *2019 IEEE Winter Conference on Applications of Computer Vision (WACV)*. IEEE. 2019, pp. 1403–1412
39. D. Dhillon and R. Chouhan. "Noise-aided Edge preserving Image Denoising using Non-Local Means with Stochastic Resonance". In: *2018 IEEE Global Conference on Signal and Information Processing (GlobalSIP)*. 2018, pp. 21–25.
40. Animesh Sengupta et al. "Edge information based image fusion metrics using fractional order differentiation and sigmoidal functions". In: *IEEE Access* 8 (2020), pp. 88385–88398.
41. Benoit Brummer and Christophe De Vleeschouwer. "Natural image noise dataset". In: *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition Workshops*. 2019, pp. 139-151.
42. Arash Akbarinia and C. Alejandro Parraga. "Feedback and surround modulated boundary detection". In: *International Journal of Computer Vision* 126.12 (2018), pp. 1367–1380
43. Y.X. Wang and J.M. Chen, "Iris edge detection algorithm based on adaptive canny operator and multi-directional Sobel operator," *Computer and Digital Engineering*, Vol. 11, No. 4, pp. 2744–2749, 2020.
44. C.W. Tian, X.C. Wang, and J.N. Yang, "Research on parallelization of kirsch operator edge detection algorithm for geological image," (in Chinese), *Journal of Xinjiang University*, vol. 38, No. 1, pp. 54–60, 2021.
45. J.H. Zeng and S.J. Huang, "Comparison and analysis on typical image edge detection operators," *Journal of Hebei Normal University (Natural Science)*, vol. 44, No. 1, pp. 295–300, 2020.
46. S.J. Chen, X.H. Wang, Y.P. Ge, C. Li, and Y.C. Li., "Application of image edge extraction algorithm in the third land survey," *Computer Technology and Development*, vol. 30, No. 10, pp. 161–166, 2020.
47. Cadena, Luis, Franklin Cadena, Nikolai Espinosa, Anna Korneeva, Alexy Kruglyakov, Alexander Legalov, Alexey Romanenko, and Alexander Zotin. (2017) "Brain's tumor image processing using shearlet transform." *Proc. SPIE 10396, Applications of Digital Image Processing XL*, 103961B, doi: 10.1117/12.2272792; in United States.
48. Yuksel M.E., Borlu M. Accurate Segmentation of Dermoscopic Images by Image Thresholding Based on Type-2 Fuzzy Logic. *IEEE Trans. Fuzzy Syst.* 2009; vol. 17, no. 4, p. 976–982.
49. M. Bobyr, A. Arkhipov, A. Yakushev, Shade recognition of the color label based on the fuzzy clustering. *Inform. Autom.* 20(2) (2021) 407–434, <http://dx.doi.org/10.15622/ia.2021.20.2.6>.
50. Bobyr M.V., Emelyanov. S.G., A nonlinear method of learning neuro-fuzzy models for dynamic control systems, *Appl. Soft Comput. J.* 88 (2020) 106030, <http://dx.doi.org/10.1016/j.asoc.2019.106030>.
51. Bobyr M.V., Milostnaya N.A., Kulabuhov S.A., A method of defuzzification based on the approach of areas' ratio, *Appl. Soft Comput.* 59 (2017) 19–32, <http://dx.doi.org/10.1016/j.asoc.2017.05.040>.

52. M.V. Bobyr, A.S. Yakushev, A.A. Dorodnykh, Fuzzy devices for cooling the cutting tool of the CNC machine implemented on FPGA, Meas.: J. Int. Meas. Confed. 152 (2020) <http://dx.doi.org/10.1016/j.measurement.2019.107378>.
53. Bobyr M.V., Milostnaya N.A., Bulatnikov V.A. The fuzzy filter based on the method of areas' ratio. Applied Soft Computing. 117 (2022) 108449, <https://doi.org/10.1016/j.asoc.2022.108449>
54. Bobyr M.V., Kulabukhov S.A. Simulation of control of temperature mode in cutting area on the basis of fuzzy logic. Journal of Machinery Manufacture and Reliability, 2017, 46(3), стр. 288–295. <http://dx.doi.org/10.3103/S1052618817030049>
55. Sala, F.A. Design of false color palettes for grayscale reproduction. Displays, 2017, 46, 9–15. <https://doi.org/10.1016/j.displa.2016.11.005>
56. Abdou, I.E., & Pratt, W.K. Quantitative design and evaluation of enhancement/thresholding edge detectors. Proceedings of the IEEE, 1979, 67(5), 753–763. doi:10.1109/proc.1979.11325

Бобырь Максим Владимирович — д-р техн. наук, профессор, кафедра вычислительной техники, Юго-Западный государственный университет. Область научных интересов: распознавание образов, стереозрение, адаптивные нейро-нечеткие системы вывода, мягкие вычисления, системы искусственного интеллекта, мобильные роботы, ПЛИС. Число научных публикаций — 398. maxbobyry@gmail.com; проезд Светлый, 1, 305046, Курск, Россия; р.т.: 8(4712)222-665.

Архипов Александр Евгеньевич — канд. техн. наук, старший научный сотрудник, институт космического приборостроения и радиоэлектронных систем, Юго-Западный государственный университет. Область научных интересов: системы искусственного интеллекта, системы технического зрения, распознавание образов, стереозрение, ПЛИС. Число научных публикаций — 102. alex.76_09@mail.ru; 50 лет Октября, 94, 305040, Курск, Россия; р.т.: 8(4712)222-665.

Горбачев Сергей Викторович — канд. техн. наук, старший научный сотрудник, факультет инновационных технологий, Томский государственный университет. Область научных интересов: нейронные сети, нечеткие множества, гибридные модели, ансамблевые вычисления. Число научных публикаций — 161. spp03@sibmail.com; проспект Ленина, 36, 634050, Томск, Россия; р.т.: 8(3822)529-498.

Цао Цзиньде — Ph.D., профессор, декан, математический факультет, Юго-Восточный университет. Область научных интересов: нейронные сети, вычислительный интеллект, сложные сети, сложные системы, интеллектуальное управление, нелинейные системы, прикладная математика. Число научных публикаций — 1546. jdciao@seu.edu.cn; трасса CEU, район Цзянин, 2, 210096, Нанкин, Китай; р.т.: +86-25-52090588.

Бхаттачарья Сиддхартха — Ph.D., ректор, вычислительная техника, Раджнагар Махалавидья - филиал Университета Бурдвана. Область научных интересов: мягкие вычисления, распознавание образов, обработка изображений, поиск мультимедийной информации, квантовые мягкие вычисления, оптимизация портфолио, социальные сети. Число научных публикаций — 300. dr.siddhartha.bhattacharyya@gmail.com; Раджнагар, 1, 731130, Бирбхум, Индия; р.т.: +919830354195.

Поддержка исследований. Работа проводится в рамках Государственного задания (грант №0851-2020-0032).

M. BOBYR, A. ARKHIPOV, S. GORBACHEV, J. CAO, S.B. BHATTACHARYYA
**FUZZY LOGIC APPROACHES IN THE TASK OF OBJECT EDGE
DETECTION**

Bobyр M., Arkhipov A., Gorbachev S., Cao J., Bhattacharyya S.B. Fuzzy Logic Approaches in the Task of Object Edge Detection.

Abstract. The task of reducing the computational complexity of contour detection in images is considered in the article. The solution to the task is achieved by modifying the Canny detector and reducing the number of passes through the original image. In the first case, two passes are excluded when determining the adjacency of the central pixel with eight adjacent ones in a frame of size 3×3 . In the second case, three passes are excluded, two as in the first case and the third one necessary to determine the angle of gradient direction. This passage is provided by a combination of fuzzy rules. The goal of the work is to increase the performance of computational operations in the process of detecting the edges of objects by reducing the number of passes through the original image. The process of edge detection is carried out by some computational operations of the Canny detector with the replacement of the most complex procedures. In the proposed methods, fuzzification of eight input variables is carried out after determining the gradient and the angle of its direction. The input variables are the gradient difference between the central and adjacent cells in a frame of size 3×3 . Then a base of fuzzy rules is built. In the first method, four fuzzy rules and one pass are excluded depending on the angle of gradient direction. In the second method, sixteen fuzzy rules themselves set the angle of the gradient direction, while eliminating two passes along the image. The gradient difference between the central cell and adjacent cells makes it possible to take into account the shape of the gradient distribution. Then, based on the center of gravity method, the resulting variable is defuzzified. Further use of fuzzy α -cut makes it possible to binarize the resulting image with the selection of object edges on it. The presented experimental results showed that the noise level depends on the value of the α -cut and the parameters of the labels of the trapezoidal membership functions. The software was developed to evaluate fuzzy edge detection methods. The limitation of the two methods is the use of piecewise-linear membership functions. Experimental studies of the performance of the proposed edge detection approaches have shown that the time of the first fuzzy method is 18% faster compared to the Canny detector and 2% faster than the second fuzzy method. However, during the visual assessment, it was found that the second fuzzy method better determines the edges of objects.

Keywords: fuzzy logic, Canny detector, boundary detection, Sobel operator, centre of gravity.

Bobyр Maksim — Ph.D., Dr.Sci., Professor, Computer science department, Southwest State University. Research interests: pattern recognition, stereo vision, adaptive neuro-fuzzy inference systems, soft computing, artificial intelligence systems, mobile robots, FPGA. The number of publications — 398. maxbobyр@gmail.com; 1, Svetly pass., 305046, Kursk, Russia; office phone: 8(4712)222-665.

Arkhipov Alexander — Ph.D., Senior researcher, Research institute of space instrumentation and radioelectronic systems, Southwest State University. Research interests: artificial intelligence systems, fuzzy logic, technical vision systems, pattern recognition, stereo vision, FPGA. The number of publications — 102. alex.76_09@mail.ru; 94, 50 years of October St., 305040, Kursk, Russia; office phone: 8(4712)222-665.

Gorbachev Sergey — Ph.D., Senior researcher, Faculty of innovative technologies, Tomsk State University. Research interests: neural networks, fuzzy sets, hybrid models, ensemble computing. The number of publications — 161. spp03@sibmail.com; 36, Lenin Ave., 634050, Tomsk, Russia; office phone: 8(3822)529-498.

Cao Jinde — Ph.D., Professor, Dean, school of mathematics, Southeast University. Research interests: neural networks, computational intelligence, complex networks, complex systems, intelligent control, nonlinear systems, applied mathematics. The number of publications — 1546. jdcao@seu.edu.cn; 2, SEU Road, Jiangning District, 210096, Nanjing, China; office phone: +86-25-52090588.

Bhattacharyya Siddhartha — Ph.D., Principal, Computer science, Rajnagar Mahalavidya affiliated to Burdwan University. Research interests: soft computing, pattern recognition, image processing, multimedia information retrieval, quantum inspired soft computing, portfolio optimization, social networks. The number of publications — 300. dr.siddhartha.bhattacharyya@gmail.com; 1, Rajnagar, 731130, Birbhum, India; office phone: +919830354195.

Acknowledgements. This research is supported by GZ (grant 0851-2020-0032).

References

1. Yuksel M.E. Edge detection in noisy images by neuro-fuzzy processing. *AEU – Int J Electron Commun* 2007; 61(2, no. 1): 8289. <http://dx.doi.org/10.1016/j.aeue.2006.02.006>.
2. Kang C.C., Wang W.J. A novel edge detection method based on the maximizing objective function. *Pattern Recogn* 2007; 40(2): 609–18. <http://dx.doi.org/10.1016/j.patcog.2006.03.016>.
3. Lopez-Molina C., De Baets B., Bustince H. Generating fuzzy edge images from gradient magnitudes. *Comput Vision Image Understand* 2011; 115(11): 1571–80. <http://dx.doi.org/10.1016/j.cviu.2011.07.003>.
4. Bovik A. *Handbook of image and video processing*. New York: Academic; 2000.
5. Canny J. A computational approach to edge detection. *IEEE Trans Pattern Anal Mach Intell* 1986; 8(6): 679–98.
6. Sobel E. *Camera Models and Machine Perception*. Ph.D thesis. Stanford University, Stanford, California; 1970.
7. Chen G., Yang Y.H.H. Edge detection by regularized cubic B-spline fitting. *IEEE Trans Syst, Man Cybern* 1995; 25: 636–43.
8. Canny J. A computational approach to edge detection. *IEEE Trans Pattern Anal Mach Intell* 1986; PAMI-8: 679–97.
9. Shen J., Castan S. An optimal linear operator for step edge detection. *Graph Models Image Process* 1992; 54(1):112–33.
10. deSouza P. Edge detection using sliding statistical tests. *Comput Vis, Graph Image Process* 1983; 23(1).
11. Bhandarkar S.M., Zhang Y., Potter W.D. An edge detection technique using genetic algorithm based optimization. *Pattern Recognit* 1994; 27(9): 1159–80.
12. Srinivasan V., Bhatia P., Ong S.H. Edge detection using neural network. *Pattern Recognit* 1995; 27(12): 1653–62.
13. Chen M.H., Lee D., Pavlidis T. Residual analysis for feature detection. *IEEE Trans Pattern Anal Mach Intell* 1991; 13: 30 – 40.
14. Hebert T.J., Malagre D. Edge detection using a priori model. *Int Conf Image Process* 1994; 94: 303–7.

15. Mejias A., Romero S., Moreno F. A new algorithm to extract the lines and edges through orthogonal projections. *Digital Signal Process.* - 2012; 22(1): 147–52.
16. Rakesh R.R., Chaudhuri P., Murthy C.A. Thresholding in edge detection: a statistical approach. *IEEE Trans Image Process* 2004; 13(7): 927–36. <http://dx.doi.org/10.1109/TIP.2004.828404>.
17. S. Uguz, U. Sahin, F. Sahin. Edge detection with fuzzy cellular automata transition function optimized by PSO. *Computers and Electrical Engineering.* - 2015; 43. – pp.180–192
18. Alexander Zotin, Konstantin Simonov, Mikhail Kurako, Yousif Hamad, Svetlana Kirillova Edge detection in MRI brain tumor images based on fuzzy C-means clustering. *Procedia Computer Science.* – 2018; 126. – pp. 1261–1270
19. Er-sen L., Shu-long Z., Bao-shan Z., Yong Z., Chao-gui X., Li-hua S. An Adaptive Edge Detection Method Based on The Canny Operator. *IEEE Int. Conf. Environmental Sci. and Inform. Applicat. Technology* 2009; p. 265–269.
20. Cho S.M., Cho J.H. Thresholding for Edge Detection using Fuzzy Reasoning Technique. *IEEE Int. Conf. Computational Sci. Proc.* 1994; p. 1121–1124.
21. Xiao W., Hui X. An Improved Canny Edge Detection Algorithm Based on Prepositional Method for Image Corrupted by Gaussian Noise. *IEEE World Automation Congr.* 2010; p. 113–116.
22. Wang H.R., Yang J.L., Sun H.J., Chen D., Liu X.L. An improved Region Growing Method for Medical Image Selection and Evaluation Based on Canny Edge Detection. *IEEE Int. Conf. Manage. and Service Sci.* 2011; p. 1–4, DOI: 10.1109/ICMSS.2011.5999180.
23. Ranita Biswas, Jaya Sil An Improved Canny Edge Detection Algorithm Based on Type-2 Fuzzy Sets. *Procedia Technology.* 2012: 4. – pp. 820 – 824
24. Shah, Hemang J. Detection of Tumor in MRI Images using Image Segmentation. *International Journal of Advance Research in Computer Science and Management Studies.* – 2014; 2(6).– pp. 53-56.
25. Zotin Alexander, Konstantin Simonov, Fedor Kapsargin, Tatyana Cherepanova, Alexey Kruglyakov, and Luis Cadena. Techniques for Medical Images Processing Using Shearlet Transform and Color Coding”, in Favorskaya M. and Jain L. (eds) *Computer Vision in Control Systems-4. Intelligent Systems.* – 2018: 136, Springer, Cham.
26. D. Xu, W. Ouyang, X. Alameda-Pineda, E. Ricci, X. Wang, and N. Sebe, “Learning deep structured multi-scale features using attention gated crfs for contour prediction,” in *Conference on Neural Information Processing Systems*, 2017, pp. 3964–3973.
27. J. He, S. Zhang, M. Yang, Y. Shan, and T. Huang, “Bi-directional cascade network for perceptual edge detection,” in *IEEE Conference on Computer Vision and Pattern Recognition*, 2019, pp. 3828–3837.
28. Jianzhong He et al. “Bi-directional cascade network for perceptual edge detection”. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. 2019, pp. 3828–3837.
29. Yang Liu, Zongwu Xie, and Hong Liu. “An Adaptive and Robust Edge Detection Method Based on Edge Proportion Statistics”. In: *IEEE Transactions on Image Processing* 29 (2020), pp. 5206–5215.
30. S. Yun, J. Choi, Y. Yoo, K. Yun, and J. Young Choi, “Action-decision networks for visual tracking with deep reinforcement learning,” in *IEEE Conference on Computer Vision and Pattern Recognition*, 2017, pp. 2711–2720.
31. K.H. Choi and J.E. Ha, “Edge detection based-on U-Net using edge classification CNN,” *Journal of Institute of Control, Robotics and Systems*, vol. 25, no. 8, pp. 684–689, 2019 (in Korean)

32. Yun Liu et al. "Richer convolutional features for edge detection". In: Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition. 2017, pp. 3000–3009.
33. Xavier Soria Poma, Edgar Riba, and Angel Sappa. "Dense extreme inception network: Towards a robust cnn model for edge detection". In: Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision. 2020, pp. 1923–1932.
34. Raihan F, Ce W. "PCB defect detection USING OPENCV with image subtraction method". In International Conference on Information Management and Technology, 2017, pp. 204-209.
35. Lee D.H., Chen P.Y., Yang F.J., et al. "High-Efficient Low-Cost VLSI Implementation for Canny Edge Detection". Journal of Information Science & Engineering, Vol. 36, no. 3, 2020, pp. 34-57.
36. François Chollet. "Xception: Deep learning with depthwise separable convolutions". In: Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition. 2017, pp. 1251–1258.
37. Debotosh Bhattacharjee and Hiranmoy Roy. "Pattern of local gravitational force (PLGF): A novel local image descriptor". In: IEEE Transactions on Pattern Analysis and Machine Intelligence 43.2 (2019), pp. 595–607.
38. Mengtian Li et al. "Photo-sketching: Inferring contour drawings from images". In: 2019 IEEE Winter Conference on Applications of Computer Vision (WACV). IEEE. 2019, pp. 1403-1412
39. D. Dhillon and R. Chouhan. "Noise-aided Edge preserving Image Denoising using Non-Local Means with Stochastic Resonance". In: 2018 IEEE Global Conference on Signal and Information Processing (GlobalSIP). 2018, pp. 21–25.
40. Animesh Sengupta et al. "Edge information based image fusion metrics using fractional order differentiation and sigmoidal functions". In: IEEE Access 8 (2020), pp. 88385–88398.
41. Benoit Brummer and Christophe De Vleeschouwer. "Natural image noise dataset". In: Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition Workshops. 2019, pp. 139-151.
42. Arash Akbarinia and C. Alejandro Parraga. "Feedback and surround modulated boundary detection". In: International Journal of Computer Vision 126.12 (2018), pp. 1367–1380
43. Y.X. Wang and J.M. Chen, "Iris edge detection algorithm based on adaptive canny operator and multi-directional Sobel operator," Computer and Digital Engineering, Vol. 11, No. 4, pp. 2744–2749, 2020.
44. C.W. Tian, X.C. Wang, and J.N. Yang, "Research on parallelization of kirsch operator edge detection algorithm for geological image," (in Chinese), Journal of Xinjiang University, Vol. 38, No. 1, pp. 54–60, 2021.
45. J.H. Zeng and S.J. Huang, "Comparison and analysis on typical image edge detection operators," Journal of Hebei Normal University (Natural Science), Vol. 44, No. 1, pp. 295–300, 2020.
46. S.J. Chen, X.H. Wang, Y.P. Ge, C. Li, and Y.C. Li., "Application of image edge extraction algorithm in the third land survey," Computer Technology and Development, Vol. 30, No. 10, pp. 161–166, 2020.
47. Cadena, Luis, Franklin Cadena, Nikolai Espinosa, Anna Korneeva, Alexy Kruglyakov, Alexander Legalov, Alexey Romanenko, and Alexander Zotin. (2017) "Brain's tumor image processing using shearlet transform." Proc. SPIE 10396, Applications of Digital Image Processing XL, 103961B, doi: 10.1117/12.2272792; in United States.

