

БЕЗПЕКА ІНФОРМАЦІЙНИХ СИСТЕМ І ТЕХНОЛОГІЙ

№ 1/2 (3/4) 2020

НАУКОВИЙ ЖУРНАЛ

ISSN 2707-1758

КИЇВСЬКИЙ НАЦІОНАЛЬНИЙ УНІВЕРСИТЕТ ІМЕНІ ТАРАСА ШЕВЧЕНКА

Представлено результати досліджень за такими напрямками: інформаційна та кібернетична безпека, комп'ютерні науки та інформаційні технології, інформаційно-комунікаційні системи, телекомунікації та радіотехніка.

Для науковців, докторантів, аспірантів, фахівців відповідних напрямів, студентів.

ГОЛОВНИЙ РЕДАКТОР	Н. В. Лукова-Чуйко , д-р техн. наук, проф.
ЗАСТУПНИК ГОЛОВНОГО РЕДАКТОРА	С. В. Толюпа , д-р техн. наук, проф.
ВІДПОВІДАЛЬНІ СЕКРЕТАРІ	О. А. Лаптєв , д-р техн. наук, ст. наук. співроб. С. Ю. Даков , канд. техн. наук А. О. Фесенко , канд. техн. наук
РЕДАКЦІЙНА КОЛЕГІЯ	Т. В. Бабенко , С. С. Бучик , А. О. Білошицький , О. Є. Іларіонов , О. Є. Колесніков , Ю. В. Кравченко , О. О. Кузнецов , В. В. Морозов , А. П. Мусієнко , В. С. Наконечний , Л. Т. Пархуць , В. Л. Плєскач , В. Г. Сайко , І. Ю. Субач , В. Є. Снитюк , В. Л. Шевченко , Ю. Н. Аміргалієв , Б. С. Ахметов , М. П. Явич , Г. П. Димитров , В. І. Гопєєнко , У. Мілош , В. Рудзьоніс
Адреса редколегії	вул. Богдана Гаврилишина, 24, Київ, 04116, Україна ☎ факс (+38044) 481 44 07, e-mail: fit@univ.net.ua
Затверджено	Вченою радою факультету інформаційних технологій Київського національного університету імені Тараса Шевченка 29.06.21 (протокол № 20)
Зареєстровано	Міністерством інформації України. Свідоцтво про державну реєстрацію КВ № 23831-13671Р від 15.03.2019
Засновник і видавець	Київський національний університет імені Тараса Шевченка, Видавничо-поліграфічний центр "Київський університет". Свідоцтво внесено до Державного реєстру ДК № 1103 від 31.10.02
Адреса видавця	ВПЦ "Київський університет", б-р Тараса Шевченка, 14, м. Київ, 01601, Україна ☎ (38044) 239 31 72, 239 32 22; факс 239 31 28

INFORMATION SYSTEMS AND TECHNOLOGIES SECURITY

№ 1/2 (3/4) 2020

SCIENTIFIC JOURNAL

ISSN 2707-1758

TARAS SHEVCHENKO NATIONAL UNIVERSITY OF KYIV

The scientific publication presents the results of research in such areas as Information and Cyber Security, Computer Science and Information Technology, Information and Communication Systems, Telecommunications and Radio Engineering.

For scientists, doctoral students, graduate students, specialists, students.

EDITOR-IN-CHIEF	N. Lukova-Chuiko , Dr Hab, Prof.
DEPUTY OF EDITOR-IN-CHIEF	S. Toliupa , Dr Hab, Prof.
EXECUTIVE SECRETARIES	O. Laptev , Dr Hab, Senior Researcher S. Dakov , PhD A. Fesenko , PhD
EDITORIAL BOARD	T. V. Babenko , S. S. Buchyk , A. O. Biloshitskyi , O. E. Ilarionov , O. E. Kolesnikov , Yu. V. Kravchenko , O. O. Kuznetsov , V. V. Morozov , A. P. Musienko , V. S. Nakonechnyi , L. T. Parkhuts , V. L. Pleskach , V. G. Saiko , I. Yu. Subach , V. E. Snityuk , V. L. Shevchenko , Y. N. Amirgaliyev , B. S. Akhmetov , M. P. Iavich , G. P. Dimitrov , V. I. Gopejenko , U. Miloš , V. Rudzionis
Editorial board address	Bohdan Hawrylyshyn str. 24, UA-04116, Kyiv, Ukraine ☎ fax (+38044) 481 44 07, e-mail: fit@univ.net.ua
Confirmed	Academic Council of the Faculty of Information Technology of Taras Shevchenko National University of Kyiv 29.06.21 (протокол № 20)
Registered	Ministry of Information of Ukraine. State registration certificate KB № 23831-13671P from 15.03.2019
Founder and Publisher	Taras Shevchenko National University of Kyiv, Publishing and Polygraphic Center "Kyiv University". The certificate is added to registry ДК № 1103 from 31.10.02
Publisher's address	Publishing and Polygraphic Center "Kyiv University" 14, Taras Shevchenko blvd., Kyiv, 01601, Ukraine ☎ (38044) 239 3172, 239 32 22; fax 239 31 28

