

**УЧРЕДИТЕЛИ**

**ФГАОУ ВО «ЮЖНО-УРАЛЬСКИЙ
ГОСУДАРСТВЕННЫЙ
УНИВЕРСИТЕТ (НИУ)»**

**ПРЕДСЕДАТЕЛЬ
РЕДАКЦИОННОГО
СОВЕТА**

ЧУВАРДИН О. П.,
руководитель Управления
Федеральной службы по
техническому и экспортному
контролю России по Уральскому
федеральному округу

**ГЛАВНЫЙ РЕДАКТОР
СОКОЛОВ А. Н.,**

к. т. н., доцент, зав. кафедрой
«Защита информации», Южно-
Уральский государственный
университет (национальный
исследовательский университет)
(г. Челябинск)

**ВЫПУСКАЮЩИЙ
РЕДАКТОР
СОГРИН Е. К.****ОТВЕТСТВЕННЫЙ
СЕКРЕТАРЬ
АНДРИАДИС Е. Ю.****ВЁРСТКА
ПЕЧЕНКИН В. А.****КОРРЕКТОР
ФЁДОРОВ В. С.**

**Подписной индекс 73852
в каталоге «Почта России»**

Журнал зарегистрирован Федераль-
ной службой по надзору в сфере
связи, информационных технологий
и массовых коммуникаций.

Свидетельство
ПИ № ФС77-65765 от 20.05.2016

Адрес редакции и издателя: Россия,
454080, г. Челябинск, пр. Ленина, д.
76. ЮУрГУ, Издательский центр
Тел./факс (351) 267-97-01.

Электронная версия
журнала в Интернете:
www.info-secur.ru,
e-mail: urvest@mail.ru

РЕДАКЦИОННЫЙ СОВЕТ:**БАРАНКОВА И. И.,**

д. т. н., профессор, зав. кафедрой
«Информатика и информаци-
онная безопасность», Магнитогор-
ский государственный техниче-
ский университет им. Г. И. Носова
(г. Магнитогорск);

ВАСИЛЬЕВ В. И.,

д. т. н., профессор, профессор
кафедры «Вычислительная
техника и защита информации»,
Уфимский государственный
авиационный технический
университет (г. Уфа);

ВОЙТОВИЧ Н. И.,

д. т. н., профессор, профессор
кафедры «Радиоэлектроника и
системы связи», Южно-Ураль-
ский государственный универ-
ситет (национальный исследо-
вательский университет)
(г. Челябинск);

ГАЙДАМАКИН Н. А.,

д.т.н., профессор, профессор
Учебно-научного центра
«Информационная безопас-
ность», Уральский федеральный
университет им. первого
президента России Б.Н. Ельцина
(г. Екатеринбург);

ДИК Д. И.,

к. т. н., доцент, зав. кафедрой
«Безопасность информаци-
онных и автоматизированных
систем», Курганский государ-
ственный университет
(г. Курган);

ЗАХАРОВ А. А.,

д.т.н., профессор, профессор
школы компьютерных наук,
Тюменский государственный
университет (г. Тюмень);

ЗЫРЯНОВА Т. Ю.,

к. т. н., доцент, зав. кафедрой
«Информационные технологии
и защита информации»,
Уральский государственный
университет путей сообщения
(г. Екатеринбург);

МЕЛЬНИКОВ А. В.,

д. т. н., профессор, директор
Югорского научно-исследова-
тельского института информа-
ционных технологий
(г. Ханты-Мансийск);

МИНБАЛЕЕВ А. В.,

д. ю. н., доцент, зав. кафедрой
«Информационное право и
цифровые технологии»,
Московский государственный
юридический университет
им. О. Е. Кутафина (МГЮА,
г. Москва);

ПОРШНЕВ С. В.,

д.т.н., профессор, профессор
Учебно-научного центра
«Информационная безопас-
ность», Уральский федеральный
университет им. первого
президента России
Б.Н. Ельцина (г. Екатеринбург);

РУЧАЙ А.Н.,

д.т.н., доцент, зав. кафедрой
«Компьютерная безопасность и
прикладная алгебра», Челяб-
инский государственный универ-
ситет (г. Челябинск);

ХОРЕВ А. А.,

д. т. н., профессор, зав. кафедрой
«Информационная безопас-
ность», Национальный исследо-
вательский университет
«Московский институт
электронной техники»
(г. Москва, г. Зеленоград);

ШАБУНИН С. Н.,

д.т.н., профессор, зав. кафедрой
«Радиоэлектроника и телеком-
муникации», Уральский
федеральный университет
им. первого президента России
Б.Н. Ельцина (г. Екатеринбург).



FOUNDER

**SOUTH URAL STATE
UNIVERSITY (NIU)**

CHAIRMAN OF THE EDITORIAL BOARD

CHUVARDIN O. P.,

Head of Department Federal
Service for Technical and Export
Control of Russia for the Urals
Federal District

CHIEF EDITOR

SOKOLOV A.N.,

Ph.D., Associate Professor, Head
of Department «Information
Protection», South Ural State
University (National Research
University) (Chelyabinsk city)

PRODUCING EDITOR

SOGRIN E. K.

LAYOUT

PECHENKIN V. A.

PROOFREADING

FEDOROV V. S.

EDITORIAL COUNCIL:

BARANKOVA I. I.,

Doctor of Technical Sciences,
Professor, Head of Department
«Informatics and Information
Security», Magnitogorsk State
Technical University named after
G.I. Nosova (Magnitogorsk city);

VASILYEV V. I.,

Doctor of Technical Sciences,
Professor, Professor of the
Department «Computer Science
and Information Protection», Ufa
State Aviation Technical
University (Ufa city);

VOITOVICH N. I.,

Doctor of Technical Sciences,
Professor, Professor of the
Department «Radioelectronics
and Communication Systems»,
South Ural State University
(National Research University)
(Chelyabinsk city);

GAYDAMAKIN N. A.,

Doctor of Technical Sciences,
Professor, Professor of the
Information Security Training and
Research Center of the Ural
Federal University named after
the first President of Russia
B.N.Yeltsin (Ekaterinburg city);

DIK D. I.,

Ph.D., Associate Professor, Head of
Department «Security of
information and automated
systems», Kurgan State University
(Kurgan city);

ZAHAROV A. A.,

Doctor of Technical Sciences,
Professor, Professor of the School
of Computer Science, Tyumen
State University (Tyumen city);

ZYRYANOVA T. Y.,

Ph.D., Associate Professor, Head of
Department «Information
Technologies and Information
Protection», Ural State
University ways of
communication (Ekaterinburg
city);

MELNIKOV A. V.,

Doctor of Technical Sciences,
Professor, Director Ugra Research
Institute of Information
Technologies (Khanty-Mansiysk
city);

MINBALEEV A.V.,

Doctor of Law, Associate
Professor, Head of Department of
«Information Law and Digital
Technologies», Moscow State Law
University. O. E.Kutafina (Moscow
city);

PORSHNEV S. V.,

Doctor of Technical Sciences,
Professor, Professor of the
Training and Scientific Center
«Information Security», Ural
Federal University named after
the first President of Russia
B.N.Yeltsin (Ekaterinburg city);

RUCHAY A.N.,

Doctor of Technical Sciences,
Associate Professor, Head
of the Department «Computer
Security and Applied Algebra»,
Chelyabinsk State University
(Chelyabinsk city);

HOREV A. A.,

Doctor of Technical Sciences,
Professor, Head of Department of
«Information Security», National
Research University «Moscow
Institute of Electronic
Technology» (Moscow, the city of
Zelenograd);

SHABUNIN S. N.,

Doctor of Technical Sciences,
Professor, Head of Department
«Radioelectronics and
Telecommunications», Ural
Federal University named after
the first President of Russia
B.N.Yeltsin (Ekaterinburg city).

**Subscription index 73852
in the «Russian Post» catalog**

The journal is registered by the Federal
service in the field of communication,
information technology and mass
communications.

Certificate

PI No. ФС77-65765 dd. 05/20/2016

Editorial and publisher address: Russia,
454080, Chelyabinsk, Lenin Avenue, 76
SUSU, Publishing Center

Phone / fax (351) 267-97-01.

Electronic version of the
magazine in the Internet:

www.info-secur.ru,

e-mail: urvest@mail.ru



**СИСТЕМНЫЙ АНАЛИЗ,
УПРАВЛЕНИЕ И ОБРАБОТКА
ИНФОРМАЦИИ**

**КУШНИР Н. В., ВЛАСЕНКО А. В.,
БУЛГАКОВ В. В., АКУЛИНИН Г. А.**
Машинное обучение и эвристические
методы в задачах классификации текстов
для информационной безопасности 5

**КУШНИР Н. В., ВЛАСЕНКО А. В.,
ТЮТЮНИК Н. П. СМИРНОВ Б. А.**
Практические аспекты применения
ИИ-систем обнаружения вторжений 14

ПОРШНЕВ С. В., РЯБКО Н. Ю.
Исследование методов
электоральной криминалистики 21

**МЕТОДЫ И СИСТЕМЫ
ЗАЩИТЫ ИНФОРМАЦИИ,
ИНФОРМАЦИОННАЯ
БЕЗОПАСНОСТЬ**

МЕЛЬНИКОВ А. В., ЯРЫШКИН И. Н.
Автоматизация оценки вредоносности
скриптов на основе больших языковых
моделей при расследовании
компьютерных инцидентов 36

МЕЛЬНИКОВ Д. Ю.
Система алгоритмов обнаружения
комплексных компьютерных атак
на основе анализа действий пользователей
и сопоставления с базой знаний
MITRE ATT&CK 44

КОРЖЕНЕВСКИЙ Д. А., ЗЫРЯНОВА Т. Ю.
Методика оценки рисков информационной
безопасности координатно-графовым
методом 53

МЕДВЕДЕВА Н. В., ТИТОВ С. С.
Аналоги совершенных шифров
в теории кодирования 64

IN THIS ISSUE

SYSTEM ANALYSIS, MANAGEMENT AND INFORMATION PROCESSING

KUSHNIR N. V., VLASENKO A. V., BULGAKOV V. V., AKULININ G. A. Machine learning and heuristic methods in text classification tasks for information security	5
KUSHNIR N. V. VLASENKO A. V., TYUTYUNIK N. P. SMIRNOV B. A. Practical aspects of using AI-based intrusion detection systems	14
PORSHNEV S. V., RYABKO N. YU. The electoral forensics methods research	21

METHODS AND SYSTEMS OF INFORMATION PROTECTION, INFORMATION SECURITY

MELNIKOV A. V., YARYSHKIN I. N. Automating the assessment of script maliciousness based on large language models when investigating computer incidents	36
MELNIKOV D. Y. System of algorithms for detecting complex computer attacks based on analysis of user actions and comparison with the MITRE ATT&CK knowledge base	44
KORZHENEVSKIY D. A., ZYRYANOVA T. Y. Methodology for assessing information security risks using the coordinate graph method	53
MEDVEDEVA N. V., TITOV S. S. Analogues of perfect ciphers in coding theory	64



МАШИННОЕ ОБУЧЕНИЕ И ЭВРИСТИЧЕСКИЕ МЕТОДЫ В ЗАДАЧАХ КЛАССИФИКАЦИИ ТЕКСТОВ ДЛЯ ИНФОРМАЦИОННОЙ БЕЗОПАСНОСТИ

Статья посвящена сравнительному анализу эвристических и обучаемых методов в задачах классификации текстов, рассматриваемых в контексте информационной безопасности. В качестве прикладного примера выбрана задача фильтрации спам-сообщений, позволяющая продемонстрировать особенности, преимущества и ограничения каждого подхода. Описана реализация двух систем: эвристической, основанной на наборе детерминированных правил, и обучаемой модели, использующей алгоритмы машинного обучения и предварительную обработку текстов на естественном языке.

Далее проведено экспериментальное сравнение качества классификации, скорости обработки и гибкости адаптации обеих систем. Особое внимание уделено критериям точности, полноты, F1-меры, а также возможностям автономного функционирования и расширяемости.

В заключении на основе полученных результатов формулируются выводы о целесообразности применения каждого из подходов в зависимости от условий эксплуатации и требований к безопасности.

Ключевые слова: классификация текстов, эвристические методы, машинное обучение, информационная безопасность, спам-фильтрация, бинарная классификация, обработка естественного языка.

MACHINE LEARNING AND HEURISTIC METHODS IN TEXT CLASSIFICATION TASKS FOR INFORMATION SECURITY

The article is devoted to a comparative analysis of heuristic and trainable methods in the tasks of text classification considered in the context of information security. The task of filtering spam messages is chosen as an applied example, which makes it possible to demonstrate the features, advantages and limitations of each approach. The implementation of two systems is described: a heuristic one based on a set of deterministic rules, and a trainable model using machine learning algorithms and preprocessing of natural language texts.

Next, an experimental comparison was made of the classification quality, processing speed, and adaptability of both systems. Special attention is paid to the criteria of accuracy, completeness, F1-measure, as well as the possibilities of autonomous operation and extensibility.

In conclusion, based on the results obtained, conclusions are drawn about the expediency of using each of the approaches, depending on the operating conditions and safety requirements.

Keywords: text classification, heuristic methods, machine learning, information security, spam filtering, binary classification, natural language processing.

Введение

Информационная безопасность в современном цифровом мире становится неотъемлемой частью стабильного функционирования частных и государственных информационных систем. Одной из ключевых задач в этой области является эффективная и быстрая обработка и классификация текстовых данных, среди которых могут находиться потенциально опасные сообщения. Тексты, содержащие признаки фишинга, мошенничества или деструктивного контента, могут угрожать конфиденциальности, целостности и доступности данных, особенно при массовом распространении, как это происходит, например, с текстовыми сообщениями. Традиционные эвристические методы фильтрации и классификации текстов, основанные на фиксированных правилах, обладают высокой интерпретируемостью и простотой настройки, но зачастую не справляются с постоянно эволюционирующими типами угроз. В то же время методы машинного обучения демонстрируют высокую адаптивность и точность за счёт способности выявлять скрытые закономерности в больших объёмах текстов. Однако их применение требует качественно размеченных данных, вычислительных ре-

сурсов и дополнительной инфраструктуры. В условиях стремительного роста объёмов цифровых сообщений и увеличения количества сложных атак особенно актуальным становится вопрос выбора подходящего инструмента для автоматической классификации текстов в рамках задач информационной безопасности. Сравнительный анализ эвристических и обучаемых подходов позволяет не только оценить их эффективность в прикладных задачах, но и сформировать понимание о целесообразности их использования в различных контекстах — от автономных систем до крупных защищённых сетей. Таким образом, изучение различий между этими подходами становится важным этапом в построении более устойчивых и интеллектуальных систем защиты информации.

Применение эвристических методов в системах защиты

Эвристические методы классификации текстов, особенно в контексте информационной безопасности, представляют собой совокупность правил и эвристик, разработанных вручную и основанных на анализе признаков, характерных для вредоносных или нежелательных сообщений. Наиболее известным

примером является система SpamAssassin, в основе которой лежит обширный набор регулярных выражений, логических правил и сигнатур, способных обнаруживать шаблоны спама. Подобные подходы широко применяются в почтовых клиентах Outlook, Mozilla Thunderbird и фильтрующих прокси-серверах вроде Squid и WebSense.

Эвристические правила также активно применяются в антифишинговых механизмах: анализируется домен отправителя, наличие в тексте подозрительных фраз, частота появления ссылок и вложений. Специальные правила позволяют выявлять сообщения, содержащие ключевые слова вроде "вы выиграли", "переведите деньги", "действуйте сейчас" и т.д. Преимущество данных методов заключается в их простоте, быstroдействии и полной автономности, что делает их пригодными для систем с ограниченными вычислительными ресурсами.

Однако у эвристических решений есть существенные недостатки. В первую очередь, это низкая гибкость: при появлении новых методов обхода фильтров злоумышленниками, правила требуют ручного обновления. Кроме того, они демонстрируют невысокую точность при работе с неоднозначными или грамматически искажёнными текстами.

Машинное обучение в задачах классификации и защиты

Методы машинного обучения (МО) предлагают альтернативу ручному проектированию правил. Они способны выявлять скрытые закономерности в больших объёмах данных, основываясь на примерах помеченных сообщений. Такие подходы стали особенно популярны в задачах классификации спама, фишинга и обнаружения вредоносных вложений [1].

Один из наиболее известных примеров — фильтрация спама в Gmail, где используются модели, обученные на огромном количестве пользовательских писем. Подходы включают классические алгоритмы, такие как наивный байесовский классификатор, который определяет спам на основе вероятностей слов; логистическая регрессия, метод статистической классификации для разделения спама и не спама; и SVM, или машина опорных векторов, которая находит оптимальную границу для разделения данных. Также применяются современные нейросетевые модели, включая BERT, модель на основе трансформеров, глубоко понимающая контекст

текста, и DistilBERT, более легкая и быстрая версия BERT, оптимизированная для задач обработки естественного языка [2] [3].

Применение МО также эффективно в системах обнаружения фишинга: модели анализируют как текст письма, так и характеристики ссылок и вложений. В других случаях — при анализе логов систем безопасности, журналов событий, сетевого трафика — машинное обучение позволяет выявлять аномалии и потенциальные угрозы без необходимости заранее формализовать правила [4] [5] [6].

Отдельно стоит отметить тренд на гибридные подходы, в которых эвристические фильтры используются как предварительный барьер, а более ресурсоёмкие МО-модели подключаются на следующих этапах анализа.

Постановка задач исследования

Цель настоящего исследования — провести сравнительный анализ эффективности эвристического и обучаемого подходов к классификации текстов в задачах информационной безопасности. В качестве примера рассматривается проблема автоматической фильтрации текстовых сообщений.

Для достижения цели поставлены следующие задачи:

1. Разработать систему фильтрации сообщений на основе набора эвристических правил;
2. Реализовать систему машинного обучения для классификации сообщений на основе размеченного корпуса данных;
3. Провести экспериментальное сравнение решений по метрикам точности, полноты, F1-меры и скорости обработки;
4. Оценить возможности масштабирования, автономности и адаптивности обоих подходов.

Система на основе эвристических правил

Фильтрация реализована в виде Python-скрипта, обрабатывающего текстовые SMS-сообщения. Для каждого сообщения применяются эвристические правила, где каждое правило добавляет по одному баллу к итоговой оценке.

Примеры правил:

- Наличие ключевых слов: "бесплатно", "подарок", "срочно";
- Более 2 ссылок в сообщении;
- Высокая доля заглавных букв;
- Подозрительная длина сообщения (меньше 20 или более 160 символов);
- Высокая частота специальных символов.

Сообщение считается спамом при достижении порогового значения (например, три балла). Порог можно настроить вручную. Для каждого сообщения формируется отчёт с указанием, какие именно правила сработали.

Листинг 1. Эвристический классификатор сообщений

```
import re
from typing import List, Dict

# Настраиваемые параметры
SPAM_KEYWORDS = [
    'бесплатно', 'выиграть', 'срочно', 'акция', 'купи',
    'free', 'win', 'urgent', 'act now', 'buy now'
]
MAX_LINKS = 2
UPPERCASE_RATIO_THRESHOLD = 0.5
SPECIAL_CHAR_THRESHOLD = 10
SPAM_SCORE_THRESHOLD = 3

def extract_links(text: str) -> List[str]:
    return re.findall(r'https?:\/\/\S+', text)

def get_uppercase_ratio(text: str) -> float:
    if not text:
        return 0
    uppercase_chars = sum(1 for c in text if c.isupper())
    return uppercase_chars / len(text)

def get_special_char_count(text: str) -> int:
    return sum(1 for c in text if not c.isalnum() and not c.isspace())

def classify_sms(text: str) -> Dict:
    score = 0
    reasons = []

    text_lower = text.lower()

    # 1. Ключевые слова
    for keyword in SPAM_KEYWORDS:
        if keyword in text_lower:
            score += 1
            reasons.append(f"Ключевое слово: '{keyword}'")
            break

    # 2. Ссылки
    links = extract_links(text)
    if len(links) > MAX_LINKS:
        score += 1
        reasons.append(f"Слишком много ссылок: {len(links)}")

    # 3. Заглавные буквы
    uppercase_ratio = get_uppercase_ratio(text)
    if uppercase_ratio > UPPERCASE_RATIO_THRESHOLD:
        score += 1
        reasons.append(f"Высокая доля ЗАГЛАВНЫХ букв: {uppercase_ratio:.2f}")

    # 4. Спецсимволы
    special_chars = get_special_char_count(text)
    if special_chars > SPECIAL_CHAR_THRESHOLD:
```



```

        score += 1
        reasons.append(f"Много специальных символов: {special_chars}")

# 5. Длина сообщения
if len(text) < 20 or len(text) > 160:
    score += 1
    reasons.append(f"Подозрительная длина сообщения: {len(text)}
символов")

is_spam = score >= SPAM_SCORE_THRESHOLD

return {
    'result': 'SPAM' if is_spam else 'NOT SPAM',
    'score': score,
    'reasons': reasons
}

# Пример использования
if __name__ == "__main__":
    test_sms = "АКЦИЯ! Выиграй миллион! https://spam.link"
    result = classify_sms(test_sms)
    print("Результат:", result['result'])
    print("Сработавшие правила:")
    for reason in result['reasons']:
        print("-", reason)

```

Система на основе машинного обучения

В качестве датасета используется SMS Spam Collection Dataset, содержащий 5572 текстовых SMS-сообщений, размеченных как спам (747 сообщений) или не спам (4825 сообщений). Датасет представляет собой общедоступный набор данных, часто используемый для задач классификации текстов. Предварительная обработка включает:

- Очистку HTML-тегов и специальных символов;
- Приведение текста к нижнему регистру;
- Удаление стоп-слов с использованием библиотеки NLTK;
- Токенизацию и векторизацию текста с применением TF-IDF.

Для классификации рассматриваются модели:

- Наивный байесовский классификатор (MultinomialNB);

- Логистическая регрессия (Logistic Regression);
- Метод опорных векторов (SVM).

Обучение производится на 80% выборки, 20% оставляется для тестирования. Модель сохраняется через joblib и может дообучаться на новых пользовательских данных.

Для оценки используются метрики:

- Accuracy (доля верно классифицированных сообщений);
- Precision и Recall;
- F1-мера как гармоническое среднее.

Также выводится вероятность классификации для каждого сообщения, что повышает доверие пользователя к системе.

Скорость обработки одного сообщения не превышает 0.2 секунды при использовании TF-IDF и логистической регрессии. Система работает в автономном режиме и поддерживает обновление модели через загрузку локального файла с новыми данными.

Листинг 2. Классификатор сообщений на основе машинного обучения

```

import pandas as pd
import numpy as np
import re
import joblib
from sklearn.feature_extraction.text import TfidfVectorizer
from sklearn.linear_model import LogisticRegression
from sklearn.model_selection import train_test_split, cross_validate
from sklearn.metrics import classification_report
from sklearn.pipeline import Pipeline

```

```

from sklearn.base import BaseEstimator, TransformerMixin
from nltk.corpus import stopwords
from nltk.stem.snowball import SnowballStemmer
import nltk
import os

nltk.download('stopwords')

# --- Кастомный препроцессор ---
class TextCleaner(BaseEstimator, TransformerMixin):
    def __init__(self, language='english'):
        self.language = language
        self.stop_words = stopwords.words(language)
        self.stemmer = SnowballStemmer(language)

    def clean(self, text):
        text = re.sub(r'<[^>+>', ' ', text) # Remove HTML tags
        text = re.sub(r'http\S+|www\S+', '', text) # Remove URLs
        text = re.sub(r'^a-zA-Za-яA-Я\s', ' ', text) # Keep only
letters
        text = text.lower()
        tokens = text.split()
        tokens = [t for t in tokens if t not in self.stop_words]
        tokens = [self.stemmer.stem(t) for t in tokens]
        return ' '.join(tokens)

    def fit(self, X, y=None):
        return self

    def transform(self, X):
        return [self.clean(text) for text in X]

# --- Загрузка и обучение модели ---
def train_model(dataset_path='spam.csv', language='english'):
    df = pd.read_csv(dataset_path, encoding='latin-1')[['v1', 'v2']]
    df.columns = ['label', 'text']
    df['label'] = df['label'].map({'ham': 0, 'spam': 1})

    X, y = df['text'], df['label']

    pipeline = Pipeline([
        ('cleaner', TextCleaner(language=language)),
        ('tfidf', TfidfVectorizer(max_features=3000)),
        ('classifier', LogisticRegression())
    ])

    # 5-кратная кросс-валидация
    scoring = ['accuracy', 'precision', 'recall', 'f1']
    scores = cross_validate(pipeline, X, y, cv=5, scoring=scoring)
    print("Результаты 5-кратной кросс-валидации:")
    print(f"Точность (Accuracy): {scores['test_accuracy'].mean():.3f}")
    print(f"Прецизия (Precision): {scores['test_precision'].mean():.3f}")
    print(f"Полнота (Recall): {scores['test_recall'].mean():.3f}")
    print(f"F1-мера: {scores['test_f1'].mean():.3f}")

    # Обучение на всей тренировочной выборке
    X_train, X_test, y_train, y_test = train_test_split(X, y, test_
size=0.2, random_state=42)
    pipeline.fit(X_train, y_train)
    y_pred = pipeline.predict(X_test)
    print("\nОтчёт по тестовой выборке:")
    print(classification_report(y_test, y_pred))

```

```
joblib.dump(pipeline, 'spam_classifier.pkl')
print("Модель сохранена в spam_classifier.pkl")

# --- Классификация новых сообщений ---
def classify_sms(text, model_path='spam_classifier.pkl'):
    if not os.path.exists(model_path):
        raise FileNotFoundError("Модель не найдена. Пожалуйста, обучите
её сначала.")

    model = joblib.load(model_path)
    pred = model.predict([text])[0]
    prob = model.predict_proba([text])[0][pred]
    return {
        'result': 'SPAM' if pred == 1 else 'NOT SPAM',
        'confidence': round(prob, 3)
    }

# --- Пример использования ---
if __name__ == '__main__':
    train_model('spam.csv')
    test_text = "Выиграй миллион прямо сейчас! Быстро, бесплатно, без
регистрации!"
    result = classify_sms(test_text)
    print("Результат:", result['result'])
    print("Уверенность:", result['confidence'])
```

**Сравнение эффективности
эвристического и машинного методов**

Эксперименты проводились на подмно- жестве SMS Spam Collection Dataset, включа- ющем 1000 текстовых сообщений (500 — спам, 500 — не спам), предварительно очи- щенных и токенизированных. Подмножество было сформировано путём случайной стра- тифицированной выборки для сохранения баланса классов. Для эвристического метода порог в 3 балла был выбран эмпирически на основе анализа ложных срабатываний на ва- лидационной выборке. Для машинного обу- чения гиперпараметры логистической ре- грессии (параметр регуляризации C в диапа- зоне [0.1, 1, 10]) подбирались с использовани- ем поиска по сетке (GridSearchCV). Параме- тры TF-IDF включали максимальный размер словаря 3000 слов и исключение слов с ча-

стотой менее 5. Время классификации изме- рялось как среднее время предсказания для одного сообщения, исключая предобработку данных (табл. 1).

Для сравнения применялись стандарт- ные метрики качества классификации:

- Accuracy (точность) — доля правильно классифицированных сообщений;
- Precision (прецизия) — доля правильно распознанных спамов среди всех поме- ченных как спам;
- Recall (полнота) — доля правильно рас- познанных спамов среди всех реаль- ных спамов;
- F1-мера — гармоническое среднее между Precision и Recall, отражает ба- ланс между ними;
- Время классификации одного сообще- ния.

Таблица 1

Результаты сравнения.

Метрика	Эвристический метод	Машинное обучение
Accuracy	84.1%	96.7%
Precision	87.4%	97.9%
Recall	81.6%	95.5%
F1-мера	84.3%	96.7%
Время на сообщение	0.014 с	0.097 с
Обновление правил	Вручную	Автоматически
Гибкость	Низкая	Высокая

Машинное обучение демонстрирует значительно лучшие результаты по всем основным метрикам качества классификации, особенно в условиях увеличенного объема данных. Эвристический подход выигрывает по скорости обработки и простоте реализации, но уступает в гибкости и способности адаптироваться к новым типам спама.

Разработанные системы классификации текстовых сообщений демонстрируют два подхода к решению задач информационной безопасности: эвристический и основанный на машинном обучении. Эвристический классификатор показывает приемлемую точность и высокую скорость обработки, что делает его подходящим для использования в условиях ограниченных ресурсов и необходимости полной автономности. В то же время система, построенная на методах машинного

обучения, продемонстрировала значительно более высокие показатели точности, полноты и F1-меры, что указывает на её эффективность в условиях сложной и динамичной структуры входящих данных. Сравнительный анализ подтвердил, что каждый из подходов имеет свои преимущества в зависимости от целей, требований и условий эксплуатации. Однако, система на основе машинного обучения оказывается более перспективной в контексте дальнейшего развития, особенно с учётом возможности адаптации к новым видам угроз. Полученные результаты подчёркивают актуальность применения современных методов анализа текста в задачах обеспечения информационной безопасности и позволяют сформировать основу для разработки интеллектуальных фильтрующих систем нового поколения.

Литература

1. Волкова, С. С. Введение в машинное обучение. Линейные модели : учебное пособие / С. С. Волкова. – Вологда: Вологод, 2023. – 76 с. – ISBN 978-5-907606-46-3. – EDN RXAMCS (дата обращения 28.10.2025).
2. Спам-фильтрация электронных почтовых сообщений на основе нейросетевой и нейронечеткой моделей / А. С. Катасев, Д. В. Катасева, А. П. Кирпичников, Я. Е. Семенов // Вестник Технологического университета. – 2015. – Т. 18, № 15. – С. 217-220. – EDN ULRPSX (дата обращения 03.11.2025).
3. Гуров, В. В. Спам-фильтры для предприятий / В. В. Гуров // Сети и системы связи. – 2007. – № 6. – С. 80-89. – EDN HZTOIL (дата обращения 03.11.2025).
4. Анализ данных и процессов / А. Барсегян, М. Куприянов, И. Холод [и др.]. – 3-е изд. – Санкт-Петербург: БХВ-Петербург, 2010. – 512 с. – ISBN 978-5-9775-0368-6. – EDN SDPQDT (дата обращения 05.11.2025).
5. Информационная безопасность и защита персональных данных. Проблемы и пути их решения: сборник материалов и докладов XVI межрегиональная научно-практическая конференция, Брянск, 29 апреля 2024 года. – Брянск: Брянский государственный технический университет, 2024. – 320 с. – ISBN 978-5-907570-87-0. – EDN LVASON (дата обращения 05.11.2025).
6. Повышение эффективности фильтрации спама на основе методов машинного обучения в сообщениях различной природы / И. А. Чижова, Е. И. Кублик, М. С. Чипчагов, А. И. Лабинцев // Нейрокомпьютеры: разработка, применение. – 2022. – Т. 24, № 5. – С. 5-18. – DOI 10.18127/j19998554-202205-01. – EDN RYRSXG (дата обращения 05.11.2025).