48. Yuksel M.E., Borlu M. Accurate Segmentation of Dermoscopic Images by Image Thresholding Based on Type-2 Fuzzy Logic. *IEEE Trans. Fuzzy Syst.* 2009; vol. 17, no. 4, p. 976–982
49. M. Bobyr, A. Arkhipov, A. Yakushev, Shade recognition of the color label based on the fuzzy clustering, *Inform. Autom.* 20 (2) (2021) 407–434, <http://dx.doi.org/10.15622/ia.2021.20.2.6>.
50. Bobyr. M.V., Emelyanov. S.G., A nonlinear method of learning neuro-fuzzy models for dynamic control systems, *Appl. Soft Comput. J.* 88 (2020) 106030, <http://dx.doi.org/10.1016/j.asoc.2019.106030>.
51. Bobyr. M.V., Milostnaya. N.A., Kulabuhov. S.A., A method of defuzzification based on the approach of areas' ratio, *Appl. Soft Comput.* 59 (2017) 19–32, <http://dx.doi.org/10.1016/j.asoc.2017.05.040>.
52. M.V. Bobyr, A.S. Yakushev, A.A. Dorodnykh, Fuzzy devices for cooling the cutting tool of the CNC machine implemented on FPGA, *Meas.: J. Int. Meas. Confed.* 152 (2020) <http://dx.doi.org/10.1016/j.measurement.2019.107378>.
53. Bobyr M.V., Milostnaya N.A., Bulatnikov V.A. The fuzzy filter based on the method of areas' ratio. *Applied Soft Computing.* 117 (2022) 108449, <https://doi.org/10.1016/j.asoc.2022.108449>
54. Bobyr, M.V., Kulabukhov, S.A. Simulation of control of temperature mode in cutting area on the basis of fuzzy logic. *Journal of Machinery Manufacture and Reliability*, 2017, 46(3), pp. 288–295. <http://dx.doi.org/10.3103/S1052618817030049>
55. Sala, F.A. Design of false color palettes for grayscale reproduction. *Displays*, 2017, 46, 9–15. <https://doi.org/10.1016/j.displa.2016.11.005>
56. Abdou, I.E., & Pratt, W.K. Quantitative design and evaluation of enhance-ment/thresholding edge detectors. *Proceedings of the IEEE*, 1979, 67(5), 753–763. doi:10.1109/proc.1979.11325

К.Н. ДУБРОВИН, А.С. СТЕПАНОВ, А.Л. ВЕРХОТУРОВ, Т.А. АСЕЕВА
**ИДЕНТИФИКАЦИЯ СЕЛЬСКОХОЗЯЙСТВЕННЫХ КУЛЬТУР С
ИСПОЛЬЗОВАНИЕМ РАДАРНЫХ ИЗОБРАЖЕНИЙ**

Дубровин К.Н., Степанов А.С., Верхотуров А.Л., Асеева Т.А. Идентификация сельскохозяйственных культур с использованием радарных изображений.

Аннотация. Одной из наиболее важных задач в практической сельскохозяйственной деятельности является идентификация сельскохозяйственных культур, произрастающих на отдельных полях в данный момент и ранее. Для снижения трудоемкости процесса идентификации в последние годы используются данные дистанционного зондирования Земли (ДЗЗ), в том числе значения индексов, рассчитываемые по ходу периода вегетации. При этом обработка оптических спутниковых снимков и получение достоверных значений индексов зачастую бывает затруднено из-за облачности во время съемки. Для решения этой проблемы в статье предложено использовать в качестве основного показателя, характеризующего сельскохозяйственную культуру, кривую сезонного хода радарного вегетационного индекса с двойной поляризацией (DpRVI). В период 2017-2020 гг. для идентификации культур на опытных полях Дальневосточного научно-исследовательского института сельского хозяйства (ДВ НИИСХ) было получено и обработано 48 радарных снимков Хабаровского муниципального района Хабаровского края со спутника Sentinel-1 (разрешение 22 м, интервал съемки – 12 дней). В качестве основных идентифицируемых культур выступали соя и овес. Также были добавлены пиксели полей, не занятых данными культурами (кормовые травы, заброшенные поля). Были получены ряды значений DpRVI как для отдельных пикселей и полей, так и аппроксимированные ряды для трех классов. Аппроксимация проводилась с использованием функции Гаусса, двойной логистической функции, квадратного и кубического полиномов. Установлено, что оптимальным алгоритмом аппроксимации является использование двойной логистической функции (средняя ошибка составила 4,6%). В среднем, ошибка аппроксимации индекса вегетации для сои не превышала 5%, для многолетних трав – 8,5%, а для овса – 11%. Для опытных полей общей площадью 303 га с известным севооборотом была проведена классификация взвешенным методом к ближайших соседей (обучающая выборка сформирована по данным 2017-2019 гг, тестовая – 2020 г.). В результате верно идентифицировано 90% полей. Общая точность классификации по пикселям составила 73%, что позволило выявить несоответствие реальных границ полей заявленным, определить заброшенные и заболоченные участки. Таким образом, установлено, что индекс DpRVI может быть использован для идентификации сельскохозяйственных культур юга Дальнего Востока и служить основой для автоматического классифицирования пахотных земель.

Ключевые слова: идентификация сельскохозяйственных культур, вегетационный индекс, дистанционное зондирование, моделирование.

1. Введение. Идентификация сельскохозяйственных культур является одной из важнейших задач в практике сельского хозяйства. Под идентификацией понимается установление тождественности неизвестной культуры известной на основании совпадения признаков. Актуальность решения этой задачи непосредственно связана как с необходимостью уточнения севооборота на отдельных полях, так и в целом оценки использования пахотных земель. При этом наземная

визуальная экспертиза – это достаточно затратное мероприятие, а в ретроспективном периоде – невозможное. Поэтому в последнее время для определения произрастающей культуры на сельскохозяйственном поле используются данные дистанционного зондирования Земли (ДЗЗ) из космоса.

Особую важность исследование этих вопросов имеет для российского Дальнего Востока: во-первых, существующие федеральные и региональные базы данных по землям сельхозназначения (ЗСН) содержат достаточное число ошибок и некорректных данных, во-вторых, арендаторы ЗСН, в том числе иностранные предприятия, зачастую предоставляют заведомо недостоверную информацию, в-третьих, для Дальнего Востока характерно наличие значительного объема заброшенных пахотных земель.

В настоящее время для решения задачи идентификации сельскохозяйственных культур применяются значения вегетационных индексов, получаемых как с использованием оптических (видимые спектры и ближний инфракрасный диапазон), так и радарных снимков. Среди индексов, получаемых первым способом, можно выделить NDVI (нормализованный разностный вегетационный индекс). Так, в работе [1] на основе 16-дневных композитов MODIS NDVI с использованием метода Decision Tree производилось выделение шаблонов севооборота и вычисление площадей, занятых сельскохозяйственными культурами в штате Мату-Гросу (Бразилия). Для заполнения пропусков в данных использовались значения соседних пикселей. В работе [2] для другого штата Бразилии уже с использованием данных более высокого разрешения (снимки Landsat-8) вычислялись средние значения NDVI для каждого поля и производилась идентификация 10 сельскохозяйственных культур методом k средних. Для определения культуры, произрастающей на поле, использовались только средние значения NDVI по этому полю. В 3 провинциях Германии была проведена автоматическая классификация пахотных земель методом Random Forest с выделением 12 сельскохозяйственных культур по значениям индекса NDVI, полученным посредством обработки изображений со спутников Landsat-8 и Sentinel-2 [3]. Использование двух различных приборов с разным разрешением и траекториями движения требует сложной калибровки и может приводить к ошибкам при вычислении вегетационного индекса. На северо-востоке Китая группа исследователей [4] изучала возможность применения различных вегетационных индексов (NDVI, PMI, NDRI и т.д.) и методов

классификации (Random Forest, Support Vector Machine, Decision Tree) для определения принадлежности пикселей полей тому или иному классу сельскохозяйственных культур. Отсутствие безоблачных снимков за август (что соответствует пику вегетации для некоторых сельскохозяйственных культур) при создании обучающей выборки заметно усложняет использование такого алгоритма на практике. В работе [5] предпринята попытка классификации пахотных земель до достижения пика вегетации путём сопоставления эвклидова расстояния между рядами NDVI отдельных пикселей и NDVI культур, сгенерированных нейронной сетью на базе рядов NDVI и EVI (Enhanced Vegetation Index) размеченных пикселей. При классификации использовались значения сразу с 3 различных спутников, то есть пиксели более высокого разрешения усреднялись и проецировались на пиксели более низкого разрешения, что вело к потере и искажению данных.

Таким образом, основной проблемой при использовании оптических снимков является наличие пропусков во временных рядах данных, вызванных влиянием атмосферной дымки, облаков и теней от них [6]. Для восстановления пропущенных значений требуется, в любом случае, либо использование данных предшествующих лет, либо обработка математическими методами (интерполяция), либо применение данных, полученных с разных спутников (что требует проецирования данных).

В качестве такого источника данных могут выступать радарные снимки. В работе [7] приведена попытка идентификации 14 культур по различным радарным данным (3 канала) как на уровне провинции Наварра в Испании, так и на уровне отдельных муниципалитетов. Для обучения модели использовались данные о границах полей и произрастаемых на них культурах. Исследователями производилась трехкратная кросс-валидация и валидация посредством полевых наблюдений. Точность классификации по данным миссии Sentinel-1 трех поляризационных каналов (VH, VV и VH/VV) не превысила 70%. Также предпринимались попытки совместного использования радарных и оптических данных для идентификации сельскохозяйственных культур [8,9]. В качестве исходных данных выступали снимки Landsat-8 (либо Sentinel-2) и Sentinel-1, а в качестве метода классификации – алгоритм Random Forest. В обоих случаях наблюдались значительные пробелы в данных, а линейная интерполяция таких пропусков сильно влияла на точность классификации.

Для идентификации культур с помощью радарных изображений чаще всего используется радарный вегетационный индекс RVI (Radar Vegetation Index [10]). RVI определяется через значение коэффициента обратного рассеяния (σ_0 , [dB]), принимаемого радиолокатором сигнала от объектов на земной поверхности. На входе используются данные изображений в проекции наклонной дальности, уровня обработки Level-1 GRD (Ground Range Detected). Во многих исследованиях показана высокая чувствительность σ_0 к динамике роста растений. Так, в работе [10] авторами отмечается, что для значений $RVI > 0.35$ биомасса кукурузного поля была выше $2,5 \text{ кг} / \text{м}^2$. Максимальные значения биомассы за весь вегетационный период составляли около $7 \text{ кг} / \text{м}^2$. Радиолокационные данные были исследованы при углах падения 35° . Аналогичные результаты показаны в работе [11]. При углах падения радиолокационного сигнала в диапазоне $35\text{--}45^\circ$ соотношение $\sigma_{0VH}/\sigma_{0VV}$ лучше коррелировало с ростом биомассы кукурузы и NDVI, в отличие от отдельных коэффициентов σ_{0VH} и σ_{0VV} . В работах [8,12] авторы указывают, что это соотношение применимо и к оценке фенологического состояния культур, характеристике растительности и их классификации. Кроме того, это позволяет разделять культуры по признакам на кукурузу, сою и подсолнух на поздних стадиях их фазы развития.

В отличие от RVI, в основе получения индекса DpRVI (Dual polarimetric Radar Vegetation Index) лежат преобразования комплексных поляриметрических радиолокационных данных уровня обработки Level-1 SLC (Single look Complex) [13]. Обработка выполняется одним из методов поляриметрической декомпозиции [14–16]. При этом происходит разложение комплексного поляриметрического отклика сигнала от объекта на составляющие, которые характеризуют вклад того или иного механизма рассеяния в общий радиолокационный сигнал (однократного, двукратного или объёмного). Информация о рассеянии рассматривается в таких терминах как: степень поляризации и мера доминирующего механизма рассеяния. Благодаря этим расчетным показателям, как показывают авторы в [13], индекс DpRVI становится более чувствительным к росту культур и применяется как относительно простой и физически интерпретируемый дескриптор растительности.

Цель работы – разработка метода идентификации сельскохозяйственных культур с использованием радарного индекса DpRVI с высоким уровнем точности. Для достижения этой цели в рамках работы решались следующие задачи: расчёт значений индексов DpRVI для пикселей полей Дальневосточного научно-

исследовательского института сельского хозяйства, аппроксимация сезонного хода индекса DpRVI различными функциями и сравнительная оценка точности, классификация пахотных земель и идентификация сельскохозяйственных культур на исследуемых полях, оценка качества классификации и идентификации.

2. Материалы и методы. В качестве исходных данных для классификации использовались аппроксимированные ряды сезонного хода DpRVI для каждого пикселя 11 опытных полей ДВ НИИСХ в период 2017-2020 гг. Опытные поля ДВ НИИСХ располагаются в Хабаровском муниципальном районе Хабаровского края между селами Мирное, Ровное и Сергеевка. Климат района характеризуется достаточно холодными зимами с малым количеством осадков и теплым, влажным летом [17]. Обилие солнечных дней и благоприятные почвенные условия делают Хабаровский район ведущим сельскохозяйственным производителем края (более 35% пахотных земель края) [18,19]. Область исследования выделена на рисунке 1.

Контуры опытных полей (shp-файлы для каждого поля) получены из Единой федеральной информационной системы о землях сельскохозяйственного назначения.

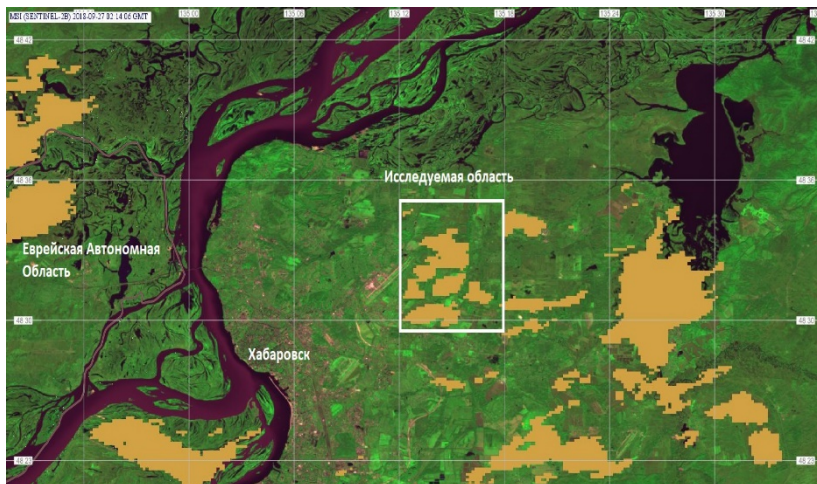


Рис. 1. Пахотные земли Хабаровского района

В качестве входных данных для расчета DpRVI использовались радиолокационные изображения спутника Sentinel-1A/B уровня обработки Level-1 SLC [20] из распределенного архива спутниковых

данных ASF DAAC (Alaska Satellite Facility Distributed Active Archive Center). На первом этапе была проведена процедура корегистрации одновременных снимков. Результатом данной процедуры является серия изображений, совмещенных с субпиксельной точностью, в которой все снимки (slave) преобразованы в геометрию заранее выбранного опорного изображения (master). Далее, ко всей серии изображений была применена операция некогерентного накопления (в англ. multilooking) [21]. Данная операция позволяет снизить уровень спекл-шума и добиться, чтобы пиксели стали «квадратного» размера. При этом пространственное разрешение для Sentinel-1 достигает 22 м. Затем, для каждой даты всего стека данных формируется ковариационная матрица C_2 . Элементы матрицы представляют собой комплексные величины, полученные на предыдущих этапах вычислений, в которых содержится вся информация о поляриметрических свойствах рассеяния исследуемой области. Все элементы матрицы дополнительно подвергаются процедуре подавления шумов при помощи адаптивного фильтра Lee [22].

$$C_2 = \begin{bmatrix} C_{11} & C_{12} \\ C_{21} & C_{22} \end{bmatrix} = \begin{bmatrix} \langle |S_{VV}|^2 \rangle & \langle |S_{VV}S_{VH}^*| \rangle \\ \langle |S_{VH}S_{VV}^*| \rangle & \langle |S_{VH}|^2 \rangle \end{bmatrix}, \quad (1)$$

где:

S_{VV} – комплексная величина радиолокационных данных канала с VV поляризацией;

S_{VH} – комплексная величина радиолокационных данных канала с VH поляризацией;

оператор $*$ – комплексное сопряжение;

оператор $\langle \rangle$ – среднее по каждому элементу матрицы.

Каждый из элементов C_2 несет в себе информацию о механизме рассеяния, а две независимых компоненты S_{VV} и S_{VH} образуют вектор рассеяния $\vec{k}$:

$$\vec{k} = [S_{VV}, S_{VH}]^T. \quad (2)$$

Далее, из элементов ковариационной матрицы индекс DpRVI был вычислен для каждой даты по формуле:

$$DpRVI = 1 - m\beta = 1 - \sqrt{\frac{4|C_2|}{(Tr(C_2))^2}} \cdot \frac{\lambda_1}{\lambda_1 + \lambda_2}, \quad (3)$$

где:

оператор Tr – сумма диагональных элементов матрицы $C_2(1)$;

$||$ – определитель матрицы C_2 ;

m – степень поляризации – отношение средней интенсивности поляризованной части волны к средней общей интенсивности волны ($0 \leq m \leq 1$);

β – мера доминирующего механизма рассеяния, которая определяется из спектрального разложения матрицы C_2 на два неотрицательных собственных значения ($\lambda_1 \geq \lambda_2 \geq 0$).

В дальнейшем для получения однородной выборки проводилась фильтрация значений DpRVI, рассчитанных для отдельных пикселей. Были удалены пиксели с аномальными значениями DpRVI (аномалии связаны либо с выходом пикселя за границу поля, либо неравномерностью посева или роста культуры на одном поле). Для этого для каждого поля на каждую дату рассчитывались средние значения индекса DpRVI. Оценка однородности выборки осуществлялась с использованием критерия “ 3σ ”.

На следующем этапе динамические ряды значений сезонного хода DpRVI, полученные для каждого пикселя, аппроксимировались с помощью квадратного (Sq) и кубического (Cube) полинома [23], двойной логистической функции (DL) [24] и функции Гаусса (Gauss) [25]. Аппроксимация применялась для сглаживания рядов DpRVI и фильтрации искажений радарных данных, возникающих при съемке.

Функция Гаусса представляет выражение вида:

$$Gauss = DpRVI_{max} e^{\frac{-(i-b)^2}{2c^2}}, \quad (4)$$

где:

i – это номер недели;

b – характеризует рост функции;

c – характеризует длительность вегетационного периода.