КОМП'ЮТЕРНІ НАУКИ Й ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ

В. В. Морозов, О. О. Мезенцева

Використання навчальних нейронних мереж
для прогнозування подій у розробленні продукції ІТ 5

Д. Коротін, С. В. Поперешняк

Аналіз RTB-платформ для автоматизації купівлі та продажу медіаконтенту 13

Г. А. Гайна

Тенденції розвитку штучного інтелекту в Україні 20

Г. М. Гнатієнко, О. Круглов, Н. П. Тмєнова

Обчислення результуючого ранжування альтернатив
на основі використання неповних експертних ранжувань 27

ІНФОРМАЦІЙНА ТА КІБЕРНЕТИЧНА БЕЗПЕКА

С. В. Толюпа, Н. В. Лукова-Чуйко,

В. С. Наконечний, В. Г. Сайко, О. М. Кулінич

Формування стратегії управління режимами роботи систем захисту
на основі моделі ігрового управління 37

О. І. Махович, Р. В. Миколайчук, В. Г. Миколайчук

Рекомендації щодо вибору способу безпечного зберігання паролів 47

В. І. Ізніска, Д. Вдовенко

Аналіз методів захисту даних 52

КОМП'ЮТЕРНА ІНЖЕНЕРІЯ ТА ПРОГРАМНЕ ЗАБЕЗПЕЧЕННЯ

Б. Бондаренко, Ю. Я. Самохвалов

Пошук мультимедійної інформації на основі нейронних мереж 57

М. Завалій, А. Завалій, Л. В. Дакова,

С. Ю. Даков, І. І. Пархоменко

Дослідження безсерверних обчислень
та їхнє використання в мережах мобільного зв'язку 63

С. В. Толюпа, В. С. Наконечний, В. Г. Сайко,

М. Котов, В. Солодовник

Використання симетричних алгоритмів шифрування
для передачі сигналів у пристроях бездротового введення даних 70

CONTENTS

COMPUTER SCIENCE AND INFORMATION TECHNOLOGIES

V. Morozov, O. Mezentseva Use Training Neural Networks for Predicting Product Development of IT Project.....	5
D. Korotin, S. Popershiak Analysis of RTB platforms for automating the purchase and sale of media content.....	13
G. Haina Trends in the development of artificial intelligence in Ukraine	20
H. Hnatienko, O. Kruglov, N. Tmenova Calculation of the resulting ranking of alternatives based on the use of incomplete expert rankings	27

INFORMATION AND CYBER SECURITY

S. Toliupa, N. Lukova-Chuiko, V. Nakonechnyi, V. Saiko, O. Kulinich Formation of a strategy for managing the modes of operation of protection systems based on the game management model.....	37
O. Makhovych, R. Mykolaichuk, V. Mykolaichuk Recommendations for choosing a method of safe storage of passwords.....	47
V. Ignisca, D. Vdovenko Analysis of methods data security	52

COMPUTER ENGINEERING AND SOFTWARE

B. Bondarenko, Yu. Samokhvalov Searching for multimedia information based on neural networks	57
M. Zavaliy, A. Zavaliy, L. Dakova , S. Dakov, I. Parkhomenko Research of serverless computing and its use in mobile communication networks	63
S. Toliupa, V. Nakonechnyi, V. Saiko, M. Kotov, V. Solodovnik Use of symmetric encryption algorithms for signal transmission in wireless data input devices	70