References

1. Volkova, S. S. Vvedeniye v mashinnoye obucheniye. Lineynyye modeli : uchebnoye posobiye / S. S. Volkova. – Vologda: Vologod, 2023. – 76 s. – ISBN 978-5-907606-46-3. – EDN RXAMCS (data obrashcheniya 28.10.2025).
2. Spam-fil'tratsiya elektronnykh pochtovykh soobshcheniy na osnove neyrosetevoy i neyronechetkoy modeley / A. S. Katasev, D. V. Kataseva, A. P. Kirpichnikov, YA. Ye. Semenov // Vestnik Tekhnologicheskogo universiteta. – 2015. – T. 18, № 15. – S. 217-220. – EDN ULRPSX (data obrashcheniya 03.11.2025).
3. Gurov, V. V. Spam-fil'try dlya predpriyatiy / V. V. Gurov // Seti i sistemy svyazi. – 2007. – № 6. – S. 80-89. – EDN HZTOIL (data obrashcheniya 03.11.2025).
4. Analiz dannykh i protsessov / A. Barsegyan, M. Kupriyanov, I. Kholod [i dr.]. – 3-ye izd. – Sankt-Peterbur: BKHV-Peterburg, 2010. – 512 s. – ISBN 978-5-9775-0368-6. – EDN SDPQDT (data obrashcheniya 05.11.2025).
5. Informatsionnaya bezopasnost' i zashchita personal'nykh dannykh. Problemy i puti ikh resheniya: sbornik materialov i dokladov KHVI mezhregional'naya nauchno-prakticheskaya konferentsiya, Bryansk, 29 aprelya 2024 goda. – Bryansk: Bryanskiy gosudarstvennyy tekhnicheskyy universitet, 2024. – 320 s. – ISBN 978-5-907570-87-0. – EDN LVASON (data obrashcheniya 05.11.2025).

КУШНИР Надежда Владимировна, старший преподаватель кафедры информационных систем и программирования ФГБОУ ВО «Кубанский государственный технологический университет». 350072, г. Краснодар, ул. Московская, д. 2. E-mail: kushnir.06@mail.ru

ВЛАСЕНКО Александра Владимировна, кандидат технических наук, доцент, Врио заведующего кафедрой информационной безопасности, профессор кафедры ФГКОУ ВО «Краснодарский университет МВД России», 350005, г. Краснодар, ул. Ярославская 128. E-mail: alex_vlasenko@list.ru

БУЛГАКОВ Владислав Витальевич, студент 3 курса направления подготовки 09.03.04 (программная инженерия) факультета информационных технологий и кибербезопасности (ФИТК) ФГБОУ ВО «Кубанский государственный технологический университет». 350072, г. Краснодар, ул. Московская, д. 2. E-mail: vlad1230943@mail.ru

АКУЛИНИН Глеб Александрович, студент 3 курса направления подготовки 09.03.04 (программная инженерия) факультета информационных технологий и кибербезопасности (ФИТК) ФГБОУ ВО «Кубанский государственный технологический университет». 350072, г. Краснодар, ул. Московская, д. 2. E-mail: akuliningleb1@mail.ru

KUSHNIR Nadezhda Vladimirovna, Senior Lecturer at the Department of Information Systems and Programming of the Federal State Budgetary Educational Institution of Higher Education «Kuban State Technological University». 350072, Krasnodar, Moskovskaya St., 2. E-mail: kushnir.06@mail.ru

VLASENKO Alexandra Vladimirovna, Candidate of Technical Sciences, Associate Professor, Acting Head of the Department of Information Security Professor of the Department, of the Krasnodar University of the Ministry of Internal Affairs of Russia. 350005, Krasnodar, Yaroslavskaya str., 128. E-mail: alex_vlasenko@list.ru

BULGAKOV Vladislav Vitalievich, Third-year year student of the training course 09.03.04 (software Engineering) of the Federal State Budgetary Educational Institution of Higher Education «Kuban State Technological University». 350072, Krasnodar, Moskovskaya St., 2. E-mail: vlad1230943@mail.ru

AKULININ Gleb Alexandrovich, Third-year year student of the training course 09.03.04 (software Engineering) of the Faculty of Information Technology and Cybersecurity (FITC) of the Federal State Budgetary Educational Institution of Higher Education «Kuban State Technological University». 350072, Krasnodar, Moskovskaya St., 2. E-mail: akuliningleb1@mail.ru

Кушнир Н. В., Власенко А. В.,
Тютюник Н. П. Смирнов Б. А.

10.14529/secur250402

ПРАКТИЧЕСКИЕ АСПЕКТЫ ПРИМЕНЕНИЯ ИИ-СИСТЕМ ОБНАРУЖЕНИЯ ВТОРЖЕНИЙ

В статье рассматриваются практические аспекты внедрения и эксплуатации ИИ систем обнаружения вторжений в реальных сетях, включая выявление новых угроз, работу в реальном времени и интеграцию с существующей инфраструктурой.

Ключевые слова: ИИ системы обнаружения вторжений, машинное обучение, глубокое обучение, сетевая безопасность, состязательные атаки.

Kushnir N. V. Vlasenko A. V., Tyutyunik N. P. Smirnov B. A.

PRACTICAL ASPECTS OF USING AI-BASED INTRUSION DETECTION SYSTEMS

The article discusses the practical aspects of implementing and operating AI-based intrusion detection systems in real networks, including the identification of new threats, real-time operation, and integration with existing infrastructure.

Keywords: AI-based intrusion detection systems, machine learning, deep learning, network security, and adversarial attacks.

Обнаружение zero-day атак и ранее неизвестных угроз

ИИ-системы IDS способны выявлять атаки, для которых ещё не существуют сигнатуры — zero-day угрозы. Особую роль тут играют методы неконтролируемого и полусконтролируемого обучения, которые фокусируются на аномалиях, отклонениях от нормы и многостадийных атаках [3].

Механизмы выявления аномалий

Применение RNN, LSTM и GAN позволяет фиксировать девиантное поведение без предварительной информации о типах угроз или шаблонов атак. Автокодировщики демонстрируют высокую эффективность в выявлении аномалий: обучаясь только на примерах нормального трафика, они фиксируют значительное увеличение ошибки восстанов-

ления для подозрительных пакетов [4]. Исследования показывают, что автокодировщики способны выявить до 95% новых типов атак, которые не встречались в обучающей выборке.

Практическое значение имеет комбинирование сигнатурных и аномалистических подходов. Система сначала проверяет известные угрозы через быструю сигнатурную базу, а затем направляет остальной трафик через глубокие нейросети для выявления новых паттернов. Такой двухэтапный процесс позволяет значительно снизить количество ложных срабатываний при сохранении высокой чувствительности. В промышленной практике такие системы снижают количество ложных срабатываний на 40-60% в сравнении с использованием только аномалистических подходов.

Анализ поведения пользователей

Профилирование пользовательского поведения (User and Entity Behavior Analytics — UEBA) позволяет выявлять скомпрометированные учётные записи и внутренние угрозы. ИИ-системы анализируют типичные паттерны доступа: рабочее время, типичные источники и целевые системы, объёмы передаваемых данных. При отклонении от этих паттернов система генерирует оповещение. Комбинация UEBA с IDS обеспечивает 360-градусное покрытие угроз как с внешней, так и с внутренней стороны.

Вызовы развёртывания и функционирования в реальном времени

В реальном мире основным ограничением внедрения ИИ является баланс между высокой производительностью и минимальной задержкой обработки информации. Такие системы требуют значительных вычислительных ресурсов: чем сложнее модель, тем выше требования к инфраструктуре. Гигабайты трафика в секунду требуют оптимизированных алгоритмов и архитектуры [5].

Обработка больших объёмов данных

Проблема ложных срабатываний остаётся актуальной и требует регулярного дообучения ИИ на свежих данных, чтобы:

- Уменьшить нагрузку на аналитиков;
- Избежать синдрома «alert fatigue» (усталости от сигналов);
- Адаптировать модель к эволюции сетевого поведения.

В крупных корпоративных сетях ежедневно генерируется от 100 ТБ до 1 ПБ данных сетевого трафика. Традиционные подходы требуют хранения всех этих данных для анализа, что приводит к экспоненциальному росту затрат на инфраструктуру. ИИ-системы позволяют проводить анализ в потоковом режиме (stream processing), выявляя угрозы в режиме реального времени без необходимости сохранения всего трафика.

Оптимизация производительности

Практические решения для оптимизации включают:

1. Кэширование и предвычисления — сохранение результатов часто встречающихся паттернов трафика в памяти высокоскоростных хранилищ (Redis, Memcached). Это может ускорить обработку повторяющихся паттернов на 50-70%.

2. Векторизация и параллелизм — использование GPU (NVIDIA Tesla V100, A100) и TPU (Google TPU) для ускорения матричных операций. Современные ускорители могут обработать трафик со скоростью до 200+ Гбит/с на одном устройстве.

3. Квантизация моделей — снижение точности представления чисел (от float32 к float16 или int8) без значительной потери точности, но с ускорением вычислений [6]. Исследования показывают, что квантизация может ускорить инференс на 3-4 раза с потерей точности менее 1%.

4. Архитектурная оптимизация — использование лёгких моделей для первоначальной фильтрации (например, Random Forest), а сложных моделей (Deep Neural Networks) только для подозрительных пакетов. Это позволяет обработать 80-90% трафика быстро с минимальными вычислительными затратами.

5. Распределённая обработка — использование Spark, Kafka и других систем для параллельной обработки данных на кластерах. Apache Spark может обработать петабайты данных за часы.

Состязательные атаки и устойчивость к обходу

ИИ-модели уязвимы для атак, при которых вредоносный трафик целенаправленно модифицируется для обхода системы обнаружения. Adversarial examples — это минимально изменённые входы, которые вызывают ошибки в классификации [7].

Типы состязательных атак

Исследователи выделяют несколько типов атак на ИИ-системы:

- Evasion attacks — модификация вредоносного трафика на этапе классификации для обхода системы.
- Poisoning attacks — внедрение вредоносных примеров в обучающую выборку для деградации производительности модели.
- Model extraction attacks — попытка скопировать или воспроизвести функциональность чёрного ящика через запросы.

Пример evasion attack: вредонос может добавить в пакет большое количество значений признаков, которые не влияют на его вредоносность, но сдвигают вектор признаков в пространство, где модель его классифицирует как нормальный трафик.

Методы защиты

Для борьбы с состязательными атаками применяются:

- Ансамблевое обучение — использование множества независимых моделей (5-10 различных архитектур) снижает вероятность обхода всех одновременно [8]. Исследования показывают, что ансамбль из 10 моделей на 90% более устойчив к состязательным примерам, чем отдельная модель.

- Состязательное обучение — обучение на специально сгенерированных adversarial примерах, которые пытаются обмануть модель, повышает её устойчивость. Процесс подобен боевой подготовке: система "тренируется" против возможных атак.

- Валидация признаков — проверка логичности значений атрибутов трафика может предотвратить использование нелогичных adversarial примеров. Например, проверка того, что значения портов находятся в диапазоне 0-65535.

- Детектирование аномалий вторичного уровня — дополнительные нейросети, обученные выявлять характеристики adversarial примеров. Это создаёт многоуровневую защиту.

Модели, натренированные с использованием состязательных примеров, демонстрируют повышение устойчивости к обходу на 15-30% в сравнении с базовыми моделями [9].

Интерпретируемость и объяснимость решений

Одна из сложнейших задач внедрения ИИ-IDS — объяснить, почему та или иная сессия классифицирована как угроза. Аналитики безопасности должны понимать, почему система сработала, чтобы принять правильное решение. В продвинутых ИИ-IDS применяются механизмы XAI (Explainable AI) [10]:

Методы объяснимости

- Карты значимости признаков (Feature Importance) — визуализация того, какие атрибуты трафика наиболее повлияли на решение. Например, если система классифицировала пакет как DDoS, она может показать, что решающими факторами были "частота пакетов" (95% вклада) и "размер пакета" (4% вклада).

- LIME (Local Interpretable Model-agnostic Explanations) — построение локальных интерпретируемых моделей вокруг конкретного предсказания [11]. LIME работает путём создания возмущений входных данных и наблюдения, как меняется выход модели.

- SHAP (SHapley Additive exPlanations) — теоретико-игровой подход к определению вклада каждого признака в предсказание модели [12]. SHAP обеспечивает более точные объяснения, чем Feature Importance, благодаря использованию теории кооперативных игр.

- Attention Visualization — для LSTM и Transformer моделей визуализация механизмов внимания показывает, на какие части входа модель обращала внимание при принятии решения. Это может помочь выявить, на какие типы событий реагирует система.

Высокая интерпретируемость критична для принятия решений аналитиками безопасности и для соответствия нормативно-правовым требованиям (GDPR требует объяснения любого решения, влияющего на человека; регулярные аудиты и комплаенс).

Метрики оценки производительности ИИ-IDS

При оценке эффективности ИИ-IDS используются различные метрики:

Основные метрики:

- True Positive Rate (TPR) / Чувствительность (Recall) — доля правильно обнаруженных атак из всех реальных атак. $TPR = TP / (TP + FN)$. Идеальное значение: 100%.

- False Positive Rate (FPR) / Ложных срабатываний — доля неправильно классифицированных нормальных пакетов из всех нормальных пакетов. $FPR = FP / (FP + TN)$. Идеальное значение: близко к 0%.

- Точность (Precision) — доля правильно обнаруженных атак из всех обнаруженных как атаки. $Precision = TP / (TP + FP)$. Высокая точность означает, что система редко выдаёт ложные срабатывания.

- F1-Score — гармоническое среднее между Precision и Recall. $F1 = 2 * (Precision * Recall) / (Precision + Recall)$. Используется для комплексной оценки баланса.

- ROC-AUC (Area Under the Receiver Operating Characteristic Curve) — показатель, объединяющий TPR и FPR. ROC-AUC = 0,5 означает случайное угадывание, ROC-AUC = 1,0 означает идеальную классификацию.

Практические требования

На практике корпорации обычно требуют: $TPR \geq 95\%$ (выявление более 95% атак); $FPR \leq 1\%$ (менее 1% ложных срабатываний); $Detection\ latency \leq 100\text{ мс}$ (обнаружение в течение 100 миллисекунд)

Результаты сравнения

Алгоритм	TPR (%)	FPR (%)	Скорость (пакетов/сек)	Требуемая память
SVM	89	3,2	50	256 МБ
Random Forest	94	2,1	120	512 МБ
KNN	92	2,8	80	1 ГБ
CNN	97	1,5	200	2 ГБ
LSTM	96	1,8	150	3 ГБ
CNN+LSTM	98	0,9	100	4 ГБ
Ансамбль (5 моделей)	99	0,5	60	6 ГБ

Сравнение различных архитектур на практике

Различные архитектуры ИИ показывают разную производительность в зависимости от типа атаки и характера сетевого трафика. Из таблицы 1 видно, что ансамблевые подходы обеспечивают лучший баланс между точностью и ложными срабатываниями, но требуют больше вычислительных ресурсов.

Стоимость внедрения и ROI

При расчёте экономической целесообразности внедрения ИИ-IDS необходимо учитывать различные факторы.

Компоненты стоимости

- 1. Оборудование и инфраструктура — GPU-ускорители (50-100 тыс. руб.), серверы хранения данных (500 тыс.-2 млн. руб.), сетевое оборудование. Типичная стоимость: 2-10 млн. рублей.
- 2. Программное обеспечение и лицензии — лицензии на ИИ-платформы (TensorFlow, PyTorch — открытые), платные платформы для IDS (Fortinet, Palo Alto Networks — от 500 тыс. до 5 млн. руб./год).
- 3. Разработка и внедрение — наём специалистов по ИИ (зарплата 200-400 тыс. руб./месяц), консалтинг, адаптация под специфику сети. Типичная стоимость: 500 тыс.-3 млн. руб.
- 4. Обучение персонала — обучение аналитиков безопасности, операционного персонала. Типичная стоимость: 100-300 тыс. руб.
- 5. Текущее обслуживание и поддержка — ежегодное обновление моделей, техническая поддержка, мониторинг. Типичная стоимость: 10-20% от стоимости внедрения в год.

Расчет ROI

Преимущества ИИ-IDS:

- Предотвращение инцидентов — средняя стоимость инцидента информационной безопасности: 2-20 млн. рублей (в зависимости от типа). Даже предотвращение одного серьёзного инцидента окупает годовые расходы на систему.
 - Снижение затрат на персонал — автоматизация снижает количество необходимых аналитиков с 5-10 до 2-3 человек. Экономия: 1,5-3 млн. руб./год.
 - Сокращение времени отклика — с нескольких часов до минут. Это снижает ущерб от инцидентов на 30-50%.
- Типичный ROI составляет 200-400% в течение первых 3 лет внедрения.

Тестирование и валидация ИИ-IDS

- Перед развёртыванием системы в продакшене необходимо провести комплексное тестирование.
- Фазы тестирования:
- 1. Lab Testing — тестирование на изолированной инфраструктуре с использованием известных датасетов (NSL-KDD, CIC-IDS2017). Целевые метрики: TPR ≥ 95%, FPR ≤ 1%.
 - 2. Staging Testing — развёртывание на копии продакшн-сети. Длительность: 2-4 недели мониторинга без принятия решений (Offline Detection Mode). Целевые метрики: те же, но с реальным трафиком организации.
 - 3. Pilot Deployment — развёртывание на 10-20% сегмента сети в режиме alert-only. Аналитики получают оповещения, но система не блокирует трафик.

4. Full Production Deployment — полное развёртывание с автоматическим блокированием подозрительного трафика.

Типичный период от Lab Testing до Full Deployment: 3-6 месяцев.

Обновление и поддержка системы

ИИ-модели деградируют со временем, так как сетевое поведение эволюционирует, появляются новые типы атак. Это явление называется Data Drift или Concept Drift.

Процесс переобучения:

1. Ежемесячное переобучение — система переучивается на свежих данных, собранных за месяц. Это позволяет учесть изменения в сетевом поведении и новые типы атак.

2. Валидация после переобучения — проверка метрик на холдаут тестовом наборе. Если метрики упали > 5%, переобучение откатывается.

3. A/B тестирование — параллельный запуск старой и новой версий моделей на 5% трафика. Сравнение метрик перед полным развёртыванием.

4. Вероятностное переобучение — использование байесовских методов для оценки неопределённости модели. Если неопределённость растёт, это сигнал к переобучению.

Интеграция с другими системами безопасности

ИИ-IDS не функционирует изолированно, а является частью экосистемы безопасности.

- SIEM (Security Information and Event Management) — корреляция событий от IDS с другими источниками (firewalls, antivirus, logs). Пример: Splunk, IBM QRadar, ArcSight.

- EDR (Endpoint Detection and Response) — интеграция с системами защиты конечных точек. IDS предоставляет контекст о сетевых атаках, EDR анализирует поведение на конечных точках.

- Threat Intelligence Platform — обогащение данных ИИ-IDS информацией о известных угрозах, ботнетах, C2 серверах.

- Firewall и WAF (Web Application Firewall) — автоматическое формирование правил на основе выявленных атак. Пример: если IDS выявил DDoS с определённого IP, firewall блокирует его.

- Ticketing System — автоматическое создание tickets в системе управления инцидентами (Jira, ServiceNow) при выявлении серьёзных атак.

На практике используются API и SIEM connectors для автоматизации этого обмена.

Заключение

ИИ-системы обнаружения вторжений позволяют перейти к проактивной защите сетевых инфраструктур, выявляя как известные, так и ранее неизвестные угрозы. Однако успешное внедрение требует комплексного подхода, включающего: оптимизацию производительности для обработки больших объёмов трафика в реальном времени, защиту от состязательных атак на саму ИИ-систему, обеспечение интерпретируемости и объяснимости решений, регулярное переобучение моделей, интеграцию с существующей инфраструктурой и постоянный мониторинг метрик производительности.

Будущее ИИ-IDS лежит в развитии гибридных моделей, которые объединяют быстрые классические алгоритмы с мощными глубокими нейросетями, в применении механизмов XAI для повышения доверия аналитиков, в разработке методов защиты от состязательных атак и в автоматизации процесса переобучения. Постоянное совершенствование алгоритмов, инвестиции в инфраструктуру и обучение персонала остаются критическими факторами успеха при внедрении таких систем в любой организации.

Литература

1. Ляо, Х. Джей, Линь, К. Х. Р., Линь, И. К., Тунг, К. Й. Система обнаружения вторжений: комплексный обзор // Журнал сетевых и компьютерных приложений. – 2013. – Т. 36, № 1. – С. 16–24. – DOI 10.1016/j.jnca.2012.09.003.
2. Шарафалдин, И., Лашкари, А. Х., Гхорбани, А. А. Создание нового набора данных для обнаружения вторжений и характеристика трафика вторжений // Материалы Международной конференции по информационным системам, безопасности и конфиденциальности (ICISSP). – 2018. – С. 108–116. – DOI 10.5220/0006639001080116.
3. Каусар, А., Жайолд, М., Джиурка, Д. К автоматизированному крупномасштабному обнаружению появляющихся аномалий, независимому от домена // Материалы Международной конференции по искусственному интеллекту и статистике (AISTATS). – 2018. – С. 800–808.
4. Гудфеллоу, И., Бенджио, Й., Курвиль, А. Глубокое обучение / пер. с англ. – М.: ООО "И.Д. Вильямс", 2016. – 800 с.
5. Ким, Дж., Ли, С., Ким, С. Подход глубокого обучения к системе обнаружения вторжений в сети // Материалы Международной конференции по большим данным и смарт-вычислениям (BigComp). – 2016. – С. 313–316. – DOI 10.1109/BIGCOMP.2016.7425916.
6. Чжоу, И., Цзян, Г., Сюй, К. Обнаружение состязательных примеров почти так же сложно, как их классификация // Материалы IEEE Европейского симпозиума по безопасности и конфиденциальности (EuroS&P). – 2019. – С. 250–265. – DOI 10.1109/EuroSP.2019.00023.
7. Гудфеллоу, И. Дж., Шленс, Дж., Сегеди, К. Объяснение и использование состязательных примеров // arXiv препринт. – 2014. – arXiv:1412.6572.
8. Чжоу, З. Х. Методы ансамблевого обучения: основы и алгоритмы / пер. с англ. – М.: CRC Press, 2012. – 290 с.
9. Паперно, Н., МакДэниел, П., Гудфеллоу, И., Джа, С., Челик, З. Б., Свами, А. Практические чёрные атаки на машинное обучение // Материалы 2016 ACM SIGSAC конференции по компьютерной и коммуникационной безопасности. – 2016. – С. 1317–1328. – DOI 10.1145/2976749.2978339.
10. Рибейро, М. Т., Сингх, С., Гвестрин, К. "Почему я должен вам доверять?": Объяснение предсказаний любого классификатора // Материалы 22-й Международной конференции ACM SIGKDD по обнаружению знаний и интеллектуальному анализу данных. – 2016. – С. 1135–1144. – DOI 10.1145/2939672.2939778.
11. Лундберг, С. М., Ли, С. И. Унифицированный подход к интерпретации предсказаний моделей // Advances in Neural Information Processing Systems. – 2017. – Т. 30. – С. 4765–4774.
12. Вавани, А., Шазеер, Н., Пармар, Н., Ушкорейт, Дж., Джонс, Л., Гомес, А. Н., Кайзер, Л., Полозункин, И. Вниманию — это всё, что вам нужно // Advances in Neural Information Processing Systems. – 2017. – Т. 30. – С. 5998–6008.

References

1. Lyao, KH. Dzhey, Lin', K. KH. R., Lin', I. K., Tung, K. Y. Sistema obnaruzeniya vtorzheniy: kompleksnyy obzor // Zhurnal setevykh i komp'yuternykh prilozheniy. – 2013. – Т. 36, № 1. – С. 16–24. – DOI 10.1016/j.jnca.2012.09.003.
2. Sharafaldin, I., Lashkari, A. KH., Gkhorbani, A. A. Sozdaniye novogo nabora dannykh dlya obnaruzeniya vtorzheniy i kharakteristika trafika vtorzheniy // Materialy Mezhdunarodnoy konferentsii po informatsionnym sistemam, bezopasnosti i konfidentsial'nosti (ICISSP). – 2018. – S. 108–116. – DOI 10.5220/0006639001080116.
3. Kausar, A., Zhayold, M., Dzhiurka, D. K avtomatizirovannomu krupnomasshtabnomu obnaruzeniyu poayavlyayushchikhsya anomalii, nezavisimomu ot domena // Materialy Mezhdunarodnoy konferentsii po iskusstvennomu intellektu i statistike (AISTATS). – 2018. – S. 800–808.
4. Gudfellow, I., Bendzhio, Y., Kurvil', A. Glubokoye obucheniye / per. s angl. – M.: ООО "I.D. Vil'yams", 2016. – 800 s.
5. Kim, Dzh., Li, S., Kim, S. Podkhod glubokogo obucheniya k sisteme obnaruzeniya vtorzheniy v seti // Materialy Mezhdunarodnoy konferentsii po bol'shim dannym i smart-vychisleniyam (BigComp). – 2016. – S. 313–316. – DOI 10.1109/BIGCOMP.2016.7425916.
6. Chzhou, I., TSzyan, G., Syuy, K. Obnaruzeniye sostyazatel'nykh primerov pochtii tak zhe slozhno, kak ikh klassifikatsiya // Materialy IEEE Yevropeyskogo simpoziuma po bezopasnosti i konfidentsial'nosti (EuroS&P). – 2019. – S. 250–265. – DOI 10.1109/EuroSP.2019.00023.
7. Gudfellow, I. Dzh., Shlens, Dzh., Segedi, K. Ob'yasneniye i ispol'zovaniye sostyazatel'nykh primerov // arXiv preprint. – 2014. – arXiv:1412.6572.

8. Chzhou, Z. KH. Metody ansamblevogo obucheniya: osnovy i algoritmy / per. s angl. – M.: CRC Press, 2012. – 290 s.
9. Paperno, N., MakDeniyel, P., Gudfellou, I., Dzha, S., Chelik, Z. B., Svami, A. Prakticheskiye chornyye ataki na mashinnoye obucheniye // Materialy 2016 ACM SIGSAC konferentsii po komp'yuternoy i kommunikatsionnoy bezopasnosti. – 2016. – S. 1317–1328. – DOI 10.1145/2976749.2978339.
10. Ribeyro, M. T., Singkh, S., Gvestrin, K. "Pochemu ya dolzhen vam doveryat'?: Ob'yasneniye predskazaniy lyubogo klassifikatora // Materialy 22-y Mezhdunarodnoy konferentsii ACM SIGKDD po obnaruzheniyu znaniy i intellektual'nomu analizu dannykh. – 2016. – S. 1135–1144. – DOI 10.1145/2939672.2939778.
11. Lundberg, S. M., Li, S. I. Unifitsirovanny podkhod k interpretatsii predskazaniy modeley // Advances in Neural Information Processing Systems. – 2017. – T. 30. – S. 4765–4774.
12. Vasvani, A., Shazeyer, N., Parmar, N., Ushkoreyt, Dzh., Dzhons, L., Gomes, A. N., Kayzer, L., Polozukhin, I. Vnimaniye — eto vso, chto vam nuzhno // Advances in Neural Information Processing Systems. – 2017. – T. 30. – S. 5998–6008.
-

КУШНИР Надежда Владимировна, старший преподаватель кафедры информационных систем и программирования ФГБОУ ВО «Кубанский государственный технологический университет». 350072, г. Краснодар, ул. Московская, д. 2. E-mail: kushnir.06@mail.ru

ВЛАСЕНКО Александра Владимировна, кандидат технических наук, доцент, Врио заведующего кафедрой информационной безопасности, профессор кафедры ФГКОУ ВО «Краснодарский университет МВД России», 350005, г. Краснодар, ул. Ярославская 128. E-mail: alex_vlasenko@list.ru

ТЮТЮНИК Николай Павлович, студент 3 курса направления подготовки 09.03.04 (программная инженерия) факультета информационных технологий и кибербезопасности (ФИТК) ФГБОУ ВО «Кубанский государственный технологический университет». 350072, г. Краснодар, ул. Московская, д. 2. E-mail: nikolat1236@gmail.com

СМИРНОВ Богдан Александрович, студент 3 курса направления подготовки 09.03.04 (программная инженерия) факультета информационных технологий и кибербезопасности (ФИТК) ФГБОУ ВО «Кубанский государственный технологический университет». 350072, г. Краснодар, ул. Московская, д. 2. E-mail: solo.logger@gmail.com

KUSHNIR Nadezhda Vladimirovna, Senior Lecturer at the Department of Information Systems and Programming of the Federal State Budgetary Educational Institution of Higher Education «Kuban State Technological University». 350072, Krasnodar, Moskovskaya St., 2. E-mail: kushnir.06@mail.ru

VLASENKO Alexandra Vladimirovna, Candidate of Technical Sciences, Associate Professor, Acting Head of the Department of Information Security Professor of the Department, of the Krasnodar University of the Ministry of Internal Affairs of Russia. 350005, Krasnodar, Yaroslavskaya str., 128. E-mail: alex_vlasenko@list.ru

TYUTYUNIK Nikolay Pavlovich, 3rd year student of the 09.03.04 (Software Engineering) Department of Information Technology and Cybersecurity (FITK) of the Federal State Budgetary Educational Institution of Higher Education «Kuban State Technological University». 350072, Krasnodar, Moskovskaya St., 2. E-mail: nikolat1236@gmail.com

SMIRNOV Bogdan Alexandrovich, 3rd year student of the 09.03.04 training course (software Engineering), Faculty of Information Technology and Cybersecurity (FITC), Federal State Budgetary Educational Institution of Higher Education «Kuban State Technological University». 350072, Krasnodar, Moskovskaya St., 2. E-mail: solo.logger@gmail.com

ИССЛЕДОВАНИЕ МЕТОДОВ ЭЛЕКТОРАЛЬНОЙ КРИМИНАЛИСТИКИ

Анализ современной истории проведения выборов показывает, что в целом ряде случаев, несогласие с результатами проведенных выборов становилось поводом к проведению массовых публичных протестов (например, в Армении (2008 г.), Белоруссии (2020 г.), Грузии и Киргизия (2018 г.), Российской Федерации (РФ) (2011–2012 гг.), Украине (2014 г.) и др.). При этом в большинстве случаев основной причиной несогласия с их результатами проигравшие участники, а также их сторонники (политики и политологи), они называют многочисленные, по их мнению, нарушения свободного волеизъявления избирателей совершенные официальными лицами, ответственными за проведение выборов, в том числе, преднамеренно внесенные аномалии в электоральные данные (ЭД) на этапе подсчета голосов избирателей (фальсификация электоральных данных).

В связи с тем, что при фальсификации ЭД оказываются измененными как, собственно, их статистические характеристики, так и вычисляемых на их основе электоральные показатели (ЭП), была выдвинута гипотеза о том, что, потенциально, преднамеренно внесенные в результаты голосования изменения можно обнаружить с помощью соответствующих математических методов, совокупность которых получила название методов электоральной криминалистики (ЭК), в том числе:

- методов ЭК, основанных на анализе плотностей распределений (ПР) и функций распределений (ФР) цифр и пар одинаковых цифр в значениях относительной явки избирателей на выборы и доли избирателей, проголосовавших за данного кандидата;

- метод, предложенный А. Собяниным и В. Суховольским;

- метод С. Шпилькина¹, основанный на постулате, что число голосов избирателей, полученных данным участником выборов, не зависит от отношения числа проголосовавших избирателей к числу зарегистрированных избирателей (явки избирателей на выборы);

- метод, предложенный А. Подлазовым, в котором для каждого участника выборов проводится анализ зависимости «число голосов избирателей, проголосовавших за победителя выборов, от числа голосов проголосовавших за остальных кандидатов», принимавших участие в выборах, отличие которой от линейной трактуется как индикатор преднамеренно внесенных аномалий в ЭД.

Однако строгих математических доказательств адекватности методов ЭК, требующих использования ЭД, собранных в ходе проведения выборов, признанных всеми их участниками абсолютно честными, справедливости данной гипотезы на сегодняшний день не представлено. В этой связи было проведено независимое исследование данных методов.

Статья состоит из двух частей: в первой части статьи описаны наиболее популярные в настоящее время методы ЭК, а также проведен критический анализ мето-

¹ С. А. Шпилькин 10.02.2023 включен Министерством юстиции Российской Федерации в реестр иностранных агентов в соответствии со статьей 7 Федерального закона от 14.07.2022 № 255-ФЗ «О контроле за деятельностью лиц, находящихся под иностранным влиянием»

да, предложенного А. Собяниным и В. Суховольским. Результаты применения методов ЭК для анализа синтезированных, заведомо нефальсифицированных ЭД и оценке адекватности полученных с их помощью результатов обсуждаются во второй части статьи, публикуемой далее.