Двойная логистическая функция FD записывается в виде [26]:

$$DL = c_1 + c_2 * \left(\frac{1}{1 + \exp\left(\frac{a_1 - t}{a_2}\right)} - \frac{1}{1 + \exp\left(\frac{a_3 - t}{a_4}\right)} \right), \quad (5)$$

где:

c_1 – минимальное из значений DpRVI;

c_2 – размах варьирования DpRVI;

a_1 – точка перегиба кривой, где она начинает расти;

a_2 – темп роста;

a_3 – точка перегиба кривой, где она начинает идти вниз;

a_4 – темп снижения значений на кривой.

Квадратический и кубический полином имеют следующий вид:

$$Sq = ai^2 + bi + c, \quad (6)$$

$$Cube = ai^3 + bi^2 + ci + d, \quad (7)$$

где a, b, c, d – неизвестные параметры модели.

Для оценки точности аппроксимации рассчитывался показатель MAPE (Mean absolute percentage error), % – средняя абсолютная ошибка аппроксимации $DpRVI$, выраженная в процентах:

$$MAPE = \frac{1}{n} \sum_{i=1}^n \frac{|DpRVI_i^{\text{аппрокс}} - DpRVI_i^{\text{факт}}|}{DpRVI_i^{\text{факт}}} 100, \quad (8)$$

где:

n – количество измерений;

$DpRVI_i^{\text{факт}}$ – фактическое значение $DpRVI$ в неделю i ;

$DpRVI_i^{\text{аппрокс}}$ – аппроксимированное значение $DpRVI$ в неделю i .

Аппроксимированные ряды вегетационных индексов пахотных земель с известным севооборотом были разбиты на обучающую и тестовую выборку. Обучающая выборка включала в себя сезонные ряды $DpRVI$ за 2017-2019 годы для пикселей 10 опытных полей общей площадью 219 га (4180 рядов). В качестве тестовой выборки рассматривались пиксели 11 полей площадью 303 гектара (1753 ряда). Каждому пикселю обучающего множества были поставлены в соответствие метка одного из трех классов: овёс, соя и многолетние травы (сюда вошли как земли с кормовыми травами, в частности, тимopheевкой луговой, так и заброшенные или заболоченные участки). В качестве метода классификации был выбран взвешенный метод k ближайших соседей. Классификация и её оценка производились с помощью инструмента Classification Learner среды Matlab ($k = 10$, метрика близости – евклидово расстояние). Для оценки точности классификации использовались данные о произрастаемых культурах за

2020 год и значения метрик: общей точности (OA), точности каждого класса (UA), полноты для каждого класса (PA) и каппы Коэна (κ) [27].

$$OA = \frac{\sum_{j=1}^r X_{jj}}{N} * 100\%, \quad (9)$$

$$UA_j = \frac{X_{jj}}{X_{+j}} * 100\%, \quad (10)$$

$$PA_j = \frac{X_{jj}}{X_{i+}} * 100\%, \quad (11)$$

$$\kappa = \frac{N \sum_{j=1}^r X_{jj} - \sum_{j=1}^r X_{j+} X_{+j}}{N^2 - \sum_{j=1}^r X_{j+} X_{+j}}, \quad (12)$$

где:

N – общее число пикселей;

r – число классов;

X_{jj} – количество правильно классифицированных пикселей в j -м классе;

X_{j+} – общее число пикселей j -го класса для исходной карты;

X_{+j} – общее число пикселей j -го класса для классифицированной карты.

3. Результаты. В период 2017-2020 гг для Хабаровского района были получены радарные снимки Sentinel-1. На рисунке 2 для некоторых календарных дат 2017 г представлены композитные RGB изображения ($R - \sigma\theta$ VV-поляризация, $G - \sigma\theta$ VH-поляризация, $B - \sigma\theta$ отношения VV/VH-поляризации), которые показывают достаточно четкое изменение цвета опытных полей во времени, что обеспечивает возможность построения кривых сезонного хода, отражающих изменения вегетационного цикла. Помимо этого, как видно из рисунка, разным полям для каждой даты соответствуют разные оттенки цвета, что свидетельствует о возможности использования радарных индексов для идентификации отдельных культур. В период с мая по сентябрь каждого года было получено по 12 радарных снимков, для которых в дальнейшем проводился расчет значений DpRVI.

На рисунке 3 представлены усредненные значения DpRVI для полей с овсом, соей и многолетними травами, а также их аппроксимированные графики в период с 2017 по 2020 год. На графике видно, что все представленные функции достаточно хорошо отображают сезонную динамику DpRVI для всех культур.

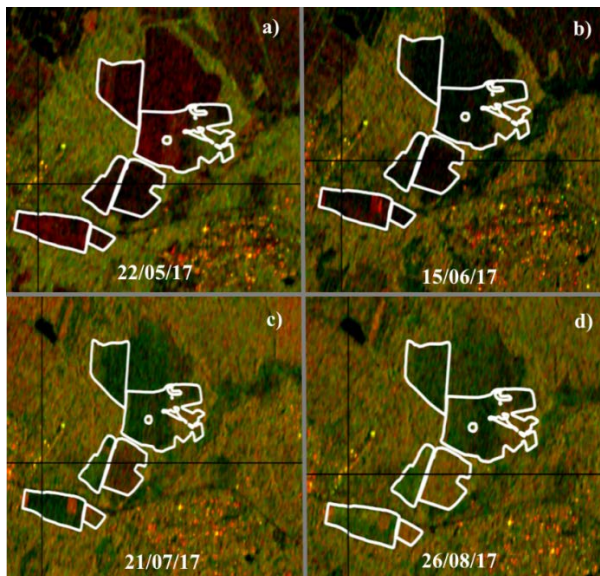


Рис. 2. Радарные композиты для области исследования (выделены контуры опытных полей), полученные в 2017 г: а) 22.05.2017, б) 15.06.2017, в) 21.07.2017, д) 26.08.2017

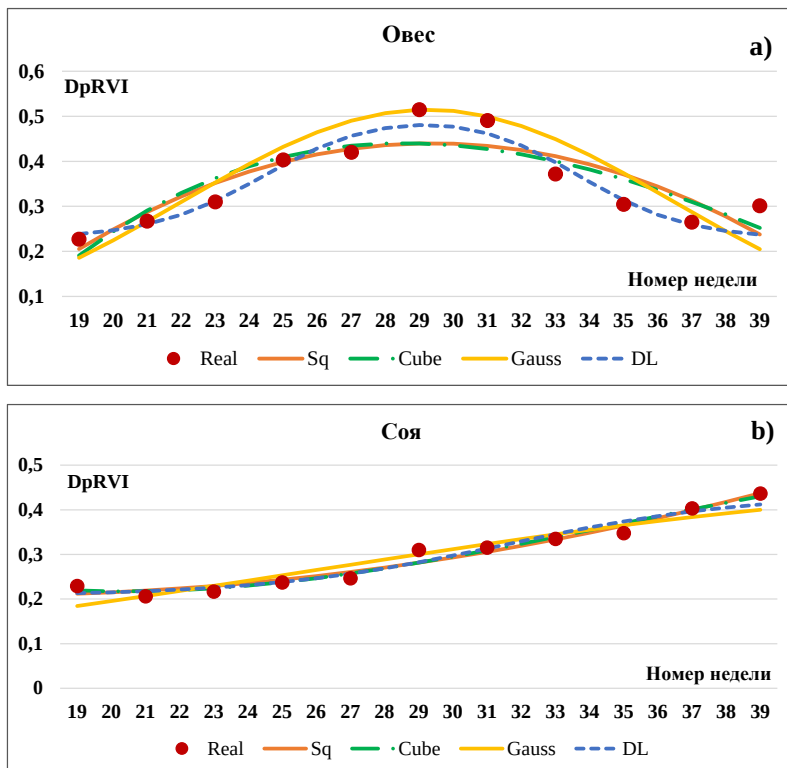
Оценка сезонного хода $DpRVI$ также позволила выявить различия в форме графиков для овса, сои и многолетних трав. Поля с соей характеризовались достаточно плавным изменением значения индекса, от 0,2 в 19-20 календарные недели, до 0,45 в 39 календарную неделю. В то же время для полей с овсом наблюдался максимум в период 29-30 календарных недель со значением 0,52, а для полей с многолетними травами – в 27-28 неделю года.

С использованием двухфакторного дисперсионного анализа достоверно установлено, что значения средней абсолютной ошибки аппроксимации значительно различаются для разных типов аппроксимирующих функций, а также для полей с разными культурами (таблица 1). Апостериорный анализ на основе критерия Фишера показал, что точность модели с использованием DL существенно выше, чем при использовании функции Гаусса и полиномов (таблица 2). Средняя абсолютная ошибка при применении DL была равной 4,6%, в то время как для функции Гаусса – 9,2%, а для квадратного и кубического полиномов – 9,3% и 8,6% соответственно.

Как видно из таблицы 3, точность аппроксимации значений $DpRVI$ для полей с соей существенно выше, чем для полей с овсом и

многолетними травами. Средняя абсолютная ошибка для полей с соей находилась на уровне 4,6%, для полей с овсом – 10,7%, многолетними травами – 8,4%.

Анализ диаграммы размаха (рисунок 4) показал, что статистически значимые значения абсолютных ошибок аппроксимации при применении DL не превышали 10% для всех культур, что соответствовало высокой точности аппроксимации. При использовании других аппроксимирующих функций для полей с соей статистически значимый диапазон также находился в пределах 0-10%, выбросов фактически не наблюдалось. Для значений DpRVI, характеризовавших поля с овсом, отмечалось увеличение статистически значимого диапазона при применении функции Гаусса и полиномов – до 20%. Аппроксимации значений DpRVI для полей с многолетними травами соответствовало большее число выбросов с абсолютным значением 30 - 50%, что могло быть вызвано неоднородностью этого класса.



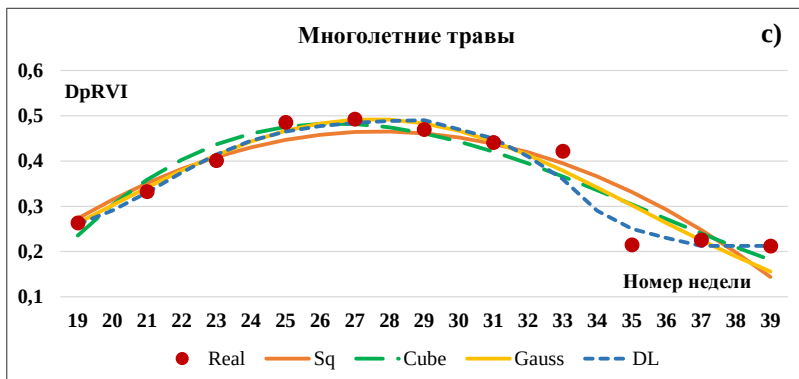


Рис. 3. Аппроксимация значений DpRVI разными функциями для пахотных земель Хабаровского района (19-39 календарные недели, 2017-2020 гг.): а) овес, б) соя, с) многолетние травы

Таблица 1. Средние значения ошибки аппроксимации разными функциями сезонного хода индекса DpRVI для сои, овса и многолетних трав (Хабаровский район, 2017-2020 гг.), %

С/х культура	Аппроксимирующая функция					Р
	Квадрат. полином	Куб. полином	Функция Гаусса	Двойная логист. функция	Всего	
Овес	11,9	12,1	12,9	6,0	10,7	P<0,05
Соя	4,2	3,5	6,5	4,4	4,6	
Многолет. травы	11,8	10,3	8,3	3,3	8,4	
Всего	9,3	8,6	9,2	4,6		
Р	P<0,05					

Таблица 2. Р-значения апостериорного критерия Фишера при попарном сравнении точности аппроксимации сезонного хода DpRVI четырьмя функциями

Аппроксимирующая функция	Квадратный полином	Кубический полином	Функция Гаусса	Двойная логистическая функция
Квадратный полином	-	0,742	0,959	0,024 (P<0,05)
Кубический полином	0,742	-	0,781	0,048 (P<0,05)
Функция Гаусса	0,959	0,781	-	0,027 (P<0,05)
Двойная логистическая функция	0,024 (P<0,05)	0,048 (P<0,05)	0,027 (P<0,05)	-

Таблица 3. Р-значения апостериорного критерия Фишера при попарном сравнении точности аппроксимации сезонного хода DpRVI для трех сельскохозяйственных культур

Сельскохозяйственная культура	Овес	Соя	Многолетние травы
Овес	-	0,001 ($P < 0,05$)	0,209
Соя	0,001 ($P < 0,05$)	-	0,036 ($P < 0,05$)
Многолетние травы	0,209	0,036 ($P < 0,05$)	-

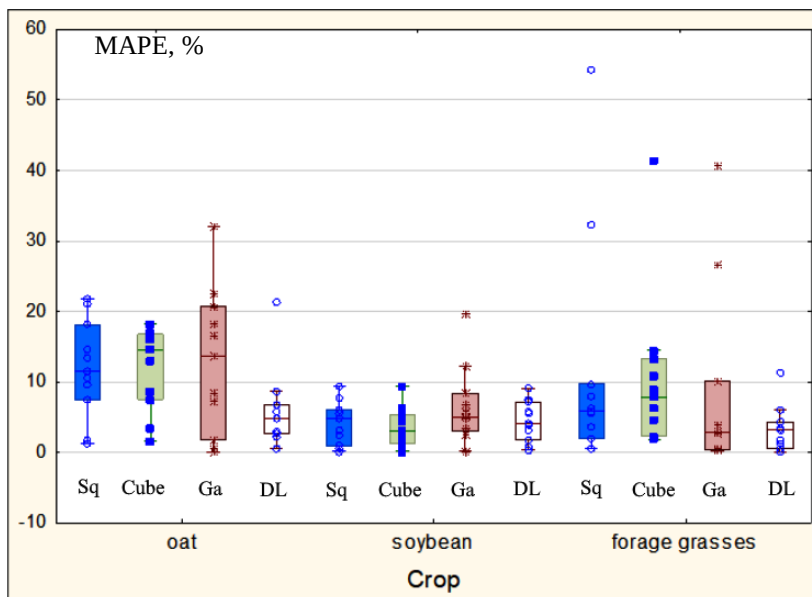
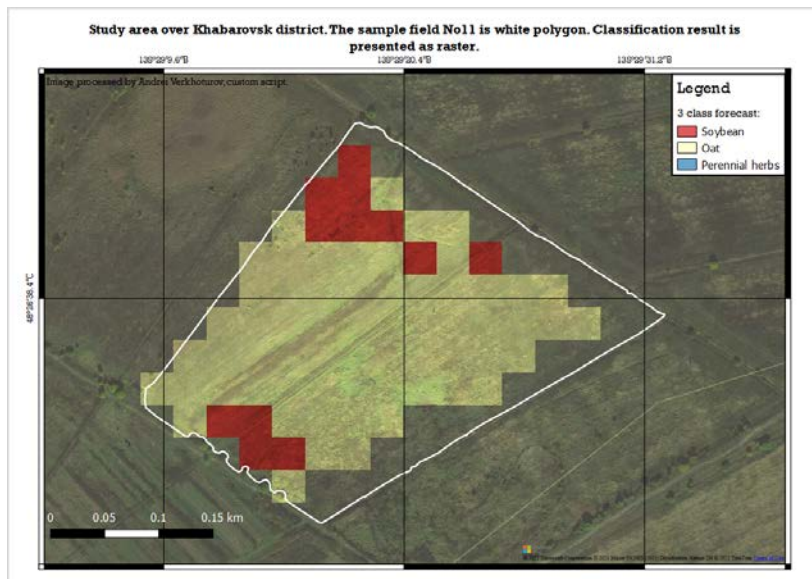


Рис. 4. Диаграмма размаха абсолютной ошибки аппроксимации значений DpRVI для сельскохозяйственных культур полей Хабаровского района в 2017-2020 гг.

На следующем этапе производилась идентификация произрастаемых культур на уровне отдельных полей. Для каждого поля подсчитывалось количество пикселей, отнесённых к тому или иному классу. Решение о том, какая культура произрастала на поле, принималось на основании того, пиксели какого класса наиболее представлены на данном поле.

На рисунке 5 представлены два поля, вошедших в тестовую выборку. Как видно из рисунка, поле с соей (рисунок 5a) идентифицировано успешно – вкраплений пикселей других классов почти нет. Поле с овсом (рисунок 5b) также идентифицировано точно, в то же время пиксели на поле с многолетними травами (рисунок 6a) были определены с меньшей точностью – связано это в том числе и с неоднородностью растительного состава внутри данного класса.

На рисунке 6b представлено тестовое поле номер 13, которое было заявлено в 2020 году как поле с овсом. Как видно, однозначно определяется как овес совокупность пикселей западной половины поля, в то время как восточная часть не является однородной. Было установлено, что восточная часть поля заболочена и фактически не возделывается. Такие же заброшенные или заболоченные участки, в основном прилегающие к очерченным границам, были выявлены и для других полей.



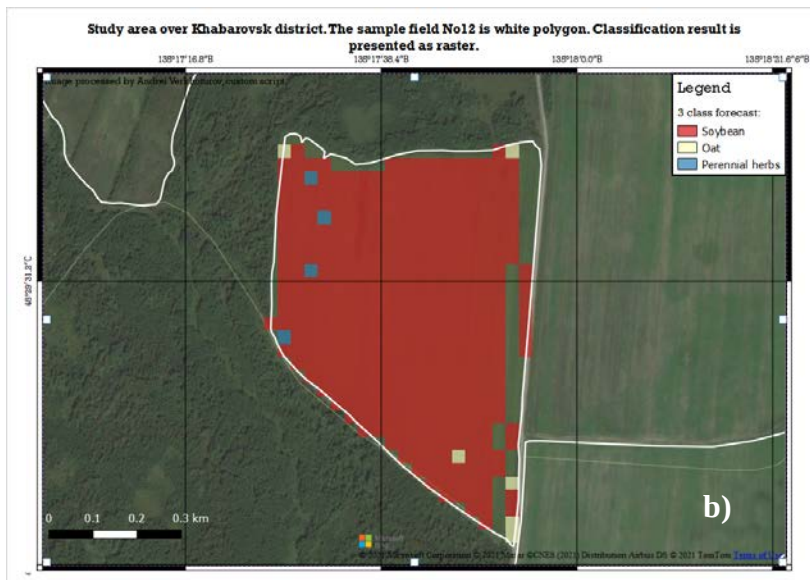
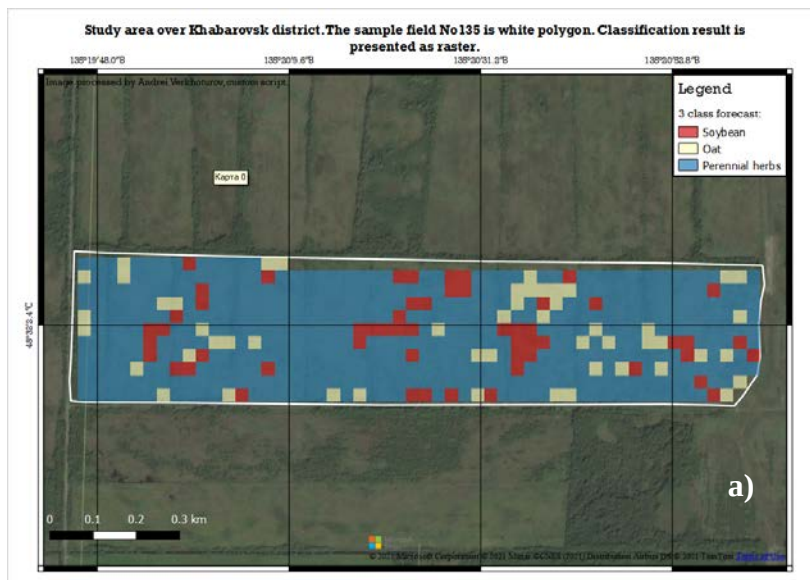


Рис. 5. Результаты классификации тестовой выборки: а) поле с соей, б) поле с овсом



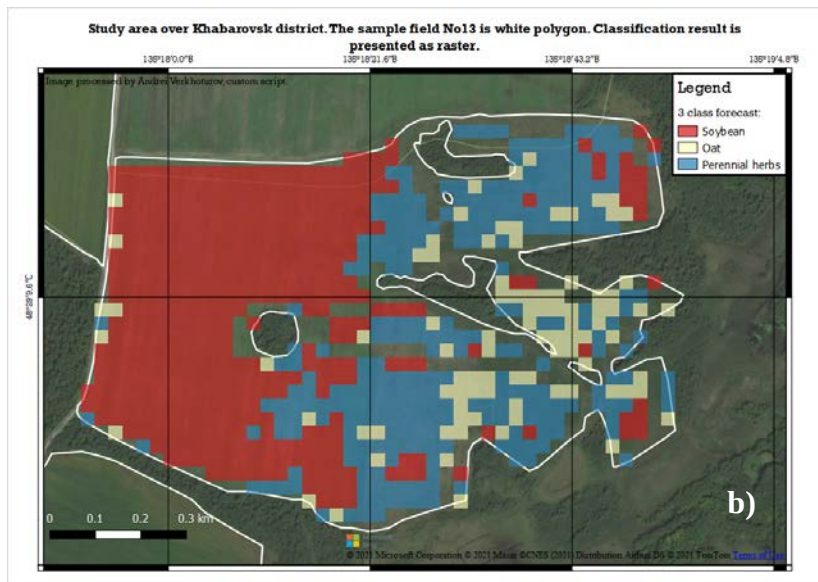


Рис. 6. Результаты классификации тестовой выборки: а) поле с многолетними травами, б) поле с овсом с заболоченным участком

Было верно идентифицировано 90% тестовых полей. Результаты попиксельной классификации ($\kappa = 0,5$, $OA = 73\%$, $UA_1 = 84\%$, $UA_2 = 61\%$) свидетельствовали: во-первых, о возможности использования предложенного метода для предварительной идентификации культур и выявления несоответствий в севообороте; во-вторых, о необходимости уточнения границ и построения новых контуров для полей перед проведением классификации ЗСН на районном или региональном уровне.

4. Заключение. Таким образом, в результате проведенных исследований установлено, что классификацию пахотных земель на Дальнем Востоке возможно проводить с использованием аппроксимированных рядов значений радарного индекса (DpRVI). В данной работе впервые для решения задачи идентификации сельскохозяйственных культур были применены сезонные ряды индекса DpRVI. Для устранения разреженности в данных, вызванных искажениями при получении снимков и обработке спутниковых данных, также была впервые предпринята аппроксимация сезонного хода индекса DpRVI. Значения DpRVI были рассчитаны по радарным снимкам спутника Sentinel-1 в период вегетации сельскохозяйственных культур Дальнего Востока (май-сентябрь

2017-2020 гг.). Для аппроксимации индексов DpRVI целесообразно использовать двойную логистическую функцию: средняя ошибка аппроксимации двойной логистической функцией составила 4,6%, ошибка аппроксимации функцией Гаусса, квадратным и кубическим полиномом превысила 8,5%. Проведенная классификация пахотных земель Хабаровского района для произрастающих культур сои, овса и многолетних трав показала точность определения 90% для тестовых полей. Вместе с тем, общая точность классификации на уровне отдельных пикселей составила 73%, что объясняется неоднородностью некоторых полей и наличием неиспользуемых участков. Проведенный анализ продемонстрировал, что для каждого из классов сезонный ход индекса DpRVI имеет отличительные особенности, и ряды значений индекса могут быть использованы для проверки соответствия реальной и заявленной культуры на отдельном поле, уточнения границ полей, поиска заболоченных и заброшенных участков.

Литература

1. Mapping croplands, cropping patterns, and crop types using MODIS time-series data / Y. Cheng [и др.] // *International Journal of Applied Earth Observation and Geoinformation*. 2018. vol. 69. pp. 133-147.
2. Improved regional-scale Brazilian cropping systems' mapping based on a semi-automatic object-based clustering approach / B. Bellon [и др.] // *International Journal of Applied Earth Observation and Geoinformation*. 2018. vol. 68. pp. 127-138.
3. Griffiths P., Nendel C., Hostert P. Intra-annual reflectance composites from Sentinel-2 and Landsat for national-scale crop and land cover mapping // *Remote Sensing of Environment*. 2019. vol. 220. pp. 135-151.
4. Accessing the temporal and spectral features in crop type mapping using multi-temporal Sentinel-2 imagery: A case study of Yi'an County, Heilongjiang province, China / H. Zhang [и др.] // *Computers and Electronics in Agriculture*. 2020. vol. 176. 105618.
5. Early-season crop type mapping using 30-m reference time series / P. Hao [и др.] // *Journal of Integrative Agriculture*. 2020. vol. 19. iss. 7. pp. 1897-1911.
6. Миклашевич Т.С., Барталев С.А., Плотников Д.Е. Интерполяционный алгоритм восстановления длинных временных рядов данных спутниковых наблюдений растительного покрова // *Современные проблемы дистанционного зондирования Земли из космоса*. 2019. Т. 16. №6. С. 143-154.
7. Arias M., Campo-Bescós M.Á, Álvarez-Mozos J. Crop Classification Based on Temporal Signatures of Sentinel-1 Observations over Navarre Province, Spain // *Remote Sensing*. 2020. vol. 12. iss. 2. 278.
8. Improved Early Crop Type Identification by Joint Use of High Temporal Resolution SAR And Optical Image Time Series / J. Inglada [и др.] // *Remote Sensing*. 2016. vol. 8. iss. 5. 362.
9. Synergistic Use of Radar Sentinel-1 and Optical Sentinel-2 Imagery for Crop Mapping: A Case Study for Belgium / van Tricht K. [и др.] // *Remote Sensing*. 2018. vol. 10. iss. 10. 1642.
10. Kim Y., van Zyl J.J. A Time-Series Approach to Estimate Soil Moisture Using Polarimetric Radar Data // *IEEE Transactions on Geoscience and Remote Sensing*. 2009. vol. 47. №8. pp. 2519-2527.

11. C-band polarimetric indexes for maize monitoring based on a validated radiative transfer model / X. Blaes [и др.] // IEEE Transactions on Geoscience and Remote Sensing. 2006. vol. 44. iss. 4. pp. 791–800.
12. Integration of optical and Synthetic Aperture Radar (SAR) imagery for delivering operational annual crop inventories / H. McNairn [и др.] // ISPRS Journal of Photogrammetry and Remote Sensing. 2009. vol. 64. iss. 5. pp. 434–449.
13. Dual polarimetric radar vegetation index for crop growth monitoring using Sentinel-1 SAR data / D. Mandal [и др.] // Remote Sensing of Environment. 2020. vol. 247. 111954.
14. Freeman A., Durden S.L. A Three-Component Scattering Model for Polarimetric SAR Data // IEEE Transactions on Geoscience and Remote Sensing. 1998. vol. 36. iss. 3. pp. 963–973.
15. Four Component Scattering Model for Polarimetric SAR Image Decomposition / Yamaguchi Y. [и др.] // IEEE Transactions on Geoscience and Remote Sensing. 2005. vol. 43. iss. 8. pp. 1699–1706.
16. Arii M., van Zyl J.J., Kim Y. Adaptive Model-Based Decomposition for Polarimetric SAR Covariance Matrices // IEEE Transactions on Geoscience and Remote Sensing. 2011. vol. 49. iss. 3. pp. 1104–1113.
17. Костенков Н.М., Ознобихин В.И. Почвы и почвенные ресурсы юга Дальнего Востока, и их оценка // Почвоведение. 2006. №5. С. 517–526.
18. Новороцкий П.В. Климатические изменения в бассейне Амура за последние 115 лет // Метеорология и гидрология. 2007. №2. С. 43–53.
19. База данных показателей муниципальных образований. URL: www.gks.ru/dbscripts/munst/ (дата обращения: 21.08.2021).
20. Sentinel-1 Mission Status / P. Potin [и др.] // 11th European Conference on Synthetic Aperture Radar. Proceedings EUSAR. 2016. pp. 59–64.
21. Intensity and phase statistics of multilook polarimetric interferometric SAR imagery / J.S. Lee [и др.] // IEEE Transactions on Geoscience and Remote Sensing. 1994. vol. 32. iss. 5. pp. 1017–1028.
22. Lee J.S., Potter E. Polarimetric SAR Radar Imaging: From Basic to Applications // Boca Raton: CRC Press. 2009. 438 p.
23. Predicting the Normalized Difference Vegetation Index (NDVI) by training a crop growth model with historical data / A. Berger [и др.] // Computers and Electronics in Agriculture. 2018. vol. 161. pp. 305–311.
24. An improved logistic method for detecting spring vegetation phenology in grasslands from MODIS EVI time-series data. / R. Cao [и др.] // Agric. For. Meteorol. 2015. vol. 200. pp. 9–20.
25. Predicting Soybean Yield at the Regional Scale Using Remote Sensing and Climatic Data / A. Stepanov [и др.] // Remote Sensing. 2020. vol. 12. iss. 12. 1936.
26. Evaluating the impacts of models, data density and irregularity on reconstructing and forecasting dense Landsat time series. / J. Zhang [и др.] // Science of Remote Sensing. 2021. №4. 100023.
27. Mapping crops within the growing season across the United States / V.S. Konduri [и др.] // Remote Sensing of Environment. 2020. vol. 251. 112048.

Дубровин Константин Николаевич — младший научный сотрудник, лаборатория численных методов математической физики, Вычислительный центр Дальневосточного отделения Российской академии наук (ВЦ ДВО РАН). Область научных интересов: математическое моделирование, машинное обучение, применение методов дистанционного зондирования в сельском хозяйстве. Число научных публикаций — 14. nobforward@gmail.com; ул. Ким Ю Чена, 65, 680000, Хабаровск, Россия; р.т.: +7(909)859-0881.

Степанов Алексей Сергеевич — д-р фармацевт. наук, ведущий научный сотрудник, лаборатория селекции зерновых и колосовых культур, Дальневосточный научно-исследовательский институт сельского хозяйства (ДВ НИИСХ). Область научных интересов: сельскохозяйственная экономика, прогнозирование урожайности сельскохозяйственных культур, математическое моделирование в сельском хозяйстве. Число научных публикаций — 80. stepanfx@mail.ru; ул. Клубная, 13, 680521, Восточное, Россия; р.т.: 8(924)210-9102.

Верхотуров Андрей Леонидович — старший научный сотрудник, лаборатория цифровых методов исследования природных и технических систем, Институт горного дела Дальневосточного отделения Российской академии наук (ИГД ДВО РАН). Область научных интересов: обработка данных дистанционного зондирования Земли, обработка данных радиолокационной интерферометрии, геоинформационные системы. Число научных публикаций — 26. andrey@ccfebras.net; ул. Тургенева, 51, 680000, Хабаровск, Россия; р.т.: +7(4212)703-913.

Асеева Татьяна Александровна — д-р с.-х. наук, член-корреспондент РАН, директор, Дальневосточный научно-исследовательский институт сельского хозяйства (ДВ НИИСХ). Область научных интересов: сельское хозяйство, агротехнологии, земледелие, селекция. Число научных публикаций — 290. aseeva59@mail.ru; ул. Клубная, 13, 680521, Восточное, Россия; р.т.: 8(421)249-7546.

Поддержка исследований. Исследования проведены с использованием ресурсов Центра коллективного пользования научным оборудованием «Центр обработки и хранения научных данных ДВО РАН», финансируемого Российской Федерацией в лице Минобрнауки России по соглашению № 075-15-2021-663.

K. DUBROVIN, A. STEPANOV, A. VERKHOTUROV, T. ASEEEVA
CROP IDENTIFICATION USING RADAR IMAGES

Dubrovin K., Stepanov A., Verkhoturov A., Aseeva T. **Crop Identification Using Radar Images.**

Abstract. One of the most important tasks in practical agricultural activity is the identification of agricultural crops, both those growing in individual fields at the moment and those that grew in these fields earlier. To reduce the complexity of the identification process in recent years, data from remote sensing of the Earth (remote sensing), including the values of vegetation indices calculated during the growing season, have been used. At the same time, processing optical satellite images and obtaining reliable index values is often difficult, which is due to cloud cover during the shooting. To solve this problem, the article suggests using the seasonal course curve of the radar vegetation index with double polarization (DpRVI) as the main indicator characterizing agricultural crops. In the period 2017-2020, 48 radar images of the Khabarovsk Municipal District of the Khabarovsk Territory from the Sentinel-1 satellite were received and processed to identify crops in the experimental fields of the Far Eastern Research Institute of Agriculture (FEARI) (resolution 22 m, shooting interval - 12 days). Soybeans and oats were the main identified crops. Pixels of fields not occupied by these crops (forage grasses, abandoned fields) were also added. The series of values of DpRVI were obtained both for individual pixels and fields, and approximated series for three classes. The approximation was carried out using the Gaussian function, the double logistic function, the square and cubic polynomials. It is established that the optimal approximation algorithm is the use of a double logistic function (the average error was 4.6%). On average, the approximation error of the vegetation index for soybeans did not exceed 5%, for perennial grasses – 8.5%, and for oats - 11%. For experimental fields with a total area of 303 hectares with a known crop rotation, the classification was carried out by the weighted method of k nearest neighbors (the training sample was formed according to the data of 2017-2019, the test sample -2020). As a result, 90% of the fields were correctly identified, and the overall pixel classification accuracy was 73%, which made it possible to identify the discrepancy between the actual boundaries of the fields declared to identify abandoned and swampy areas. Thus, it is established that the DpRVI index can be used to identify agricultural crops in the south of the Far East and serve as the basis for the automatic classification of arable land.

Keywords: crop identification, vegetation index, remote sensing, modelling.

Dubrovin Konstantin — Junior researcher, Laboratory of numerical methods of mathematical physics, Computing Center of the Far Eastern Branch of the Russian Academy of Sciences. Research interests: mathematical modeling, machine learning, application of remote sensing methods in agriculture. The number of publications — 14. nobforward@gmail.com; 65, Kim Yu Chen St., 680000, Khabarovsk, Russia; office phone: +7(909)859-0881.

Stepanov Alexey — Ph.D., Dr.Sci., Leading researcher, laboratory of grain and ear crops breeding, Far Eastern Agriculture Research Institute of the Russian Academy of Sciences (FEARI). Research interests: agricultural economics, crop yield forecasting, mathematical modeling in agriculture. The number of publications — 80. stepanfx@mail.ru; 13, Klubnaya St., 680521, Vostochnoe, Russia; office phone: 8(924)210-9102.

Verkhoturov Andrey — Senior researcher, Laboratory of digital methods for the study of natural and technical systems, Mining Institute of the Far Eastern Branch of the Russian Academy of Sciences (MI FEB RAS). Research interests: processing of Earth remote sensing

data, processing of radar interferometry data, geoinformation systems. The number of publications — 26. andrey@ccfebras.net; 51, Turgeneva St., 680000, Khabarovsk, Russia; office phone: +7(4212)703-913.

Aseeva Tatiana — Ph.D., Dr.Sci., Corresponding member, Director, Far Eastern Agriculture Research Institute of the Russian Academy of Sciences (FEARI). Research interests: agriculture, agrotechnology, agriculture, breeding. The number of publications — 290. aseeva59@mail.ru; 13, Klubnaya St., 680521, Vostochnoe, Russia; office phone: 8(421)249-7546.

Acknowledgements. The studies were carried out using the resources of the Center for Shared Use of Scientific Equipment «Center for Processing and Storage of Scientific Data of the Far Eastern Branch of the Russian Academy of Sciences», funded by the Russian Federation represented by the Ministry of Science and Higher Education of the Russian Federation under project No. 075-15-2021-663.

References

1. Mapping croplands, cropping patterns, and crop types using MODIS time-series data / Y. Cheng [et al.] // *International Journal of Applied Earth Observation and Geoinformation*. 2018. vol. 69. pp. 133-147.
2. Improved regional-scale Brazilian cropping systems' mapping based on a semi-automatic object-based clustering approach / B. Bellon [et al.] // *International Journal of Applied Earth Observation and Geoinformation*. 2018. vol. 68. pp. 127-138.
3. Griffiths P., Nendel C., Hostert P. Intra-annual reflectance composites from Sentinel-2 and Landsat for national-scale crop and land cover mapping // *Remote Sensing of Environment*. 2019. vol. 220. pp. 135-151.
4. Accessing the temporal and spectral features in crop type mapping using multi-temporal Sentinel-2 imagery: A case study of Yi'an County, Heilongjiang province, China / H. Zhang [et al.] // *Computers and Electronics in Agriculture*. 2020. vol. 176. 105618.
5. Early-season crop type mapping using 30-m reference time series / P. Hao [et al.] // *Journal of Integrative Agriculture*. 2020. vol. 19. iss. 7. pp. 1897-1911.
6. Miklashevich, T., Bartalev S., Plotnikov D. [Interpolation algorithm for the recovery of long satellite data time series of vegetation cover observation]. *Sovremennye problemy distantsionnogo zondirovaniya Zemli iz kosmosa* - Current problems in remote sensing of the Earth from space. 2019. vol. 16. pp. 143-154. (in Russ.).
7. Arias M., Campo-Bescós MÁ, Álvarez-Mozos J. Crop Classification Based on Temporal Signatures of Sentinel-1 Observations over Navarre Province, Spain // *Remote Sensing*. 2020. vol. 12. iss. 2. 278.
8. Improved Early Crop Type Identification by Joint Use of High Temporal Resolution SAR And Optical Image Time Series / J. Inglada [et al.] // *Remote Sensing*. 2016. vol. 8. iss. 5. 362.
9. Synergistic Use of Radar Sentinel-1 and Optical Sentinel-2 Imagery for Crop Mapping: A Case Study for Belgium / van Tricht K. [et al.] // *Remote Sensing*. 2018. vol. 10. iss. 10. 1642.
10. Kim Y., van Zyl J.J. A Time-Series Approach to Estimate Soil Moisture Using Polarimetric Radar Data // *IEEE Transactions on Geoscience and Remote Sensing*. 2009. vol. 47. №8. pp. 2519-2527.
11. C-band polarimetric indexes for maize monitoring based on a validated radiative transfer model / X. Blaes [et al.] // *IEEE Transactions on Geoscience and Remote Sensing*. 2006. vol. 44. iss. 4. pp. 791-800.