Ключевые слова: электоральная криминалистика, электоральные данные, электоральные показатели, фальсификация электоральных данных, избирательная комиссия, выборы, случайная величина с ограниченной областью рассеяния, случайные блуждания, плотность распределения, функция распределения

Porshnev S. V., Ryabko N. Yu.

THE ELECTORAL FORENSICS METHODS RESEARCH

The modern history of elections to government bodies of various levels and local self-government, held both in the Russian Federation and abroad analysis shows that in a number of cases, disagreement with the elections results became the reason for mass public protests (for example, in Armenia (2008), Belarus (2020 Georgia and Kyrgyzstan (2018), the Russian Federation (2011–2012), Ukraine (2014), etc.). However, in most cases, the main reason for disagreement with the election results was the participants who lost the elections, as well as their supporters (politicians and political scientists), who claimed to have committed numerous violations of the free will of voters by officials responsible for the conduct of the elections. Traditionally, the lost the election participants consider intentionally introduced anomalies in the electoral data at the stage of vote counting (falsification of electoral data) to be the main reason for their defeat. Due to the fact that when electoral data is falsified, certain statistical characteristics of electoral data, as well as electoral indicators calculated on their basis, are changed, it was hypothesized that, potentially, they can be detected using appropriate mathematical methods called electoral forensics (EF) methods. However, no rigorous mathematical proofs of this hypothesis validity have been presented to date.

The conducted research used synthesized deliberately unfalsified electoral data. Well-known EF methods (EF methods that based on analysis the functions distributions and the densities distributions of digits and pairs of identical digits in the values of relative voter turnout and the proportion of voters who voted for a given candidate; the method proposed by A. Sobyenin and V. Sukhovolsky; S. Shpilkin's² method that based on the postulate that in the fraud absence, the given election participant result does not depend on voter turnout; the proposed A. Podlazov method in which analyzes the dependences of the votes cast for the winner and of the dependence of the election number the votes cast for other election participants on the turnout, which, according to the method author, in the absence of fraud will be linear, in the fraud presence – different from linear) were applied to the synthesized electoral data and then the adequacy of the results was assessed.

To prove the EF methods adequacy, electoral data collected during "absolutely fair" elections is needed. However, such data does not exist today. In this regard, an independent study of these methods was conducted, which used the analytical method and synthesized data.

² S. A. Shpilkin was included in the Register of Foreign Agents by the Ministry of Justice of the Russian Federation on 10.02.2023 in accordance with Article 7 of Federal Law No. 255-FZ dated 14.07.2022 «On Control over the Activities of Persons under Foreign Influence».

The article consists of two parts: the first part of the article describes the currently most popular methods of electoral criminology, as well as a critical analysis of the A. Sobyenin and V. Sukhovolsky method. The results of applying the EF methods to analyze synthesized, deliberately unfalsified electoral data and assess the adequacy of the results obtained with their help are discussed in the second part of the article published below.

Keywords: electoral forensics, electoral data, electoral indicators, falsification of electoral data, election commission, elections, random variable with a limited scattering area, random walks, distribution density, distribution function

ВВЕДЕНИЕ

На современном этапе развития общественных отношений выборы в органы государственной власти и местного самоуправления, призванные обеспечить легитимное делегирование избирателями властных полномочий их представителям, является одним из основополагающих условий существования и развития демократического правового государства.

В демократических обществах выборы считаются способом мирного и легального урегулирования противоречий и конфликтов между участниками политического процесса. Однако анализ опыта проведения избирательных кампаний в различных странах (США, Молдова, Румыния и др.), показывает, что сами выборы и их результаты во многих случаях оказывались причиной конфликтов, напротив, активизируя противоборство политических сил. Это подтверждает, в том числе, история проведения выборов в РФ и странах постсоветского пространства в конце XX – начале XXI вв., в которых проигравшие кандидаты/партии и их сторонники ставили под сомнение честность проведения выборов, фактически, любого уровня, а также оспаривали результаты голосования, обвиняя сотрудников государственных органов и членов Центральной избирательной комиссии (ЦИК) РФ, избирательных комиссий субъектов РФ и территорий, находящихся за пределами РФ, (ИК), территориальных избирательных комиссий (ТИК), а также участковых избирательных комиссий (УИК), в фальсификации выборов. В данном случае под термином «фальсификация» понимается преднамеренно внесенные искажения в ЭД. Одним из наиболее распространенных видов фальсификаций, состоит, по мнению ЭК, в увеличении на ΔN^V официально публикуемого числа проголосовавших избирателей $N_{official}^V = N_{real}^V + \Delta N_{real}^V$, N_{real}^V – число избирателей, действительно, принявших участие в выборах, которые, в первую очередь, передаются заранее определенному победителю выборов ($[N_{official}^V]^1 = [N_{real}^V]^1 + \Delta N^V$).

При этом наличие законодательства, регулирующего избирательный процесс, а также комплекса мер, реализуемых организаторами выборов для обеспечения их прозрачности и легитимности, по мнению экспертов, поддерживающих проигравших кандидатов/партии, в отнюдь не гарантируют отсутствия масштабных фальсификаций ЭД в пользу победителя выборов, выявить которые, по их мнению, можно с помощью соответствующих математических методов, которые получили название «методы электоральной криминалистики» (методы ЭК).

По классификации, предложенной Н. Шалаевым [1], методы ЭК могут быть объединены в следующие группы:

1) методы, основанные на анализе функций распределений (ФР) и плотностей распределений (ПР) цифр, находящихся в выбранных разрядах электоральных показателей (например,

$$[\alpha_{Level_k}^{\%}]_j = \frac{[N_{Level_k}^V]_j}{[N_{Level_k}^{ER}]_j} \cdot 100\%, \quad k = \overline{1, K}, K -$$

количество уровней избирательной системы³, $[N_{Level_k}^{ER}]_j$ – число избирателей, зарегистрированных в j -ой избирательной комиссии $Level_k$ -го; $[N_{Level_k}^V]_j$ – число избирателей,

принявших участие в выборах в j -ой избирательной комиссии $Level_k$ -го уровня;

$$[\gamma_{Level_k}^{\%}]_j^{(i)} = \frac{[N_{Level_k}^V]_j^{(i)}}{[N_{Level_k}^V]_j} \cdot 100\%, \quad [N_{Level_k}^V]_j^{(i)} -$$

число избирателей, проголосовавших в j -ой избирательной комиссии $Level_k$ -го за i -го кандидата);

³ Выделение уровней избирательных комиссий, участвующих в проведении выборов в органы государственной власти и местного самоуправления, обусловлено иерархической структурой организации систем избирательных комиссий, как в РФ, так и других странах, например, в РФ $K = 4$.

2) методы, основанные на анализе функций, описывающих зависимости между выбранными электоральными показателями

(например, функции $\left[\gamma_{Level_k}^{\%} \right]_j^{(i)} = f\left(\left[\alpha_{Level_k}^{\%} \right]_j \right)$).

К первой группе № 1 методов относятся.

1) Метод проверки соответствия ПР выборов, составленных из первых цифр ЭД, закону Бенфорда [2].

2) Метод анализа ПР выборов, составленных из цифр, находящихся во втором после десятичной запятой значащем разряде анализируемого набора чисел [3–6].

3) Методы, основанные на аксиоме, утверждающей, что при отсутствии фальсификаций ЭД появление цифр 0,1,...,9 в последнем разряде числа равновероятно (эквивалентно, случайная последовательность, составленная из цифр, находящихся в последнем разряде ЭД, имеет равномерную ПР), а при наличии преднамеренно внесенных фальсификаций в ЭД ПР отличается от равномерного закона (поскольку при внесении придуманных результатов голосования в ито-

говых протоколах, по мнению авторов данных методов, преобладают «психологически привлекательные» цифры (для целых чисел в последнем разряде – цифры 0 и 5, для процентных чисел – значения, у которых присутствует только одна цифра после запятой)) [7].

4) Метод, основанный на изучении встречаемости в наборах процентных чисел парных цифр в младших разрядах (например, 11, 22 и т.д.), которая (как постулировано, но не доказано), при отсутствии фальсификаций должна быть одинаковой для всех пар цифр [8].

5) Метод, основанный на законе Стиглера, описывающем распределение цифр в первом разряде набора чисел [9].

Данная группа методов была использована экспертами при анализе выборов в странах Европы, Северной, Южной Америки (см., например, [10–12]). Также отметим, что различные модификации методов, относящихся к данной группе, использовались У. Мебейном У. и К. Калининым [13], а также А.В. Подлазовым для анализа результатов российских выборов (см., например, [14,15]).

В Российской Федерации также получила распространение вторая группа методов, к которой относятся.

1) Метод, предложенный А. Собяниным и В. Суховольским [16], суть которого состоит в проверке гипотезы о наличии (отсутствии) детерминированной связи между ЭП

$$\left[\alpha_{Level} \right]_j^{(i)} = \left[N_{Level}^V \right]_j^{(i)} / \left[N_{Level_k}^V \right]_j, \left[\alpha_{Level_k}^{\%} \right]_j^{(i)} = \left[\alpha_{Level_k} \right]_j^{(i)} \cdot 100\%,$$

где $Level_k$ – уровень избирательной комиссии. (Наличие детерминированной связи между этими ЭП, по мнению авторов, метода является маркером фальсификации ЭД, отсутствие детерминированной связи, наоборот, индикатором, свидетельствующем об отсутствии преднамеренно внесенных изменений в ЭД).

2) Метод С. Шпилькина, основанный на постулате о том, что при отсутствии фальсификаций ЭД число избирателей, проголосовавших за i -го кандидата $\left[N_{Level_k}^V \right]_j^{(i)}$ в j -ой избирательной

комиссии $Level_k$ -го уровня, не должно зависеть от $\left[\alpha_{Level_k} \right]_j = \frac{\left[N_{Level_k}^V \right]_j}{\left[N_{Level_k}^{ER} \right]_j}$ явки избирателей и,

более того, ПР случайных последовательностей, составленных из числа голосов избирателей проголосовавших за различных кандидатов в избирательных комиссиях выбранного уровня, с точностью до масштабирующего коэффициента должны быть подобными друг другу, и наоборот, при наличии фальсификаций обсуждаемое подобие ПР должно отсутствовать в ЭД, представленных избирательными комиссиями выбранного уровня, на которых явкой избирателей на выборы была высокой. (см., например, [17–21]).

3) Метод А. Подлазова [22], основанный на постулате о том, что при отсутствии преднамеренных фальсификаций ЭД, представленных избирательными комиссиями $Level_k$ - го уровня, зависимость:

$$\left[N_{Level_k}^V \right]^{(i)} = \sum_{j=1}^{n_{Level}} \left[N_{Level_k}^V \right]_j^{(i)} = f \left(\sum_{i=2}^m \left[N_{Level_k}^V \right]^{(i)} \right) = f \left(\sum_{j=1}^{n_{Level}} \sum_{i=2}^m \left[N_{Level_k}^V \right]_j^{(i)} \right),$$

где n_{Level_k} – число избирательных комиссий уровня $Level$, является линейной:

$$\begin{aligned} \left[N_{Level_k}^V \right]^{(i)} &= \sum_{j=1}^{n_{Level}} \left[N_{Level_k}^V \right]_j^{(i)} = a_{Level_k}^{(i)} \left(\sum_{i=2}^m \left[N_{Level_k}^V \right]^{(i)} \right) + b_{Level_k}^{(i)} = \\ &= a_{Level_k}^{(i)} \left(\sum_{j=1}^{n_{Level}} \sum_{i=2}^m \left[N_{Level_k}^V \right]_j^{(i)} \right) + b_{Level_k}^{(i)}, \end{aligned}$$

при наличии целенаправленно внесенных изменений, напротив, нелинейной.

Следует отметить, что в обеих группах методов выявление фальсификаций основано на сравнении результатов анализа реальных ЭД с соответствующими оценками статистических характеристик ЭД, полученных в ходе «идеальных» выборов, которые постулированы, но не доказаны их авторам. Данная ситуация обусловлена тем, для оценки адекватности представлений, на которых основаны методы ЭК, необходимо использовать и ЭД, опубликованные после проведения «абсолютно честных» выборов, однако, при этом сначала требуется доказать, что выборы, действительно, были проведены «абсолютно честно».

В этой связи оказывается актуальным анализ математического базиса обсуждаемых методов ЭК, а также результатов их применения для обработки массивов синтезированных, а потому заведомо несфальсифицированных, ЭД.

Статья состоит из пяти разделов и заключения. В первом разделе рассматриваются и анализируются наиболее распространенные методы ЭК. Во втором разделе проводится критический анализ метода С-С. В третьем разделе представлен алгоритм генерации модельных ЭД, подобных реальным ЭД, и результаты анализа их статистических свойств. В четвертом разделе описана математическая модель плотности распределения (ПР) и функции распределения (ФР) ЭД и ЭП. В пятом разделе обсуждаются результаты анализа для модельных ЭД с помощью методов ЭК.

1. МЕТОДЫ ЭЛЕКТОРАЛЬНОЙ КРИМИНАЛИСТИКИ

1.1. МЕТОД ПРОВЕРКИ СООТВЕТСТВИЯ ПЛОТНОСТЕЙ РАСПРЕДЕЛЕНИЙ ВЫБОРОК, СОСТАВЛЕННЫХ ИЗ ПЕРВЫХ ЦИФР ЭЛЕКТОРАЛЬНЫХ ДАННЫХ, ЗАКОНУ БЕНФОРДА

Впервые опыт применения обсуждаемого далее метода с целью выявления возможных фальсификаций ЭД был описан в [2]. Напомним, что в соответствии с законом Бенфорда вероятность появления цифры d в первом разряде (для краткости d_1) наборов значений измеренных количественных характеристик некоторых реальных процессов $p(d_1)$ вычисляется по известной формуле [2]:

$$p(d_1) = \log_N (1 + 1/d_1)$$

где N – основание системы счисления (например, последовательность степеней натурального числа 2, экспоненциальные последовательности, факториалы и т.п.).

Отметим, что справедливость закона Бенфорда постулируется, но не доказывается для закона распределения второй значащей цифры выбранного электорального показателя. В этой связи, с научной точки зрения закон Бенфорда является не универсальной, но некоторой феноменологической закономерностью (правилом), выполняющейся для некоторого конечного числа наборов количе-

ственных показателей. Однако на сегодняшний день правомерность отнесения ЭД к наборам данных, распределение цифр в которых описывается законом Бенфорда, остается недоказанной. Более того, известны публикации, авторы которых утверждают, что ПР случайных выборок, составленных из первых и вторых цифр ЭД, опубликованных по результатам абсолютно честных выборов, отнюдь, не обязаны соответствовать закону Бенфорда [4]. Данный вывод, косвенно, подтверждается результатами, опубликованными в [5,6]. Здесь авторы применили обсуждаемый метод для анализа ЭД, которые по косвенным признакам были разделены на два класса: фальсифицированные (класс № 1) и нефальсифицированные (класс № 2) ЭД. Результаты анализа показали, что данный метод, в ряде случаев, не смог обнаружить признака фальсификации в данных класса № 1 (ошибка первого рода), а также обнаружил фальсификации в электоральных данных, отнесенных к классу № 2 (ошибка второго рода).

Также известны методы выявления электоральных аномалий, основанные на анализе распределений последних цифр в значениях выбранных показателей электорального поведения, округленных с точностью до заданного числа цифр после запятой, в основу которых не строго доказанная лемма и/или теорема, но аксиома, в соответствии с которой, априори, принимается, что при отсутствии фальсификаций появление цифр 0,1,..., 9 в последнем разряде целочисленного ЭП равновероятно, а при наличии фальсификаций закон распределения данных цифр будет отличаться от равномерного закона. Логика подобных методов основывается на априорном предположении, что при внесении придуманных результатов голосования в протоколы подсчета голосов преобладают психологически привлекательные, для вносящего изменения в ЭД, цифры «0» и «5» [7], которые и указываются в целочисленных ЭП, а также в первом после десятичной запятой разряде ЭП, измеряемых в процентах. Аналогичный подход использовался в [8] для анализа встречаемости парных цифр в младших разрядах в целочисленных ЭП (метод Бебера и Скакко).

К данной группе методов также относятся методы, основанные на проверке гипотезы о соответствии ФР случайных последовательностей, составленных из цифр, находящихся в выбранных разрядах в ЭД, закону

Стиглера [9], в соответствие с которым вероятность обнаружить одну из цифр из диапазона в первом разряде ЭД P_1^d вычисляется по формуле:

$$P_1^d = \frac{d \ln(d) - (d+1) \ln(d+1) + \left(1 + \frac{10}{9} \ln(10)\right)}{9}.$$

Модификация метода, основанного на использовании закона Стиглера, была предложена Л. Лееманом и Д. Бохлером в [23], заметивших, что в проанализированных ими наборах целых чисел с одинаковым числом разрядов частота встречаемости данной цифры оказалась зависящей от числа разрядов (длины числа). В этой связи они предложили вычислять вероятность появления данной цифры в наборе чисел, имеющих разной длину, с учетом длины каждого из чисел набора. Модифицированный метод вычисления вероятности появления данной цифры в последнем разряде чисел используемого набора чисел $\{N_i\}, i=1, I$ реализуется выполнением следующей последовательности действий.

1. Определение длины каждого из чисел используемого набора чисел $\{L_i\} = \{\text{len}(N_i)\}, i=1, I$, $\text{len}(\)$ – функция, возвращающая длину числа.

2. Извлечение из набора длин чисел $\{L_i\}$ набора, содержащего уникальные (без дубликатов) значения длин чисел, $\{L_k\}, k=1, K$.

3. Подсчет вероятности встречаемости цифры d в последнем разряде чисел длиной $\{L_k\}, k=1, K$ P_k .

4. Вычисление вероятности появления в последнем разряде цифры d по формуле:

$$P_d = \frac{1}{I} \sum_{k=1}^K L_k P_k.$$

Л. Перикки и Д. Торрес заметили, что эмпирические ФР случайных последовательностей, составленных из цифр, находящихся в первом разряде целочисленных ЭД, отличаются от ФР, предсказываемых законом Бенфорда [24]. Обнаруженное отличие они объяснили тем, что ЭД представляют собой целочисленные случайные последовательности, области значений членов которых являются ограниченными (например, число избирателей, зарегистрированных в j -ой избирательной комиссии уровня $Level$

$$\left[N_{Level}^{ER} \right]_j \in \left[\min \left(\left[N_{Level}^{ER} \right]_j \right), \max \left(\left[N_{Level}^{ER} \right]_j \right) \right];$$

число избирателей, принявших участие в вы-

борах в j -ой избирательной комиссии уровня Level –

$$\left[N_{Level}^V \right]_j \in \left[\min \left(\left[N_{Level}^V \right]_j \right), \max \left(\left[N_{Level}^V \right]_j \right) \right]$$

и т.д.). Для учета обнаруженных отличий от закона Бенфорда они использовать для расчета вероятности появления цифры d в i -ом разряде (далее для краткости d_i) целого числа $N \leq K$ $P_i(d|N \leq K)$ следующую формулу:

$$P_i(d|N \leq K) = \frac{P_i^B(d) \# \{d_i \leq K\}}{\sum_{\forall d} [P_i^B(d) \# \{d_i \leq K\}]},$$

где $P_i^B(d)$ – вероятность обнаружения в соответствие с законом Бенфорда в i -ом разряде числа N цифры d ; $\# \{d_i \leq K\}$ – число чисел $d_i \leq K$, а сумма $\sum_{\forall d_i} [P_i^B(d) \# \{d_i \leq K\}]$ вычисляется для всех d_i .

Отметим, что в предложенном Л. Перикки и Д. Торресом способе вычисления $P_i(d|N \leq K)$, принято во внимание только ограничение значений целочисленных ЭД справа. Для учета ограничений области значений ЭД с обеих сторон в [1], используя подход Л. Леemannоса и Д. Бохлером [15], были предложены и апробированы следующие модификации метода Перикки-Торреса.

Модификация № 1, реализующаяся выполнением следующей последовательности действий.

1) Кластеризация избирательных участков на уровне $Level_k$ по степени близости значений показателей $\left[N_{Level_k}^{ER} \right]_j$, $\left[N_{Level_k}^V \right]_j$ – формирование кластеров

$$\left[Name_{Level_k} \right]_m = \left\{ \left[Name_{Level_k} \right]_{1,m}, \left[Name_{Level_k} \right]_m, \dots, \left[Name_{Level_k} \right]_m \right\},$$

$m = \overline{1, M}$, M – число выделенных кластеров, $\left[Name_{Level_k} \right]_m$ – множество, составленное из названий избирательных комиссий $Level_k$ -го уровня, отнесенных к m -му кластеру.

2) Вычисление в соответствие с методом Перикки-Торреса $P_{\left[Name_{Level_k} \right]_m}(d_i)$ для ЭД, представленных избирательными комиссиями $Level_k$ -го уровня, отнесенных к m -ому кластеру.

3) Вычисление $P(d_i)$, как взвешенной суммы $P_{\left[Name_{Level_k} \right]_m}(d_i)$:

$$P(d_i) = \frac{1}{N} \sum_{m=1}^M P_{\left[Name_{Level_k} \right]_m}(d_i) \left[\left[Name_{Level_k} \right]_m \right],$$

где $\left[\left[Name_{Level_k} \right]_m \right]$ – число избирательных комиссий выбранного уровня, отнесенных к m -му кластеру.

Модификация № 2, основанная на аналитическом вычислении вероятности для каждого встречающегося значения списочного числа избирателей. При этом в обоих способах установлен порог отсеечения избирательных участков с малым числом избирателей. В результате используются только те группы (кластеры), которые кластеры объясняют не менее 90% дисперсии значений электоральных показателей. Выбор данного значения критерия для отбора избирательных участков, используемых в дальнейшем анализе, обусловлен тем, что начиная с данного уровня прирост объяснённой дисперсии резко замедляется. При этом порог в 95% объяснённой дисперсии достигается только существенным (~ удвоением относительно 90% уровня) увеличением числа кластеров.

1.2. МЕТОДЫ, ОСНОВАННЫЕ НА АНАЛИЗЕ СВЯЗЕЙ МЕЖДУ КОЛИЧЕСТВЕННЫМИ ПОКАЗАТЕЛЯМИ ЭЛЕКТОРАЛЬНЫХ ПРОЦЕССОВ

Начало разработки данной группы методов было положено А. Собяниным и В. Суховольским [16]. Методы, относящиеся к данной группе, основаны на постулате о том, в случае «абсолютно честных» выборов ЭП

$\left[\gamma_{Level_k}^{(i)} \right]_j$, $i = \overline{1, m}$, $j = \overline{1, n}$, не зависит от значения ЭП $\left[\alpha_{Level_k}^{(i)} \right]_j$. Это означает, что значения

ЭП $\left[\gamma_{Level_k}^{(i)} \right]_j$ должны быть примерно одинаковыми в каждой j -ой избирательной комиссии каждого из уровней $Level$ независимо от ЭП

$\left[\alpha_{Level_k}^{(i)} \right]_j$ а также ЭП $\left[\alpha_{Level_k}^{(i)} \right]_j$.

Для проверки гипотезу о наличии/отсутствии детерминированной связи между ЭП

$\left[\alpha_{Level_k}^{(i)} \right]_j$ и ЭП $\left[\gamma_{Level_k}^{(i)} \right]_j$ с помощью метода

наименьших квадратов вычислялись оценки значений коэффициенты линейной регрессии:

$$\left[\gamma_{Level_k}^{(i)} \right]_j = a_{Level_k}^{(i)} \left[\alpha_{Level_k}^{(i)} \right]_j + b_{Level_k}^{(i)}$$

и далее проверяется гипотеза о статистически значимом отличии коэффициентов линейной регрессии $a_{Level_k}^{(i)}$ от нуля. Факт принятия данной статистической гипотезы $\left[\alpha_{Level_k}^{\%}\right]_j$

интерпретируется как подтверждение целенаправленного внесения аномалий в ЭД

$\left[N_{Level_k}^{ER}\right]_j$, $\left[N_{Level_k}^V\right]_j$, и $\left[N_{Level_k}^V\right]_j^{(j)}$, использованные для расчета ЭП $\left[\alpha_{Level_k}^{\%}\right]_j$, $\left[\gamma_{Level_k}^{\%}\right]_j^{(i)}$.

Для наглядной иллюстрации полученных подтверждений внесения целенаправленных изменений в ЭД визуализируют точки

с координатами $\left(\left[\alpha_{Level_k}^{\%}\right]_j, \left[\gamma_{Level_k}^{\%}\right]_j^{(i=1)}\right)$ и $\left(\left[\alpha_{Level_k}^{\%}\right]_j, \left[\gamma_{Level_k}^{\%}\right]_j^{(i=2)}\right)$ (здесь значение индекса i соответствует месту занятому кандидатом по результатам выборов), а также соответствующие линейные регрессии, которые при отсутствии фальсификаций ЭД должны быть параллельными оси абсцисс. Критика метода С-С приведена в разделе 2.

Метод Шпилькина – еще один известный метод выявления целенаправленно внесенных аномалий в ЭД основан на гипотезе об отсутствии детерминированной связи между ЭП $\left[\gamma_{Level_k}^{\%}\right]_j^{(i)}$ и ЭП $\left[\alpha_{Level_k}^{\%}\right]_j$ явкой избирателей

на выборы (см. [17–21]), традиционно, не доказанной математически. В этой связи эмпирические

ПР $P_{Level}^{(i)}\left(\left[\gamma_{Level_k}^{\%}\right]_j^{(i)}\right)$ и ФР $F_{Level}^{(i)}\left(\left[\gamma_{Level_k}^{\%}\right]_j^{(i)}\right)$, j – число избирательных ко-

миссий уровня $Level$, $i = \overline{1, m}$, m – число кандидатов, принявших участие в выборах, с точностью до масштабирующего коэффициента должны быть подобными друг другу:

$P_{Level}^{(1)}\left(\left[\gamma_{Level_k}^{\%}\right]_j^{(1)}\right) \square P_{Level}^{(2)}\left(\left[\gamma_{Level_k}^{\%}\right]_j^{(2)}\right) \square \dots \square P_{Level}^{(m)}\left(\left[\gamma_{Level_k}^{\%}\right]_j^{(m)}\right)$,

$F_{Level}^{(1)}\left(\left[\gamma_{Level_k}^{\%}\right]_j^{(1)}\right) \square F_{Level}^{(2)}\left(\left[\gamma_{Level_k}^{\%}\right]_j^{(2)}\right) \square \dots \square F_{Level}^{(m)}\left(\left[\gamma_{Level_k}^{\%}\right]_j^{(m)}\right)$.

В случае фальсификации результатов выборов в пользу определенного участника выборов, указанное выше подобие эмпирических

ПР $P_{Level}^{(1)}\left(\left[\gamma_{Level_k}^{\%}\right]_j^{(1)}\right)$, $P_{Level}^{(2)}\left(\left[\gamma_{Level_k}^{\%}\right]_j^{(2)}\right)$, ...,

и ФР $F_{Level}^{(1)}\left(\left[\gamma_{Level_k}^{\%}\right]_j^{(1)}\right)$, $F_{Level}^{(2)}\left(\left[\gamma_{Level_k}^{\%}\right]_j^{(2)}\right)$, ...,

отсутствует (см. [17–21]). В этой связи автор, предлагает, задавшись некоторым произвольно выбираемым пороговым значением ЭП

$\left[\tilde{\alpha}_{Level_k}^{\%}\right]_j$, отсеять ЭД, представленные избирательными комиссиями $Level_k$ -го уровня (и, соответственно, вычисленные по ним ЭП), в которых

$\left[\alpha_{Level_k}^{\%}\right]_j \geq \left[\tilde{\alpha}_{Level_k}^{\%}\right]_j$ и далее оценивать «истинные» значениями явки и результаты кандидатов, участвовавших в выборах, используя в качестве таковых координаты центра тяжести усеченного кластера. Затем на основе сравнения значений координат центров тяжести не усеченного и усеченного наборов ЭД автор предлагает вычислять оценки (по его утверждению «достоверные») «согласованности ЭД» и «объемов разных типов фальсификаций».

$P_{Level_k}^{(m)}\left(\left[\gamma_{Level_k}^{\%}\right]_j^{(m)}\right)$ и ФР $F_{Level_k}^{(1)}\left(\left[\gamma_{Level_k}^{\%}\right]_j^{(1)}\right)$,

$F_{Level}^{(2)}\left(\left[\gamma_{Level_k}^{\%}\right]_j^{(2)}\right), \dots, F_{Level_k}^{(m)}\left(\left[\gamma_{Level_k}^{\%}\right]_j^{(m)}\right)$ отсут-

ствует (см. [17–21]). В этой связи автор, предлагает, задавшись некоторым произвольно выбираемым пороговым значением ЭП

$\left[\tilde{\alpha}_{Level_k}^{\%}\right]_j$, отсеять ЭД, представленные избирательными комиссиями $Level_k$ -го уровня (и, соответственно, вычисленные по ним ЭП), в ко-

торых $\left[\alpha_{Level_k}^{\%}\right]_j \geq \left[\tilde{\alpha}_{Level_k}^{\%}\right]_j$ и далее оценивать «истинные» значениями явки и результаты кандидатов, участвовавших в выборах, используя в качестве таковых координаты центра тяжести усеченного кластера. Затем на основе сравнения значений координат центров тяжести не усеченного и усеченного наборов ЭД автор предлагает вычислять оценки (по его утверждению «достоверные») «согласованности ЭД» и «объемов разных типов фальсификаций».

В последствие метод С. Шпилькина был дополнен А. Подлазовым, предложившим

анализировать зависимость ЭП $\left[N_{Level_k}^V\right]_j^{(1)}$ от

ЭП $\sum_{i=2}^m \left[N_{Level_k}^V\right]_j^{(i)}$ который, традиционно, постулировал, что при отсутствии фальсификаций ЭД

$\left[N_{Level_k}^V\right]_j^{(1)} = k \sum_{i=2}^m \left[N_{Level_k}^V\right]_j^{(i)}$,

где угловой коэффициент k , которой можно вычислить в соответствии с регрессией Деминга [22]:

$k = p + (1 + p^2)^{1/2}$,

$p = \frac{\left\langle \left(\left[N_{Level_k}^V \right]_j^{(1)} \right)^2 - \left(\sum_{i=2}^m \left[N_{Level_k}^V \right]_j^{(i)} \right)^2 \right\rangle}{2 \left\langle \left[N_{Level_k}^V \right]_j^{(1)} \cdot \sum_{i=2}^m \left[N_{Level_k}^V \right]_j^{(i)} \right\rangle}$,

здесь $\langle \rangle$ означает операцию усреднения по всем точкам анализируемых случайных выборок, а качество линейной аппроксимации оценить величиной

$$\frac{\sigma_{\square}}{\sigma_{\perp}} = \frac{\left\langle \left(k \cdot [N_{Level_k}^V]_j^{(1)} + \sum_{i=2}^m [N_{Level_k}^V]_j^{(i)} \right)^2 \right\rangle}{\left\langle \left(k \cdot [N_{Level_k}^V]_j^{(1)} - \sum_{i=2}^m [N_{Level_k}^V]_j^{(i)} \right)^2 \right\rangle},$$

где $\sigma_{\square}$ – средний квадрат проекций векторов остатков моделей на координатную ось, сонаправленную с аппроксимирующей прямой, $\sigma_{\perp}$ – средний квадрат проекции векторов остатков моделей на координатную ось, перпендикулярную аппроксимирующей прямой. При наличии преднамеренно внесенных аномалий в ЭД, обсуждаемая зависимость, по мнению А. Подлазова, будет отличаться от линейной зависимости.

Таким образом, рассмотренным выше методам ЭК присущ ряд недостатков, основным из которых является отсутствием доказательств адекватности постулатов, положенных в основу этих методов, а также субъективный подбор «эталонных» характеристик, присущих, по представлению авторов того или иного метода, нефальсифицированным ЭД.