12. Integration of optical and Synthetic Aperture Radar (SAR) imagery for delivering operational annual crop inventories / H. McNairn [et al.] // ISPRS Journal of Photogrammetry and Remote Sensing. 2009. vol. 64. iss. 5. pp. 434–449.
13. Dual polarimetric radar vegetation index for crop growth monitoring using Sentinel-1 SAR data / D. Mandal [et al.] // Remote Sensing of Environment. 2020. vol. 247. 111954.
14. Freeman A., Durden S.L. A Three-Component Scattering Model for Polarimetric SAR Data // IEEE Transactions on Geoscience and Remote Sensing. 1998. vol. 36. iss. 3. pp. 963-973.
15. Four Component Scattering Model for Polarimetric SAR Image Decomposition / Yamaguchi Y. [et al.] // IEEE Transactions on Geoscience and Remote Sensing. 2005. vol. 43. iss. 8. pp. 1699-1706.
16. Ariei M., van Zyl J.J., Kim Y. Adaptive Model-Based Decomposition for Polarimetric SAR Covariance Matrices // IEEE Transactions on Geoscience and Remote Sensing. 2011. vol. 49. iss. 3. pp. 1104-1113.
17. Kostenkov N.M., Oznobikhin V.I. Soils and soil resources in the southern Far East and their assessment // Eurasian Soil Sc. 2006 vol. 39. pp. 461–469.
18. Novorotskii P.V. Climate changes in the Amur River basin in the last 115 years // Russian Meteorology and Hydrology. 2007. vol. 32. №2. pp. 102-109.
19. Baza dannyh pokazatelej municipal'nyh obrazovaniy [Database of indicators of municipalities]. Available at www.gks.ru/dbscripts/munst/ (acssesed 21.08.2021). (in Russ.)
20. Sentinel-1 Mission Status / P. Potin [et al.] // 11th European Conference on Synthetic Aperture Radar. Proceedings EUSAR. 2016. pp. 59–64.
21. Intensity and phase statistics of multilook polarimetric interferometric SAR imagery / J.S. Lee [et al.] // IEEE Transactions on Geoscience and Remote Sensing. 1994. vol. 32. iss. 5. pp. 1017-1028.
22. Lee J.S., Pottier E. Polarimetric SAR Radar Imaging: From Basic to Applications // Boca Raton: CRC Press. 2009. 438 p.
23. Predicting the Normalized Difference Vegetation Index (NDVI) by training a crop growth model with historical data / A. Berger [et al.] // Computers and Electronics in Agriculture. 2018. vol. 161. pp. 305-311.
24. An improved logistic method for detecting spring vegetation phenology in grasslands from MODIS EVI time-series data. / R. Cao [et al.] // Agric. For. Meteorol. 2015. vol. 200. pp. 9–20.
25. Predicting Soybean Yield at the Regional Scale Using Remote Sensing and Climatic Data / A. Stepanov [et al.] // Remote Sensing. 2020. vol. 12. iss.12. 1936.
26. Evaluating the impacts of models, data density and irregularity on reconstructing and forecasting dense Landsat time series. / J. Zhang [et al.] // Science of Remote Sensing. 2021. №4. 100023.
27. Mapping crops within the growing season across the United States / V.S. Konduri [et al.] // Remote Sensing of Environment. 2020. vol. 251. 112048.

А.С. ШАУРА, А.Г. ЗЛОБИНА, И.В. ЖУРБИН, А.И. БАЖЕНОВА
**АНАЛИЗ ДАННЫХ РАЗНОВРЕМЕННОЙ
МУЛЬТИСПЕКТРАЛЬНОЙ АЭРОФОТОСЪЕМКИ ДЛЯ
ОБНАРУЖЕНИЯ ГРАНИЦ ИСТОРИЧЕСКОГО
АНТРОПОГЕННОГО ВОЗДЕЙСТВИЯ**

Шаура А.С., Злобина А.Г., Журбин И.В., Баженова А.И. Анализ данных разновременной мультиспектральной аэрофотосъемки для обнаружения границ исторического антропогенного воздействия.

Аннотация. В работе представлено применение алгоритма статистического анализа данных разновременной мультиспектральной аэрофотосъемки с целью выявления участков исторического антропогенного воздействия на природную среду. Исследуемый участок расположен на окраине поселка городского типа Знаменка (Знаменский район Тамбовской области) в лесостепной зоне с типичными черноземными почвами, где во второй половине XIX – начале XX вв. были расположены пашни. Признаком для выявления следов исторического антропогенного воздействия может быть растительность, возникшая в результате вторичной сукцессии на заброшенных участках. Отличительной особенностью такой растительности от окружающей природной среды является ее тип, возраст и плотность произрастания. Таким образом, задача обнаружения границ антропогенного воздействия по мультиспектральным изображениям сводится к задаче классификации растительности. Исходными данными являлись результаты разновременной мультиспектральной съемки в зеленом (Green), красном (Red), краевом красном (RedEdge) и ближнем инфракрасном (NIR) спектральных диапазонах. На первом этапе алгоритма предполагается вычисление текстурных признаков Харалика по данным мультиспектральной съемки, на втором этапе – уменьшение количества признаков методом главных компонент, на третьем – сегментация изображений на основе полученных признаков методом k-means. Эффективность предложенного алгоритма показана при сопоставлении результатов сегментации с эталонными данными исторических картографических материалов. Полученный результат сегментации отражает не только конфигурацию участков антропогенно-преобразованной природной среды, но и особенности зарастания заброшенной пашни, поскольку исследование разновременных мультиспектральных снимков позволяет более полно охарактеризовать и учесть динамику наращивания фитомассы в разные периоды вегетации.

Ключевые слова: мультиспектральная съемка, текстурная сегментация, признаки Харалика, метод главных компонент, кластеризация, k-means, разновременные данные, период вегетации, вторичная сукцессия.

1. Введение. Исследование природной среды регионов, где современные ландшафты сложились под влиянием хозяйственного воздействия, становится заметным явлением в современной науке. Различные виды исторического природопользования вызывают долговременные изменения растительности, почв и гидрографии (например, [1–3]). Эти проблемы в губерниях средней полосы России привлекли пристальное внимание еще в конце XIX – начале XX в. [4]. В настоящее время для некоторых регионов, в частности для Тамбовской области (бывш. Тамбовской губернии), на основании разнообразных исто-

рических источников конкретизированы некоторые причины сложившейся экологической ситуации [5–7]. В связи с этим, поиск участков сплошных рубок леса, пашен и сенокосов второй половины XIX – начала XX вв. необходим для реконструкции системы хозяйства в исторической ретроспективе и использования исторического опыта для определения перспективных направлений в освоении территории.

Необходимо отметить, что следы исторического антропогенного воздействия неоднозначно выражены в современном ландшафте. Одним из признаков заброшенных сельскохозяйственных угодий могут являться участки растительности, сформировавшейся в результате вторичной сукцессии [8, с. 225–257; 9].

Апробация метода выявления участков антропогенно-преобразованной природной среды проведена на материалах мультиспектральной съемки ныне заброшенной пашни имения Знаменское (совр. поселок городского типа Знаменка, Знаменский район Тамбовской области) [10]. Усадебная часть имения расположена в лесостепной зоне с типичными черноземными почвами. Известно, что в лесостепи восстановление профиля чернозема происходит в течение 30–40 лет после завершения эксплуатации возделываемых полей [9, 11]. Далее, при отсутствии блокирующих процессов (пожары, выпасы скота и др.), происходит полное восстановление растительности, типичной для данной местности, – в большинстве случаев, разнотравья и кустарников. Этот процесс занимает в среднем 50–60 лет с момента прекращения хозяйственной деятельности. В некоторых случаях продолжительное отсутствие антропогенного воздействия на природную среду (более 100 лет) способствует формированию участков достаточно плотного произрастания древовидных пород, преобладающих в лесостепной зоне [12]. Таким образом, признаком для выявления следов сплошных рубок леса и пашен второй половины XIX – начала XX вв. может быть растительность, возникшая в результате вторичной сукцессии на заброшенных участках. Отличительной особенностью такой растительности от характеристик окружающей природной среды является ее тип, возраст и плотность произрастания [8, с. 250–257], что отображается в текстуре соответствующих фрагментов мультиспектральных изображений. Таким образом, задача обнаружения границ антропогенного воздействия по мультиспектральным изображениям сводится к задаче классификации растительности на территории обследования.

Использование мультиспектральных данных в задачах сегментации изображений поверхности Земли обусловлено тем, что разные ландшафтные объекты (в том числе и растительность) отличаются

спектральными характеристиками. Кроме того, отражательная способность растительности существенно изменяется на разных этапах вегетационного периода. Эта особенность наиболее ярко проявляется у лиственных деревьев и кустарников, а также у травянистой растительности. Поэтому анализ разновременных снимков позволяет выявлять изменения растительного покрова и используется для мониторинга состояния природной среды и сельскохозяйственных угодий [13–16]. Не менее важно, комплексный анализ мультиспектральных данных, полученных в разные периоды вегетации, обеспечивает повышение достоверности классификации участков растительности [17, 18]. Высокая эффективность такого подхода доказана для различных природно-климатических зон. В частности, при распознавании различных типов пойменной растительности в южной части Нидерландов анализ комплекса мультиспектральных снимков позволил достигнуть высоких показателей точности классификации (6 классов с общей точностью более 90%) [19]. Принципиально важно, для схожих классов растительности (луга и травянистая растительность) было доказано существенное повышение точности классификации в сравнении с результатами, полученными без учета динамики процесса вегетации. Схожие результаты были получены при классификации участков хозяйственной деятельности на юге Португалии [20]. Показано, что совместный анализ материалов разновременной мультиспектральной съемки необходим для достижения удовлетворительной точности картографирования участков открытых почв, кустарников, полей, засеянных зерновыми культурами, а также лесных массивов и лугов (общая точность классификации – 74,5%). Необходимость использования набора мультиспектральных снимков также продемонстрирована для тропических и засушливых регионов центральной и южной Африки, где особенно ярко проявляется сезонность растительности [21]. Для 12 классов ландшафтных объектов, включающих древесную, кустарниковую и травянистую растительность с различной плотностью произрастания, а также открытые пространства, была достигнута общая точность классификации более 90%.

Таким образом, применение разновременной мультиспектральной съемки позволяет учесть динамику изменения спектральных отражательных свойств растительности. Это играет существенную роль при классификации схожих ландшафтных объектов: луга и заброшенные поля, покрытые травянистой растительностью; участки разноразрастной кустарниковой растительности. Для выявления участков хозяйственного использования второй половины XIX – начала XX вв.

исходными данными являются мультиспектральные изображения в период ранней вегетации (май) и развитой растительности (июль).

2. Исходные и эталонные данные. Участок обследования расположен на окраине поселка городского типа Знаменка (Знаменский район Тамбовской области). В начале XVIII века, здесь, у слияния рек Кариана и Цны, возникла усадьба Кариан-Загряжское (поздние названия – Кариан-Строганово, Кариан-Знаменка). Зоны жилой застройки и парка усадьбы являются центром современного поселка. Северо-восточная часть имения, где были расположены пашни, луга и лесные угодья усадьбы, в настоящее время выведена из хозяйственного оборота.

Согласно исторической карте 1913-1914 гг. «План Тамбовской губернии Тамбовского уезда усадебной части Знаменского имения владения князя Г.А. Щербатова» [10], фрагмент которой приведен на рисунок 1а, значительную часть участка обследования занимала пашня (области со штриховкой). На северо-востоке участка располагались луга, а его западная часть покрыта кустарником. Анализ ортофотоплана и натурные обследования показали, что участок достаточно равномерно покрыт плотной кустарниковой растительностью, следы пашни и луга начала XX в. визуальнo не фиксируются.

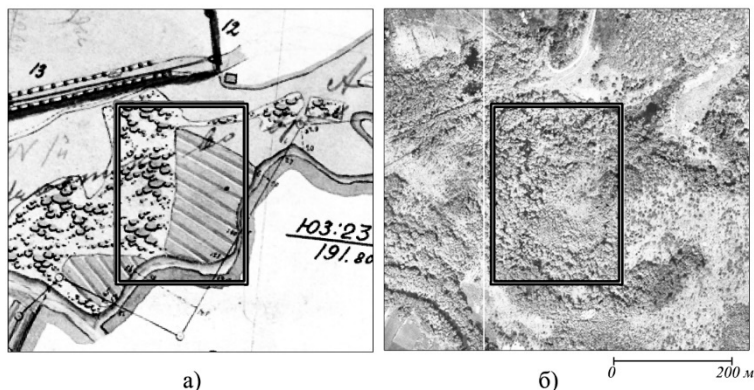


Рис. 1. Исследуемый участок: фрагменты исторической карты (а) и ортофотоплана (б). Черная рамка – граница участка обследования

Вероятно, этот участок обрабатывался короткий промежуток времени. На ранних картографических материалах – «Топографический межевой атлас Тамбовской губернии» 1862 г., выполненный под руководством А.И. Менде [22], – вся территория имения на северном берегу р. Цны еще обозначена как «мокрый луг с кустарником всякого

рода». На поздних картографических источниках – карта РККА 1930-х гг. [23] – это участок уже обозначен пустующим. Предположительно, изучаемый участок был распахан в пик малоземелья в конце XIX в., так как в ходе реформы 1861 г. крестьяне Знаменки получили очень маленькие наделы. Также распашка могла быть связана с сильными засухами 1880-х и, особенно, начала 1890-х гг. Известно, что заливные участки периодически распахивались в годы засух 1905-1906, 1920-1921, 1932-1933 гг. Здесь могли сеять озимую пшеницу. Помимо этого, распашка могла быть связана с посевами кормовых трав, для устойчивого роста которых необходима достаточно высокая влажность почвы. Из описания хозяйства графа П.С. Строганова, владельца Знаменки до 1911 г., известно, что тимopheевку в имении сеяли несколько десятилетий [24]. Поэтому следы пашни могли сохраниться на долгое время. Вероятно, регулярная распашка на участке обследования прекратилась в 1918 г., после социализации помещичьей земли. В результате вторичной сукцессии, продолжавшейся около 100 лет, растительность постепенно вернулась к состоянию непашаного заливного луга.

Современные беспилотные летательные аппараты (БПЛА) позволяют выполнить крупномасштабную аэрофотосъемку с высоким пространственным разрешением. Очевидным преимуществом аэрофотосъемки с БПЛА перед космической съемкой и пилотируемой аэрофотосъемкой является ее доступность и возможность целенаправленного обследования локальных участков (единицы-десятки гектар) при выборе оптимальных условий (высоты полета, времени суток, погодных условий и т.п.). Это обеспечивает контрастное выявление теневых, почвенных и растительных признаков разноплановых ландшафтных объектов поиска.

Аэрофотосъемка в видимом диапазоне поселка Знаменка и прилегающей территории проведена в период ранней вегетации (май 2019 г.). Съемка выполнена с БПЛА самолетного типа Supercam S350-F (ООО «Финко», Ижевск) при высоте 250 м над средним уровнем рельефа местности. В результате фотограмметрической обработки (PHOTOMOD UAS) был построен ортофотоплан в системе координат МСК-68 с разрешением 0,05 м на пиксель (рисунок 16).

Для локализации участка заброшенной пашни использованы материалы мультиспектральной съемки, полученные на этапе ранней вегетации в мае 2019 г. и при развитой растительности в июле 2019 г. (рисунок 2). Съемка выполнена мультиспектральной камерой Parrot Sequoia. Применяемая камера имеет следующие характеристики: динамический диапазон – 10 бит, пространственное разрешение мультиспектральной съемки – 0,05 м на пиксель.

тиспектральных снимков – 0,3 м, спектральное разрешение – 4 канала (видимый зеленый Green, 530 – 570 нм; видимый красный Red, 640 – 680 нм; граничный красный RedEdge, 730 – 740 нм и ближний инфракрасный NIR, 770 – 810 нм). Фотограмметрическая обработка мультиспектральных аэрофотоснимков выполнена в программе Pix4d. Исходя из известных особенностей вегетационного периода лиственной и травянистой растительности, на этапе ранней вегетации (активация жизненных процессов) объем зеленой массы существенно меньше, в сравнении с этапом развитой растительности, когда в полной мере проявляется зеленая биомасса лиственных лесов, кустарников и трав. Эти отличия наглядно прослеживаются при визуальном анализе исходных снимков территории обследования (рисунок 2)

Спектральная отражательная способность растительного покрова имеет характерные особенности: высокое значение в канале Green, резкий подъем в каналах RedEdge и NIR [25]. Так как на канал Red приходится минимум отражательной способности зеленой фитомассы, данный канал не используется при дальнейшем анализе результатов мультиспектральной съемки.

3. Алгоритм обработки комплекса мультиспектральных данных. Традиционно, в качестве признаков разноплановых объектов на мультиспектральных изображениях, принято использовать отличия спектрально-яркостных характеристик отдельных областей снимка [26–29]. Однако яркостные характеристики в значительной степени подвержены влиянию условий съемки: неравномерная освещенность отдельного снимка, вызванная низкой облачностью, или отличия в освещенности смежных снимков и др. Поэтому, для исключения искажений исходных изображений, вызванных условиями съемки, используются методы фильтрации [30–32]. Сглаживание и подавление шумов, с одной стороны, устраняет артефакты съемки, но с другой – снижает контрастность проявления ландшафтных объектов, которые незначительно отличаются по спектрально-яркостным характеристикам. Именно к таким объектам относятся участки растительности, возникшие в результате вторичной сукцессии. Другим известным подходом является вычисление вегетационных индексов, учитывающих взаимное изменение изображений в различных парах спектральных каналов мультиспектральной съемки [33, 34] Обычно такое преобразование также выполняется после фильтрации исходных данных. При этом очевидно, что предварительная фильтрация не гарантирует однородность исходных данных в разных спектральных каналах. Поэтому анализ вегетационных индексов не исключает проблему, возникающую при выявлении малоконтрастных областей растительности.

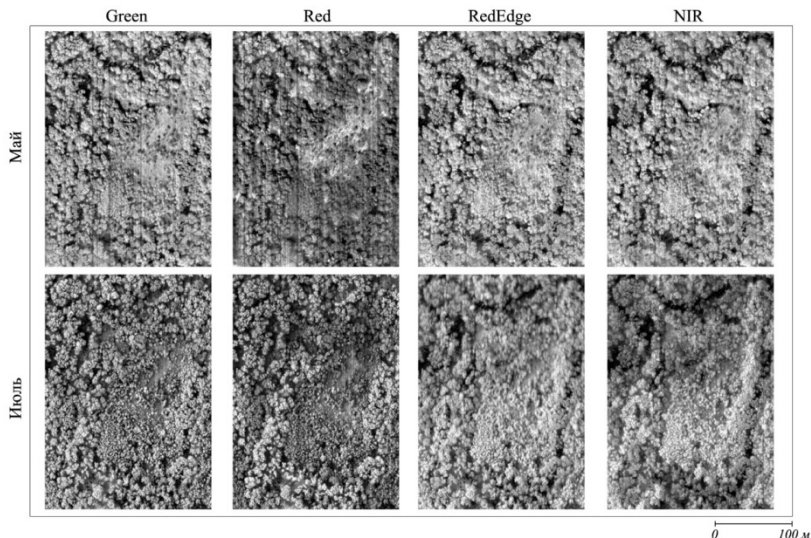


Рис. 2. Разновременные снимки в спектральных каналах Green, Red, RedEdge и NIR

Следовательно, для контрастного проявления ландшафтных объектов более эффективны подходы, которые нечувствительны к спектральным шумам. Примером такого преобразования является вычисление текстурных признаков на основе матрицы смежности [35]. Текстура ландшафта на мультиспектральных изображениях является более стабильной характеристикой, чем интенсивность отдельных пикселей: для определенных типов ландшафта текстура достаточно постоянна и в меньшей степени зависит от условий съемки. Кроме того, текстура разнородных областей в пределах анализируемого участка существенно различна [36–38]. Таким образом, анализ текстурных особенностей ландшафта делает возможным выявление «малоконтрастных» отличий растительного покрова на мультиспектральных изображениях. Поэтому в настоящей работе в качестве признаков используются текстурные признаки Харалика, которые рассчитываются по всем каналам мультиспектрального изображения, что позволяет учесть и спектральные особенности объектов.

Основным способом выявления разноплановых объектов на преобразованном изображении является сегментация. Для выделения наиболее информативных признаков зачастую используется метод главных компонент. Такой подход нередко применяется при анализе гипер- и мультиспектральных изображений для классификации расти-

тельного покрова [39–41]. Сегментация изображения по текстурным признакам, вычисленным по данным одномоментной мультиспектральной съемки, позволяет локализовать следы участков антропогенного преобразования природной среды, что было показано при поиске пашен первой половины XIX в., расположенных на участках сплошной вырубki лесных массивов, в условиях вторичной сукцессии елово-пихтово-южнотаетжных лесов [42]. Необходимо отметить, что одномоментный снимок позволяет с некоторой достоверностью констатировать лишь доминирование определенного типа растительности и оценить особенности ее состояния (здоровая растительность; растительность, испытывающая угнетение или стресс). Однако применение разновременной мультиспектральной съемки позволяет оценить этапы процесса вегетации. Именно на разновременных снимках наглядно проявляется влияние характеристик природной среды, определяющих скорость вегетации (тип и состав почв, мощность почвенного горизонта, увлажненность и экспозиция склона и пр.) [15, 18]. Отличие этих характеристик на участках природной среды, не подвергавшейся антропогенному воздействию, и участках антропогенно-преобразованной природной среды дает принципиально новую дополнительную информацию для выявления растительности, появившейся в результате вторичной сукцессии.

Для сегментации мультиспектральных изображений используется классическая схема анализа данных feature extraction → feature selection → classification [43, 44], включающая в себя этапы расчета признаков, сокращение признакового пространства и распознавания. В связи с этим алгоритм обнаружения следов антропогенного преобразования природной среды принимает вид, представленный на рисунок 3.



Рис. 3. Алгоритм обнаружения следов антропогенного преобразования природной среды

3.1. Текстурные признаки. Задача состоит в разделении мультиспектрального изображения территории обследования на множество площадных объектов: участков природной среды, сохранившихся в

естественном состоянии, и участков антропогенного воздействия на природную среду. Поэтому принципиальную роль играет текстура ландшафтных объектов и ее изменение при переходе от одной области к другой. Для описания текстуры в количественной шкале применяются текстурные признаки Харалика, которые рассчитываются по матрицам смежности полутоновых изображений и содержат информацию о различных текстурных характеристиках: однородности, линейной зависимости тона (линейная структура), контрасте, количестве и природе границ, сложности изображения [35]. Поскольку не все текстурные признаки Харалика являются взаимно независимыми [45], было принято решение ограничиться рассмотрением основных из них: Contrast, Correlation, Energy, Entropy, Homogeneity, Variance (таблица 1).

Для построения матрицы смежности выполняется квантование исходного изображения I , что позволяет преобразовать его из полутонового 10-битного изображения в квантованное изображение I_q с N уровнями градаций серого в интервале $[1, \dots, N]$. По квантованному изображению I_q для каждого положения скользящего окна размером $W \times W$ строится матрица X_θ – матрица смежности уровня серого (*Gray-Level Co-Occurrence Matrix, GLCM*):

$$X_\theta(i, j) = \sum_{m=1}^W \sum_{k=1}^W \delta_{mk} \begin{cases} 1, & \text{если } I_q(m, k) = i \text{ и } I_q(m + d_x, k + d_y) = j, \\ 0, & \text{в противном случае.} \end{cases}$$

Таблица 1. Признаки Харалика

Contrast	$\sum_{i=1}^N \sum_{j=1}^N (i - j)^2 p_{ij}$
Correlation	$\sum_{i=1}^N \sum_{j=1}^N \left(\frac{i - \mu_x}{\sigma_x} \right) \left(\frac{j - \mu_y}{\sigma_y} \right) p_{ij},$ $\mu_x = \sum_{i=1}^N i \cdot \sum_{j=1}^N p_{ij}, \quad \mu_y = \sum_{j=1}^N j \cdot \sum_{i=1}^N p_{ij},$ $\sigma_x^2 = \sum_{i=1}^N (i - \mu_x)^2 \cdot \sum_{j=1}^N p_{ij},$ $\sigma_y^2 = \sum_{j=1}^N (j - \mu_y)^2 \cdot \sum_{i=1}^N p_{ij}$

Продолжение Таблицы 1

Energy	$\sum_{i=1}^N \sum_{j=1}^N p_{ij}^2$
Entropy	$-\sum_{i=1}^N \sum_{j=1}^N p_{ij} \log p_{ij}$
Homogeneity	$\sum_{i=1}^N \sum_{j=1}^N \frac{p_{ij}}{1 + (i - j)^2}$
Variance	$\sum_{i=1}^N \sum_{j=1}^N (i - \mu_x)^2 p_{ij}$

Матрица смежности показывает, как часто пара пикселей с интенсивностями i и j оказывается смежной. Понятие смежности в направлении θ определяется вектором $r_\theta = (d_x, d_y)$, где d_x и d_y – расстояние между пикселями по горизонтали и по вертикали соответственно. *GLCM* рассчитывается для каждого из направлений:

$$r_0 = (d_x, d_y), r_{\frac{\pi}{4}} = (d_x, d_y), r_{\frac{\pi}{2}} = (d_x, d_y) \text{ и } r_{\frac{3\pi}{4}} = (d_x, d_y),$$

после чего усредняется:

$$X = \frac{1}{4} \left(X_0 + X_{\frac{\pi}{4}} + X_{\frac{\pi}{2}} + X_{\frac{3\pi}{4}} \right),$$

и нормируется:

$$P = \frac{X}{\sum_{i=1}^N \sum_{j=1}^N X(i, j)}.$$

Нормированная матрица смежности выступает в роли функции плотности распределения вероятностей проявления смежных пар пикселей в различных областях изображения. По нормированной матрице

смежности $P = (p_{ij})$ рассчитываются текстурные признаки Харалика (таблица 1).

Контраст (Contrast) характеризует резкость изображения и глубину «борозд» текстуры. Низкая контрастность соответствует размытым текстурам.

Корреляция (Correlation) оценивает линейность зависимости уровня серого от соответствующих значений для смежных пикселей.

Энергия (Energy) оценивает однородность и грубость текстуры, она максимальна для однотонных областей.

Энтропия (Entropy) характеризует случайность и неравномерность. Максимальные значения признака соответствуют случайному распределению значений яркости пикселей.

Однородность (Homogeneity) противоположна контрасту, характеризует «сглаженность» области изображения. Большие значения признака соответствуют однородным областям (с небольшой разницей в уровне серого), а близкие к 0 – наоборот.

Вариация (Variance) – соответствует квадрату стандартного отклонения, вычисленному по матрице смежности *GLCM*. Однотонное серое изображение имеет вариацию, равную 0.

Текстурные признаки Харалика, рассчитанные для каждого спектрального канала (видимый зеленый Green, граничный красный RedEdge и ближний инфракрасный NIR) по снимкам, сделанным в мае и июле, представлены на рисунке 4 и рисунке 5 соответственно.

Полученные изображения более наглядно, в сравнении с исходными мультиспектральными снимками (рисунок 2), выявляют участки растительности схожей текстуры, сглаживая малоинформативные детали ландшафтных объектов. Визуально выделяемые разноплановые фрагменты изображения могут соответствовать участкам, покрытым растительностью различных типов и плотности. Кроме того, очевидное преимущество дает сравнительный анализ разновременных снимков. Так изображение на «фоновых» участках – вне зоны исторического антропогенного воздействия (западная и северо-восточная части исследуемого участка, рисунок 1а) – практически неизменно по структуре и интенсивности. Данная особенность наблюдается для всех пар преобразованных изображений (май – июль) в каждом из трех каналов съемки. С другой стороны, на участке возможной исторической пашни (центральная и восточная части исследуемого участка, рисунок 1а), характер вариации изображений текстурных признаков существенно иной. Здесь, в зависимости от времени съемки, существенно изменяется конфигурация выделяемых областей и значение текстурных признаков (вплоть обратного соотношения их интенсивности). Очевидным

образом разделяется северная и южная части участка предполагаемой пашни. Эти наблюдения позволяют предположить принципиально различный характер современного растительного покрова, как минимум, в трех областях исследуемого участка: в «фоновой» области, в южной и северной частях исторической пашни. Последующее определение количества и конфигурации участков разноплановой растительности основано на статистическом анализе полученных 36 массивов текстурных признаков.

3.2. Снижение размерности признакового пространства. Для исключения корреляционных связей между полученными признаками проводится процедура снижения размерности методом главных компонент [46]. Идея метода заключается в формировании системы новых признаков и выборе из них наиболее информативных, объясняющих большую часть изменчивости данных в целом. Новые признаки, которые называют главными компонентами, независимы и представляют собой линейную комбинацию исходных признаков.

Основной способ отбора главных компонент основан на оценке кумулятивной суммы их дисперсий. Как правило, эта сумма (в процентах) должна быть не менее установленного порога 80-95% [47, 48].

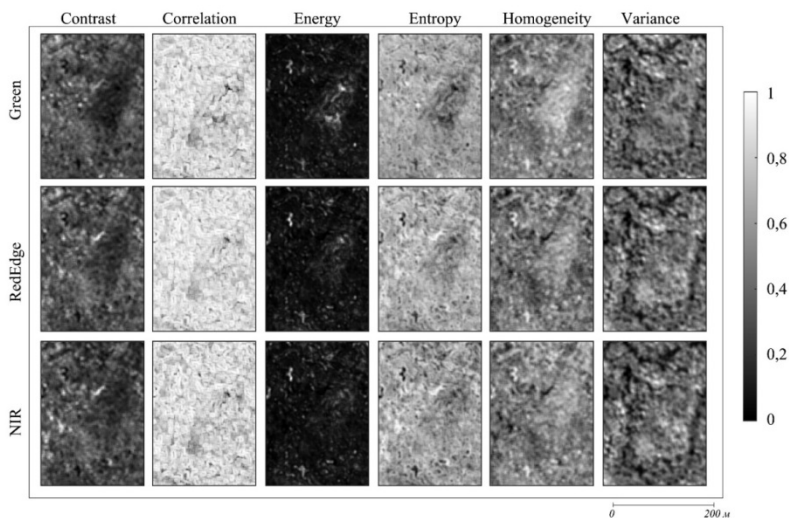


Рис. 4. Текстурные признаки Харалика (май)

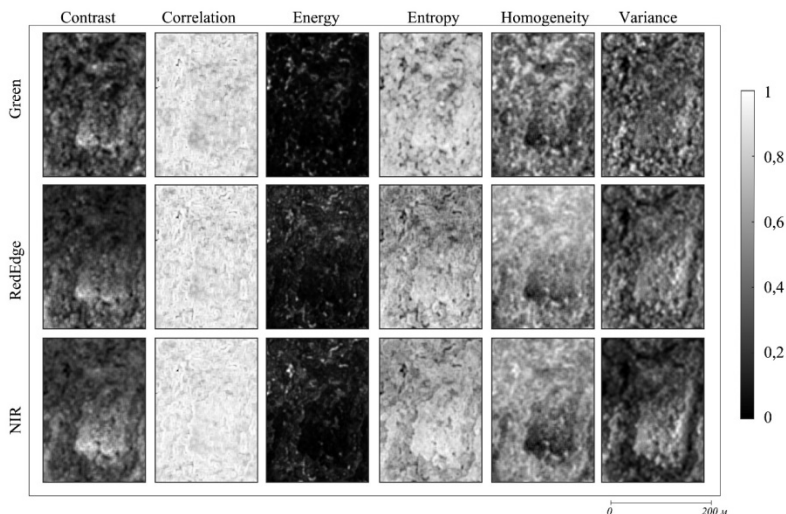


Рис. 5. Текстурные признаки Харалика (июль)

Для каждого из наборов текстурных признаков Харалика выделено по четыре главные компоненты (рисунок 6). По наибольшему вкладу признаков (таблица 2, серая заливка) можно сделать следующие предположения. Первая главная компонента PC1, описывающая большую часть общей дисперсии, отображает наиболее контрастные объекты – участки растительности, значительно отличающиеся по своей текстуре от окружающей территории. Конфигурация областей максимальных значений компоненты близка к линейной (рисунок 6), в большинстве случаев они приурочены к «переходной зоне» между природной средой, не подвергавшейся антропогенному воздействию, и участками антропогенно-преобразованной природной среды (рисунок 1a). Вторая главная компонента PC2 характеризует степень однородности: чем меньше значение, тем более хаотичная текстура. Именно поэтому в период ранней вегетации область максимальных значений компоненты охватывает участок непреобразованной природной среды, а на этапе развитой растительности ситуация на всем участке обследования становится однородной (рисунок 6).

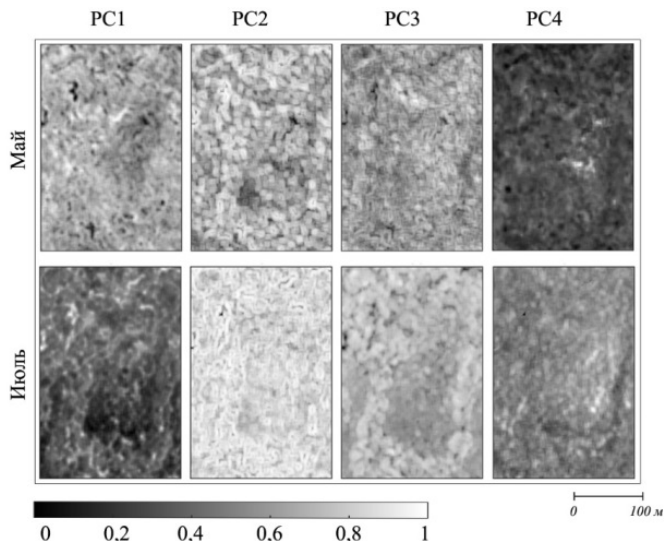


Рис. 6. Визуализация главных компонент

Таблица 2. Коэффициенты корреляции между исходными признаками и главными компонентами

		Май				Июль			
		PC1	PC2	PC3	PC4	PC1	PC2	PC3	PC4
Green	Contrast	0,27	-0,05	-0,24	-0,27	-0,26	-0,17	0,09	-0,20
	Correlation	0,10	0,41	0,22	-0,14	0,06	0,47	0,22	-0,08
	Energy	-0,26	-0,08	-0,07	0,48	0,23	-0,02	-0,44	0,08
	Entropy	0,28	0,13	0,01	-0,37	-0,26	0,01	0,36	-0,16
	Homogeneity	-0,28	0,09	0,18	0,35	0,27	0,16	-0,23	0,13
	Variance	0,04	-0,36	0,39	-0,24	-0,19	-0,12	0,47	0,29
RedEdge	Contrast	0,28	-0,08	-0,23	0,17	-0,26	0,00	-0,23	-0,31
	Correlation	0,11	0,40	0,29	0,11	-0,04	0,53	0,05	0,10
	Energy	-0,28	0,00	-0,17	0,01	0,26	-0,23	-0,07	0,12
	Entropy	0,31	0,09	0,06	0,12	-0,26	0,24	-0,09	-0,17
	Homogeneity	-0,29	0,14	0,16	-0,09	0,28	0,00	0,17	0,26
	Variance	0,07	-0,39	0,39	0,02	-0,24	-0,13	0,04	0,55
NIR	Contrast	0,27	-0,07	-0,21	0,31	-0,27	-0,08	-0,30	-0,04
	Correlation	0,10	0,39	0,31	0,17	-0,05	0,50	-0,13	0,25
	Energy	-0,28	0,02	-0,19	-0,13	0,26	-0,14	-0,01	-0,11
	Entropy	0,30	0,07	0,08	0,28	-0,28	0,13	-0,20	0,11
	Homogeneity	-0,28	0,14	0,13	-0,25	0,28	0,08	0,24	0,00
	Variance	0,07	-0,38	0,39	0,09	-0,25	-0,12	-0,20	0,46
Кумулятивная сумма дисперсий		54%	77%	88%	93%	60%	78%	86%	91%

Визуализация первых главных компонент PC1 и PC2, объясняющих большую часть дисперсии исходных данных, даёт обобщённое представление о структуре участка по всем спектральным каналам. На этих компонентах выделяются все природные объекты, которые устойчиво отображаются на исходных снимках. Последние главные компоненты PC3 и PC4 содержат в себе сведения о скрытых закономерностях – локальных особенностях участка, которые явно не фиксируются в исходных данных. Эти особенности не имеют выраженного контраста, поэтому могут быть интерпретированы как следы антропогенного воздействия, которые существенно «сглажены» вторичной сукцессией.

Построенные главные компоненты содержат практически всю информацию об исходных данных (более 90%, таблица 2). Поэтому переход от текстурных признаков (рисунок 4, 5) к новым переменным PC1 – PC4 (рисунок 6) существенно упрощает сравнительный анализ разновременных снимков.

3.3. Сегментация. Сегментация изображения заключается в разбиении его на непересекающиеся области с однородными свойствами и, вероятно, близкими значениями рассматриваемых характеристик природной среды. Поэтому на последнем этапе алгоритма может быть получена «карта» распределения типов растительности и следов антропогенного преобразования природной среды.

Сегментированное изображение (рисунок 7) получается в результате кластеризации методом *k-means* в пространстве главных компонент. Выбор метода обусловлен, прежде всего, высокой скоростью и эффективностью обработки большого набора данных [49]: матрица главных компонент содержит 746496×4 элементов (≈ 3 млн.). Количество выделяемых классов задано, исходя из визуального анализа ортофотоплана (рисунок 1б), который позволил выделить 4 вида ландшафтных объектов: древесная влаголюбивая растительность, плотная и редкая поросль кустарника, участки травянистой растительности.

В работе использована реализация алгоритма *k-means*, приведенная в [50]. Интерпретация выделенных классов (таблица 3 и 4) выполнена путем сопоставления результата сегментации (рисунок 7) с исторической картой (рисунок 1а) и ортофотопланом (рисунок 1б).

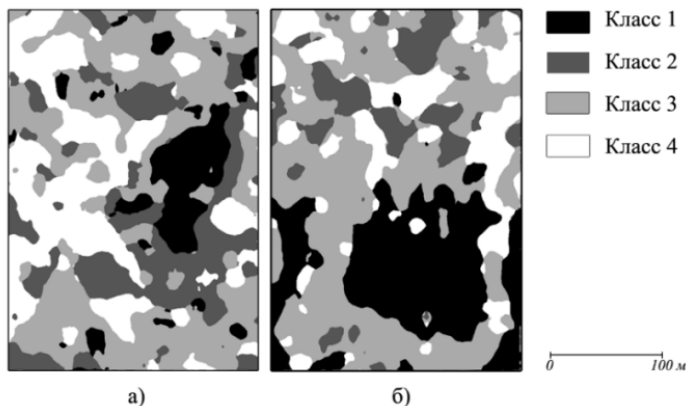


Рис. 7. Результат сегментации: а) май, б) июль

Таблица 3. Интерпретация результатов сегментации (май)

Класс 1	Поросль восстановившегося кустарника, прогалины. Северная часть пашни
Класс 2	Восстановившийся кустарник. Южная часть пашни
Класс 3	Крупный кустарник, древесная влаголюбивая растительность
Класс 4	Редкий кустарник, преимущественно вдоль русла реки

Таблица 4. Интерпретация результатов сегментации (июль)

Класс 1	Плотный восстановившийся кустарник. Южный участок пашни
Класс 2	Плотный крупный кустарник
Класс 3	Крупный кустарник, древесная влаголюбивая растительность
Класс 4	Редкий кустарник, преимущественно вдоль русла реки

Очевидны визуальные отличия между полученными сегментированными изображениями (рисунок 7а и 7б), описывающие один и тот же участок, но в разные периоды вегетации. Эти отличия объяснимы разной плотностью растительного покрова и, как следствие, – текстурой на разновременных снимках. Условное разделение заброшенной пашни на две части: северную (рисунок 7а, Класс 1) и южную (рисунок 7б, Класс 1), – можно связать с неравномерным процессом восстановления природного кустарника. Северная часть зафиксирована по результатам сегментации майских данных, южная – июльских. Выявленные участки антропогенного воздействия соответствуют положению пашни на исторической карте (рисунок 8).