2. КРИТИКА МЕТОДА СУХОВОЛЬСКОГО-СОБЯНИНА

Рассмотрим результаты выборов в гипотетическом регионе, где проживают $N_{РИК}^{ER}$ официально зарегистрированных избирателей, из которых приняло участие в выборах $N_{РИК}^V \leq N_{РИК}^{ER}$ избирателей. Будем полагать, что для проведения выборов использовалась двухуровневая избирательная система:

$$K = 2, \quad k = 1, 2,$$

$$Level_1 = РИК, \quad Level_2 = РИК$$

Тогда на уровне региональной избирательной комиссии ($Level_1 = РИК$) ЭП «относительная явка избирателей на выборы» $\alpha_{РИК}^{\%}$ составил:

$$\alpha_{РИК}^{\%} = N_{РИК}^V / N_{РИК}^{ER} \cdot 100\%$$

Будем считать, что выборы в данном регионе проводили две «территориальных избирательных комиссий» ($ТИК \text{ № } 1, ТИК \text{ № } 2$), поэтому $Name_{Level_2} = \{ТИК_1, ТИК_2\}$. В $ТИК \text{ № } 1$ было зарегистрировано $[N_{ТИК}^{ER}]_1$ избирателей, из которых в выборах приняло участие $[N_{ТИК}^V]_1$ избирателей. Аналогичные ЭД, представленные $ТИК \text{ № } 2$ в РИК, обозначим,

$[N_{ТИК}^{ER}]_2, [N_{ТИК}^V]_2$ соответственно. Тогда общее число избирателей, зарегистрированных в РИК, составило:

$$N_{РИК}^{ER} = [N_{ТИК}^{ER}]_1 + [N_{ТИК}^{ER}]_2,$$

общее число избирателей, принявших участие в выборах, в РИК, составило:

$$N_{РИК}^V = [N_{ТИК}^V]_1 + [N_{ТИК}^V]_2.$$

Будем считать, что.

1) В $ТИК \text{ № } 1$ за первого кандидата проголосовало $[N_{ТИК}^V]_1^{(1)}$ избирателей, за второго кандидата – $[N_{ТИК}^V]_1^{(2)}$ избирателей, соответственно.

2) В $ТИК \text{ № } 2$ за первого кандидата проголосовало $[N_{ТИК}^V]_2^{(1)}$ избирателей, за второго кандидата – $[N_{ТИК}^V]_2^{(2)}$ избирателей.

Тогда значения ЭП $[\alpha_{ТИК}]_1, [\alpha_{ТИК}]_2$ вычисляются по формулам:

$$[\alpha_{ТИК}]_1 = \frac{[N_{ТИК}^V]_1}{[N_{ТИК}^{ER}]_1} = \frac{[N_{ТИК}^V]_1^{(1)} + [N_{ТИК}^V]_1^{(2)}}{[N_{ТИК}^{ER}]_1},$$

$$[\alpha_{ТИК}]_2 = \frac{[N_{ТИК}^V]_2}{[N_{ТИК}^{ER}]_2} = \frac{[N_{ТИК}^V]_2^{(1)} + [N_{ТИК}^V]_2^{(2)}}{[N_{ТИК}^{ER}]_2},$$

соответственно, значения ЭП $[\gamma_{ТИК}]_1^{(1)}, [\gamma_{ТИК}]_1^{(2)}$ – по формулам:

$$[\gamma_{ТИК}]_1^{(1)} = \frac{[N_{ТИК}^V]_1^{(1)}}{[N_{ТИК}^V]_1} = \frac{[N_{ТИК}^V]_1^{(1)}}{[N_{ТИК}^V]_1^{(1)} + [N_{ТИК}^V]_1^{(2)}},$$

$$[\gamma_{ТИК}]_1^{(2)} = \frac{[N_{ТИК}^V]_1^{(2)}}{[N_{ТИК}^V]_1} = \frac{[N_{ТИК}^V]_1^{(2)}}{[N_{ТИК}^V]_1^{(1)} + [N_{ТИК}^V]_1^{(2)}},$$

соответственно, значение ЭП $[\gamma_{ТИК}]_2^{(1)}, [\gamma_{ТИК}]_2^{(2)}$ – по формулам:

$$[\gamma_{ТИК}]_2^{(1)} = \frac{[N_{ТИК}^V]_2^{(1)}}{[N_{ТИК}^V]_2} = \frac{[N_{ТИК}^V]_2^{(1)}}{[N_{ТИК}^V]_2^{(1)} + [N_{ТИК}^V]_2^{(2)}},$$

$$[\gamma_{ТИК}]_2^{(2)} = \frac{[N_{ТИК}^V]_2^{(2)}}{[N_{ТИК}^V]_2} = \frac{[N_{ТИК}^V]_2^{(2)}}{[N_{ТИК}^V]_2^{(1)} + [N_{ТИК}^V]_2^{(2)}},$$

соответственно, значение ЭП $[\gamma_{РИК}]^{(1)}$, $[\gamma_{РИК}]^{(2)}$ – по формулам:

$$[\gamma_{РИК}]^{(1)} = \frac{[N_{РИК}^V]}{[N_{РИК}^V]} = \frac{[N_{ТИК}^V]_1 + [N_{ТИК}^V]_2}{[N_{ТИК}^V]_1 + [N_{ТИК}^V]_2},$$

$$[\gamma_{РИК}]^{(2)} = \frac{[N_{РИК}^V]}{[N_{РИК}^V]} = \frac{[N_{ТИК}^V]_1 + [N_{ТИК}^V]_2}{[N_{ТИК}^V]_1 + [N_{ТИК}^V]_2},$$

соответственно.

Используя приведенные выше формулы для вычисления значений ЭП, легко получить выражения, связывающее ЭП $[\gamma_{РИК}]^{(1)}$ с другими ЭП:

$$\begin{aligned} [\gamma_{РИК}]^1 &= \\ &= ([\gamma_{ТИК}]_1^1 + [\gamma_{ТИК}]_1^2) \frac{[N_{ТИК}^V]_1 + [N_{ТИК}^V]_2}{[N_{ТИК}^V]_1 + [N_{ТИК}^V]_2} = \\ &= ([\gamma_{ТИК}]_1^1 + [\gamma_{ТИК}]_1^2) \frac{[N_{ТИК}^V]_1}{[N_{ТИК}^V]_1 + [N_{ТИК}^V]_2} = (1) \\ &= ([\gamma_{ТИК}]_1^1 + [\gamma_{ТИК}]_1^2) \frac{[N_{ТИК}^V]_1}{N_{РИК}^V} = \\ &= ([\gamma_{ТИК}]_1^1 + [\gamma_{ТИК}]_1^2) \frac{[\alpha_{ТИК}]_1}{\alpha_{РИК}} \frac{[N_{ТИК}^{ER}]_1}{N_{РИК}^{ER}}. \end{aligned}$$

Поступая аналогично предыдущему, получаем выражение, связывающее ЭП $[\gamma_{РИК}]^{(1)}$ с другими ЭП:

$$\begin{aligned} [\gamma_{РИК}]^{(2)} &= ([\gamma_{ТИК}]_2^{(1)} + [\gamma_{ТИК}]_2^{(2)}) \frac{[N_{ТИК}^V]_2}{N_{РИК}^V} = \\ &= ([\gamma_{ТИК}]_2^{(1)} + [\gamma_{ТИК}]_2^{(2)}) \frac{[\alpha_{ТИК}]_2}{\alpha_{РИК}} \frac{[N_{ТИК}^{ER}]_2}{N_{РИК}^{ER}} \end{aligned} \quad (2)$$

Из (1), (2) видно, что ЭП $[\gamma_{РИК}]^{(1)}$ зависит от суммы ЭП $[\gamma_{ТИК}]_1^{(1)}$, $[\gamma_{ТИК}]_1^{(2)}$ умноженной на весовой коэффициент $\frac{[\alpha_{ТИК}]_1}{\alpha_{РИК}} \frac{[N_{ТИК}^{ER}]_1}{N_{РИК}^{ER}}$, ЭП $[\gamma_{РИК}]^{(2)}$ – от суммы ЭП $[\gamma_{ТИК}]_2^{(1)}$, $[\gamma_{ТИК}]_2^{(2)}$ умноженной на весовой коэффициент $\frac{[\alpha_{ТИК}]_2}{\alpha_{РИК}} \frac{[N_{ТИК}^{ER}]_2}{N_{РИК}^{ER}}$. Отметим, что в весовые ко-

эффициенты в (1) представляют собой произведение отношения ЭП $[\alpha_{ТИК}]_1$ к ЭП $\alpha_{РИК}$ и отношения $[N_{ТИК}^{ER}]_1$ к ЭП $N_{РИК}^{ER}$; весовые коэффи-

циенты в (2) – произведение отношения ЭП $[N_{ТИК}^{ER}]_2$ к ЭП $N_{РИК}^{ER}$. При этом, очевидно, что $[\alpha_{ТИК}]_1/\alpha_{РИК} \square 1$, $[\alpha_{ТИК}]_2/\alpha_{РИК} \square 1$. Таким образом, можно считать, что при вычислении значений ЭП $[\gamma_{РИК}]^{(1)}$, $[\gamma_{РИК}]^{(1)}$ используются весовые коэффициенты $[N_{ТИК}^{ER}]_1/N_{РИК}^{ER}$, $[N_{ТИК}^{ER}]_2/N_{РИК}^{ER}$ которые позволяют учесть степень влияния численности избирателей, зарегистрированных в данной ТИК, на значения значения ЭП $[\gamma_{ТИК}]_1^{(1)}$, $[\gamma_{ТИК}]_1^{(2)}$. Соответственно, в обсуждаемом методе С-С, должен проводится анализ зависимостей взвешенных ЭП $([N_{ТИК}^{ER}]_1/N_{РИК}^{ER})[\gamma_{ТИК}]_1^{(1)}$, $([N_{ТИК}^{ER}]_2/N_{РИК}^{ER})[\gamma_{ТИК}]_1^{(2)}$ от ЭП «явка избирателей в ТИК № 1», «явка избирателей в ТИК № 2». Очевидно, что данная рекомендация относится к случаю большего числа ТИК и большего числа кандидатов (партий), участвующих в выборах.

Для подтверждения того, что при использовании взвешенных значений ЭП $([N_{ТИК}^{ER}]_1/N_{РИК}^{ER})[\gamma_{ТИК}]_1^{(1)}$, $([N_{ТИК}^{ER}]_2/N_{РИК}^{ER})[\gamma_{ТИК}]_1^{(2)}$ их линейные связи, соответственно, с ЭП $[\alpha_{ТИК}^%]_1$, $[\alpha_{ТИК}^%]_2$ авторы провели самостоятельные подсчеты, используя официальные ЭД, опубликованные Центральной избирательной комиссией (ЦИК) РФ на сайте ЦИК, по итогам выборов в 2018 г. Президента РФ (далее – Выборы-2018) [26], структура которых подробно описанная в [27], соответствовала иерархической архитектуре избирательной системы, проводившей Выборы-2018:

- 1) вершина дерева – ЦИК РФ;
- 2) избирательные комиссии (ИК) субъектов РФ, г. Байконур (Республика Казахстан), ИК, объединившая участковые избирательные комиссии (УИК), созданные пределами РФ (всего – 87);
- 3) ТИК (всего – 2 776);
- 4) УИК (всего – 97 314).

Для автоматической выгрузки ЭД с сайта ЦИК РФ в файл CVS-формата, содержащий таблицу, состоящую из 23 столбцов и 100579 строк, на основе использования информационных технологий, выбор которых обоснован в [27], был создан специализированный программный инструмент, подробно описанный в [28]. Результаты многомерного статистического анализа ЭД, выгруженных в автомати-

ческом режиме с сайта ЦИК, которые подтверждают отсутствие в них целенаправленно внесенных аномалий, опубликованы авторами в [29].

На основе использования ЭД, находящихся в обсуждаемой таблице, были вычислены значения ЭП $\left[\gamma_{ИК}^{\%}\right]_j^{(1)}$, $\left[\tilde{\gamma}_{ИК}^{\%}\right]_j^{(1)}$, $\left[\gamma_{ИК}^{\%}\right]_j^{(2)}$, $\left[\tilde{\gamma}_{ИК}^{\%}\right]_j^{(2)}$, $\left[\alpha_{ИК}^{\%}\right]_j$, $j = \overline{1,87}$ (1 – ЭП Путина Вла-

димира Владимировича, занявшего второе место). Визуализации зависимостей, заданных таблицами:

$$1) \left\langle \left[\alpha_{ИК}^{\%}\right]_j, \left[\gamma_{ИК}^{\%}\right]_j^{(1)} \right\rangle, \left\langle \left[\alpha_{ИК}^{\%}\right]_j, \left[\tilde{\gamma}_{ИК}^{\%}\right]_j^{(2)} \right\rangle, \\ j = \overline{1,87} \text{ отсчетов;}$$

$$2) \left\langle \left[\alpha_{ТИК}^{\%}\right]_j, \left[\gamma_{ТИК}^{\%}\right]_j^{(1)} \right\rangle, \left\langle \left[\alpha_{ТИК}^{\%}\right]_j, \left[\tilde{\gamma}_{ТИК}^{\%}\right]_j^{(2)} \right\rangle, \\ j = \overline{1,97314},$$

и их линейные аппроксимации представлены на рисунках 1,2 [26].

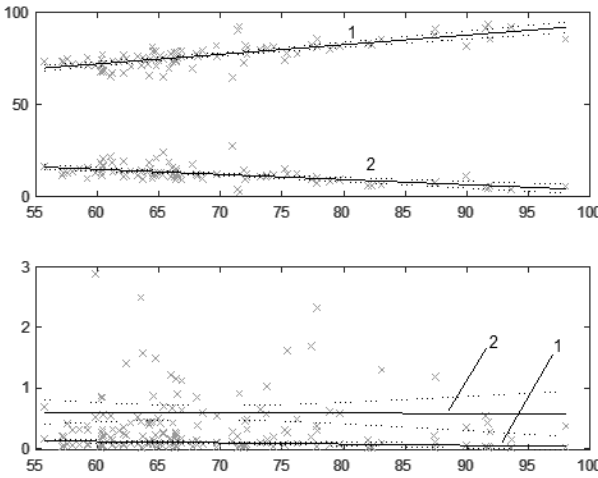


Рис. 1. Сверху: графики зависимостей $\left\langle \left[\alpha_{ИК}^{\%}\right]_j, \left[\gamma_{ИК}^{\%}\right]_j^{(1)} \right\rangle$ (1), $\left\langle \left[\alpha_{ИК}^{\%}\right]_j, \left[\tilde{\gamma}_{ИК}^{\%}\right]_j^{(2)} \right\rangle$ (2), $j = \overline{1,87}$ и их линейные аппроксимации.
Снизу: графики зависимостей $\left\langle \left[\alpha_{ИК}^{\%}\right]_j, \left[\tilde{\gamma}_{ИК}^{\%}\right]_j^{(1)} \right\rangle$ (1), $\left\langle \left[\alpha_{ИК}^{\%}\right]_j, \left[\gamma_{ИК}^{\%}\right]_j^{(2)} \right\rangle$ (2) и их линейной аппроксимации

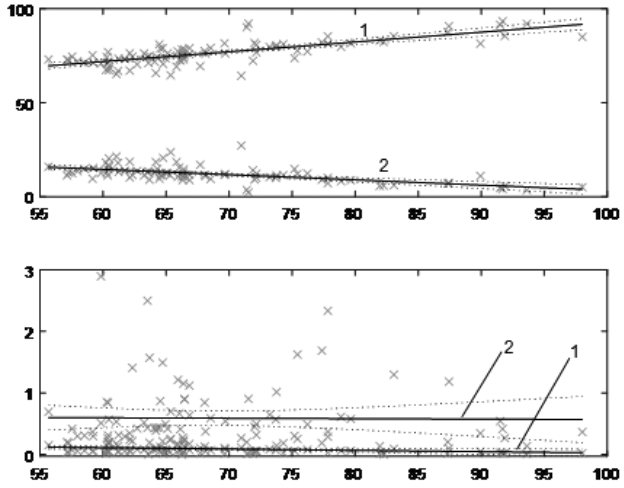


Рис 2. Сверху: графики зависимостей $\left\langle \left[\alpha_{ТИК}^{\%}\right]_j, \left[\gamma_{ТИК}^{\%}\right]_j^{(1)} \right\rangle$ (1), $\left\langle \left[\alpha_{ТИК}^{\%}\right]_j, \left[\tilde{\gamma}_{ТИК}^{\%}\right]_j^{(2)} \right\rangle$ (2), $j = \overline{1,97314}$ и их линейные аппроксимации.
Снизу: графики зависимостей $\left\langle \left[\alpha_{ТИК}^{\%}\right]_j, \left[\tilde{\gamma}_{ТИК}^{\%}\right]_j^{(1)} \right\rangle$ (1), $\left\langle \left[\alpha_{ТИК}^{\%}\right]_j, \left[\gamma_{ТИК}^{\%}\right]_j^{(2)} \right\rangle$ (2) и их линейной аппроксимации

Таким образом, результаты критического анализа метода С-С позволяют сделать обоснованный вывод о неправомерности его использования для анализа ЭД, а также публикации полученных с его помощью результатов. Необходимо отметить, что известен опыт использования метода С-С в его изначальном варианте для случая «честные выборы, неоднородные ТИК», а также для построения модели, описывающей фальсифицированные ЭД за счет уменьшения числа голосов, поданных за данного кандидата в пользу другого кандидата (переток голосов), и обоснования выбора критериев для проверки гипотезы об отсутствии нерегулярностей в ЭД [25]. Однако результаты, изложенные выше, опровергают выводы авторов [25].

ЗАКЛЮЧЕНИЕ

В первой части статьи обоснована актуальность изучения методов ЭК, которые по мнению авторов данных методов и их сторонников обеспечивают выявление целенаправленно внесенных аномалий в электоральные данные. Приведено описание методов ЭК, наиболее активно в постсоветский период истории России использовавшихся представителями проигравшей стороны для обнаружения преднамеренно внесенных искажений в электоральные данные, публикуемые по результатам выборов в органы государственной власти РФ, местного самоуправления и референдумов, с целью делигитимации их результатов. Данные методы, как показали результаты их анализа, основаны на сравнении тех или иных статистических характеристик ЭД и вычисляемых на их основе ЭП с некоторыми эталонами, которые, зачастую, выбираются их авторами, субъективно, без должного научного обоснования. не имеют под собой строго обоснованной научной базы.

В этой ситуации возникает бесконечная рекурсия, так для доказательства факта честного проведения выборов нужно использовать методы ЭК, для доказательства адекватности которых, в свою очередь, нужно использовать ЭД, опубликованные по результатам честно проведенных выборов, и т.д. (аналогично, известной истории о слуге культа и его любимой собаке). По мнению авторов возможный выход из описанной ситуации состоит в использовании аналитических подходов к анализу методов ЭК, а также их применение к анализу синтезированных заведомо нефальсифицированным ЭД, статистические свойства которых подобны статистическим свойствам реальных электоральных данных.

В первой части статьи подробно рассмотрен метод Собянина-Суходольского и аналитических показано, что авторы данного метода при анализе зависимости ЭП «доля голосов избирателей, проголосовавших за победителя выборов» от ЭП «явка избирателей на выборы», игнорируют необходимость использования весовых множителей пропорциональных отношению числа избирателей, зарегистрированных в избирательной комиссии данного уровня, к числу избирателей, имевших право участия в выборах. Без их учета вклад избирательных комиссий с небольшим числом избирателей, на которых, как очевидно, высокую явку обеспечить легче, чем на участках с большим числом зарегистрированных избирателей, оказывается аналогичной вкладу последних в анализируемую зависимость.

Следовательно, выводы, сделанные на основе результатов применения метода Собянина-Суходольского и его модификаций, требуют критического переосмысления.

ЛИТЕРАТУРА

1. Шалаев Н.Е. Электоральные аномалии в постсоциалистическом пространстве: опыт политологического анализа. Дисс...канд. политических наук по специальности 23.00.02 – Политические институты, процессы и технологии. – Санкт-Петербург: СПбГУ, 2016 г. – 191 с.
2. Benford F. The Law of Anomalous Numbers// *Proceedings of the American Philosophical Society.* – 1938. – Mar. 31. – Vol. 78. – № 4. – P. 551–572.
3. Deckert J., M. Myagkov, Ordeshook P.C. The Irrelevance of Benford's Law for Detecting Fraud in Elections [Электронный ресурс] // Caltech/MIT Voting Technology Project Working Paper. №. 9. 2010. URL: <http://vote.caltech.edu/content/irrelevance-benford-s-law-detecting-fraud-elections> (Дата обращения: 4.02.2023).
4. Diekmann A. Benford's Law and Fraud Detection: Facts and Legends/ Diekmann, Andreas, Ben Jann// *German Economic Review.* – 2010. – Vol. 11 (3). – P. 397–401.
5. Shikano S. When Does the Second-Digit Benford's Law-Test Signal an Election Fraud? Facts or Misleading Test Results // *Jahrbucher f. Nationalokonomie u. Statistik (Lucius & Lucius, Stuttgart 2011).* – Bd. (Vol.) 231/5+6. – P. 719–732.
6. Breunig C. Searching for electoral irregularities in an established democracy: Applying Benford's Law tests to Bundestag elections in Unified Germany// *Electoral Studies.* – 2011. – Vol. 30. – P. 534–545.
7. Kobak D., Shpilkin S., Pshenichnikov M.S. Integer Percentages as Electoral Falsification Fingerprints// *Ann. Appl. Stat.* – 2016. – Vol. 10. – № 1. – P. 54–73.
8. Beber B., Scacco A. What the Numbers Say: A Digit-Based Test for Election Fraud//*Political Analysis.* – 2012. – Vol. 20. – P. 211–234.
9. Lee J., Tam Cho W.K., Judge G.G. Stigler's approach to recovering the distribution of first significant digits in natural data sets// *Statistics & Probability Letters.* – 2010. – Vol. 80. Issue 2. – P. 82–88.
10. 10. Walter R. Mebane Jr. Fraud in the 2009 Presidential Election in Iran? // *CHANCE.* – 2010. – Vol. 23. – № 1. – P. 9.
11. Walter R. Mebane Jr. Comment on «Benford's Law and the Detection of Election Fraud». // *Political Analysis.* – 2011. – Vol. 19. – P. 270.
12. Walter R. Meban, Jr. A Layman's Guide to Statistical Election Forensics. [Электронный ресурс] // *ElectionGuide Digest.* – 2010. URL: <http://www.electionguide.org/digest/post/271/>
13. Мебейн, У. Электоральные фальсификации в России: комплексная диагностика выборов 2003–2004, 2007–2008 гг.// *Российское Электоральное Обозрение.* – 2009. – №2. – С 61.
14. Подлазов А.В. Формальные методы выявления масштабных электоральных фальсификаций на материале федеральных выборов 1999–2018 гг. / Препринты ИПМ им. М.В. Келдыша. – 2019. – № 2. – 28 с. doi:<http://doi.org/10.20948/prepr-2019-2>. URL: <http://library.keldysh.ru/preprint.asp?id=2019-2> (Дата обращения: 25.11.2025).
15. Подлазов А.В. Исследование статистических методов выявления вымышленных результатов выборов: Часть 1. Круглые числа // Препринты ИПМ им. М.В. Келдыша. – 2019. – № 147. – 28 с. DOI: <http://doi.org/10.20948/prepr-2019-147>. URL: <http://library.keldysh.ru/preprint.asp?id=2019-147> (Дата обращения: 25.11.2025).
16. Собянин А.А., Суховольский В.Г. Демократия, ограниченная фальсификациями. – М.: ИНТУ, 1995. – 263 с.
17. Шпилькин С. Статистическое исследование результатов российских выборов 2007–2009 гг. // *Троицкий вариант – наука.* – 2009. – № 21(40). – С. 2–4.
18. Шпилькин С. Математика выборов – 2011 // *Троицкий вариант – наука.* – 2011. – № 25(94). – С. 2–4.
19. Шпилькин С. Двугорбая Россия // *Троицкий вариант – наука.* – 2016. № 20(214). – С. 1–3.
20. Шпилькин С. Выборы 2018 года: Фактор X и «пила Чурова» // *Троицкий вариант – наука.* – 2018. – № 8(252). – С. 8–10.
21. Шпилькин С. Поправки на 27 миллионов // *Троицкий вариант – наука.* – 2020. – № 14(308). – С. 4–5.
22. Подлазов А.В. Реконструкция фальсифицированных результатов выборов с помощью интегрального метода Шпилькина // Проектирование будущего. Проблемы цифровой реальности: труды 4-й Международной конференции (4–5 февраля 2021 г., Москва). – М.: ИПМ им. М.В. Келдыша, 2021. – С. 193–208// URL: <https://keldysh.ru/future/2021/18.pdf> (Дата обращения: 25.11.2025).

23. Leemann L., Bochsler D. A systematic approach to study electoral fraud// *Electoral Studies*. – 2014. – Vol. 35. – P. 33–47.
24. Pericchi L., Torres D. Quick Anomaly Detection by the Newcomb-Benford Law, with Applications to Electoral Processes Data from the USA, Puerto Rico and Venezuela// *Statistical Science*. – 2011. – Vol. 26. – №. 4. – P. 502–516.
25. Кунов А., Мягков М., Ситников А., Шакин Д. Россия и Украина: нерегулярные результаты регулярных выборов. Аналитический доклад Института открытой экономики. – М.: Институт открытой экономики, 2005. – 37 с. URL: http://sergey.shulgin.ru/files/paper_elections.pdf
26. URL:<http://www.vybory.izbirkom.ru/region/izbirkom?action=show&global=1&vrn=100100084849062®ion=0&prver=0&pronetvd=null> (Дата обращения: 25.11.2025)
27. Мирвода С.Г., Поршнев С.В., Рябко Н.Ю. Автоматизация процедуры доступа к электоральным данным, размещенным на сайте Центральной избирательной комиссии Российской Федерации// *Вестник УрФО. Безопасность в информационной сфере*. – 2022. – № 1(43). – С. 28– 34. DOI: 10.14529/secur220104
28. Инструмент для загрузки официальных данных, публикуемых Центральной избирательной комиссией Российской Федерации// Мирвода С.Г., Поршнев С.В., Рябко Н.Ю., Стойчин К.Л., Тимошенко-ва Ю.С./ Свидетельство о государственной регистрации программы для ЭВМ № 2023683261 (Заявка № 2023681765. Дата поступления 23.10.2023. Дата государственной регистрации в Реестре программ для ЭВМ 07.11.2023.)
29. Поршнев С.В., Рябко Н.Ю. Опыт многомерного статистического анализа электоральных данных// *Вестник УрФО. Безопасность в информационной сфере*. – 2023. – № 2(48). – С. 5– 29. DOI: 10.14529/secur230201

References

1. Shalaev N.E. Elektoral'nye anomalii v postsotsialisticheskom prostranstve: opyt politologicheskogo analiza. Diss...kand. politicheskikh nauk po spetsial'nosti 23.00.02 – Politicheskie instituty, protsessy i tehnologii. – Sankt-Peterburg: SPbGU, 2016. – 191 s.
2. Benford F. The Law of Anomalous Numbers// *Proceedings of the American Philosophical Society*. – 1938. – Mar. 31. – Vol. 78. – № 4. – P. 551– 572.
3. Deckert J., M. Myagkov, Ordeshook P.C. The Irrelevance of Benford's Law for Detecting Fraud in Elections// *Caltech/MIT Voting Technology Project Working Paper*. №. 9. 2010. URL: <http://vote.caltech.edu/content/irrelevance-benfords-law-detecting-fraud-elections> (Data obrascheniya: 4.02.2023).
4. Diekmann A. Benford's Law and Fraud Detection: Facts and Legends/ Diekmann, Andreas, Ben Jann// *German Economic Review*. – 2010. – Vol. 11 (3). – P. 397–401.
5. Shikano S. When Does the Second-Digit Benford's Law-Test Signal an Election Fraud? Facts or Misleading Test Results// *Jahrbucher f. Nationalokonomie u. Statistik (Lucius & Lucius, Stuttgart 2011)*. – Bd. (Vol.) 231/5+6. – P. 719–732.
6. Breunig C. Searching for electoral irregularities in an established democracy: Applying Benford's Law tests to Bundestag elections in Unified Germany// *Electoral Studies*. – 2011. – Vol. 30. – P. 534–545.
7. Kobak D., Shpilkin S., Pshenichnikov M.S. Integer Percentages as Electoral Falsification Fingerprints// *Ann. Appl. Stat.* – 2016. – Vol. 10. – № 1. – P. 54– 73.
8. Beber B., Scacco A. What the Numbers Say: A Digit-Based Test for Election Fraud// *Political Analysis*. – 2012. – Vol. 20. – P. 211–234.
9. Lee J., Tam Cho W.K., Judge G.G. Stigler's approach to recovering the distribution of first significant digits in natural data sets// *Statistics & Probability Letters*. – 2010. – Vol. 80. Issue 2. – P. 82–88.
10. Walter R. Mebane Jr. Fraud in the 2009 Presidential Election in Iran?// *CHANCE*. – 2010. – Vol. 23. – № 1. – P. 9.
11. Walter R. Mebane Jr. Comment on «Benford's Law and the Detection of Election Fraud»// *Political Analysis*. – 2011. – Vol. 19. – P. 270.
12. Walter R. Meban, Jr. A Layman's Guide to Statistical Election Forensics// *ElectionGuide Digest*. – 2010. URL: <http://www.electionguide.org/digest/post/271/> (Data obrascheniya 20.10.2023)
13. Mebejn, U. Elektoral'nye fal'sifikatsii v Rossii: kompleksnaya diagnostika vyborov 2003– 2004, 2007– 2008 gg.// *Rossiiskoe Elektoral'noe Obozrenie*. – 2009. – № 2. – S 61.
14. Podlazov A.V. Formal'nye metody vyjavleniya masshtabnyh `elektoral'nyh fal'sifikatsij na materiale federal'nyh vyborov 1999–2018 gg. / Preprinty IPM im. M.V. Keldysha. – 2019. – № 2. – 28 s. doi:<http://doi.org/10.20948/prepr-2019-2>. URL: <http://library.keldysh.ru/preprint.asp?id=2019-2> (Data obrascheniya: 25.11.2025).