Таким образом, анализ разновременных снимков дает взаимодополняющую информацию о текстуре растительного покрова. Такой анализ обеспечивает возможность выявить следы антропогенного воздействия в случае неравномерной вторичной сукцессии на заброшенном участке.

4. Результаты и обсуждение. Снимки в периоды ранней вегетации и развитой растительности выделяют разные участки заброшенной пашни второй половины XIX – начала XX вв. (рисунок 7, таблица 3 и 4). Следовательно, совместный анализ разновременных снимков позволяет выделить зону исторического хозяйственного освоения на фоне природной среды, не подвергавшейся антропогенному воздействию, а также дает возможность оценить ее возможную конфигурацию. В частности, результаты сегментации по майским (рисунок 8а) и июльским (рисунок 8б) снимкам, сопоставленные с фрагментом исторической карты (эталонные данные), согласуются с расположением пашни на участке обследования. Кроме того, объединение выделенных сегментов мультиспектральных изображений позволяет восстановить конфигурацию зоны антропогенного воздействия (рисунок 8в). Независимая сегментация разновременных снимков выделяет разные участки исторической пашни, именно поэтому, только совместный анализ позволяет восстановить границы области пашни наиболее полно.

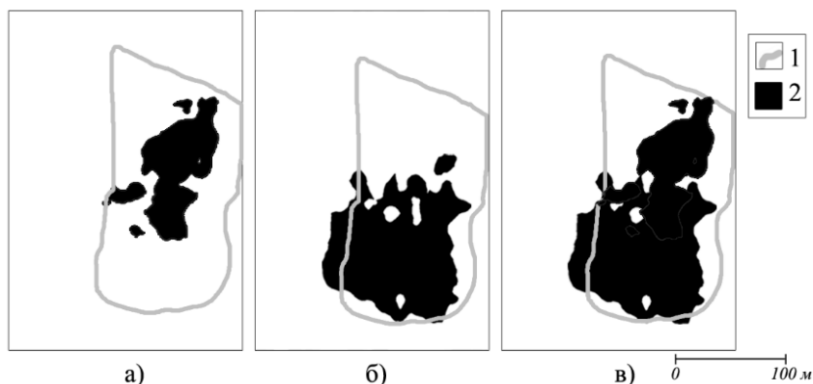


Рис. 8. Сопоставление результата сегментации с исторической картой: а) участок вторичной сукцессии в период ранней вегетации; б) участок вторичной сукцессии в период развитой растительности; в) конфигурация территории заброшенной пашни по результатам сегментации; 1 – граница пашни на исторической карте; 2 – участок пашни по результатам сегментации

Полученный результат сегментации отражает неравномерность вторичной сукцессии на участке заброшенной пашни. С точки зрения восстановившейся растительности территория пашни разделяется на две условные части: северная (рисунок 8а) и южная (рисунок 8б). Растительность северной части (первая надпойменная терраса р. Цны) – поросль кустарника с травой – отчетливо проявляется по текстуре на майских снимках (рисунок 7а, Класс 1). В период зрелой вегетации (июльские снимки) – «сливается» с окружающей средой (рисунок 7б, Класс 3). Отличия растительности в южной части исторической пашни от окружающей территории фиксируются в разные периоды вегетации. Здесь, в пойме р. Цны, наблюдается более контрастная текстура плотного молодого кустарника. На майских снимках восстановившаяся растительность выделяется как «переходная зона» (рисунок 7а, Класс 2) от участка вторичной сукцессии (рисунок 7а, Класс 1) к территории природной среды, не подвергавшейся антропогенному воздействию (рисунок 7а, Классы 3 и 4). На июльских снимках южный участок пашни определяется по компактной зоне значительной площади (рисунок 7б, Класс 1) на фоне достаточно однородной растительности природной среды (рисунок 7б, Класс 3). Разделение территории заброшенной пашни на две зоны связано с тем, что исходное состояние непаханого заливного луга восстанавливается неравномерно. Вероятной причиной этого являются особенности рельефа. Во время весенних паводков затопливается южная часть территории, расположенная в пойме р. Цны. Вследствие этого она менее контрастно проявляется в период ранней вегетации, но явно выделяется более активным наращиванием фитомассы в период зрелой растительности.

Для анализа комплекса мультиспектральных данных успешно использован классический подход *feature extraction* → *feature selection* → *classification*, адаптированный для применения статистических текстурных признаков Харалика на этапе извлечения признаков (*feature extraction*). Необходимость использования признаков Харалика вызвана тем, что текстура ландшафта на мультиспектральных изображениях является более стабильной характеристикой площадных объектов чем интенсивность отдельных пикселей мультиспектрального изображения. Именно это дает возможность выявления сегментов с «малоконтрастными» отличиями растительного покрова, возникшего в результате вторичной сукцессии. Другой отличительной особенностью предложенного подхода является анализ набора снимков, сделанных в разные периоды вегетации. В условиях неравномерности вторичной сукцессии в лесостепной зоне это позволяет полнее охарактеризовать изменение интенсивности лиственной и травянистой растительности.

Объединение результатов сегментации обеспечивает выявление участков восстановившейся растительности, соответствующих антропогенно-преобразованным областям природного ландшафта.

Авторы благодарят д.и.н. В.В. Канищева и к.и.н. К.С. Кунавина, коллег из Тамбовского государственного университета им. Г.Р. Державина, за конструктивное обсуждение и предоставленные материалы.

Литература

1. Бешенцев А.Н. Геоинформационное обеспечение мониторинга трансформации природных ландшафтов в бассейне оз. Байкал на основе ретроспективных картографических материалов // Аридные экосистемы. 2011. Т. 17. № 4 (49). С. 53–62.
2. Черепанова Е.С. Исторические аспекты освоения лесных территорий бассейна Верхней Камы и их последствия // Интерэкспо ГЕО-Сибирь. 2012. Т. 4. С. 37–41.
3. Чернов С.З. Рекомендуемые форматы исторических карт земельных дач средневековой России XIII–XVII вв. (по материалам древнего Радонежа) // Актуальные проблемы аграрной истории Восточной Европы X–XXI вв.: источники и методы исследования: материалы XXXII сессии симпозиума по аграрной истории Восточной Европы. Рязань: Ряз. гос. ун-т им. С.А. Есенина, 2012. С. 58–78.
4. Россия. Полное географическое описание нашего Отечества: настольная и дорожная книга для русских людей. Т.2: Среднерусская Черноземная область: Курская, Орловская, Тульская, Рязанская, Тамбовская, Воронежская и Пензенская губернии // СПб.: А.Ф. Девриен. 1902. 717 с.
5. Аврех А.Л., Канищев В.В. Естественно-исторические условия модернизации аграрного общества. Тамбовская губерния, XIX–XX вв. // Социальная история российской провинции в контексте модернизации аграрного общества в XVIII–XX вв. Тамбов. 2002. С. 3–17.
6. Канищев В.В. Экономика, демография, экология в контексте модернизации аграрного общества (Тамбовская губерния в XIX – начале XX в.) // Экономическая история: Ежегодник. 2002. М.: РОССПЭН. 2003. С. 513–530.
7. Цинцадзе Н.С. Демографические и экологические проблемы развития аграрного общества России во второй половине XIX – начале XX века в восприятии современников. Монография // Тамбов: Изд. дом ТГУ им. Г.Р. Державина. 2012. 286 с.
8. Мониторинг биологического разнообразия лесов России: методология и методы // М.: Наука. 2008. 453 с.
9. Kalinina O., Goryachkin S.V., Lyuri D.I., et.al. Post-agrogenic development of vegetation, soils, and carbon stocks under self-restoration in different climatic zones of European Russia // Catena. 2015. vol. 129. pp. 18–29. DOI: 10.1016/j.catena.2015.02.016.
10. План Знаменской усадьбы и окрестных территорий // Государственный архив Тамбовской области. Ф. 29. Оп. 4. Д. 10052.
11. Kalinina O., Krause S.-E., Goryachkin S.V., et. al. Self-restoration of post-agrogenic chernozems of Russia: Soil development, carbon stocks, and dynamics of carbon pools // Geoderma. 2011. vol. 162. pp. 196–206. DOI: 10.1016/j.geoderma.2011.02.005.

12. Агроэкологическое состояние и перспективы использования земель России, выбывших из активного сельскохозяйственного оборота / под ред. Г.А. Романенко // М.: ФГНУ «Росинформагротех». 2008. 64 с.
13. Cao Y., Li G.L., Luo Y.K., et al. Monitoring of sugar beet growth indicators using wide-dynamic-range vegetation index (WDRVI) derived from UAV multispectral images // *Computers and Electronics in Agriculture*. 2020. vol. 171: 105331. DOI: 10.1016/j.compag.2020.105331.
14. Eng L.S., Ismail R., Hashim W., et al. Vegetation monitoring using UAV: a preliminary study // *International Journal of Engineering & Technology*. 2018. no.7. pp. 223–227. DOI: 10.14419/ijet.v7i4.35.22736.
15. Liu S., Marinelli D., Bruzzone L., et al. A review of change detection in multitemporal hyperspectral images: current techniques, applications, and challenges // *IEEE Geo science and Remote Sensing Magazine*. 2019. vol. 7. no. 2. pp. 140–158. DOI: 10.1109/MGRS.2019.2898520.
16. Wei Z., Gu X., Sun Q., et al. Analysis of the spatial and temporal pattern of changes in abandoned farmland based on long time series of remote sensing data // *Remote Sensing*. 2021. vol. 13(13): 2549. DOI: 10.3390/rs13132549.
17. Alonso L., Picos J., Armesto J. Forest Land Cover Mapping at a Regional Scale Using Multi-Temporal Sentinel-2 Imagery and RF Models // *Remote Sensing*. 2021. vol. 13(12): 2237. DOI: 10.3390/rs13122237.
18. Possoch M., Bieker S., Hoffmeister D., et al. Multi-temporal crop surface models combined with the rgb vegetation index from UAV-based images for forage monitoring in grassland // *The International Archives of the Photogrammetry, Remote Sensing and Spatial Information Sciences*. 2016. vol. XLI-B1. pp. 991–998. DOI:10.5194/isprsarchives-XLI-B1-991-2016.
19. Iersel W., Straatsma M., Middelkoop H., et al. Multitemporal classification of river floodplain vegetation using time series of UAV Images // *Remote Sensing*. 2018. vol. 10(7): 1144. DOI:10.3390/rs10071144.
20. Senf C., Leitão P.J., Pflugmacher D., et al. Mapping land cover in complex mediterranean landscapes using Landsat: improved classification accuracies from integrating multi-seasonal and synthetic imagery // *Remote Sensing of Environment*. 2015. vol. 156. pp. 527–536. DOI: 10.1016/j.rse.2014.10.018.
21. Simonetti D., Simonetti E., Szantoi Z., et al. First results from the phenology-based synthesis classifier using Landsat 8 imagery // *IEEE Geoscience and Remote Sensing Letters*. 2015. vol. 12. no. 7. pp. 1496–1500 DOI: 10.1109/LGRS.2015.2409982.
22. Топографический межевой атлас Тамбовской губернии 1:84000 (1862) [Электронный ресурс] / Это Место. URL: <http://www.etomesto.ru/karta5623> (дата обращения 12.07.2021 г.).
23. Карта РККА 1:100000 (1935-1941) [Электронный ресурс] / Это Место. URL: <http://www.etomesto.ru/karta2027> (дата обращения 12.07.2021 г.).
24. Самодуров И. Краткое описание имения графа Павла Сергеевича Строганова, Тамбовской губернии и уезда, при селе Знаменском-Кариан. М.: типо-лит. т-ва И.Н. Кушнерев и К°. 1895. 19 с.
25. Xie Q., Dash J., Huang W., et al. Vegetation indices combining the red and red-edge spectral information for leaf area index retrieval // *IEEE Journal of Selected Topics in Applied Earth Observations and Remote Sensing*. 2018. vol. 11. no. 5. pp. 1482–1493. DOI: 10.1109/JSTARS.2018.2813281.
26. Кондратьев К.Я., Федченко П.П. Спектральная отражательная способность и распознавание растительности // Л.: Гидрометеиздат. 1982. 215 с.
27. Roy P.S. Spectral reflectance characteristics of vegetation and their use in estimating productive potential // *Proceedings of the Indian Academy of Sciences (Plant Science)*. 1989. vol. 99. no 1. P. 59 – 81. DOI: 10.1007/BF03053419.

28. Govender M., Chetty K., Bulcock H. A review of hyperspectral remote sensing and its application in vegetation and water resource studies // *Water SA*. 2007. vol. 33. no. 2. pp. 145–151. DOI:10.4314/wsa.v33i2.49049.
29. Sharma A., Kaur D., Gupta A., et. al. Application and Analysis of Hyperspectral Imaging // *IEEE International Conference on Signal Processing, Computing and Control (ISPCC 2k19)*. 2019. pp. 30–35. DOI: 10.1109/ISPCC48220.2019.8988436.
30. Гонсалес Р., Вудс Р. Цифровая обработка изображений // М.: Техносфера. 2006. 1072 с.
31. Gao J., Li H., Han Z., et. al. Aircraft Detection in Remote Sensing Images Based on Background Filtering and Scale Prediction // *PRICAI 2018: Trends in Artificial Intelligence. Lecture Notes in Computer Science*. 2018. vol. 11012. pp. 604–616. DOI: /10.1007/978-3-319-97304-3_46.
32. Andriyanov N.A., Vasiliev K.K., Dement'ev V.E. Analysis of the efficiency of satellite image sequences filtering // *Journal of Physics: Conference Series*. 2018. vol.1096: 012036. DOI: 10.1088/1742- 6596/1096/1/012036.
33. Bannari, A., Morin, D., Bonn, F, et. al. A review of vegetation indices // *Remote Sensing Reviews*. 1995. vol. 13. no.1. pp. 95–120. DOI: 10.1080/02757259509532298.
34. Черепанов А.С. Вегетационные индексы // *Geomatics*. 2011. №2. С. 98 – 102.
35. Haralick R.M., Shanmugam K., Dinstein I. Textural features for image classification // *IEEE Transactions on Systems, Man, and Cybernetics*. 1973. vol. SMC-3. no. 6. pp. 610–621. DOI: 10.1109/TSMC.1973.4309314.
36. Feng Q., Liu J., Gong J. UAV Remote Sensing for Urban Vegetation Mapping Using Random Forest and Texture Analysis // *Remote Sensing*. 2015. vol. 7. pp. 1074–1094. DOI: 10.3390/rs70101074.
37. Kwak G.-H., Park N.-W. Impact of Texture Information on Crop Classification with Machine Learning and UAV Images // *Applied Science*. 2019. vol. 9: 643. DOI: 10.3390/app9040643.
38. Dian Y., Li Z., Pang Y. Spectral and texture features combined for forest tree species classification with airborne hyperspectral imagery // *Journal of the Indian Society of Remote Sensing*. 2015. vol. 43. pp. 101–107. DOI: 10.1007/s12524-014-0392-6.
39. Bekkari A., Idbraim S., Elhassouny A., et. al. SVM and Haralick features for classification of high resolution satellite images from urban areas // *ICISP 2012: Image and Signal Processing. Lecture Notes in Computer Science*. 2012. vol. 7340. pp. 17–26. DOI: 10.1007/978-3-642-31254-0_3.
40. Rodarmel C., Shan J. Principal component analysis for hyperspectral image classification // *Surveying and Land Information Systems*. 2002. vol. 62. no. 2, pp. 115–122.
41. Rejichi S., Chaabane F. Feature extraction using PCA for VHR satellite image time series spatial-temporal classification // *IEEE International Geoscience and Remote Sensing Symposium (IGARSS)*. 2015. pp. 485–488. DOI: 10.1109/IGARSS.2015.7325806.
42. Zlobina A.G., Shaura A.S., Zhurbin I.V., et al. Algorithm for statistical analysis of multispectral survey data to identify the anthropogenic impact of the 19th century on the natural environment // *Pattern Recognition and Image Analysis*. 2021. vol. 31. no. 2. pp. 344–354. DOI: 10.1134/S1054661821020176.
43. Khalid S., Khalil T., Nasreen S. A survey of feature selection and feature extraction techniques in machine learning // *Science and Information Conference*. 2014. pp. 372–378. DOI: 10.1109/SAL.2014.6918213.
44. Popescu M.C., Sasu L.M. Feature extraction, feature selection and machine learning for image classification: A case study // *2014 International Conference on*

- Optimization of Electrical and Electronic Equipment (OPTIM). 2014. pp. 968–973. DOI: 10.1109/OPTIM.2014.6850925.
45. Ulaby F.T., Kouyate F., Brisco B., et al. Textural Information in SAR Images // IEEE Transactions on Geoscience and Remote Sensing. 1986. vol. GE-24. no. 2. pp. 235–245. DOI:10.1109/TGRS.1986.289643.
46. Jolliffe I.T. Principal Components Analysis. 2nd ed. // N.Y.: Springer-Verlag New York Inc. 2002. 487 p.
47. Kang, B., Jung H., Jeong H., et al. Characterization of three-dimensional channel reservoirs using ensemble Kalman filter assisted by principal component analysis // Petroleum Science. 2020. vol. 17. pp. 182–195. DOI: 10.1007/s12182-019-00362-8.
48. Artoni F., Delorme A., Makeig S. Applying dimension reduction to EEG data by principal component analysis reduces the quality of its subsequent independent component decomposition // NeuroImage. 2018. vol. 175. pp. 176–187. DOI: 10.1016/j.neuroimage.2018.03.016.
49. Huang Z. Extensions to the k-means algorithm for clustering large data sets with categorical values // Data mining and knowledge discovery. 1998. vol. 2. pp. 283–304.
50. Jain A.K. Data clustering: 50 years beyond k-means // Pattern Recognition Letters. 2010. vol. 31. pp. 651–666. DOI: 10.1016/j.patrec.2009.09.011.

Шаура Александр Сергеевич — канд. техн. наук, старший научный сотрудник, Удмуртский федеральный исследовательский центр Уральского отделения Российской академии наук. Область научных интересов: искусственный интеллект, методы оптимизации и оптимальное управление, интеллектуальный анализ данных. Число научных публикаций — 23. shauraa@mail.ru; ул. Татьяны Барамзиной, 34, 426067, Ижевск, Россия; р.т.: 8(912)467-8365.

Злобина Анна Григорьевна — канд. техн. наук, научный сотрудник, Удмуртский федеральный исследовательский центр Уральского отделения Российской академии наук. Область научных интересов: информационные технологии, анализ изображений, цифровая обработка сигналов. Число научных публикаций — 25. elf54@yandex.ru; ул. Татьяны Барамзиной, 34, 426067, Ижевск, Россия; р.т.: +7(952)401-8129.

Журбин Игорь Витальевич — д-р ист. наук, главный научный сотрудник, Удмуртский федеральный исследовательский центр Уральского отделения Российской академии наук. Область научных интересов: информационные технологии, анализ изображений, междисциплинарные исследования. Число научных публикаций — 120. zhurb@udm.ru; ул. Татьяны Барамзиной, 34, 426067, Ижевск, Россия; р.т.: 8(912)853-0973.

Баженова Айгуль Илсуровна — канд. техн. наук, научный сотрудник, Удмуртский федеральный исследовательский центр Уральского отделения Российской академии наук. Область научных интересов: цифровая обработка сигналов, распознавание образов, вейвлет-анализ. Число научных публикаций — 25. aigul_bazh@udman.ru; ул. Татьяны Барамзиной, 34, 426067, Ижевск, Россия; р.т.: 8(950)174-3743.