15. Podlazov A.V. Issledovanie statisticheskikh metodov vyjavleniya vydumannykh rezul'tatov vyborov: Chast' 1. Kruglye chisla // Preprinty IPM im. M.V. Keldysha. – 2019. – № 147. – 28 s. DOI: <http://doi.org/10.20948/prepr-2019-147>. URL: <http://library.keldysh.ru/preprint.asp?id=2019-147> (Data obrascheniya: 25.11.2025).
16. Sobjanin A.A., Suhovol'skij V.G. Demokratiya, ogranichennaya fal'sifikatsijami. – M.: INTU, 1995. – 263 s.
17. Shpil'kin S. Statisticheskoe issledovanie rezul'tatov rossijskikh vyborov 2007–2009 gg.// Troitskij variant – nauka. – 2009. – № 21(40). – S. 2– 4.
18. Shpil'kin S. Matematika vyborov – 2011 // Troitskij variant – nauka. – 2011. – № 25(94). – S. 2– 4.
19. Shpil'kin S. Dvugorbaja Rossiya // Troitskij variant – nauka. – 2016. № 20(214). – S. 1– 3.
20. Shpil'kin S. Vybory 2018 goda: Faktor H i «pila Churova» // Troitskij variant – nauka. – 2018. – № 8(252). – S. 8– 10.
21. Shpil'kin S. Popravki na 27 millionov // Troitskij variant – nauka. – 2020. – № 14(308). – S. 4– 5.
22. Podlazov A.V. Rekonstruktsiya fal'sifitsirovannykh rezul'tatov vyborov s pomosh'ju integral'nogo metoda Shpil'kina // Proektirovanie buduschego. Problemy tsifrovoy real'nosti: trudy 4-j Mezhdunarodnoj konferentsii (4– 5 fevralja 2021 g., Moskva). – M.: IPM im. M.V. Keldysha, 2021. – S. 193– 208. URL: <https://keldysh.ru/future/2021/18.pdf> (Data obrascheniya: 25.11.2025)
23. Leemann L., Bochsler D. A systematic approach to study electoral fraud// Electoral Studies. – 2014. – Vol. 35. – P. 33–47.
24. Pericchi L., Torres D. Quick Anomaly Detection by the Newcomb-Benford Law, with Applications to Electoral Processes Data from the USA, Puerto Rico and Venezuela// Statistical Science. – 2011. – Vol. 26. – №. 4. – P. 502–516.
25. Kunov A., Mjagkov M., Sitnikov A., Shakin D. Rossiya i Ukraina: nereguljarnye rezul'taty reguljarnykh vyborov. Analiticheskij doklad Instituta otkrytoj `ekonomiki. – M.: Institut otkrytoj ekonomiki, 2005. – 37 s. URL: http://sergey.shulgin.ru/files/paper_elections.pdf (Data obrascheniya: 25.11.2025).
26. URL:<http://www.vybory.izbirkom.ru/region/izbirkom?action=show&global=1&vrn=100100084849062®ion=0&prver=0&pronetvd=null> (Data obrascheniya: 25.11.2025)
27. Mirvoda S.G., Porshnev S.V., Rjabko N.Ju. Avtomatizatsiya protsedury dostupa k `elektoral'nykh dannym, razmeshchennym na sajte Tsentral'noj izbiratel'noj komissii Rossijskoj Federatsii// Vestnik UrFO. Bezopasnost' v informatsionnoj sfere. – 2022. – № 1(43). – S. 28– 34. DOI: 10.14529/secur220104
28. Instrument dlja vygruzki ofitsial'nykh dannyh, publikuemykh Tsentral'noj izbiratel'noj komissiej Rossijskoj Federatsii// Mirvoda S.G., Porshnev S.V., Rjabko N.Ju., Stojchin K.L., Timoshenkova Ju.S/ Svidetel'stvo o gosudarstvennoj registratsii programmy dlja `EVM № 2023683261 (Zajavka № 2023681765. Data postuplenija 23.10.2023. Data gosudarstvennoj registratsii v Reestre programm dlja `EVM 07.11.2023.)28.
29. Porshnev S.V., Rjabko N.Ju. Opyt mnogomernogo statisticheskogo analiza `elektoral'nykh dannyh// Vestnik UrFO. Bezopasnost' v informatsionnoj sfere. – 2023. – № 2(48). – S. 5– 29. DOI: 10.14529/secur230201

ПОРШНЕВ Сергей Владимирович, доктор технических наук, профессор, профессор Учебно-научного центра «Информационная безопасность» ФГАОУ ВО «Уральский федеральный университет им. первого Президента России Б.Н. Ельцина». 620002, г. Екатеринбург, ул. Мира, 32.; ФГБОУ ВО «Уральский государственный экономический университет». 620144, г. Екатеринбург, ул. 8 Марта/Народной Воли, 62/45. E-mail: s.v.porshnev@urfu.ru

РЯБКО Николай Юрьевич, аспирант, ФГАОУ ВО «Уральский федеральный университет им. первого Президента России Б.Н. Ельцина». 620002, г. Екатеринбург, ул. Мира, 32. E-mail: N.Yu.Ryabko@urfu.ru

PORSHNEV Sergey Vladimirovich, Doctor of Technical Sciences, Professor, Director of the Educational and Scientific Center «Information Security» of the Federal State Autonomous Educational Institution of Higher Education «Ural Federal University named after the first President of Russia B.N. Yeltsin». 620002, Yekaterinburg, st. Mira, 32.; Ural State University of Economics. 620144, Yekaterinburg, 8 Marta/Narodnoy Voli St., 62/45. E-mail: N.Yu.Ryabko@urfu.ru

RYABKO Nikolay Yurievich, post-graduate student of the Federal State Autonomous Educational Institution of Higher Education «Ural Federal University named after the first President of Russia B.N. Yeltsin». 620002, Yekaterinburg, st. Mira, 32. E-mail: N.Yu.Ryabko@urfu.ru



АВТОМАТИЗАЦИЯ ОЦЕНКИ ВРЕДОНОСНОСТИ СКРИПТОВ НА ОСНОВЕ БОЛЬШИХ ЯЗЫКОВЫХ МОДЕЛЕЙ ПРИ РАССЛЕДОВАНИИ КОМПЬЮТЕРНЫХ ИНЦИДЕНТОВ

В статье рассматривается возможность автоматизации одного из самых трудоёмких этапов расследования компьютерных инцидентов — анализа скриптов на предмет вредоносности. Для этого нами разработан специализированный датасет, включающий 98 скриптов на языках Python, PowerShell и Bash, размеченных экспертами на «вредные» и «безвредные». На основе этого датасета проведено сравнительное исследование проприетарной модели gpt-4-turbo и открытой LLaMA-3.1-70B с использованием подхода Chain-of-Thought (CoT) для повышения интерпретируемости выводов. Эксперимент показал, что LLaMA-3.1-70B достигла 100% точности ($F1=1,0$), хотя такой результат, вероятно, связан с ограниченным составом датасета. Модель gpt-4-turbo продемонстрировала более умеренные, но реалистичные метрики ($F1=0,72$), однако с высоким числом ложноположительных срабатываний. Полученные данные подтверждают перспективность интеграции больших языковых моделей в процессы цифровой криминалистики, при этом подчёркивают критическую важность учёта контекста выполнения скриптов.

Ключевые слова: защита информации, кибербезопасность, большие языковые модели, анализ скриптов, расследование инцидентов, датасет, машинное обучение в кибербезопасности.

AUTOMATING THE ASSESSMENT OF SCRIPT MALICIOUSNESS BASED ON LARGE LANGUAGE MODELS WHEN INVESTIGATING COMPUTER INCIDENTS

This paper explores the automation of one of the most labor-intensive stages of computer incident investigation—the analysis of script maliciousness. To this end, we developed a specialized dataset comprising 98 expert-labeled scripts in Python, PowerShell, and Bash, categorized as 'malicious' or 'benign'. Using this dataset, a comparative study was conducted between the proprietary gpt-4-turbo model and the open-source LLaMA-3.1-70B, employing a Chain-of-Thought (CoT) approach to enhance the interpretability of the conclusions. The experiment showed that LLaMA-3.1-70B achieved 100% accuracy ($F1=1.0$), a result likely attributed to the limited composition of the dataset. The gpt-4-turbo model demonstrated more moderate yet realistic metrics ($F1=0.72$), but with a high number of false positives. The findings confirm the promise of integrating Large Language Models into digital forensics processes while highlighting the critical importance of considering script execution context.

Keywords: information security, cybersecurity, large language models (LLMs), script analysis, incident response, dataset, machine learning.

Введение. При расследовании компьютерных инцидентов особое внимание уделяется цифровым следам, оставленным злоумышленниками. Среди них особенно ценными оказываются скрипты — последовательности команд на языках Python, PowerShell или Bash, которые атакующие используют для автоматизации своих действий: от первоначального проникновения до перемещения по инфраструктуре и выполнения вредоносных операций. Анализ таких скриптов позволяет не только выявить факт компрометации, но и восстановить тактики, техники и процедуры (TTPs), применённые злоумышленниками, а также оценить потенциальный ущерб. Однако в условиях роста числа и сложности инцидентов ручной анализ становится всё менее жизнеспособным. Это объективно подталкивает к поиску решений на основе автоматизации — и здесь большие языковые модели (Large Language Models, LLM) выглядят многообещающе. Они способны не только классифицировать скрипты по степени вредоносности, но и выделять индикаторы ком-

прометации (IoC) и формировать аналитические заключения для специалистов по кибербезопасности.

Эффективность применения LLM напрямую зависит от качества обучающих и тестовых данных. В нашей работе мы предложили оригинальный подход к формированию датасета, включающего как скрипты из реальных инцидентов, так и легитимные примеры, используемые в системном администрировании, с последующей экспертной разметкой.

Важно отметить, что использование проприетарных моделей, таких как gpt-4-turbo или Claude 3, несмотря на их высокую точность, сопряжено с рисками утечки конфиденциальной информации, поскольку запросы обрабатываются на стороне облачного провайдера. В задачах информационной безопасности предпочтение целесообразно отдавать открытым LLM — таким как LLaMA [1], Mistral [2] и Qwen [3], — которые могут быть развёрнуты в изолированных средах и поддерживают fine-tuning под специфические сценарии [4].

На сегодняшний день применение LLM в кибербезопасности активно изучается, но результаты остаются неоднозначными. Например, GPT-4 показывает высокую эффективность в исправлении уязвимого кода, снижая их количество на 90% [5]. В работе [6] отмечается, что LLM успешно идентифицируют такие уязвимости, как SQL-инъекции и XSS, и дают рекомендации по их устранению. Интеграция LLM с формальными методами верификации позволяет устранить до 80% уязвимостей [7], хотя такой подход требует значительных вычислительных ресурсов. Вместе с тем существуют и ограничения. Исследование [8] показало низкую точность LLM при обнаружении сложных уязвимостей типа zero-day из-за недостатка контекста выполнения. Кроме того, даже при превосходстве GPT-4 над специализированными моделями, такими как CodeBERT [9], её использование ограничено в сценариях, где критична конфиденциальность.

Нами выделены три ключевых пробела в современных исследованиях. Во-первых, сохраняется дефицит репрезентативных датасетов, сочетающих вредоносные и легитимные скрипты из реальных инцидентов. Во-вторых, недостаточно изучена эффективность открытых LLM, таких как LLaMA, в задачах автоматизации расследований. В-третьих, практически не исследовано влияние контекста выполнения на достоверность классификации скриптов.

Целью настоящей работы является экспериментальная оценка достоверности больших языковых моделей в анализе скриптов и обоснование целесообразности их применения для автоматизации этапов расследования компьютерных инцидентов. В отличие от предыдущих работ, мы сосредоточились именно на скриптах (а не на произвольном коде), сравнили проприетарные и открытые модели, а также проанализировали влияние контекста выполнения — фактор, критически важный при анализе реальных инцидентов.

Описание датасета. Для проведения экспериментов нами был создан специализированный датасет, включающий 98 скриптов на языках Python, PowerShell и Bash. Каждый скрипт сохранён в виде отдельного текстового файла и отражает реальные сценарии, с которыми сталкиваются специалисты при расследовании инцидентов. При формировании датасета мы руководствовались практическими потребностями анализа: в него вошли как легитимные скрипты, используемые в повседневной работе системных администраторов и DevOps-инженеров, так и примеры, характерные для действий злоумышленников. Все скрипты были размечены экспертами и распределены по двум категориям:

«Безвредные» — легитимные сценарии (например, резервное копирование, развёртывание приложений, мониторинг ресурсов);

«Вредные» — команды, направленные на компрометацию систем (сбор учётных данных, обход защиты, эксфильтрация информации).

По функциональному происхождению скрипты делятся на три группы:

DevOps-скрипты [10] — автоматизация CI/CD-процессов, управление контейнерами, развёртывание инфраструктуры;

Скрипты системного администрирования — настройка ОС, управление пользователями, диагностика;

Скрипты, используемые в атаках — примеры из реальных инцидентов, включая техники, описанные в MITRE ATT&CK.

Итоговое распределение представлено в Таблице 1.

Важно отметить одно ограничение: в текущей версии датасета отсутствует контекст выполнения скриптов (окружение, права пользователя, сетевые условия). Это сознательное упрощение, моделирующее начальный этап расследования, когда аналитик имеет доступ только к самому коду.

Пример вредоносного скрипта на рисунке 1 демонстрирует типичную цепочку дей-

Таблица 1

Распределение скриптов по категориям

№ п/п	Категория	Количество скриптов	Процент от общего числа
1	Безвредный	58	59%
2	Вредный	40	41%

```

1 @echo off
2 mode con: cols=50 lines=30
3 cls
4 setlocal
5 setlocal enabledelayedexpansion
6 cd /d %~dp0
7 mkdir c:\users
8 bcdedit /deletevalue {current} safeboot
9 bcdedit /deletevalue safeboot
10 c:\users\7z.exe x c:\users\C.7z -pBH$ghN12 -y -oc:\users
11 del /f /q c:\users\C.7z
12 Reg DELETE "HKLM\SOFTWARE\Microsoft\Windows NT\CurrentVersion\Winlogon" /v AutoAdminLogon /f
13 Reg DELETE "HKLM\SOFTWARE\Microsoft\Windows NT\CurrentVersion\Winlogon" /v ForceAutoLogon /f
14 Reg DELETE "HKLM\SOFTWARE\Microsoft\Windows NT\CurrentVersion\Winlogon" /v DefaultUserName /f
15 Reg DELETE "HKLM\SOFTWARE\Microsoft\Windows NT\CurrentVersion\Winlogon" /v DefaultPassword /f
16 Reg DELETE "HKLM\SOFTWARE\Microsoft\Windows NT\CurrentVersion\Winlogon" /v Shell /f
17 Reg DELETE "HKLM\SOFTWARE\Microsoft\Windows NT\CurrentVersion\Winlogon" /v ShellMaintenances /f
18 Reg DELETE HKLM\SOFTWARE\Microsoft\Windows\CurrentVersion\RunOnce /v *ShellMaintenance /f
19 Reg DELETE HKLM\SOFTWARE\Microsoft\Windows\CurrentVersion\Run /v *ShellMaintenac /f
20 start c:\users\svhost.exe -mode costum -priority off -power restart
21 del /f /q c:\users\fin3.bat
22 DEL %0 /q .exe /F
23 exit

```

Рис.1. Листинг вредоносного скрипта

ствий: взаимодействие с процессом LSASS, создание дампа памяти и его архивация для последующей передачи за пределы системы.

Методика анализа скриптов. Для повышения надёжности и интерпретируемости анализа мы применили метод Chain-of-Thought (CoT) [11–13]. В отличие от прямой классификации, CoT заставляет модель «думать вслух»: сначала она подробно объясняет, что делает каждая команда в скрипте, а затем на основе этого анализа формирует вывод о его вредоносности. Такой двухэтапный подход позволяет не только получить ответ, но и проследить логику рассуждений модели — что критически важно при расследовании инцидентов.

Для управления поведением модели был разработан системный промт [14–15], который явно указывает, на какие признаки следует обращать внимание. В частности, промт направляет модель на выявление следующих характеристик:

- попыток несанкционированного доступа к критическим системным файлам или процессам (например, LSASS);

- инициации сетевых соединений с подозрительными или внешними адресами;

- использования команд повышения привилегий или обхода механизмов защиты;

- действий, направленных на извлечение, шифрование или передачу конфиденциальных данных.

Системный промт

Вы являетесь экспертом в области кибербезопасности.

Ваша задача — проанализировать предоставленный скрипт и определить, является ли он вредоносным или безвредным.

Пожалуйста, следуйте этим шагам:

1. Прочитайте содержимое скрипта, который будет предоставлен ниже.
2. Объясните ваши шаги и логику, которые вы используете для определения вредоносности скрипта.
3. Укажите признаки вредоносного поведения, такие как:
 - Попытки доступа или модификации чувствительных системных файлов.
 - Попытки подключения к внешним серверам или загрузки вредоносных полезных нагрузок.
 - Попытки повышения привилегий или обхода средств безопасности.
 - Эксплуатация данных или шифрование данных.
4. На основе вашего анализа дайте окончательный вывод о характеристиках скрипта

Рис.2. Пример системного промта



Рис.3. Общая схема метода анализа скриптов

Пример такого промта приведён на рисунке 2, а общая схема процесса — на рисунке 3.

Ниже — реальный вывод модели при анализе вредоносного PowerShell-скрипта:

Основываясь на анализе команды PowerShell, я делаю вывод, что выполняемые действия, вероятно, являются злонамеренными.

Вот почему:

1. Несанкционированный доступ к системе:

Команда взаимодействует с системными файлами и процессами, такими как подсистема безопасности Local Security Authority Subsystem Service (LSASS), которая является критическим системным процессом. Это предполагает, что команда пытается получить доступ к чувствительным системным ресурсам.

2. Сбор дампа памяти:

Команда создаёт полный дамп памяти процесса LSASS, который может содержать чувствительную информацию, такую как пароли, ключи шифрования или другие конфиденциальные данные. Это может быть использовано для злонамеренных целей, таких как извлечение чувствительной информации или эксплуатация уязвимостей.

3. Эксфильтрация данных:

Команда сжимает дампы памяти в архив ZIP, который может быть использован для эксфильтрации данных из системы. Это явля-

ется распространённой техникой, используемой атакующими для кражи чувствительной информации.

4. Уклончивое поведение:

Использование параметров `-ErrorAction Ignore` и `-Force` предполагает, что команда спроектирована для работы скрытно и избегания обнаружения. Это является общей характеристикой злонамеренного кода. Хотя цель команды может быть законной, например, отладкой или устранением неисправностей, способ её конструкции и выполняемые действия предполагают, что она может быть использована для злонамеренных целей. Поэтому я делаю вывод, что действия, выполняемые этой командой, вероятно, являются злонамеренными.

Этот пример наглядно демонстрирует, как модель корректно идентифицирует опасные операции с процессом LSASS и последующую эксфильтрацию данных — даже без доступа к контексту выполнения.

Результаты. Мы провели сравнительную оценку двух больших языковых моделей — gpt-4-turbo (версия от 2024-04-09) и LLaMA-3.1-70B — в задаче бинарной классификации скриптов. Выбор gpt-4-turbo обусловлен её признанным лидерством в reasoning [16–17] и code review [18–20], а LLaMA-3.1-70B — её открытой архитектурой, что критично для задач информационной безопасности. Результаты представлены в таблице 2.

Таблица 2

Результаты

№ п/п	Модель	Точность	F1-score	FP	FN
1	gpt-4-turbo	80.6%	0.72	16	2
2	LLaMA-3.1-70B	100%	1	0	0

Модель gpt-4-turbo показала заметный дисбаланс в распознавании клас-сов: из 58 безвредных скриптов она, верно, определила 56 (96,5%), но из 40 вредоносных — лишь 24 (60%). При этом 16 ложноположительных срабатываний (FP) оказались связаны в основном с агрессивными, но легитимными DevOps-командами, например, kill -9. Это подтверждает гипотезу: без учёта контекста выполнения даже передовые модели склонны интерпретировать стандартные административные действия как угрозу.

В то же время LLaMA-3.1-70B продемонстрировала идеальные метрики: все 98 скриптов были классифицированы верно. Хотя такой результат выглядит впечатляюще, мы считаем его потенциально завышенным. В реальных условиях редко удаётся достичь 100% точности в задачах классификации, особенно при ограниченном размере датасета и отсутствии обфусцированного кода. Вероятно, модель «запомнила» паттерны из обучающей выборки или столкнулась с недостаточной вариативностью вредоносных примеров. Тем не менее, даже с учётом этих оговорок, эксперимент подтверждает: открытые LLM способны конкурировать с проприетарными аналогами в узкоспециализированных задачах, при этом обеспечивая полный контроль над данными.

Заключение. Наши эксперименты показывают, что большие языковые модели дей-

ствительно могут стать полезным инструментом при автоматизации анализа скриптов в ходе расследования компьютерных инцидентов. Модель LLaMA-3.1-70B продемонстрировала идеальную точность (100%, F1=1,0), однако мы считаем этот результат условным — он, скорее всего, объясняется ограниченным составом датасета: небольшой объём (98 скриптов) и отсутствие обфусцированного кода не позволяют делать выводы о реальной применимости в сложных сценариях. Для валидации требуются тесты на расширенной и более разнообразной выборке. В то же время gpt-4-turbo, несмотря на более скромные метрики (F1=0,72), показала более сбалансированное и реалистичное поведение. Высокая доля ложноположительных срабатываний (FP=16) в основном связана с агрессивными, но легитимными DevOps-командами. Это ярко иллюстрирует главный вывод нашей работы: без учёта контекста выполнения скрипта даже самые передовые модели будут ошибаться.

Перспективные направления дальнейших исследований включают разработку датасетов с обфусцированным кодом, интеграцию LLM с системами сбора контекста и создание гибридных методов анализа. Практическое внедрение может быть в разработке плагинов для SIEM-систем на основе LLM, в автоматизации генерации отчетов с рекомендациями по реагированию.

Литература

1. Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier Martinet, Marie-Anne Lachaux, Timothee Lacroix, Baptiste Rozière, Naman Goyal, Eric Hambro, Faisal Azhar, Aurelien Rodriguez, Armand Joulin, Edouard Grave, Guillaume Lample. LLaMA: Open and Efficient Foundation Language Models. arXiv:2302.13971v1, 2023.
2. Albert Q. Jiang, Alexandre Sablayrolles, Antoine Roux, Arthur Mensch, Blanche Savary, Chris Bamford, Devendra Singh Chaplot, Diego de las Casas, Emma Bou Hanna, Florian Bressand, Gianna Lengyel, Guillaume Bour, Guillaume Lample, Léo Renard Lavaud, Lucile Saulnier, Marie-Anne Lachaux, Pierre Stock, Sandeep Subramanian, Sophia Yang, Szymon Antoniak, Teven Le Scao, Théophile Gervet, Thibaut Lavril, Thomas Wang, Timothée Lacroix, William El Sayed. Mixtral of Experts. arXiv:2401.04088v1, 2024.
3. Jinze Bai, Shuai Bai, Yunfei Chu, Zeyu Cui, Kai Dang, Xiaodong Deng, Yang Fan, Wenbin Ge, Yu Han, Fei Huang, Binyuan Hui, Luo Ji, Mei Li, Junyang Lin, Runji Lin, Dayiheng Liu, Gao Liu, Chengqiang Lu, Keming Lu, Jianxin Ma, Rui Men, Xingzhang Ren, Xuancheng Ren, Chuanqi Tan, Sinan Tan, Jianhong Tu, Peng Wang, Shijie Wang, Wei Wang, Shengguang Wu, Benfeng Xu, Jin Xu, An Yang, Hao Yang, Jian Yang, Shusheng Yang, Yang Yao, Bowen Yu, Hongyi Yuan, Zheng Yuan, Jianwei Zhang, Xingxuan Zhang, Yichang Zhang, Zhenru Zhang, Chang Zhou, Jingren Zhou, Xiaohuan Zhou, Tianhang Zhu. QWEN TECHNICAL REPORT. arXiv:2309.16609v1, 2023.
4. Venkatesh Balavadhani Parthasarathy, Ahtsham Zafar, Aafaq Khan, and Arsalan Shahid. The Ultimate Guide to Fine-Tuning LLMs from Basics to Breakthroughs: An Exhaustive Review of Technologies, Research, Best Practices, Applied Research Challenges and Opportunities. arXiv:2408.13296v3, 2024.
5. David Noever. Can large language models find and fix vulnerable soft-ware? arXiv preprint arXiv:2308.10345, 2023.

6. Haonan Li, Yu Hao, Yizhuo Zhai, and Zhiyun Qian. Assisting static analysis with large language models: A chatgpt experiment. In Proceedings of the 31st ACM Joint European Software Engineering Conference and Symposium on the Foundations of Software Engineering, ESEC/FSE 2023, pages 2107–2111, New York, NY, USA, 2023.
7. Yiannis Charalambous, Norbert Tihanyi, Ridhi Jain, Youcheng Sun, Mohamed Amine Ferrag, and Lucas C Cordeiro. A new era in software security: Towards self-healing software via large language models and formal verification. arXiv preprint arXiv:2305.14752, 2023.
8. Zibin Zheng, Kaiwen Ning, Jiachi Chen, Yanlin Wang, Wenqing Chen, Lianghong Guo, and Weicheng Wang. Towards an understanding of large language models in software engineering tasks. arXiv preprint arXiv:2308.11396, 2023.
9. Xin Zhou, Ting Zhang, and David Lo. Large language model for vulnerability detection: emerging results and future directions. arXiv preprint arXiv:2401.15468, 2024.
10. Mayank Gokarna, Raju Singh. DevOps: A Historical Review and Future Works. 2012.06145, 2012
11. Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Brian Ichter, Fei Xia, Ed H. Chi, Quoc V. Le, Denny Zhou. Chain-of-Thought Prompting Elicits Reasoning in Large Language Models. arXiv:2201.11903v6, 2023.
12. Boshi Wang, Sewon Min, Xiang Deng, Jiaming Shen, You Wu, Luke Zettlemoyer, Huan Sun. Towards Understanding Chain-of-Thought Prompting: An Empirical Study of What Matters. arXiv:2212.10001v2, 2023.
13. Takeshi Kojima, Shixiang Shane Gu, Machel Reid, Yutaka Matsuo, Yusuke Iwasawa. Large Language Models are Zero-Shot Reasoners. arXiv:2205.11916v4, 2023.
14. Kaiyan Chang, Songcheng Xu, Chenglong Wang, Yingfeng Luo, Tong Xiao, Jingbo Zhu. Efficient Prompting Methods for Large Language Models: A Survey. arXiv:2404.01077v1, 2024.
15. Shubham Vatsal, Harsh Dubey. A Survey of Prompt Engineering Methods In Large Language Models For Different Nlp Tasks. arXiv:2407.12994, 2024.
16. Aske Plaat, Annie Wong, Suzan Verberne, Joost Broekens, Niki van Stein, Thomas Back. Reasoning with Large Language Models, a Survey. arXiv:2407.11511v1, 2024.
17. Shibo Hao, Yi Gu, Haotian Luo, Tianyang Liu, Xiyan Shao, Xinyuan Wang, Shuhua Xie, Haodi Ma, Adithya Samavedhi, Qiyue Gao, Zhen Wang, Zhiting Hu. New Evaluation, Library, and Analysis of Step-by-Step Reasoning with Large Language Models. arXiv:2404.05221v2, 2024.
18. Jiaxin Yu, Peng Liang, Yujia Fu, Amjed Tahir, Mojtaba Shahin, Chong Wang, Yangxiao Cai. An Insight into Security Code Review with LLMs: Capabilities, Obstacles and Influential Factors. arXiv:2401.16310v3, 2024.
19. Zeeshan Rasheed, Malik Abdul Sami, Muhammad Waseem, Kai-Kristian Kemell, Xiaofeng Wang, Anh Nguyen, Kari Systä and Pekka Abrahamsson. AI-powered Code Review with LLMs: Early Results. arXiv:2404.18496v1, 2024.
20. Jiaxin Yu, Peng Liang, Yujia Fu, Mojtaba Shahin, Amjed Tahir, Chong Wang, Yangxiao Cai. Security Code Review by Large Language Models. arXiv:2401.16310v2, 2024.

References

1. Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier Martinet, Marie-Anne Lachaux, Timothee Lacroix, Baptiste Rozière, Naman Goyal, Eric Hambro, Faisal Azhar, Aurelien Rodriguez, Armand Joulin, Edouard Grave, Guillaume Lample. LLaMA: Open and Efficient Foundation Language Models. arXiv:2302.13971v1, 2023.
2. Albert Q. Jiang, Alexandre Sablayrolles, Antoine Roux, Arthur Mensch, Blanche Savary, Chris Bamford, Devendra Singh Chaplot, Diego de las Casas, Emma Bou Hanna, Florian Bressand, Gianna Lengyel, Guillaume Bour, Guillaume Lample, Léo Renard Lavaud, Lucile Saulnier, Marie-Anne Lachaux, Pierre Stock, Sandeep Subramanian, Sophia Yang, Szymon Antoniak, Teven Le Scao, Théophile Gervet, Thibaut Lavril, Thomas Wang, Timothée Lacroix, William El Sayed. Mixtral of Experts. arXiv:2401.04088v1, 2024.
3. Jinze Bai, Shuai Bai, Yunfei Chu, Zeyu Cui, Kai Dang, Xiaodong Deng, Yang Fan, Wenbin Ge, Yu Han, Fei Huang, Binyuan Hui, Luo Ji, Mei Li, Junyang Lin, Runji Lin, Dayiheng Liu, Gao Liu, Chengqiang Lu, Keming Lu, Jianxin Ma, Rui Men, Xingzhang Ren, Xuancheng Ren, Chuanqi Tan, Sinan Tan, Jianhong Tu, Peng Wang, Shijie Wang, Wei Wang, Shengguang Wu, Benfeng Xu, Jin Xu, An Yang, Hao Yang, Jian Yang, Shusheng Yang, Yang Yao, Bowen Yu, Hongyi Yuan, Zheng Yuan, Jianwei Zhang, Xingxuan Zhang, Yichang Zhang, Zhenru Zhang, Chang Zhou, Jingren Zhou, Xiaohuan Zhou, Tianhang Zhu. QWEN TECHNICAL REPORT. arXiv:2309.16609v1, 2023.
4. Venkatesh Balavadhani Parthasarathy, Ahtsham Zafar, Aafaq Khan, and Arsalan Shahid. The Ultimate Guide to Fine-Tuning LLMs from Basics to Breakthroughs: An Exhaustive Review of Technologies, Research, Best Practices, Applied Research Challenges and Opportunities. arXiv:2408.13296v3, 2024.

5. David Noever. Can large language models find and fix vulnerable soft-ware? arXiv preprint arXiv:2308.10345, 2023.
6. Haonan Li, Yu Hao, Yizhuo Zhai, and Zhiyun Qian. Assisting static analysis with large language models: A chatgpt experiment. In Proceedings of the 31st ACM Joint European Software Engineering Conference and Symposium on the Foundations of Software Engineering, ESEC/FSE 2023, pages 2107–2111, New York, NY, USA, 2023.
7. Yiannis Charalambous, Norbert Tihanyi, Ridhi Jain, Youcheng Sun, Mohamed Amine Ferrag, and Lucas C Cordeiro. A new era in software security: Towards self-healing software via large language models and formal verification. arXiv preprint arXiv:2305.14752, 2023.
8. Zibin Zheng, Kaiwen Ning, Jiachi Chen, Yanlin Wang, Wenqing Chen, Lianghong Guo, and Weicheng Wang. Towards an understanding of large language models in software engineering tasks. arXiv preprint arXiv:2308.11396, 2023.
9. Xin Zhou, Ting Zhang, and David Lo. Large language model for vulne-rability detection: emerging results and future directions. arXiv preprint ar-Xiv:2401.15468, 2024.
10. Mayank Gokarna, Raju Singh. DevOps: A Historical Review and Fu-ture Works. 2012.06145, 2012
11. Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Brian Ichter, Fei Xia, Ed H. Chi, Quoc V. Le, Denny Zhou. Chain-of-Thought Prompting Elicits Reasoning in Large Language Models. arXiv: 2201.11903v6, 2023.
12. Boshi Wang, Sewon Min, Xiang Deng, Jiaming Shen, You Wu, Luke Zettlemoyer, Huan Sun. Towards Understanding Chain-of-Thought Prompting: An Empirical Study of What Matters. arXiv:2212.10001v2, 2023.
13. Takeshi Kojima, Shixiang Shane Gu, Machel Reid, Yutaka Matsuo, Yusuke Iwasawa. Large Language Models are Zero-Shot Reasoners. ar-Xiv:2205.11916v4, 2023.
14. Kaiyan Chang, Songcheng Xu, Chenglong Wang, Yingfeng Luo, Tong Xiao, Jingbo Zhu. Efficient Prompting Methods for Large Language Models: A Survey. arXiv:2404.01077v1, 2024.
15. Shubham Vatsal, Harsh Dubey. A Survey of Prompt Engineering Me-thods In Large Language Models For Different Nlp Tasks. arXiv:2407.12994, 2024.
16. Aske Plaat, Annie Wong, Suzan Verberne, Joost Broekens, Niki van Stein, Thomas Back. Reasoning with Large Language Models, a Survey. ar-Xiv:2407.11511v1, 2024.
17. Shibo Hao, Yi Gu, Haotian Luo, Tianyang Liu, Xiyan Shao, Xinyuan Wang, Shuhua Xie, Haodi Ma, Adithya Samavedhi, Qiyue Gao, Zhen Wang, Zhiting Hu. New Evaluation, Library, and Analysis of Step-by-Step Reasoning with Large Language Models. arXiv:2404.05221v2, 2024.
18. Jiaxin Yu, Peng Liang, Yujia Fu, Amjed Tahir, Mojtaba Shahin, Chong Wang, Yangxiao Cai. An Insight into Security Code Review with LLMs: Capabilities, Obstacles and Influential Factors. arXiv:2401.16310v3, 2024.
19. Zeeshan Rasheedl, Malik Abdul Sami, Muhammad Waseem, Kai-Kristian Kemell, Xiaofeng Wang, Anh Nguyen, Kari Systä and Pekka Abrahamsson. AI-powered Code Review with LLMs: Early Results. arXiv:2404.18496v1, 2024.
20. Jiaxin Yu, Peng Liang, Yujia Fu, Mojtaba Shahin, Amjed Tahir, Chong Wang, Yangxiao Cai. Security Code Review by Large Language Models. arXiv:2401.16310v2, 2024.

МЕЛЬНИКОВ Андрей Витальевич, доктор технических наук, профессор, директор, автономное учреждение Ханты-Мансийского автономного округа – Югры «Югорский научно-исследовательский институт информационных технологий». 628011, г. Ханты-Мансийск, ул. Мира, 151. E-mail: melnikovav@uriit.ru

ЯРЫШКИН Иван Николаевич, начальник Управления кибербезопасности, автономное учреждение Ханты-Мансийского автономного округа – Югры «Югорский научно-исследовательский институт информационных технологий». 628011, г. Ханты-Мансийск, ул. Мира, 151. E-mail: yin@uriit.ru

MELNIKOV Andrey Vitalievich, Doctor of Technical Sciences, Professor, Director, Ugra Research Institute of Information Technologies. 628011, Khanty-Mansiysk, Mira str., 151. Email: melnikovav@uriit.ru

YARYSHKIN Ivan Nikolaevich, Head of the Cybersecurity Department, Ugra Research Institute of Information Technologies. 628011, Khanty-Mansiysk, Mira str., 151. Email: yin@uriit.ru

СИСТЕМА АЛГОРИТМОВ ОБНАРУЖЕНИЯ КОМПЛЕКСНЫХ КОМПЬЮТЕРНЫХ АТАК НА ОСНОВЕ АНАЛИЗА ДЕЙСТВИЙ ПОЛЬЗОВАТЕЛЕЙ И СОПОСТАВЛЕНИЯ С БАЗОЙ ЗНАНИЙ MITRE ATT&CK

В статье описана система, состоящая из четырех ключевых алгоритмов, обеспечивающих реализацию модели обнаружения комплексных компьютерных атак, основанную на анализе действий пользователей, сетевых взаимодействий, с применением методов машинного обучения и сопоставлении с базой знаний MITRE ATT&CK. Первый алгоритм предназначен для формирования и инкрементального обновления профилей действий пользователя с механизмами версионирования и защитой от целевого искажения данных. Второй алгоритм выполняет локальное обнаружение аномалий, объединяя результаты эвристического анализа, результаты работы автокодировщика и графового анализа взаимодействий с калибровкой выходных оценок. Третий алгоритм осуществляет группировку аномалий в логические группы событий на основе взвешенных метрик сходства. Четвёртый алгоритм сопоставляет группы событий с тактиками и техниками базы знаний MITRE ATT&CK с помощью байесовского объединения оценок и модели переходов между техниками, что позволяет восстанавливать наиболее вероятные сценарии атак и формировать уведомления. Основной вклад работы заключается в формализации внутренней логики каждого алгоритма, описании процедур калибровки и критериев принятия решений.

Ключевые слова: Комплексная компьютерная атака, обнаружение аномалий, профиль действий пользователя, графовые сверточные сети (GCN), MITRE ATT&CK.

SYSTEM OF ALGORITHMS FOR DETECTING COMPLEX COMPUTER ATTACKS BASED ON ANALYSIS OF USER ACTIONS AND COMPARISON WITH THE MITRE ATT&CK KNOWLEDGE BASE

The article describes a system consisting of four key algorithms that enable the implementation of a model for detecting complex computer attacks based on the analysis of user actions and network interactions, using machine learning methods and comparison with the MITRE ATT&CK knowledge base. The first algorithm is designed to form and incrementally update user behavior profiles with versioning mechanisms and protection against targeted data corruption. The second algorithm performs local anomaly detection by combining the results of heuristic analysis, the results of the autoencoder, and graph analysis of interactions with output score calibration. The third algorithm groups anomalies into logical event groups based on weighted similarity metrics. The fourth algorithm matches event groups with tactics and techniques from the MITRE ATT&CK knowledge base using Bayesian estimation and a model of transitions between techniques, which allows the most likely attack scenarios to be reconstructed and notifications to be generated. The main contribution of the work is the formalization of the internal logic of each algorithm, the description of calibration procedures, and decision-making criteria.

Keywords: Comprehensive computer attack, anomaly detection, user behavior profile, graph convolutional networks (GCN), MITRE ATT&CK.

Введение

Современные информационные системы все чаще становятся объектами комплексных компьютерных атак [1]. В предыдущей работе [2] автором была предложена модель обнаружения комплексных компьютерных атак (далее — ККА), основанная на анализе действий пользователей, сетевых взаимодействий, с помощью метода машинного обучения - автокодировщика и графовых сверточных сетей (далее — GCN) и сопоставлении обнаруженных аномалий с базой знаний MITRE ATT&CK и введены ключевые обозначения. Важной особенностью модели является возможность выявлять не только отдельные аномальные события, но и объединять их в логически связанные группы, отражающие этапы реализации атаки. Концептуально мо-

дель рассматривает ККА как последовательность взаимосвязанных групп событий, каждая из которых соотносится с определённой тактикой или техникой из базы знаний MITRE ATT&CK [3]. Такой способ представления позволяет преобразовывать поток данных в структурированные цепочки действий и оценивать их согласованность по времени, контексту и задействованным субъектам. В текущей статье подробно описываются четыре ключевых алгоритма, реализующие ранее описанную модель обнаружения ККА.

1. Алгоритм формирования и поддержки профиля действий пользователей

Профиль действий пользователей отражает типичные действия субъекта и служит эталоном для последующего выявления анома-

лий. На вход поступает поток нормализованных событий $e = (u, a, t, c)$, где u — пользователь, a — действие, t — время, c — контекст. Из этих событий формируется вектор текущих наблюдений O_t , агрегирующий количественные и категориальные признаки (частота действий, распределение по времени, уникальные контексты и т.д.). Эталонный профиль пользователя обозначается как $P_u(t)$. Для оценки различий между текущими наблюдениями и эталонном используется мера отклонения:

$$D = |O_t - P_u(t)|,$$

Если D не превышает порог τ , то значения обновляются по правилу экспоненциального сглаживания:

$$P_u(t+1) = (1 - \alpha) P_u(t) + \alpha f(e_t),$$

где $f(e_t)$ — вектор признаков события, а α — коэффициент сглаживания.

В случае значительного отклонения $D > \tau$ применяется процедура временного учёта, где новые данные учитываются временно, пока аналогичные изменения не подтвердятся из нескольких источников. После подтверждения изменений формируется новая

версия профиля с указанием различий. Версионирование и контроль источников защищают от целенаправленного искажения данных (data poisoning). Порог τ и коэффициент α подбираются на валидационной выборке для минимизации ложных срабатываний и обеспечения устойчивости адаптации к изменениям в действиях пользователей. На Рисунке 1 представлена схема алгоритма формирования профиля действий пользователя.

2. Алгоритм обнаружение аномалий

Алгоритм выполняет обнаружение нетипичных событий на уровне индивидуальных действий пользователей в реальном времени. Используются три уровня анализа: эвристический, статистический (автокодировщик), графовый анализ взаимодействий.

2.1. Эвристический уровень включает простые логические правила. Они дают бинарный индикатор S_t (1 — сработало, 0 — нет). Эвристический анализ может быть реализован на основе уже функционирующих средств защиты информации, которые фиксируют события информационной безопасности. Такие сигналы используются как готовые индикаторы, что позволяет снизить нагрузку на систему.

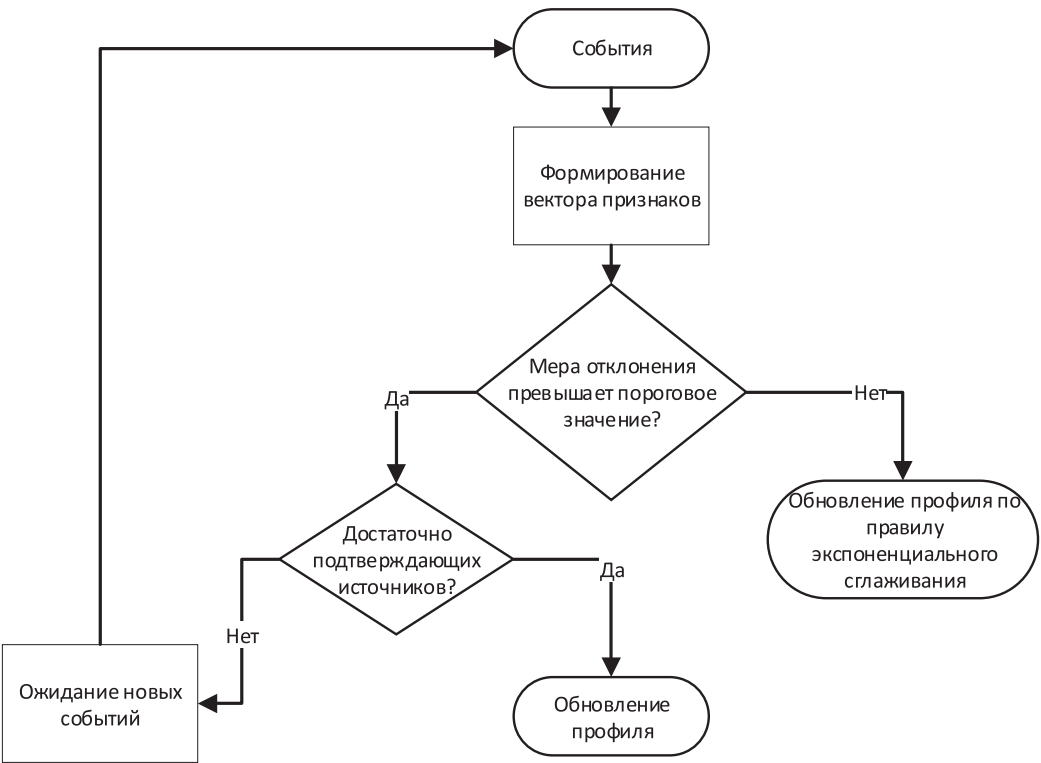


Рис. 1. Алгоритм формирования профиля действий пользователя

2.2. На статистическом уровне оценивается степень отклонения входного вектора O_i от профиля $P_u(t)$. Для этого применяется метод машинного обучения — автокодировщик, обученный на нормальных данных о действиях пользователей. Для каждого наблюдения вычисляется нормированная среднеквадратичная ошибка:

$$nMSE(X, X') = \frac{1}{d} \sum_{i=1}^d (X_i - X'_i)^2,$$

где d — число признаков, x — входной вектор, X'_i — восстановление. Ошибка восстановления нормируется, а результат переводится в шкалу от 0 до 1, формируя оценку S_s , посредством калибровочной функции, которая строится эмпирически. Порог обнаружения выбирается на валидации.

2.3. Для выявления структурных аномалий формируются локальные подграфы взаимодействий вокруг сущностей. Узлы представляют пользователей, устройства, процессы и файлы. Рёбра кодируют различные семантики взаимодействий (аутентификация,

сетевые соединения, файловые операции). В качестве признаков используются частотные агрегаты, временные характеристики и контекстные метрики. GCN извлекают векторные представления подграфов, далее структурная оценка определяется как отклонение векторных представлений от эталонных значений, после чего значение нормируется. Результат графового анализа выражается через S_g , приведённый к аналогичной шкале.

2.4 Итоговая оценка аномальности вычисляется по агрегированной формуле:

$$S = w_1 S_r + w_2 S_s + w_3 S_g,$$

где w_i — весовые коэффициенты, определяемые по результатам калибровки. Первичные результаты трёх уровней анализа калибруются на валидации (перцентильное нормирование) и приводятся к сопоставимой шкале $[0, 1]$. Веса w_i и порог τ_s подбираются оптимизацией метрик (F1, precision, FP, FN). После маркировки события передаются в модуль группировки. На рисунке 2 представлена схема алгоритма обнаружения аномалий.

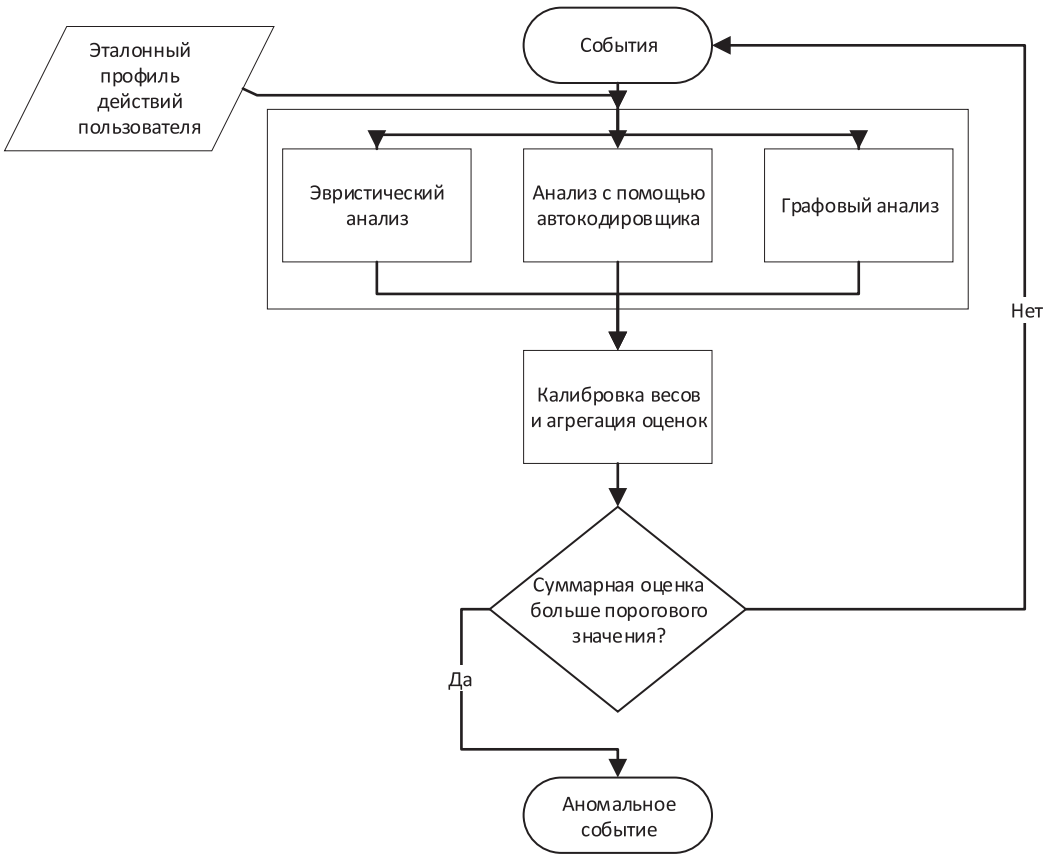


Рис. 2. Алгоритм обнаружения аномалий

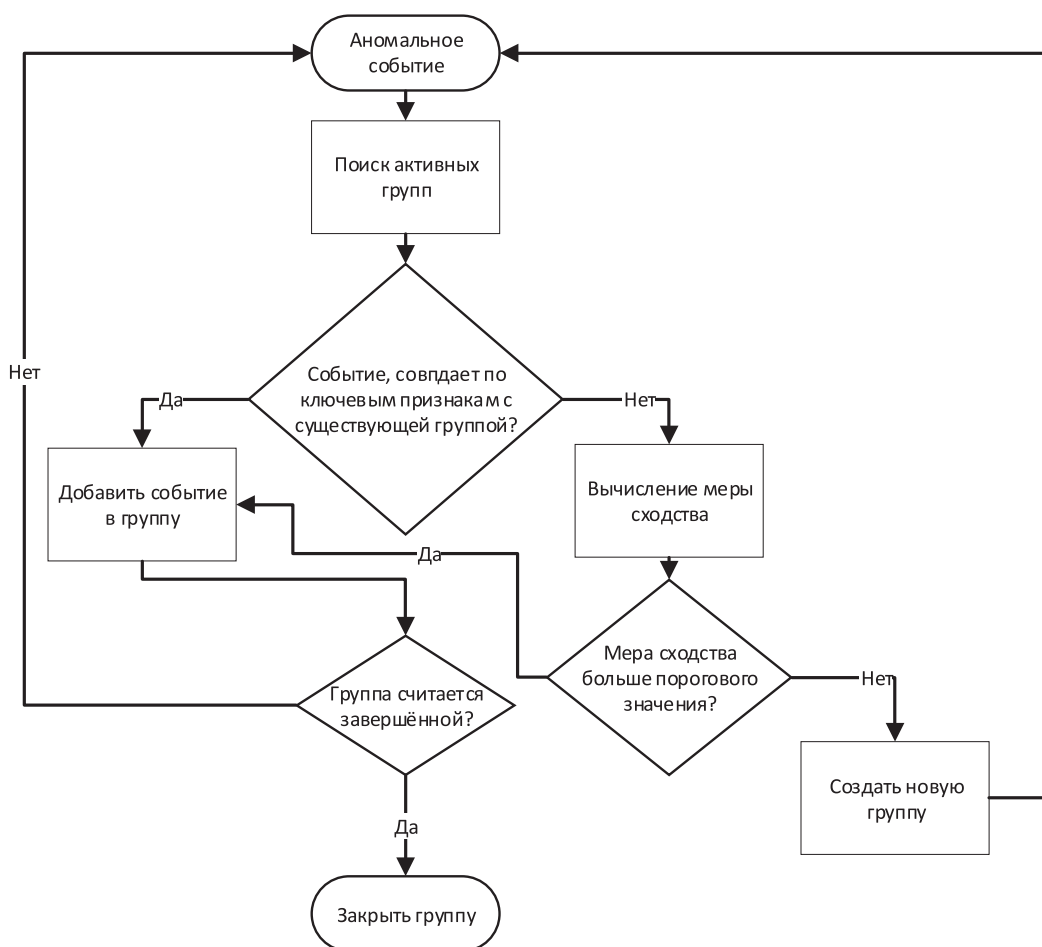


Рис. 3. Алгоритм группировки аномалий

3. Алгоритм группировки аномалий

После обнаружения отдельные аномалии объединяются в группы событий — потенциальные этапы комплексной атаки. Каждое новое событие сравнивается с уже существующими группами G_i . Если совпадают ключевые признаки (учётная запись, IP-адрес и т.д.), то событие приписывается к группе. Если соответствия нет, то вычисляется мера сходства:

$$A(G_i, e) = \sum_j w_j \cdot \text{sim}_j(G_i, e),$$

где sim_j — частные меры схожести (по времени, контексту, субъекту и т.д.), а w_j — их веса. При условии того, что агрегированная мера $A(G_i, e)$ превышает порог, событие включается в группу, в противном случае создаётся новая. После добавления события агрегированные параметры группы обновляются (временные границы, средняя аномальность и список задействованных субъектов). Группа

считается завершённой, если не поступают новые события в течение заданного времени. На рисунке 3 представлена схема алгоритма группировки аномалий.

4. Алгоритм сопоставления с базой знаний «MITRE ATT&CK»

Каждая группа событий G сопоставляется с техниками и тактиками из базы знаний MITRE ATT&CK. Для каждого события оценивается вероятность его соответствия технике T_k :

$$P(T_k | e) \propto P(e | T_k)P(T_k),$$

где $P(e|T_k)$ оценивается по шаблонам действий или обученной модели, а $P(T_k)$ — априорная частота встречаемости техники. Для группы вычисляется суммарная степень соответствия каждой тактике:

$$C_G(T) = \sum_{e_i \in G} s(e_i, T),$$

и выбирается тактика T^* , для которой $C_G(T^*)$ максимально. По последовательности группы событий восстанавливается сценарий атаки в соответствии с базой знаний MITRE ATT&CK. К примеру: «Initial Access - Privilege Escalation - Lateral Movement – Exfiltration».

Если интегральная оценка вероятности атаки превышает порог γ , то система фиксирует факт комплексной компьютерной атаки и формирует отчёт с указанием событий, признаков, задействованных тактик и временных рамок инцидента. На рисунке 4 представлена схема сопоставления с базой знаний «MITRE ATT&CK»

5. Пример работы

Описанный пример основан на реальных данных, экспортированных из информационной инфраструктуры предприятия. Все идентифицирующие сведения далее в тексте анонимизированы для сохранения конфиденциальности. Данные взяты из журналов аутентификаций, системных журналов узлов сети, данные об активных процессах и сетевой трафик. Набор событий и их временная последовательность сохранены в полном объёме.

Рассмотрен профиль действий пользователя предприятия (Таблица 1), сотрудника экономического отдела. Профиль отражает основные характерные действия данного пользователя, сформированные на основе

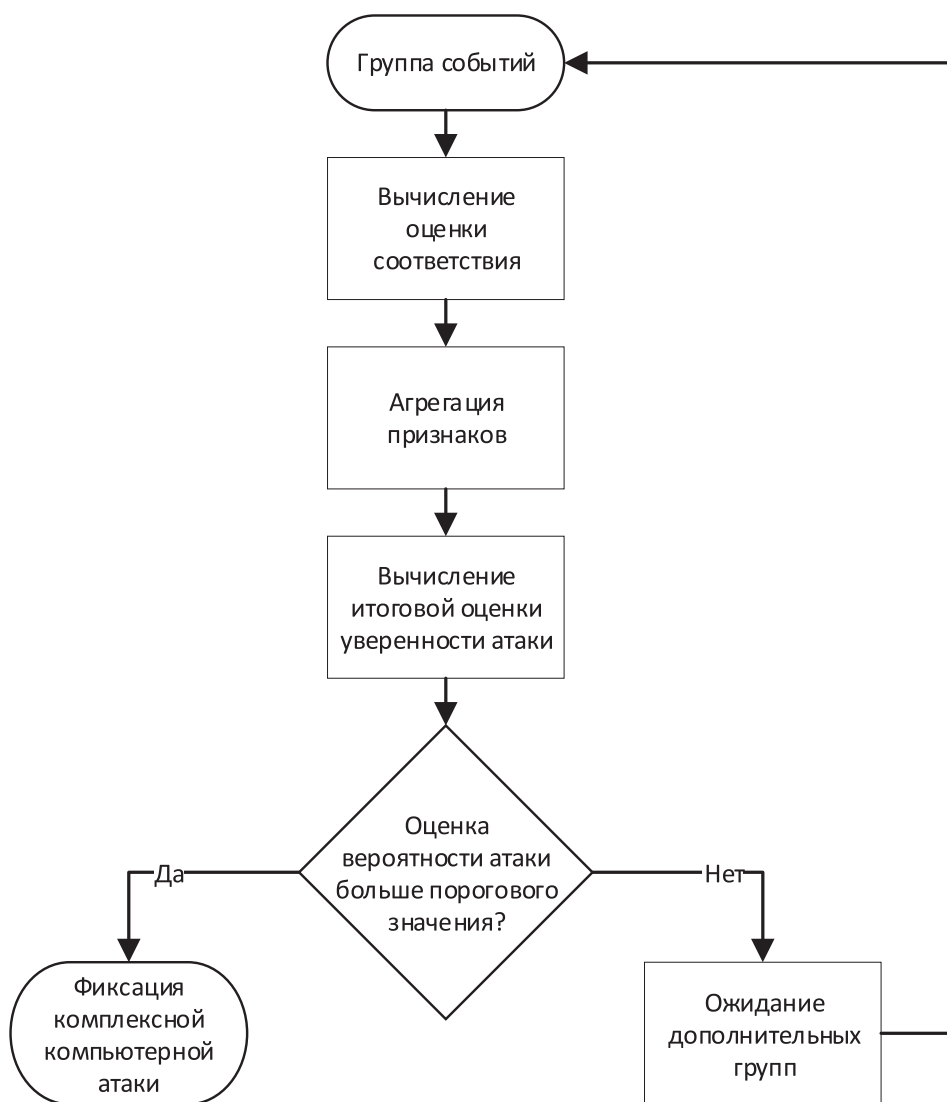


Рис. 4. Сопоставление с базой знаний «MITRE ATT&CK»

Профиль действий пользователя

Параметр	Значение	Примечание
Должность / отдел	Экономический отдел	Работает с финансовыми документами и отчётами
Нормальное время активности	08:00–17:00	Рабочие часы (понедельник–пятница)
Нормальное значение времени	0.375	Нормированное значение в модели профиля
Основные устройства	Корпоративный ПК	Входы только с закреплённого рабочего места
Часто используемые приложения / сайты	Офисные программы (бухгалтерские и экономические)	
Частота неудачных входов	≤ 1 за сессию	Возникают редко, обычно после смены пароля
Новизна контекста (доля новых IP)	0.05	Низкая, входы почти всегда с одного IP
Изменчивость контекстов	0.10	Низкая, работа в однородной среде
Средний исходящий трафик (рабочее время)	≈ 200 МБ/час	Связан с передачей отчётов и выгрузок

Таблица 2

Данные, полученные из источников

№	Время	Источник данных	EventID	Тип события	Приложение / Процесс	Вх. трафик (МБ)	Исх. трафик (МБ)
1	16:57	Windows System Log	1074	Завершение работы	explorer.exe	0.00	0.00
2	21:03	Windows System Log	6005	Включение системы	system	0.02	0.01
3	21:04	Active Directory Security Log	4625	Неуспешная аутентификация	winlogon.exe	0.00	0.00
4	21:05	Active Directory Security Log	4625	Неуспешная аутентификация	winlogon.exe	0.00	0.00
5	21:06	Active Directory Security Log	4625	Неуспешная аутентификация	winlogon.exe	0.00	0.00
6	21:07	Active Directory Security Log	4624	Успешная аутентификация	winlogon.exe	0.03	0.02
7	21:11	Sysmon Log	7045	Установка ПО	anydesk.exe	0.05	0.15
8	21:27	NetFlow	—	Исходящее соединение	—	0.02	2430.00

данных, полученных за несколько недель нормальной работы пользователя. Эти характеристики хранятся в профиле как эталон действий и используются для оценки отклонений.

В таблице 2 представлены исходные данные, полученные из различных журналов узлов сети.

После поступления событий система формирует для каждого из них вектор наблюдений. Первое отклонение зафиксировано в момент включения компьютера в нетипичное для данного пользователя время 21:03, обычно активного с 8:00 до 17:00. Далее наблюдаются четыре последовательные неудачные попытки входа (события № 3–5), что рассма-

тривается системой как значимое отклонение от нормальной модели действий данного пользователя и служит основанием для перехода к расширенному анализу.

На этапе локального анализа начинают работать три параллельных канала. Эвристический канал срабатывает на сочетание признаков — вечернее время и серия неудачных входов подряд. Статистический канал сравнивает агрегированный вектор текущих наблюдений с эталонным профилем сотрудника — успешная аутентификация в 21:07 (событие № 6) существенно отличается от нормального распределения активности, и степень отклонения превышает порог. Структурный канал формирует локальный граф событий и фиксирует аномальную последовательность действий — сразу после успешного входа происходит установка программы удалённого доступа «AnyDesk» (событие №7), а затем фиксируется значительный объём исходящего трафика (событие №8). Такая цепочка событий для одного пользователя и одного узла сети нетипична и интерпретируется как возможный сценарий компрометации учётных данных и утечки данных за пределы информационной инфраструктуры.

Результаты трёх каналов приводятся к единой шкале и взвешиваются в соответствии с параметрами системы. Эвристический канал предоставил бинарный сигнал «аномалия», статистический канал зафиксировал значительное отклонение от профиля, а структурный подтвердил наличие нетипич-

ной последовательности действий. Совокупная оценка превысила порог маркировки, и события №3–8 были отмечены как аномальные и переданы в модуль корреляции.

Модуль корреляции проверил временную и контекстную связанность событий. Все отмеченные события принадлежат одному пользователю и одному автоматизированному рабочему месту, а временные интервалы между ними не превышают нескольких минут. Вышеуказанные факты обеспечили высокую агрегированную меру сходства, превышающую порог приписывания, что позволило объединить записи в группу событий. В таблице 3 представлены значения используемых параметров в алгоритмах.

На этапе сопоставления каждая запись группы получила оценки соответствия техникам базы знаний MITRE ATT&CK. Серия неудачных аутентификаций и последующий успешный вход интерпретируются как техники «Credential Access» и «Initial Access», установка «AnyDesk» — как возможное исполнение с использованием внешнего инструмента «Remote Services», а исходящая передача данных — как техника «Exfiltration». Агрегированная поддержка по группе показала последовательность техник «Credential Access (T1110) — Initial Access (T1078) — Remote Services (T1021) — Exfiltration (T1041)». Суммарная уверенность превысила порог, и система создала уведомление, в котором указаны состав группы событий, временные рамки, перечислены ключевые признаки и техники базы знаний MITRE ATT&CK.

Таблица 3

Значения параметров

Параметр	Обозначение	Значение
Окно агрегации признаков	—	30 мин
Порог отклонения профиля	τ	0.5
Коэффициент экспоненциального сглаживания	α	0.2
Вес эвристического канала	w_r	0.2
Вес статистического канала	w_s	0.5
Вес графового канала	w_g	0.3
Порог итоговой аномалии	τ_s	0.6
Порог приписывания к группе событий	τ_d	0.6
Порог уверенности для фиксирования ККА	γ	1.5

Проведённое внутренне расследование показало, что злоумышленник, который является сотрудником организации получил несанкционированный доступ, посредством перебора, к учётной записи сотрудника экономического отдела. В вечернее время, когда сотрудник отсутствовал на рабочем месте злоумышленник, после нескольких неудачных попыток, смог авторизоваться на автоматизированном рабочем месте сотрудника и установить программу для удалённого доступа с целью передачи информации за периметр информационной инфраструктуры предприятия.

Заключение

В работе представлены четыре взаимосвязанных алгоритма обеспечивающие функционирование модели обнаружения комплексных компьютерных атак: формирование и инкрементальную поддержку профиля действий пользователей, локальное обнаружение аномалий, группировку аномалий в логические группы событий и сопоставление этих групп с тактиками и техниками из базы знаний MITRE ATT&CK. Для каждого этапа подробно описаны внутренняя логика, ключевые критерии принятия решений и процедуры калибровки оценок.

Элементы решения включают адаптивное обновление профиля действия пользователя с применением экспоненциального сглаживания, политику «мягкого слияния» профилей с версионированием и проверкой изменений по нескольким источникам, гибридную многоуровневую схему обнаружения сочетающую эвристический анализ, автокодировщик и графовый анализ при которой выходные оценки каждого уровня приводятся к единой шкале и объединяются с помощью весов, а также байесовский подход к сопоставлению групп событий с техниками базы знаний MITRE ATT&CK с учётом условных вероятностей.

Стоит отметить, что для корректной работы алгоритмов необходимы разнородные, синхронизированные данные, корректные временные метки и достаточная плотность нормальных данных для обучения автокодировщика.

В дальнейшем, планируется практическая реализация описанных алгоритмов и их экспериментальная валидация, с автоматической настройкой весов и повышением устойчивости к целенаправленным атакам на профили.