Поддержка исследований. Исследование выполнено при финансовой поддержке РФ, проект № 19-18-00322 «Сравнительно-историческое изучение антропогенных ландшафтов различных регионов средствами беспилотных летательных аппаратов (Тамбовская область и Удмуртия, середина XVIII — начало XX вв.)».

A. SHAURA, A. ZLOBINA, I. ZHURBIN, A. BAZHENOVA
**ANALYSIS OF MULTI-TEMPORAL MULTISPECTRAL AERIAL
PHOTOGRAPHY DATA TO DETECT THE BOUNDARIES OF
HISTORICAL ANTHROPOGENIC IMPACT**

Shauro A., Zlobina A., Zhurbin I., Bazhenova A. Analysis of Multi-Temporal Multispectral Aerial Photography Data to Detect the Boundaries of Historical Anthropogenic Impact.

Abstract. The article presents the application of a statistical analysis algorithm for multi-temporal multispectral aerial photography data to identify areas of historical anthropogenic impact on the natural environment. The investigated site is located on the outskirts of the urban-type village of Znamenka (Znamensky District, Tambov Region) in a forest-steppe zone with typical chernozem soils, where arable lands were located in the second half of the 19th - early 20th centuries. Grown vegetation as a result of secondary succession in abandoned areas can be a sign for identifying traces of historical anthropogenic impact. Distinctive signs of such vegetation from the surrounding natural environment are its type, age and growth density. Thus, the problem of detecting the boundaries of anthropogenic impact on multispectral images is reduced to the problem of vegetation classification. The initial data were the results of multi-temporal multispectral imaging in green (Green), red (Red), edge of red (RedEdge) and near-infrared (NIR) spectral ranges. The first stage of the algorithm is the calculation of the Haralick texture features on multispectral images, the second stage – reduction in the number of features by the principal component analysis, the third stage – the segmentation of images based on the obtained features by the k-means method. The effectiveness of the proposed algorithm is shown by comparing the segmentation results with the reference data of historical cartographic materials. The study of multi-temporal multispectral images makes it possible to more fully characterize and take into account the dynamics of phytomass growth in different periods of the growing season. Therefore, the obtained segmentation result reflects not only the configuration of areas of an anthropogenic transformed natural environment, but also the features of overgrowth of abandoned arable land.

Keywords: multispectral survey, texture segmentation, Haralick texture features, principal component analysis, clustering, k-means, multi-temporal data, growing season, secondary succession.

Shauro Alexander — Ph.D., Senior researcher, Udmurt Federal Research Center of the Ural Branch of the Russian Academy of Sciences. Research interests: artificial intelligence, optimization methods and optimal control, data mining. The number of publications — 23. shauroa@mail.ru; 34, Tatyana Baramzina St., 426067, Izhevsk, Russia; office phone: 8(912)467-8365.

Zlobina Anna — Ph.D., Researcher, Udmurt Federal Research Center of the Ural Branch of the Russian Academy of Sciences. Research interests: information technology, image analysis, digital signal processing. The number of publications — 25. elf54@yandex.ru; 34, Tatyana Baramzina St., 426067, Izhevsk, Russia; office phone: +7(952)401-8129.

Zhurbin Igor — Ph.D., Dr.Sci., Chief researcher, Udmurt Federal Research Center of the Ural Branch of the Russian Academy of Sciences. Research interests: information technologies, image analysis, interdisciplinary research. The number of publications — 120. zhurbin@udm.ru; 34, Tatyana Baramzina St., 426067, Izhevsk, Russia; office phone: 8(912)853-0973.

Bazhenova Aigul — Ph.D., Researcher, Udmurt Federal Research Center of the Ural Branch of the Russian Academy of Sciences. Research interests: digital signal processing, pattern recognition, wavelet analysis. The number of publications — 25. aigul_bazh@udman.ru; 34, Tatyana Baramzina St., 426067, Izhevsk, Russia; office phone: 8(950)174-3743.

Acknowledgements. The research was supported by the Russian Science Foundation under grant 19-18-00322 «Comparative historical research of anthropogenic landscapes in different regions by means of unmanned aerial vehicles (the Tambov region and Udmurtia, mid-XVIII - early XX centuries)».

References

1. Beshencev A.N. [Geoinformation support for monitoring the transformation of natural landscapes in the Lake Baikal basin based on retrospective cartographic materials]. *Aridnye ekosistemy*. 2011. vol. 17. no. 4 (49). pp. 53–62. (In Russ.).
2. Cherepanova E.S. [Historical aspects of the development of forest territories of the Upper Kama basin and their consequences]. *Interespo GEO-Sibir'*. 2012. vol. 4. pp. 37–41. (In Russ.).
3. Chernov S.Z. [Recommended formats of historical maps of land dachas in medieval Russia of the XIII-XVII centuries. (based on the materials of ancient Radonezh)]. *Aktual'nye problemy agrarnoj istorii Vostochnoj Evropy X–XXI vv.: istochniki i metody issledovaniya: materialy XXXII sessii simpoziuma po agrarnoj istorii Vostochnoj Evropy*. [Actual problems of the agrarian history of Eastern Europe of the X-XXI centuries: sources and research methods. Collected papers]. Ryazan': Ryaz. gos. un-t im. S.A. Esenina. 2012. pp. 58–78. (In Russ.).
4. Rossiya. *Polnoe geograficheskoe opisanie nashego Otechestva: nastol'naya i dorozhnaya kniga dlya russkikh lyudej*. T.2: *Srednerusskaya Chernozemnaya oblast': Kurskaya, Orlovskaya, Tul'skaya, Ryazanskaya, Tambovskaya, Voronezhskaya i Penzenskaya gubernii*. [Russia. A complete geographical description of our Fatherland: a desktop and travel book for Russian people. Vol. 2: Central Russian Chernozem region: Kursk, Orel, Tula, Ryazan, Tambov, Voronezh and Penza provinces]. SPb.: A.F. Devrien. 1902. 717 p. (In Russ.).
5. Avrekh A.L., Kanishchev V.V. [Natural-historical conditions of modernization of agrarian society. Tambov province, XIX-XX centuries]. *Sotsial'naya istoriya rossijskoj provincii v kontekste modernizacii agrarnogo obshchestva v XVIII–XX vv.* [The social history of the Russian province in the context of the modernization of agrarian society in the XVIII-XX centuries. Collected papers]. Tambov. 2002. pp. 3–17. (In Russ.).
6. Kanishchev V.V. [Economics, demography, ecology in the context of modernization of agrarian society (Tambov province in the XIX-early XX century)]. *Ekonomicheskaya istoriya: Ezhegodnik*. [Economic History: Yearbook]. 2002. M.: ROSSPEN. 2003. p. 513–530. (In Russ.).
7. Cincadze N.S. *Demograficheskie i ekologicheskie problemy razvitiya agrarnogo obshchestva Rossii vo vtoroj polovine XIX – nachale XX veka v vospriyatii sovremennikov*. Monografiya. [Demographic and environmental problems of the development of the agrarian society of Russia in the second half of the XIX-early XX century in the perception of contemporaries. Monograph]. Tambov: Izd. dom TGU im. G.R. Derzhavina. 2012. 286 p. (In Russ.).
8. *Monitoring biologicheskogo raznoobrazija lesov Rossii: metodologiya i metody* [Monitoring of the biological diversity of Russian forests: methodology and methods]. Moscow: Nauka. 2008. 453 p. (In Russ.).
9. Kalinina O., Goryachkin S.V., Lyuri D.I., et.al. Post-agrogenic development of vegetation, soils, and carbon stocks under self-restoration in different climatic zones

- of European Russia. Catena. 2015. vol. 129. pp. 18–29. DOI: 10.1016/j.catena.2015.02.016.
10. Plan Znamenskoy usad'by i okrestnyh territorij [Plan of the Znamenskaya estate and surrounding areas]. Gosudarstvennyj arhiv Tambovskoj oblasti [State Archive of the Tambov Region]. F. 29. Op. 4. D. 10052. (In Russ.).
11. Kalinina O., Krause S.-E., Goryachkin S.V., et. al. Self-restoration of post-agrogenic chernozems of Russia: Soil development, carbon stocks, and dynamics of carbon pools. Geoderma. 2011. vol. 162. pp. 196–206. DOI: 10.1016/j.geoderma.2011.02.005.
12. Agrojekologicheskoe sostojanie i perspektivy ispol'zovanija zemel' Rossii, vybyvshih iz aktivnogo sel'skohozjajstvennogo oborota [Agroecological state and prospects for the use of Russian lands that have been eliminated from active agricultural turnover]. Ed. by G.A. Romanenko. Moscow.: FGNU "Rosinformagroteh". 2008. 64 p. (In Russ.).
13. Cao Y., Li G.L., Luo Y.K., et al. Monitoring of sugar beet growth indicators using wide-dynamic-range vegetation index (WDRVI) derived from UAV multispectral images. Computers and Electronics in Agriculture. 2020. vol. 171: 105331. DOI: 10.1016/j.compag.2020.105331.
14. Eng L.S., Ismail R., Hashim W., et al. Vegetation monitoring using UAV: a preliminary study. International Journal of Engineering & Technology. 2018. no. 7. pp. 223–227. DOI: 10.14419/ijet.v7i4.35.22736.
15. Liu S., Marinelli D., Bruzzone L., et al. A review of change detection in multitemporal hyperspectral images: current techniques, applications, and challenges. IEEE Geoscience and Remote Sensing Magazine. 2019. vol. 7. no. 2. pp. 140–158. DOI: 10.1109/MGRS.2019.2898520.
16. Wei Z., Gu X., Sun Q., et al. Analysis of the spatial and temporal pattern of changes in abandoned farmland based on long time series of remote sensing data. Remote Sensing. 2021. vol. 13(13): 2549. DOI: 10.3390/rs13132549.
17. Alonso L., Picos J., Armesto J. Forest Land Cover Mapping at a Regional Scale Using Multi-Temporal Sentinel-2 Imagery and RF Models. Remote sensing. 2021. vol. 13(12): 2237. DOI: 10.3390/rs13122237.
18. Possoch M., Bieker S., Hoffmeister D., et al. Multi-temporal crop surface models com-bined with the rgb vegetation index from UAV-based images for forage monitoring in grassland. The International Archives of the Photogrammetry, Remote Sensing and Spatial Information Sciences. 2016. vol. XLI-B1. pp. 991–998. DOI:10.5194/isprsarchives-XLI-B1-991-2016.
19. Iersel W., Straatsma M., Middelkoop H., et al. Multitemporal classification of river floodplain vegetation using time series of UAV Images. Remote Sensing. 2018. vol. 10(7): 1144. DOI:10.3390/rs10071144.
20. Senf C., Leitão P.J., Pflugmacher D., et al. Mapping land cover in complex mediterranean landscapes using Landsat: improved classification accuracies from integrating multi-seasonal and synthetic imagery. Remote Sensing of Environment. 2015. vol. 156. pp. 527–536. DOI: 10.1016/j.rse.2014.10.018.
21. Simonetti D., Simonetti E., Szantoi Z., et. al. First results from the phenology-based synthesis classifier using Landsat 8 imagery. IEEE Geoscience and Remote Sensing Letters. 2015. vol. 12. no. 7. pp. 1496–1500 DOI: 10.1109/LGRS.2015.2409982.
22. Topograficheskij mezhevoj atlas Tambovskoj gubernii 1:84000 (1862) [Topographic boundary atlas of the Tambov province 1:84000 (1862)]. Jeto Mesto [This place]. Available at: <http://www.etomesto.ru/karta5623> (accessed 12.07.2021). (In Russ.)
23. Karta RKKA 1:100000 (1935-1941) [Red Army map 1:100000 (1935-1941)]. Jeto Mesto [This place]. Available at: <http://www.etomesto.ru/karta2027> (accessed 12.07.2021). (In Russ.).

24. Samodurov I. Kratkoe opisanie imenija grafa Pavla Sergeevicha Stroganova, Tambovskoj gubernii i uezda, pri sele Znamenskom-Karian [A brief description of the estate of Count Pavel Sergejevich Stroganov, Tambov province and district, near the village of Znamenskoye-Karian]. Moscow: tipo-lit. t-va I.N. Kushnerev i K°. 1895. 19 p. (In Russ.).
25. Xie Q., Dash J., Huang W., et al. Vegetation indices combining the red and red-edge spectral information for leaf area index retrieval. *IEEE Journal of Selected Topics in Applied Earth Observations and Remote Sensing*. 2018. vol. 11. no. 5. pp. 1482–1493. DOI: 10.1109/JSTARS.2018.2813281.
26. Kondrat'ev K.Y., Fedchenko P.P. Spektral'naya otrazhatel'naya sposobnost' i raspoznavanie rastitel'nosti. [Spectral reflectivity and vegetation recognition]. L.: Gidrometeoizdat. 1982. 215 p. (In Russ.).
27. Roy P.S. Spectral reflectance characteristics of vegetation and their use in estimating productive potential. *Proceedings of the Indian Academy of Sciences (Plant Science)*. 1989. vol. 99. no 1. P. 59 – 81. DOI: 10.1007/BF03053419.
28. Govender M., Chetty K., Bulcock H. A review of hyperspectral remote sensing and its application in vegetation and water resource studies. *Water SA*. 2007. vol. 33. no. 2. pp. 145–151. DOI:10.4314/wsa.v33i2.49049.
29. Sharma A., Kaur D., Gupta A., et. al. Application and Analysis of Hyperspectral Imaging. *IEEE International Conference on Signal Processing, Computing and Control (ISPCC 2k19)*. 2019. pp. 30–35. DOI: 10.1109/ISPCC48220.2019.8988436.
30. Gonsales R., Vuds R. Cifrovaya obrabotka izobrazhenij. [Digital image processing]. M.: Tekhnosfera. 2006. 1072 p. (In Russ.).
31. Gao J., Li H., Han Z., et. al. Aircraft Detection in Remote Sensing Images Based on Background Filtering and Scale Prediction. *PRICAI 2018: Trends in Artificial Intelligence. Lecture Notes in Computer Science*. 2018. vol. 11012. pp. 604–616. DOI: /10.1007/978-3-319-97304-3_46.
32. Andriyanov N.A., Vasiliev K.K., Dement'ev V.E. Analysis of the efficiency of satellite image sequences filtering. *Journal of Physics: Conference Series*. 2018. vol.1096: 012036. DOI: 10.1088/1742- 6596/1096/1/012036.
33. Bannari, A., Morin, D., Bonn, F, et. al. A review of vegetation indices. *Remote Sensing Reviews*. 1995. vol. 13. no.1. pp. 95–120. DOI: 10.1080/02757259509532298.
34. Cherepanov A.S. [Vegetation indices]. *Geomatics*. 2011. №2. p. 98–102. (In Russ.).
35. Haralick R.M., Shanmugam K., Dinstein I. Textural features for image classification. *IEEE Transactions on Systems, Man, and Cybernetics*. 1973. vol. SMC-3. no. 6. pp. 610–621. DOI: 10.1109/TSMC.1973.4309314.
36. Feng Q., Liu J., Gong J. UAV Remote Sensing for Urban Vegetation Mapping Using Random Forest and Texture Analysis. *Remote Sensing*. 2015. vol. 7. pp. 1074–1094. DOI: 10.3390/rs70101074.
37. Kwak G.-H., Park N.-W. Impact of Texture Information on Crop Classification with Machine Learning and UAV Images. *Applied Science*. 2019. vol. 9: 643. DOI: 10.3390/app9040643.
38. Dian Y., Li Z., Pang Y. Spectral and texture features combined for forest tree species classification with airborne hyperspectral imagery. *Journal of the Indian Society of Remote Sensing*. 2015. vol. 43. pp. 101–107. DOI: 10.1007/s12524-014-0392-6.
39. Bekkari A., Idbraim S., Elhassouny A., et. al. SVM and Haralick features for classification of high resolution satellite images from urban areas. *ICISP 2012: Image and Signal Processing. Lecture Notes in Computer Science*. 2012. vol. 7340. pp. 17–26. DOI: 10.1007/978-3-642-31254-0_3.

40. Rodarmel C., Shan J. Principal component analysis for hyperspectral image classification. *Surveying and Land Information Systems*. 2002. vol. 62. no. 2, pp. 115–122.
41. Rejichi S., Chaabane F. Feature extraction using PCA for VHR satellite image time series spatial-temporal classification. *IEEE International Geoscience and Remote Sensing Symposium (IGARSS)*. 2015. pp. 485–488. DOI: 10.1109/IGARSS.2015.7325806.
42. Zlobina A.G., Shaura A.S., Zhurbin I.V., et al. Algorithm for statistical analysis of multispectral survey data to identify the anthropogenic impact of the 19th century on the natural environment. *Pattern Recognition and Image Analysis*. 2021. vol. 31. no. 2, pp. 344–354. DOI: 10.1134/S1054661821020176.
43. Khalid S., Khalil T., Nasreen S. A survey of feature selection and feature extraction techniques in machine learning. *Science and Information Conference*. 2014. pp. 372–378. DOI: 10.1109/SAI.2014.6918213.
44. Popescu M.C., Sasu L.M. Feature extraction, feature selection and machine learning for image classification: A case study. *2014 International Conference on Optimization of Electrical and Electronic Equipment (OPTIM)*. 2014. pp. 968–973. DOI: 10.1109/OPTIM.2014.6850925.
45. Ulaby F.T., Kouyate F., Brisco B., et al. Textural Information in SAR Images. *IEEE Transactions on Geoscience and Remote Sensing*. 1986. vol. GE-24. no. 2, pp. 235–245. DOI:10.1109/TGRS.1986.289643.
46. Jolliffe I.T. *Principal Components Analysis*. 2nd ed. N.Y.: Springer-Verlag New York Inc. 2002. 487 p.
47. Kang, B., Jung H., Jeong H., et al. Characterization of three-dimensional channel reservoirs using ensemble Kalman filter assisted by principal component analysis. *Petroleum Science*. 2020. vol. 17, pp. 182–195. DOI: 10.1007/s12182-019-00362-8.
48. Artoni F., Delorme A., Makeig S. Applying dimension reduction to EEG data by principal component analysis reduces the quality of its subsequent independent component decomposition. *NeuroImage*. 2018. vol. 175, pp. 176–187. DOI: 10.1016/j.neuroimage.2018.03.016.
49. Huang Z. Extensions to the k-means algorithm for clustering large data sets with categorical values. *Data mining and knowledge discovery*. 1998. vol. 2, pp. 283–304.
50. Jain A.K. Data clustering: 50 years beyond k-means. *Pattern Recognition Letters*. 2010. vol. 31, pp. 651–666. DOI: 10.1016/j.patrec.2009.09.011.

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 15.03.2022

Printed in Publishing center GUAP, 190000, St. Petersburg, B. Morskaya 67, litera A, Russia

The journal is registered in the Russian Federal Agency for Communications and Mass-Media Supervision, certificate ПИ № ФС77-79228 dated September 25, 2020
Subscription Index П5513, Russian Post Catalog

Подписано к печати 15.03.2022. Формат 60×90 1/16. Усл. печ. л. 13,9. Заказ № 98.

Тираж 300 экз., цена свободная.

Отпечатано в Редакционно-издательском центре ГУАП, 190000,
г. Санкт-Петербург, ул. Б. Морская, д. 67, лит. А

Журнал зарегистрирован Федеральной службой по надзору в сфере связи и массовых коммуникаций, свидетельство ПИ № ФС77-79228 от 25 сентября 2020 г.

Подписной индекс П5513 по каталогу «Почта России»