Литература

1. Актуальные киберугрозы: I-II кварталы 2025 года – URL: <https://ptsecurity.com/research/analytics/aktual-nye-kiberugrozy-i-ii-kvartaly-2025-goda/> (дата обращения: 20.11.2025).
2. D. Melnikov, «A Model for Detecting Complex Computer Attacks Based on the Analysis of User Actions, Network Interactions and Comparison with the MITRE ATT&CK Knowledge Base,» 2025 IEEE Ural-Siberian Conference on Biomedical Engineering, Radioelectronics and Information Technology (USBEREIT), Yekaterinburg, Russian Federation, 2025, pp. 221-224, doi: 10.1109/USBEREIT65494.2025.11054096.
3. MITRE. «ATT&CK® Design and Philosophy», March 2020. URL: https://attack.mitre.org/docs/ATTACK_Design_and_Philosophy_March_2020.pdf (дата обращения: 20.10.2025).

References

1. Aktualnyye kiberugrozy: I-II kvartaly 2025 goda – URL: <https://ptsecurity.com/research/analytics/aktual-nye-kiberugrozy-i-ii-kvartaly-2025-goda/> (accessed: - 20.11.2025).
2. D. Melnikov, «A Model for Detecting Complex Computer Attacks Based on the Analysis of User Actions, Network Interactions and Comparison with the MITRE ATT&CK Knowledge Base,» 2025 IEEE Ural-Siberian Conference on Biomedical Engineering, Radioelectronics and Information Technology (USBEREIT), Yekaterinburg, Russian Federation, 2025, pp. 221-224, doi: 10.1109/USBEREIT65494.2025.11054096.
3. MITRE. «ATT&CK® Design and Philosophy», March 2020. URL: https://attack.mitre.org/docs/ATTACK_Design_and_Philosophy_March_2020.pdf (accessed: 20.10.2025).

МЕЛЬНИКОВ Дмитрий Юрьевич, аспирант Учебно-научного центра «Информационная безопасность» федерального государственного автономного образовательного учреждения высшего образования «Уральский федеральный университет им. первого Президента России Б.Н. Ельцина». 620002, г. Екатеринбург, ул. Мира, 32. E-mail: melnikovdima517@yandex.ru

MELNIKOV Dmitry Yuryevich, graduate student at the Educational and Scientific Center «Information Security» of the Federal State Autonomous Educational Institution of Higher Education «Ural Federal University named after the first President of Russia B.N. Yeltsin». 620002, Yekaterinburg, st. Mira, 32. E-mail: melnikovdima517@yandex.ru

МЕТОДИКА ОЦЕНКИ РИСКОВ ИНФОРМАЦИОННОЙ БЕЗОПАСНОСТИ КООРДИНАТНО-ГРАФОВЫМ МЕТОДОМ

В статье рассматривается связь оценки риска информационной безопасности и бизнес-процессов, которые определяют шаги для достижения бизнес-цели. Рассматриваются составляющие, необходимые для оценки риска информационной безопасности при учете потребности организации как выстроенного менеджмента. Рассмотрены различные подходы к оценке информационной безопасности и определен приоритетный в данной работе.

Предложен метод оценки информационной безопасности, использующий количественный подход. Метод отражает взаимосвязь активов организации и выполняемых бизнес-операций и ущерба, который является составляющей риска. Описываются шаги, в рамках предлагаемого метода оценки рисков информационной безопасности, результатом выполнения которых является показатель риска информационной безопасности для рассматриваемой системы.

Ключевые слова: *риск информационной безопасности, оценка риска информационной безопасности, бизнес-процесс, бизнес-операция, аппроксимация.*

METHODOLOGY FOR ASSESSING INFORMATION SECURITY RISKS USING THE COORDINATE GRAPH METHOD

The article examines the relationship between information security risk assessment and business processes that determine the steps to achieve a business goal. The components necessary for assessing information security risk are considered, taking into account the needs of the organization as a structured management. Various approaches to information security assessment are considered and a priority one is identified in this work.

A method for assessing information security using a quantitative approach is proposed. The method reflects the relationship between the assets of the organization and the business operations performed and the damage that is part of the risk. The steps described within the framework of the proposed method for assessing information security risks are described, the result of which is an indicator of information security risk for the system in question.

Keywords: *information security risk, information security risk assessment, business process, business operation, approximation.*

Деятельность любой организации направлена на достижение бизнес-целей, для обеспечения которых могут применяться средства автоматизации. Средства автоматизации используются в бизнес-процессах, а каждому бизнес-процессу соответствует некоторый набор бизнес-операций. В данной работе под бизнес-операцией подразумевается совершение совокупности действий, которые являются составляющими бизнес-процесса. Бизнес-операция, в данной работе, – это минимальная единица совершения действий и минимальная составляющая бизнес-процесса. Определение термина «бизнес-процесс» содержится в ГОСТ Р ИСО 19439-2022 «Интеграция предприятия. Основа моделирования предприятия»: бизнес-процесс – частично установленный набор видов деятельности предприятия, который может быть выполнен для достижения определенного желаемого конечного результата во исполнение данной цели предприятия или части предприятия [1].

Применение средств автоматизации для выполнения и ускорения бизнес-операций, сопряжено с рисками информационной безопасности. В соответствии с текущим законодательством (ГОСТ Р ИСО/МЭК 27005-2010),

риск информационной безопасности – возможность того, что данная угроза сможет воспользоваться уязвимостью актива или группы активов и тем самым нанесет ущерб организации [2]. В случае реализации риска информационной безопасности, возникает остановка выполнения бизнес-операций, как следствие невозможно выполнить бизнес-процесс или выполнение бизнес-процесса осуществляется с неудовлетворительными результатами. Остановленные бизнес-процессы не позволяют достичь бизнес-целей. Конечным результатом реализации риска информационной безопасности является ущерб, который может быть выражен во временном или денежном эквиваленте. Временной показатель ущерба оценивает время простоя актива или предприятия в целом, а время необходимое на возобновление выполнения бизнес-операций.

В качестве определения термина «ущерб» будет использоваться пояснение Банка России: ущерб – утрата активов, повреждение (утрата свойств) активов и (или) инфраструктуры организации или другой вред активам и (или) инфраструктуре организации БС РФ, наступивший в результате реализации угроз ИБ через уязвимости ИБ [3].

Определение «риск информационной безопасности» коррелирует с определением «ущерб». Согласно обоим определениям, существуют активы организации, которые содержат уязвимости, и эксплуатация которых приводит к реализации угроз информационной безопасности. Под активами организации в данной работе подразумевается следующее определение из ГОСТ Р 53114-2008: активы организации: все, что имеет ценность для организации в интересах достижения целей деятельности и находится в ее распоряжении [4]. Актив организации в данной работе – минимальная единица имущества, обладающая самостоятельной ценностью и функциональностью.

Исходя из вышеизложенного, риск информационной безопасности напрямую связан с активами организации, которые участвуют в обеспечении бизнес-операций.

Оценивать риски информационной безопасности можно с помощью нескольких подходов:

1. Качественная оценка риска информационной безопасности;
2. Количественная оценка риска информационной безопасности;
3. Метод, комбинирующий элементы качественной и количественной оценки рисков информационной безопасности.

Каждый подход обладает своими достоинствами и недостатками, что обуславливает необходимость выбора в зависимости от специфики информационной системы или сферы применения. Качественная оценка рисков информационной безопасности предназначена для общей идентификации риска с категоризацией выявленных рисков посредством присвоения им одного или нескольких атрибутов на основе оценочных шкал. Ключевым преимуществом данного подхода является его относительная простота и оперативность проведения анализа, что позволяет обеспечить быстрое получение интерпретируемых результатов. Также данный подход отличается высокой степенью доступности для широкого круга пользователей, включая лиц, не обладающих специализированными знаниями в области информационной безопасности. Однако ключевыми недостатками метода, которые ставят под сомнение результаты оценивания, являются:

1. Субъективность при оценивании риска ИБ – эксперт может субъективно подходить к оценке конкретного ри-

ска в зависимости от большого количества обстоятельств: уровень компетенций, опыт работы, различный уровень информированности о рассматриваемой системе и т.д.

2. Ограниченный спектр категорий – распределение оценки по небольшому количеству категорий приводит к ухудшению качества оценки риска и накоплению погрешности. Увеличение количества категорий может ввести в заблуждение экспертов при оценивании, так и непрофессионалов при принятии решения.
3. Трудности интерпретации результатов с существенной разницей в оценке – в случае получения от разных экспертов диаметральных оценок, сложно принять итоговое решение. Предположим, что эксперты дали оценки «низкий риск» и «высокий риск». Использование среднего значения может еще сильнее увеличить погрешность оценки и привести к необоснованно низкой оценке. А принятие версии одной из сторон может также привести к неверной оценке риска, так как эксперт выставляет субъективную оценку, которая может быть завышена.
4. Необходимость регулярного повторного оценивания – качественный подход не позволяет осуществлять точное прогнозирование изменения риска информационной безопасности. Также учитывая вышеописанные недостатки, повторная оценка необходима для корректировки произведенных ранее оценок.

Количественная оценка риска информационной безопасности представляет собой присвоение некоторого значения вероятности реализации риска в числовом формате. В ГОСТ Р ИСО/МЭК 27005-2010 определяется, что, помимо значения вероятности, необходимо рассматривать некоторую количественную оценку последствий риска [2]. Последствием риска будем считать воздействие на бизнес-операции и бизнес-процессы организации.

Количественная оценка риска информационной безопасности обладает рядом преимуществ, которые способствуют достижению высокой степени точности и объективности в анализе. Одним из ключевых преимуществ данного подхода является отсутствие

шкал оценивания, что позволяет минимизировать влияние человеческого фактора и повысить объективность результатов. Оценка с помощью количественного подхода позволяет получить точное числовое значение риска информационной безопасности. Данный подход позволяет проводить прогнозирование изменения риска информационной безопасности.

Недостатки, которые можно выделить для количественного подхода оценки рисков информационной безопасности:

1. Необходимость наличия достоверных данных – искаженные входные параметры приводят к неверному определению итогового результата, а также могут привести к ошибкам при расчете формул в автоматизированном режиме;
2. Сложность интерпретации итоговой оценки риска информационной безопасности – полученное значение может сложно интерпретироваться и непонятно человеку, не вовлеченному в процесс оценки рисков информационной безопасности.
3. Приведение входных параметров к определенной форме – подход использует числовые значения, которые в зависимости от метода подсчета, должны быть представлены в той или иной форме, а также возможно не все входные параметры поддаются конвертации в числовую величину.

Учитывая ключевые преимущества метода – точная оценка риска и возможность прогнозирования изменения риска информационной безопасности, будем считать метод количественной оценки приоритетным.

Предлагается новый метод оценки рисков информационной безопасности, основное внимание которого сконцентрировано на взаимосвязях: активов предприятия, бизнес-операциях (выполняемых с помощью активов), ущерба от прекращения бизнес-процессов, актуальной экспертной оценки опасности фактора риска. Перечисленные составляющие также являются входными данными метода.

Основные преимущества метода: возможность выполнения расчетов для любого количества входных параметров, что означает масштабируемость метода; возможность прогнозирования изменения рисков информационной безопасности с течением времени.

Для выполнения оценки риска информационной безопасности, необходимо последовательно выполнить следующие шаги:

1. Сформировать перечень бизнес-процессов и соответствующих им перечень бизнес-операций.
2. Сформировать перечень активов организации, которые обеспечивают выполнение бизнес-операций.
3. Составить структуру в виде графа, которая отображала связь активов и бизнес-операций.
4. Оценить минимальное количество активов для выполнения бизнес-операций.
5. Оценить необходимое время на восстановление выполнения бизнес-операции для каждой группы объектов.
6. Оценить с помощью экспертов существенность влияния факторов на указанные активы.
7. Построить в пространстве точки, которые отражали состояние активов по трем координатам: коэффициент показателя бизнес-операций (КПБО), ущерб, совокупная экспертная оценка.
8. Провести аппроксимацию указанных точек с целью получения полиномиальной функции.
9. Выполнить аналогичные действия для показателей с минимальным количеством активов, с целью определения приемлемого риска.
10. Определить соотношение объемов фигур и вычислить риск.

Для первого этапа необходимо, чтобы специалисты организации, где оценивается риск информационной безопасности, определили перечень рассматриваемых бизнес-процессов. Именно набор бизнес-процессов задает первый параметр масштабирования. Затем необходимо проделать декомпозицию для каждого бизнес-процесса с целью определения бизнес-операций. Бизнес-операция в данном методе, как уже упоминалось ранее, минимальная единица действия. После получения перечня бизнес-операций необходимо перейти ко второму этапу.

Необходимо определить перечень активов предприятия, которые участвуют в выполнении бизнес-операций. Организация сама принимает решение, что включить в перечень, так как согласно ГОСТ Р 53114-2008, только организация определяет поставленные цели и соответствующую ценность. Вели-

чина перечня, которая является вторым параметром масштаба оценивания, задает объем последующих вычислений.

Третий этап — это построение графа, где для каждого актива определяется связь с бизнес-операцией. При этом один актив может быть связан с несколькими бизнес-операциями. Для каждой связи – ребра графа, необходимо определить важность (вовлеченность) актива в выполнении бизнес-операции. Актив может обеспечивать выполнение нескольких бизнес-операций с разной степенью вовлеченности. Результатом данного этапа является коэффициент показателя бизнес-операции (КПБО), и значение данного коэффициента будет располагаться на оси Ох, при выполнении шага под номером семь.

Дополнительно введем показатель сложности актива. Под сложностью актива понимается коэффициент, отражающий совокупность свойств и признаков актива, а также степень их влияния. Числовое значение характеризует насколько составляющей элемент отличается по некоторым критериям от других составляющих. Рекомендуется рассматривать элементы, которые непосредственно задействованы при осуществлении бизнес-операции. Предлагается использовать следующие критерии для оценки показателя сложности актива:

1. Стоимость актива;
2. Аппаратная сложность;
3. Количество обеспечиваемых операций и функций;
4. Отказоустойчивость и время беспрерывной работы актива;
5. Сложность администрирования актива;
6. Порог компетенций для работы с активом;
7. Ценность актива с точки зрения нарушителя.

Вышеперечисленные показатели не являются конечными и могут быть дополнены организацией, выполняющей оценку риска, дополнительными критериями, которые отвечали бы интересам организации и сфере деятельности.

В общем виде показатель сложности рассчитывается по формуле:

$$D = k \prod_{i=1}^n x_i^{w_i},$$

где k – константа нормализации, x_i – признаки и их значения, w_i – весовой коэффициент признака.

Например, для некоторого автоматизированного рабочего места (АРМ) бухгалтера, которое участвует в бизнес-операции «отправка файла xml во внешнюю систему», необходимо рассмотреть следующие составляющие элементы: операционная система, клиент базы данных, браузер, подключение к внешней системе, аппаратная составляющая АРМ.

Итоговая координата для каждой бизнес-операции находится по формуле:

$$\text{КПБО}_n = (e_1 * D_1) + \dots + (e_i * D_k),$$

где e_i – вес ребра, D_k – сложность актива.

Четвертый этап предполагает нахождение показателя $\text{КПБО}_{\min}$, который отражает необходимое минимальное количество активов, для выполнения бизнес-операции с наименьшими возможными и допустимыми результатами, но без прерывания операции. В отдельных случаях минимальное количество активов может совпадать с количеством активов в штатном режиме работы, в этом случае $\text{КПБО}_{\min}$ будет равен 0. Показатель $\text{КПБО}_{\min}$ необходим для оценивания приемлемого (допустимого) риска, по отношению к которому оценивается недопустимый (фактический) риск. Под допустимым риском понимается риск, который в данной ситуации считают приемлемым при существующих общественных ценностях [5].

Пятый этап предполагает определение показателя ущерба. Результатом данного этапа является коэффициент ущерба – располагается на оси Оу, при выполнении этапа под номером семь. Согласно определению, описанному в Стандарте Банка России СТО БР ИББС-1.0-2014, утрата актива ведет к ущербу. В предлагаемом методе утрата актива ведет к прекращению выполнения бизнес-операции, а следовательно, наступает ущерб, который можно измерить в денежном эквиваленте или временном эквиваленте. Ущерб будет оцениваться по времени, затраченному на восстановление выполнения БО с наименьшими допустимыми результатами, т.е. с минимальным количеством активов. Среди перечня возможных временных значений выбирается максимальное.

Единицы измерения определяются заранее и дальнейшие расчеты конвертируются в выбранную величину, например, часы. Действия совершаемые с активами для восстановления бизнес-операции:

1. Устранение нарушения работоспособности (например, время восстановления функционирования тонкого клиента 1С занимает 1 час);
2. Замена актива на резервный (например, в случае выхода из строя АРМ, замена займет 0,6 часа).

Стоит отметить, что замена на резервный актив не всегда является лучшим выбором, так как резервное устройство может нуждаться в настройке, которое займет больше времени, чем устранение причин неработоспособности основного устройства.

В последующем теоретически возможно рассчитать ущерб в денежном эквиваленте, основываясь на временном ущербе: время, затраченное персоналом на восстановление работоспособности; время простоя персонала, из-за невозможности работать с активом; стоимость актива (в случае полного уничтожения); стоимость компенсаций контрагентам и т.д.

Шестой этап направлен на получение экспертной оценки. Показатель экспертной оценки – представляет собой числовое значение, которое описывает величину фактора риска. Под фактором риска понимается фактор, который оказывает существенное влияние на риск [6]. Эксперт, руководствуясь своим профессиональным опытом, актуальным перечнем угроз и уязвимостей для активов, формирует некоторое положительное действительное число, отражающее величину существенности влияния фактора риска. Эксперту доступны сведения о предыдущих этапах, с целью осведомления обо всех влияющих факторах подлежащих оцениванию. Сле-

дует отметить, что количество экспертов, производящих оценку, должно быть не менее двух, чтобы избежать некорректных расчетов и субъективности мнения.

После получения данных от экспертов, необходимо произвести некоторые проверки: нахождение коэффициента конкордации, стандартизацию данных для поиска аномалий.

Для нахождения коэффициента конкордации, который используется для оценки согласованности мнений нескольких экспертов относительно ранжирования объектов, необходимо:

1. Составить таблицу, где для каждого фактора риска будут соотнесены экспертные оценки. Пример выполнения данного шага представлен в таблице 1, где m – количество экспертов, n – количество факторов риска. Фактором риска в данных вычислениях является комбинация сведений об активах, коэффициентах показателя бизнес-операции, коэффициент сложности каждого актива и показатель ущерба. Предполагается, что наименьшим значениям соответствуют наименьшие ранги, т.е. минимальное значение в столбце эксперта получит 1, а наибольшее будет равняться значению количества факторов.

2. Найти сумму ранжированных оценок экспертов для каждого фактора

$$R_i = \sum_{j=1}^n r_{mj},$$

где r_{mj} – оценка каждого эксперта по каждому фактору

Данные удобно привести к виду таблицы (таблица 2).

Таблица 1

Составление таблицы ранжирования экспертных оценок

Фактор риска	Оценка эксперта 1	Оценка эксперта 2	Оценка эксперта 3	...	Оценка эксперта m
Фактор 1	$r_{1,1}$	$r_{2,1}$	$r_{3,1}$	...	$r_{m,1}$
Фактор 2	$r_{1,2}$	$r_{2,2}$	$r_{3,2}$	...	$r_{m,2}$
Фактор 3	$r_{1,3}$	$r_{2,3}$	$r_{3,3}$	...	$r_{m,3}$
...	...	...	...	...	...
Фактор n	$r_{1,n}$	$r_{2,n}$	$r_{3,n}$	...	$r_{m,n}$

Сумма ранжированных оценок

Фактор риска	Сумма оценок экспертов
Фактор 1	R_1
Фактор 2	R_2
...	...
Фактор n	R_n

3. Необходимо оценить среднюю сумму ранжированных оценок и сумму квадратов отклонений.

$$\bar{R} = \frac{\sum_{n=1}^n R_n}{n};$$

$$SS = \sum_{n=1}^n (R_n - \bar{R})^2.$$

4. Определение величины коэффициента конкордации Кендалла:

$$W = \frac{12 * SS}{m^2(n^3 - n)}$$

Величина изменяется от 0 (полная несогласованность) до 1 (полная согласованность). При значениях коэффициента, приближенных к 1, должно быть выдвинуто предположение о сговорности экспертов, а значит, оценки могут быть некорректными. Оптимальным диапазоном величины коэффициента является 0,7...0,9. Если коэффициент определяется ниже порогового уровня в 0,7, то необходимо провести проверку экспертных оценок и определить аномалии.

Для поиска аномалий используется метод стандартизации данных. Основываясь на выполнении расчетов при нахождении коэффициента конкордации, определяется стандартное отклонение.

$$\sigma = \sqrt{\frac{SS}{n - 1}}$$

Затем для каждого фактора и соответствующего ему экспертной оценки находится показатель z , который при превышении установленного порога означает аномалию. Установленный порог может изменяться по решению эксперта, проводящего оценку, однако рекомендуемый порог $|z_i| > 3$.

$$z_i = \frac{x_i - \mu}{\sigma}$$

Аномальные оценки экспертов должны быть исключены из построения точек в трехмерном пространстве. Результатом данного этапа является получение экспертной оценки R_i для каждой комбинации КПБО и ущерба. Оценка R_i располагается по оси Oz .

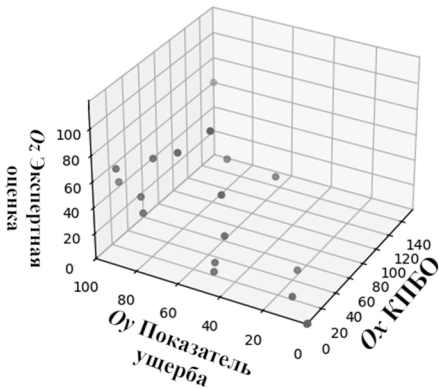


Рис. 1. Визуализация точек в пространстве

Седьмой этап предполагает, что полученные значения (связка КПБО-ущерб-экспертная оценка) являются координатами точек, которые располагаются в трехмерном пространстве: коэффициент показателя бизнес-операции (КПБО) – располагается на оси Ox , коэффициент ущерба – располагается на оси Oy , коэффициент экспертной оценки R_i – располагается на оси Oz . Визуализация данного этапа представлена на рисунке 1.

Восьмой этап предполагает работу с точками, расположенными в пространстве. После построения точек необходимо рассчитать аппроксимирующую поверхность. Аппроксимирующей функцией в данной работе будет являться полиномиальная функция.

Ключевой показатель полиномиальной функции – степень. Именно от данного показателя зависит точность аппроксимации, а также способность не учитывать некоторые значения «шума».

Для поиска оптимального значения степени полинома необходимо сравнить критерии Akaike Information Criterion (AIC) и Bayesian Information Criterion (BIC) для найденных значений точек в пространстве.

Нахождение критерия Акаике (AIC) выполняется по формуле:

$$AIC = n \ln \left(\frac{RSS}{n} \right) + 2k,$$

где RSS – сумма квадратов остатков, n – общее количество наблюдений, k – количество параметров модели (включая свободный член).

Нахождение критерия Байеса (BIC) выполняется по формуле:

$$BIC = n \ln \left(\frac{RSS}{n} \right) + k \ln(n).$$

Полученные значения для разных степеней полинома необходимо сравнить и выбрать наименьшие значения. Наименьшее значение определяет оптимальное соотношение точности аппроксимации и сложности модели, но без переобучения модели. Переобучением модели считается излишнее воздействие данных (шум) на полином.

Для нахождения RSS необходимо для каждой степени найти:

$$RSS = \sum_{i=1}^n [z_i - f(x_i, y_i)]^2,$$

где $f(x_i, y_i)$ значение полинома в точке (x_i, y_i) .

Для нахождения AIC и BIC необходимо решение уравнения полиномов с первой до восьмой степени:

$$z = a_0 + a_1x + a_2y;$$

$$z = a_0 + a_1x + a_2y + a_3xy + a_4x^2 + a_5y^2;$$

$$z = a_0 + a_1x + a_2y + a_3xy + a_4x^2 + a_5y^2 + a_6x^2y + a_7xy^2 + a_8x^3 + a_9y^3;$$

$$z = a_0 + a_1x + a_2y + a_3xy + a_4x^2 + a_5y^2 + a_6x^2y + a_7xy^2 + a_8x^3 + a_9y^3 + a_{10}x^3y + a_{11}xy^3 + a_{12}x^4 + a_{13}y^4;$$

$$z = a_0 + a_1x + a_2y + a_3xy + a_4x^2 + a_5y^2 + a_6x^2y + a_7xy^2 + a_8x^3 + a_9y^3 + a_{10}x^3y + a_{11}xy^3 + a_{12}x^4 + a_{13}y^4 + a_{14}x^4y + a_{15}xy^4 + a_{16}x^5 + a_{17}y^5;$$

$$z = a_0 + a_1x + a_2y + a_3xy + a_4x^2 + a_5y^2 + a_6x^2y + a_7xy^2 + a_8x^3 + a_9y^3 + a_{10}x^3y + a_{11}xy^3 + a_{12}x^4 + a_{13}y^4 + a_{14}x^4y + a_{15}xy^4 + a_{16}x^5 + a_{17}y^5 + a_{18}x^5y + a_{19}xy^5 + a_{20}x^6 + a_{21}y^6;$$

$$z = a_0 + a_1x + a_2y + a_3xy + a_4x^2 + a_5y^2 + a_6x^2y + a_7xy^2 + a_8x^3 + a_9y^3 + a_{10}x^3y + a_{11}xy^3 + a_{12}x^4 + a_{13}y^4 + a_{14}x^4y + a_{15}xy^4 + a_{16}x^5 + a_{17}y^5 + a_{18}x^5y + a_{19}xy^5 + a_{20}x^6 + a_{21}y^6 + a_{22}x^6y + a_{23}xy^6 + a_{24}x^7 + a_{25}y^7;$$

$$\begin{aligned}
z = & a_0 + a_1x + a_2y + a_3xy + a_4x^2 + a_5y^2 + a_6x^2y + a_7xy^2 + a_8x^3 + a_9y^3 \\
& + a_{10}x^3y + a_{11}xy^3 + a_{12}x^4 + a_{13}y^4 + a_{14}x^4y + a_{15}xy^4 + a_{16}x^5 \\
& + a_{17}y^5 + a_{18}x^5y + a_{19}xy^5 + a_{20}x^6 + a_{21}y^6 + a_{22}x^6y + a_{23}xy^6 \\
& + a_{24}x^7 + a_{25}y^7 + a_{26}x^7y + a_{27}xy^7 + a_{28}x^8 + a_{29}y^8;
\end{aligned}$$

Ввиду того, что значения функции не проходят через координаты $(0;0;z_0)$, следовательно, $a_0 \neq z_0$. В связи с этим необходимо найти a_0 , а затем найти решение уравнения для оставшихся членов.

В общем виде решение предлагается строить в матричной форме, вида:

X

$$= \begin{pmatrix} 1 & x_{11} & x_{12} & \dots & x_{1m} & x_{11}^2 & x_{12}^2 & \dots & x_{11}x_{12} & \dots & x_{11}^d & x_{12}^d & \dots & x_{11}^d x_{12}^d \\ 1 & x_{21} & x_{22} & \dots & x_{2m} & x_{21}^2 & x_{22}^2 & \dots & x_{21}x_{22} & \dots & x_{21}^d & x_{22}^d & \dots & x_{21}^d x_{22}^d \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \ddots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots \\ 1 & x_{n1} & x_{n2} & \dots & x_{nm} & x_{n1}^2 & x_{n2}^2 & \dots & x_{n1}x_{n2} & \dots & x_{n1}^d & x_{n2}^d & \dots & x_{n1}^d x_{n2}^d \end{pmatrix};$$

$$Z = \begin{pmatrix} z_1 \\ z_2 \\ \vdots \\ z_n \end{pmatrix}.$$

Для решения уравнения с применением метода Гаусса необходимо: перемножить обе стороны матричного уравнения слева на транспонированную матрицу $X^T X a = X^T Z$. Найдем матрицу $C = X^T X$ и вектор $b = X^T Z$, а затем решим систему $Ca = b$. Решение системы $Ca = b$ выполнить методом Гаусса с приведением к верхней треугольной форме.

При наличии информации о числе наблюдений (n), количестве параметров (k), и величинах ошибок (RSS) допустимо вычислить критерии — AIC и BIC, а затем составить сравнительную таблицу. Итоговым результатом вычислений является степень полинома, оптимальная для аппроксимации построенных точек в пространстве.

Построение полинома удобнее выполнять автоматизированным способом. Визуализация аппроксимации для некоторых точек в пространстве представлена на рисунке 2.

Девятый этап выполняется для определения аппроксимирующей поверхности для $KПБ O_{min}$ и ничем не отличается от выполняемых ранее вычислений. Задача данного этапа – построение аппроксимирующей поверхности допустимого риска.

Десятый этап заключается в получении значения объема фигуры, которую образует: полиномиальная поверхность (ограничение сверху), плоскость $z=0$ (ограничение снизу), плоскости, образуемые нормальными к координатным плоскостям крайних точек (боковые

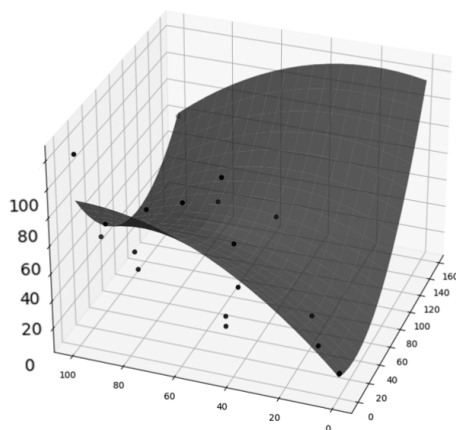


Рис. 2. Визуализация аппроксимирующей поверхности для точек в пространстве

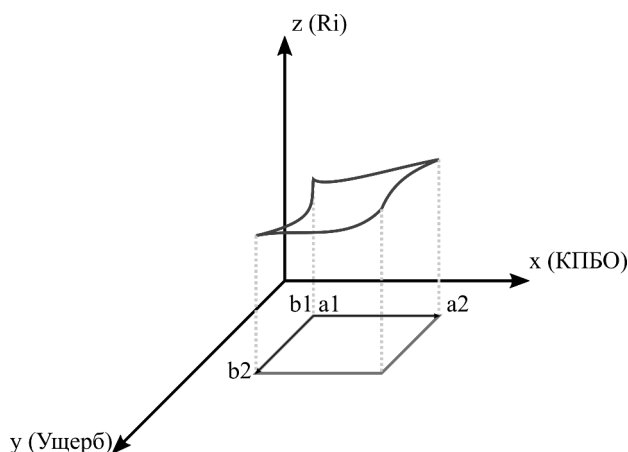


Рис. 3. Графическое представление по нахождению объема фигуры

грани). Для нахождения объема фигуры используется двойной интеграл, интегрируемый по координатам x и y , при фиксированном значении z .

Пусть границы фигуры заданы отрезками $[a1, a2]$ вдоль оси x , $[b1, b2]$ вдоль оси y , а поверхность, задаваемая полиномом, расположена целиком выше плоскости z . Начальной и конечной точками отрезков $a1, a2, b1, b2$ должны стать крайние точки, спроецированные на Oxy . Тогда нахождение объема производится по формуле:

$$V = \int_{a1}^{a2} \int_{b1}^{b2} f(x, y) dx dy.$$

Графическое представление по нахождению объема представлено на рисунке 3.

Заключительным этапом является определение объема фигуры оцениваемого риска и объема фигуры допустимого риска. Соотношение объемов фигур будет являться расчи-

танным показателем риска, относительно приемлемого риска.

Предлагаемый метод позволяет проводить прогнозирование изменения рисков информационной безопасности с течением времени. Ввиду наличия математической функции, если предположить, что технические средства организации могут выходить из строя согласно теории надежности, то возможно произвести расчет рисков информационной безопасности через некоторый промежуток времени.

Предлагаемый метод имеет перспективы для дальнейшего исследования. Предполагается, что точки в координатном пространстве позволят построить векторы, что будет являться дополнительным фактором для анализа данных. Также предполагается, что полученные данные позволят построить градиент, что будет отображать тенденцию развития риска информационной безопасности и позволит проводить действия для минимизации риска.

Литература

1. ГОСТ Р ИСО 19439-2022 // Федеральное агентство по техническому регулированию и метрологии – [сайт] – URL: <https://protect.gost.ru/document1.aspx?control=31&baseC=6&page=7&month=12&year=2022&search=&id=248845> (дата обращения 02.10.2025). – Режим доступа: свободный.
2. ГОСТ Р ИСО/МЭК 27005-2010 // РосГОСТ – [сайт] – URL: https://rosgosts.ru/file/gost/35/040/gost_r_iso/mek_27005-2010.pdf (дата обращения: 10.10.2025). – Режим доступа: свободный.
3. СТО БР ИББС-1.0-2014 «Обеспечение информационной безопасности организаций банковской системы Российской Федерации» // Центральный банк Российской Федерации URL: <https://cbr.ru/Crosscut/LawActs/File/446> (дата обращения: 18.10.2025). – Режим доступа: свободный.
4. ГОСТ Р 53114-2008 // Федеральное агентство по техническому регулированию и метрологии – [сайт] – URL: <https://protect.gost.ru/document.aspx?control=7&id=174974> (дата обращения 02.11.2025). – Режим доступа: свободный.
5. ГОСТ Р 51898—2002 «Аспекты безопасности. Правила включения в стандарты» Федеральное агентство по техническому регулированию и метрологии. – 2002. 8 с.
6. ГОСТ Р 58771—2019 «Менеджмент риска. Технологии оценки риска» // Электронный фонд правовых и нормативно-технических документов URL: <https://docs.cntd.ru/document/1200170253> (дата обращения: 20.10.2025).

References

1. GOST R ISO 19439-2022 // Federal'noye agentstvo po tekhnicheskomu regulirovaniyu i metrologii – [sayt] – URL: <https://protect.gost.ru/document1.aspx?control=31&baseC=6&page=7&month=12&year=2022&search=&id=248845> (data obrashcheniya 02.10.2025). – Rezhim dostupa: svobodnyy.
2. GOST R ISO/MEK 27005-2010 // RosGOST – [sayt] – URL: https://rosgosts.ru/file/gost/35/040/gost_r_iso/mek_27005-2010.pdf (data obra-shcheniya: 10.10.2025). – Rezhim dostupa: svobodnyy.
3. STO BR IBBS-1.0-2014 «Obespecheniye informatsionnoy bezopasno-sti organizatsiy bankovskoy sistemy Rossiyskoy Federatsii» // Tsentral'-nyy bank Ros-siyskoy Federatsii URL: <https://cbr.ru/Crosscut/LawActs/File/446> (data obrashche-niya: 18.10.2025). – Re-zhim dostupa: svobodnyy.
4. GOST R 53114-2008 // Federal'noye agentstvo po tekhnicheskomu regulirovaniyu i metrologii – [sayt] – URL: <https://protect.gost.ru/document.aspx?control=7&id=174974> (data obrashche-niya 02.11.2025). – Rezhim dostupa: svobodnyy.
5. GOST R 51898—2002 «Aspekty bezopasnosti. Pravila vkluycheniya v standarty» Federal'noye agentstvo po tekhnicheskomu regulirovaniyu i met-rologii. – 2002. 8 s.
6. GOST R 58771—2019 «Menedzhment riska. Tekhnologii otsenki ris-ka» // Elektronnyy fond pravovykh i normativno-tekhnicheskikh dokumentov URL: <https://docs.cntd.ru/document/1200170253> (data obrashcheniya: 20.10.2025).

КОРЖЕНЕВСКИЙ Денис Алексеевич, аспирант, старший преподаватель кафедры «Информационные технологии и защита информации» ФГБОУ ВО «Уральский государственный университет путей сообщения». 650034, г. Екатеринбург, ул. Колмогорова, 66. E-mail: DKorzhenevskiy@usurt.ru

ЗЫРЯНОВА Татьяна Юрьевна, кандидат технических наук, доцент, заведующий кафедрой «Информационные технологии и защита информации» ФГБОУ ВО «Уральский государственный университет путей сообщения». 650034, г. Екатеринбург, ул. Колмогорова, 66. E-mail: TZyryanova@usurt.ru

KORZHENEVSKIY Denis Alekseevich, post-graduate, Senior Lecturer at the Department of Information Technology and Information Security of the Federal State Budgetary Educational Institution of Higher Education «Ural State University of Railway Transport». 620034, Yekaterinburg, str. Kolmogorova, 66. E-mail: DKorzhenevskiy@usurt.ru

ZYRYANOVA Tatyana Yuryevna, Candidate of Technical Sciences, Associate Professor, Head of the Department of Information Technology and Information Security of the Federal State Budgetary Educational Institution of Higher Education «Ural State University of Railway Transport». 620034, Yekaterinburg, str. Kolmogorova, 66. E-mail: TZyryanova@usurt.ru

АНАЛОГИ СОВЕРШЕННЫХ ШИФРОВ В ТЕОРИИ КОДИРОВАНИЯ

В работе предложена кодовая конструкция, которая оказывается связанной с задачами теории информации, распределения вероятностей многомерных случайных величин и многомерной задаче о системах представителей. Доказана теорема об аналогах совершенных шифров, согласно которой свойство «совершенности» (отсутствия потерь информации) является глобальным свойством канала. Построен пример аналога шифра с тремя шифрвеличинами, иллюстрирующий предложенную кодовую конструкцию. Полученные результаты отражают универсальность совершенных систем шифрования.

Ключевые слова: защита информации, совершенные шифры, совершенные коды, многомерные случайные величины, кодовые конструкции.

Medvedeva N. V., Titov S. S.

ANALOGUES OF PERFECT CIPHERS IN CODING THEORY

This work proposes a code construction that is related to problems in information theory, probability distributions of multivariate random variables, and the multidimensional system of representatives problem. A theorem on the analogues of perfect ciphers is proved, stating that the property of «perfection» (absence of information loss) is a global property of the channel. An example of a cipher analogue with three cipher values is constructed to illustrate the proposed code structure. The results reflect the universality of perfect encryption systems.

Keywords: information protection, perfect ciphers, perfect codes, multidimensional random variables, code constructions.

В современных условиях, когда объемы передаваемых и хранимых данных растут экспоненциально, а требования к надежности и безопасности информационных систем становятся критически важными, особую значимость приобретают фундаментальные исследования на стыке теории информации и криптографии. Изучение аналогов совершенных шифров в теории кодирования представляет не только теоретический интерес, но и имеет практическое значение для решения

актуальных задач в области защиты информации и обеспечения достоверности данных.

Концепция совершенного шифра, впервые строго сформулированная К. Шенноном [1], устанавливает теоретический предел в криптографии – абсолютную стойкость к криптоатакам. Построение абсолютно стойких систем передачи данных исследовалось в работах [2 – 7].

Поиск и изучение аналогичных предельно эффективных конструкций в теории коди-

рования позволяет выявить универсальные математические принципы, лежащие в основе как защиты от злоумышленника, так и обеспечения целостности данных при передаче через зашумленные каналы связи.

В настоящей работе исследуются аналоги совершенных (по Шеннону [1]) шифров при построении систем передачи данных. Здесь используется вероятностная модель $\sum_B$ [1] передачи данных с шифрованием и кодированием, с соответствующими обозначениями: $X = \{x_1, x_2, \dots, x_\lambda\} = \{1, 2, \dots, \lambda\}$ – множество шифрвеличин (алфавит открытых текстов); $Y = \{y_1, y_2, \dots, y_\mu\} = \{1, 2, \dots, \mu\}$ – множество шифробозначений (алфавит закрытых текстов), с которыми оперирует некоторый шифр замены; $K = \{k_1, k_2, \dots, k_\pi\}$ – множество ключей. По условию: $|X| = \lambda > 1$, $|Y| = \mu \geq \lambda$, $|K| = \pi \geq \mu$. Тексты представляют собою слова (ℓ -граммы, т.е. биграммы, триграммы и так далее) в соответствующих алфавитах. Под шифром $\sum_B$ понимается совокупность множеств правил зашифрования и правил расшифрования с заданными распределениями вероятностей на множествах открытых текстов и ключей [8, 9].

Абсолютная стойкость вероятностной модели $\sum_B$ системы передачи данных, согласно определению совершенного по Шеннону шифра, эквивалентна равенству апостериорных вероятностей $p(x|y)$, $x \in X$, $y \in Y$ открытых текстов и соответствующих им априорных вероятностей $p(x)$ [8, 9]. В работе [2] показано, что построение абсолютно стойких систем передачи данных приводит к постановке задачи описания произвола в построении совместно распределенных случайных величин, трактуемых как посылаемые и принимаемые сигналы.

Кодовыми аналогами совершенных шифров можно считать генератор символов кода, которые зависят от двух источников – детерминированного, определяемого передаваемым сообщением (аналог символа открытого текста), и случайного, выбор которого происходит в соответствии с заданной вероятностью (аналог ключа); при этом априорные и апостериорные вероятности символов совпадают (аналог совершенности шифра). Так, если генерируемые символы кода – биты, то условная вероятность появления на выходе бита, равного единице, должна быть одинаковой для любого входного символа сообщения.

Детерминированность источников сообщений должна пониматься доста-точно отно-

сительно, поскольку она связана с логикой управления, реагирования на реальные ситуации и поэтому не входит в данную кодовую модель. По-этому можно расширить эту модель, считать появление символов источников сообщений как случайные события с некоторым распределением их вероятности. Однако это оказывается излишним, потому что имеет место

Теорема. Пусть имеется канал связи с дискретными входами $x \in X$ и выходами $y \in Y$. Если существует распределение $p_0(x)$, для которого выполняется **условие аналога совершенности**: $I_{p_0}(x; y) = C$ и $H_{p_0}(x|y) = 0$, где C – пропускная способность канала, $I(x; y)$ – взаимная информация между двумя случайными величинами x и y , $H(x|y)$ – условная энтропия x при известном y , то для любого другого распределения $p(x)$ выполняется **расширенное условие совершенности**: $I_p(x; y) = C$ и $H_p(x|y) = 0$.

Доказательство. Полагаем, что матрица $p(Y|X)$ переходных вероятностей канала фиксирована. Условие $H_{p_0}(x|y) = 0$ означает, что при наблюдении y вход x восстанавливается однозначно. Формально: $\forall y \in Y \exists! x \in X : p_0(x|y) = 1$. Это эквивалентно тому, что отображение $X \rightarrow Y$ инъективно на носителе распределения $p_0(X)$, т.е. для любого y существует не более одного x , для которого $p(y|x) > 0$ и $p_0(x) > 0$.

Из условия $H_{p_0}(x|y) = 0$ следует, что канал не создает неопределенности при использовании $p_0(X)$. Это возможно только, если для каждого y существует не более одного x с $p(y|x) > 0$. Таким образом, канал ведет себя как детерминированное отображение на подмножестве входов, где $p_0(x) > 0$. Обозначим через $X_0 = \{x \in X : p_0(x) > 0\}$. Тогда для $x_1, x_2 \in X_0$, $x_1 \neq x_2$ имеем: $p(y|x_1) \cdot p(y|x_2) = 0$ для любого $y \in Y$. Это означает, что выходы для разных входов из X_0 не пересекаются.

Условие $I_{p_0}(x; y) = C$, где C – пропускная способность канала, определяемая равенством $C = \max_{p(x)} I_p(x; y)$, означает, что распре-

деление $p_0(x)$ достигает пропускной способности. Так как $H_{p_0}(x|y) = 0$, то $I_{p_0}(x; y) = H_{p_0}(x) = C$. Следовательно, пропускная способность канала равна максимальной энтропии входа при условии инъективности отображения $X \rightarrow Y$.

Рассмотрим произвольное распределение $p(x)$. Покажем, что $I_p(x; y) = C$ и $H_p(x|y) = 0$. Для любого y с $p(y) > 0$ существует един-

ственный $x \in X_0$, для которого $p(y | x) > 0$. Это следует из структурного свойства канала, которое не зависит от распределения $p(x)$. Поэтому при наблюдении y всегда можно однозначно определить x даже, если $p(x) \neq p_0(x)$. Следовательно, $H_p(x | y) = 0$ для любого $p(x)$.

Из $H_p(x | y)$ следует, что $I_p(x; y) = H_p(x) - H_p(x | y) = H_p(x)$. Однако, $I_p(x; y) \leq C$ по определению пропускной способности. Для того, чтобы достичь C необходимо максимизировать $H_p(x)$ при условии, что отображение $X \rightarrow Y$ остается инъективным. Максимальная энтропия достигается на равномерном распределении на множестве X_0 и $H(X_0) = C$. Для произвольного $p(x)$, если $\text{supp } p \subset X_0$, то $H_p(x) \leq C$, и равенство достигается только для равномерного распределения на X_0 .

Условие теоремы гарантирует, что структура канала позволяет передавать без потерь C бит. Для произвольного $p(x)$, если $\text{supp } p \not\subset X_0$, то $H_p(x)$ может быть меньше C , но условие совершенности $H_p(x | y) = 0$ сохраняется. Следовательно, теорема утверждает, что свойство нулевой условной энтропии сохраняется, а не то, что пропускная способность достигается при любом распределении.

Выше показано, что:

1. Условие $H_{p_0}(x | y) = 0$ влечет структурное свойство канала – инъективность отображения $X \rightarrow Y$ на $\text{supp } p_0$.

2. Это структурное свойство не зависит от распределения $p(x)$, поэтому $H_p(x | y) = 0$ выполняется для любого $p(x)$.

3. Пропускная способность C достигается при равномерном распределении на X_0 , а для других распределений $I_p(x; y) = H_p(x) \leq C$.

Таким образом, условие аналога совершенности выполняется в расширенном смысле: свойство нулевой условной энтропии сохраняется для любого распределения входных символов, что и требовалось доказать.

Другими словами, если условие аналога совершенности выполняется для одного конкретного случая (чаще всего равномерного распределения), оно будет верным и для всех

остальных распределений вероятностей, т.е. схема шифрования не зависит от характера распределения исходных данных. Это обстоятельство указывает на универсальность данного подхода к созданию абсолютно стойких кодов.

В простейшем случае эта конструкция может быть представлена матрицей размерами M на N совместного распределения двух независимых случайных величин, представляющей соответствующее дискретное вероятностное пространство. В этой матрице представлены биты, задающие характеристическую функцию подмножества множества её ячеек и, таким образом, случайное событие, которое должно обладать следующим свойством: его проекция на пространство N есть пространство с тем же (исходным) распределением вероятности. При этом пространство N столбцов может быть интерпретировано как множество источников сообщений, а пространство M строк – как множество кодировок (ключей). При равновероятных распределениях такая модель задаётся $(0,1)$ -матрицей, у которой одинаковое количество единиц в каждом столбце и одинаковое количество единиц в каждой строке.

Такая матричная модель связана не только с теорией кодирования, но также и с задачей представителей: столбцы – вероятные ответы источника сообщений, представлятеля, а строки – варианты ответов всей группы. При этом выбор случайного источника и ответа не раскрывает, какой именно представитель был выбран.

Построим и проверим на совершенную секретность кодовую конструкцию с тремя шифрвеличинами.

Пример. Пусть $X = \{x_1, x_2, x_3\}$, где $x_i \in \{0,1\}$; $K = \{k_1, k_2, k_3\}$, где $k_i \in \{0,1\}$, k_i равномерно распределены, независимы; $Y = \{y_1, y_2, y_3\}$, где $y_i = x_i \oplus k_i$ (сложение по модулю 2).

Случай 1. Равномерное распределение открытого текста. Пусть $p(X) = 1/8$ для всех восьми возможных сообщений, указанных в таблице 1.

Таблица 1

Равномерное распределение входных символов

X	000	001	010	011	100	101	110	111
$p(X)$	1/8	1/8	1/8	1/8	1/8	1/8	1/8	1/8

Тогда для любого фиксированного Y получаем, что $p(Y|X) = 1/8$. Например, для $Y = 000$ и $X = 010$ имеем: $p(Y|X = 010) = p(K = 010) = 1/8$. При этом вероятность $p(Y) = \sum p(X) \cdot p(Y|X) = 1/8$ и условная вероятность $p(X|Y) = 1/8$. Следовательно, $p(X|Y) = p(X) = 1/8$ для всех X , т.е. шифр является совершенным для равномерного распределения открытых текстов.

В этом случае все элементы матрицы размера 8×8 совместного распределения $p(X|Y)$ будут равны $1/64$. Здесь для любых X и Y : $p(X|Y) = p(X) \cdot p(Y|X) = p(X) \cdot p(K = Y \oplus X) = 1/8 \cdot 1/8 = 1/64$. Дискретное пространство элементарных событий Ω содержит 64 элементарных событий, вероятность каждого равна $1/64$.

Случай 2. Неравномерное распределение (естественный язык). Пусть распределение открытых текстов неравномерно и задано таблицей 2.

Тогда для любого фиксированного Y сравнивая, априорное распределение: $p(X) = (0,4, 0,1, 0,1, 0,1, 0,1, 0,05, 0,05, 0,1)$ с апостериорным

распределением: $p(X|Y) = (0,4, 0,1, 0,1, 0,1, 0,1, 0,05, 0,05, 0,1)$, получим равенство $p(X|Y) = p(X)$.

Конструкция матрицы размера 8×8 совместного распределения $p(X|Y)$ представлена в таблице 3.

Случай 3. Вырожденное распределение. Пусть всегда передается сообщение $X = 000$ (таблица 4).

В этом случае равенство также сохраняется: $p(X|Y) = p(X)$. Например, для $Y = 000$ имеем: $p(Y = 000) = 1 \cdot 1/8 = 1/8$; $p(X = 000 | Y = 000) = 1 \cdot (1/8) / (1/8) = 1$.

Для других X : $p(X|Y = 000) = 0$.

Здесь матрица совместного распределения $p(X|Y)$ будет также размера 8×8 . При этом все элементы первой строки (соответствующей $X = 000$) равны $1/64$, а элементы остальных строк – равны 0.

Случай 4. Коррелированные сообщения. Рассмотрим случай, где биты сообщения коррелированы: $p(X = 000) = 0,3$; $p(X = 111) = 0,3$; $p(X = 001) = p(X = 010) = \dots = p(X = 110) = 0,4/6 = 0,2/3$.

Таблица 2

Неравномерное распределение входных символов

X	000	001	010	011	100	101	110	111
$p(X)$	0,4	0,1	0,1	0,1	0,1	0,05	0,05	0,1

Таблица 3

Конструкция матрицы совместного неравномерного распределения

XY	000	001	010	011	100	101	110	111
000	0,4/64	0,4/64	0,4/64	0,4/64	0,4/64	0,4/64	0,4/64	0,4/64
001	0,1/64	0,1/64	0,1/64	0,1/64	0,1/64	0,1/64	0,1/64	0,1/64
010	0,1/64	0,1/64	0,1/64	0,1/64	0,1/64	0,1/64	0,1/64	0,1/64
011	0,1/64	0,1/64	0,1/64	0,1/64	0,1/64	0,1/64	0,1/64	0,1/64
100	0,1/64	0,1/64	0,1/64	0,1/64	0,1/64	0,1/64	0,1/64	0,1/64
101	0,05/64	0,05/64	0,05/64	0,05/64	0,05/64	0,05/64	0,05/64	0,05/64
110	0,05/64	0,05/64	0,05/64	0,05/64	0,05/64	0,05/64	0,05/64	0,05/64
111	0,1/64	0,1/64	0,1/64	0,1/64	0,1/64	0,1/64	0,1/64	0,1/64

Таблица 4

Вырожденное распределение входных символов

X	000	001	010	011	100	101	110	111
$p(X)$	1	0	0	0	0	0	0	0

Конструкция матрицы совместного распределения коррелированных битов

$X \backslash Y$	000	001	010	011	100	101	110	111
000	0,3/64	0,3/64	0,3/64	0,3/64	0,3/64	0,3/64	0,3/64	0,3/64
001	0,2/192	0,2/192	0,2/192	0,2/192	0,2/192	0,2/192	0,2/192	0,2/192
010	0,2/192	0,2/192	0,2/192	0,2/192	0,2/192	0,2/192	0,2/192	0,2/192
011	0,2/192	0,2/192	0,2/192	0,2/192	0,2/192	0,2/192	0,2/192	0,2/192
100	0,2/192	0,2/192	0,2/192	0,2/192	0,2/192	0,2/192	0,2/192	0,2/192
101	0,2/192	0,2/192	0,2/192	0,2/192	0,2/192	0,2/192	0,2/192	0,2/192
110	0,2/192	0,2/192	0,2/192	0,2/192	0,2/192	0,2/192	0,2/192	0,2/192
111	0,3/64	0,3/64	0,3/64	0,3/64	0,3/64	0,3/64	0,3/64	0,3/64

Таблица 6

Конструкция битовой матрицы совместного распределения

$X \backslash Y$	000	001	010	011	100	101	110	111
000	000	001	010	011	100	101	110	111
001	001	000	011	010	101	10	111	110
010	010	011	000	001	110	111	100	101
011	011	010	001	000	111	110	101	100
100	100	101	110	111	000	001	011	011
101	101	100	111	110	001	000	011	010
110	110	111	100	101	010	011	000	001
111	111	110	101	111	011	010	001	000

Проверим для $Y = 101$: $p(Y = 101) = 0,3 \cdot 1/8 + 0,3 \cdot 1/8 + 6 \cdot 1/15 \cdot 1/8 = 1/8$; $p(X = 000 | Y = 101) = 0,3 \cdot 1/8 / (1/8) = 0,3$; $p(X = 111 | Y = 101) = 0,3 \cdot 1/8 / (1/8) = 0,3$.

Для других X : $p(X | Y = 101) = (1/8) \cdot (1/15) / (1/8) = 1/15$, т.е. равенство $p(X | Y) = p(X)$ сохраняется.

Конструкция матрицы размера 8×8 совместного распределения $p(X | Y)$ представлена в таблице 5.

Конструкция битовой матрицы размера $M \times N = 8 \times 8$ совместного распределения (для фиксированного K и X получаем определенный шифртекст Y) для шифра с тремя шифрвеличинами представлена в таблице 6.

Эта матрица универсальна и не зависит от $p(X)$ распределения вероятностей открытых текстов. При этом матрица совместного распределения $p(X | Y)$ выводится из битовой матрицы. Например, каждая ячейка (X, Y) в вероятностной матри-

це равномерного распределения (случай 1) получает вероятность $1/64$, потому что для каждой пары (K, X) в битовой матрице существует ровно один K , дающий нужный шифртекст Y .

Данный пример аналога шифра с тремя шифрвеличинами показывает, что условие совершенной секретности, выполненное для одного распределения открытых текстов (равномерного), выполняется и в расширенном смысле для любого распределения вероятностей входных символов.

Таким образом, в работе предложена кодовая конструкция, которая оказывается связанной с задачами теории информации, распределения вероятностей многомерных случайных величин и многомерной задаче о системах представителей. Доказана теорема об аналогах совершенных шифров, согласно которой свойство «совершенности» (отсутствия потерь информации) является глобальным свойством

канала. Если можно найти хотя бы одно распределение входных символов, при котором канал работает на пределе своих возможностей и без потерь, то этот канал по своей природе лишен внутренних помех, вызывающих потери информации, для любого способа его использования. Построен пример аналога шифра с тремя шифрвеличинами, иллюстрирующий предложенную кодовую конструкцию.

Полученные результаты отражают универсальность совершенных систем шифрования и делают теорию идеальной секретности пригодной для широкого спектра криптографических систем. Это, в свою очередь, дает возможность создания адаптивных систем связи, сохраняющих свои оптимальные свойства при изменении статистических характеристик передаваемых данных.

Литература

1. Шеннон К. Теория связи в секретных системах // Работы по теории информации и кибернетики. М.: Наука, 1963. – С. 333–402.
2. Медведева Н.В. Об аналогах теоремы Шеннона для совершенных шифров // CEUR Workshop Proceedings. 2016. Т. 1825. С. 232–239. ISSN 1613–0073.
3. Медведева Н. В. К разработке абсолютно стойких систем передачи данных / Н. В. Медведева, С. С. Титов // Вестник УрГУПС. – 2019. – № 2(42). – С. 34–43. – DOI 10.20291/2079-0392-2019-2-34-43. – EDN TVECM1.
4. Medvedeva N. V. Construction of minimal (with respect to inclusion) perfect ciphers / N. V. Medvedeva, S. S. Titov // International Scientific and Practical Conference "Railway Transport and Technologies" (RTT-2021): Collection of conference materials. Volume 2624, Ekaterinburg, November 24–25, 2021. Vol. 2624, Issue 1. – USA: AIP PUBLISHING, 2023. – P. 050018. – DOI 10.1063/5.0132399. – EDN LANBAV.
5. Medvedeva N. V. Construction of Absolutely Failure-free Minimal Data Transmission Systems on Railway Transport / N. V. Medvedeva, S. Titov // Transport: logistics, construction, Maintenance, management: Proceedings of the 1st International Scientific and Practical Conference on Transport: Logistics, Construction, Maintenance, Management, Yekaterinburg, March 17–30, 2022. – Portugal: SciTe-Press Digital Library (Science and Technology Publications, Lda), 2023. – P. 73–77. – EDN PHDLQE.
6. Медведева Н. В. Применение многократно стохастических матриц для построения абсолютно стойких систем передачи данных / Н. В. Медведева, С. С. Титов // Вестник УрГУПС. – 2023. – № 3(59). – С. 4–13. – DOI 10.20291/2079-0392-2023-3-4-13. – EDN GWKLMZ.
7. Медведева, Н. В. О совместных распределениях случайных величин при построении стойких кодов / Н. В. Медведева, С. С. Титов // Железнодорожный транспорт и технологии: сборник трудов Всероссийской научно-практической конференции с международным участием, Екатеринбург, 27–28 ноября 2024 года. – Екатеринбург: УрГУПС, 2025. – С. 410–413. – EDN XAWWUN.
8. Алферов А.П., Зубов А.Ю., Кузьмин А.С., Черемушкин А.В. Основы криптографии. М.: Гелиос АРБ, 2001. – 479 с.
9. Зубов А.Ю. Совершенные шифры. М.: Гелиос АРБ, 2003. – 160 с.

References

1. Shannon K. Teoriya svyazi v sekretnykh sistemakh // Raboty po teorii in-formatsii i kibernetiki. M.: Nauka, 1963. – S. 333–402.
2. Medvedeva N.V. Ob analogakh teoremy Shennona dlya sovershennykh shifrov // CEUR Workshop Proceedings. 2016. T. 1825. S. 232–239. ISSN 1613–0073.
3. Medvedeva N. V. K razrabotke absolyutno stoykikh sistem peredachi dan-nykh / N. V. Medvedeva, S. S. Titov // Vestnik UrGUPS. – 2019. – № 2(42). – S. 34–43. – DOI 10.20291/2079-0392-2019-2-34-43. – EDN TVECM1.
4. Medvedeva N. V. Construction of minimal (with respect to inclusion) perfect ciphers / N. V. Medvedeva, S. S. Titov // International Scientific and Practical Conference "Railway Transport and Technologies" (RTT-2021): Collection of conference materials. Volume 2624, Ekaterinburg, November 24–25, 2021. Vol. 2624, Issue 1. – USA: AIP PUBLISHING, 2023. – P. 050018. – DOI 10.1063/5.0132399. – EDN LANBAV.
5. Medvedeva N. V. Construction of Absolutely Failure-free Minimal Data Transmission Systems on Railway Transport / N. V. Medvedeva, S. Titov // Transport: logistics, construction, Maintenance, management: Proceedings of the 1st International Scientific and Practical Conference on Transport: Logistics, Construction, Maintenance, Management, Yekaterinburg, March 17–30, 2022. – Portugal: SciTe-Press Digital Library (Science and Technology Publications, Lda), 2023. – P. 73–77. – EDN PHDLQE.

6. Medvedeva N. V. Primeneniye mnogokratno stokhasticheskikh matrix dlya postroyeniya absolyutno stoykikh sistem peredachi dannykh / N. V. Medvedeva, S. S. Titov // Vestnik UrGUPS. – 2023. – № 3(59). – S. 4–13. – DOI 10.20291/2079-0392-2023-3-4-13. – EDN GWKLMZ.

7. Medvedeva, N. V. O sovместnykh raspredeleniyakh sluchaynykh velichin pri postroyenii stoykikh kodov / N. V. Medvedeva, S. S. Titov // Zheleznodo-rozhnyy transport i tekhnologii: sbornik trudov Vserossiyskoy nauchno-prakticheskoy konferentsii s mezhdunarodnym uchastiyem, Yekaterinburg, 27–28 noyabrya 2024 goda. – Yekaterinburg: UrGUPS, 2025. – S. 410–413. – EDN XAWWUN.

8. Alferov A.P., Zubov A.YU., Kuz'min A.S., Cheremushkin A.V. Osnovy kriptografii. M.: Gelios ARV, 2001. – 479 s.

9. Zubov A.YU. Sovershennyye shifry. M.: Gelios ARV, 2003. – 160 s.

МЕДВЕДЕВА Наталья Валерьевна, кандидат физико-математических наук, доцент, доцент кафедры «Естественнонаучные дисциплины» ФГБОУ ВО «Уральский государственный университет путей сообщения». 650034, г. Екатеринбург, ул. Колмогорова, 66. E-mail: medvedeva_n_v@mail.ru.

ТИТОВ Сергей Сергеевич, доктор физико-математических наук, профессор, профессор кафедры «Естественнонаучные дисциплины» ФГБОУ ВО «Уральский государственный университет путей сообщения». 650034, г. Екатеринбург, ул. Колмогорова, 66. 620034, г. Екатеринбург, ул. Колмогорова, д. 66. E-mail: stitov@usaaa.ru.

MEDVEDEVA Natalya Valeryevna, Candidate of Physico-Mathematical Sciences, Associate Professor, Associate Professor of the Department of Natural Sciences of the Federal State Budgetary Educational Institution of Higher Education «Ural State University of Railway Transport». 620034, Yekaterinburg, str. Kolmogorova, 66. E-mail: medvedeva_n_v@mail.ru.

TITOV Sergey Sergeevich, Doctor of Physico-Mathematical Sciences, Professor, Professor of the Department of Natural Sciences of the Federal State Budgetary Educational Institution of Higher Education «Ural State University of Railway Transport». 620034, Yekaterinburg, str. Kolmogorova, 66. E-mail: stitov@usaaa.ru.

***Материалы к публикации отправлять по адресу E-mail: urvest@mail.ru
в редакцию журнала «Вестник УрФО. Безопасность в информационной сфере».***

***Или по почте по адресу: Россия, 454080, г. Челябинск, пр. им. Ленина, д. 76, ЮУрГУ,
Издательский центр***

ВЕСТНИК УрФО

Безопасность в информационной сфере № 4(58) / 2025

Подписано в печать 29.12.2025. Дата выхода в свет 29.12.2025.

Формат 70×108 1/16. Печать цифровая. Усл.-печ. л. 7,79. Тираж 50 экз.

Заказ XXX/XXX.

Цена свободная.

Отпечатано в типографии Издательского центра ФГАОУ ВО «ЮУрГУ (НИУ)».
454080, г. Челябинск, пр. им. В. И. Ленина, 76, ЮУрГУ, Издательский центр.

Bulletin of the Ural Federal District

Security in the Sphere of Information No. 4(58) / 2025

Signed to print 29.12.2025. Date of publication of the 29.12.2025.

Format 70×108 1/16. Screen printing. Conventional printed sheet 7,79. Circulation – 50 issues.

Order XXX/XXX.

Open price.

Printed in the printing house of the Publishing Center of FGAOU VO «SUSU (NIU)».
SUSU, Publishing Center, 76, Lenina Str., Chelyabinsk, 454080