КОМП'ЮТЕРНІ НАУКИ Й ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ

УДК 005.8:316.422

DOI <https://doi.org/10.17721/ISTS.2020.4.3-10>V. V. Morozov, orcid.org/0000-0001-7946-0832,
knumvv@gmail.comO. O. Mezentseva, orcid.org/0000-0002-8430-4022,
olga.mezentseva.fit@gmail.com

Taras Shevchenko National University of Kyiv, Kyiv, Ukraine

USE TRAINING NEURAL NETWORKS FOR PREDICTING PRODUCT DEVELOPMENT OF IT PROJECT

The state of development of innovations in Ukraine is characterized by an increase in development on the basis of start-up projects with the use as a project product of information systems of varying complexity. The article analyzes the weak survivability of the results of start-up projects. The conclusion on the need to predict the stages of development of IT project products based on the analysis of the processes of interaction of users (customers) with the information system (product). In this article, components of the model of forecasting of IT products development of innovative start-up projects are considered based on the analysis of formed datasets of the interactions of prospective clients. We offered the algorithm of formation of initial datasets based on Customer Journey Map (CJM), which are the tool of fixing of events of the interaction of clients with the system. Examples of models of analogues of clients' travel maps are given, which are the basis for recording and analyzing interactions. This fact is the basis for the formation of appropriate data sets of large dimension. As a mechanism for processing big data sets and building strategies for IT products development, it is proposed to use a learning neural network. Mathematical models for further modeling and analysis of the obtained results are built. We used a simple linear regression analysis to model the relationship between a single explanatory variable and a continuous response variable (dependent variable). An exploratory data analysis method was applied to the available data to find repetitive patterns and anomalies. In the course of the research, we constructed a model of linear regression implementation using the gradient optimisation approach. The linear models of the scikit-learn library for the regression task were also applied, and the stabilisation regression method was implemented. Modelling and analysis of the obtained results were carried out, which showed greater efficiency over the extended life cycle of IT project products.

Keywords: start-up; information interaction; customer journey map; forecasting.

1. INTRODUCTION

At the current stage of development of information systems engineering technologies, innovations in the form of start-up projects are becoming increasingly important [1]. Such projects are often formalised as separate commercial enterprises to obtain funding and financial profitability. Often, such companies are represented with the implementation of SaaS (Software as a Service) distribution model [2] or B2B (Business to Business) transactions [3]. Such business models have problems with long sales cycles, as the decision of the client is collective and depends on many factors. The success of the product development and,

accordingly, the efficiency of the enterprise depends on the indicator of the quality of customer service at information interaction with the IT system. However, to define directions of development of such products, it is not simple. For this purpose, it is necessary to form datasets during some period, characterising points of preferable information interaction of different kinds of clients, and also to spend analytics and to predict directions of such development.

The method of information interaction that has proved to be effective in many projects can be used to solve these problems. It is based on the analysis of the "journey" (interaction with IT product) of the

© Morozov V. V., Mezentseva O. O., 2020



prospective client (Customer Journey Map – CJM) [4, 5]. Also, it is necessary to combine millions of events, which will provide the necessary analysis of the impact on the customers' journeys and will determine the future content of projects to develop such IT products.

By analysing [6] millions of real-time interaction data points, it will be possible to find the most critical customer movement events in the system [7] and prioritise those opportunities that have a significant impact on business objectives, such as increasing revenue, reducing customer churn, improving customer service [8] and developing innovations in the company.

However, sufficient constructive and technological complexity of construction of programs of development of complex IT products demands the use of the project approach [9–12] wherein processes of creation and development of such systems methods and information technologies of project management are applied.

Experience shows that for analytical processing of considerable volumes of the information of interaction of clients with IT system use of methods of artificial intelligence (AI) will have a significant influence on the efficiency of development programs (of projects) creation.

The search for optimal variants of distribution of resources at the preparation of development programs in innovative projects can reduce terms of performance of project tasks and, as a result, decrease their cost. At the same time, predicting the impact of interaction with clients [13] on the variants of development programs in such projects is a multidimensional task that can be solved using technologies of modern artificial neural networks [14]. Besides, it is necessary to take into account numerous changes [15] that affect various parameters of innovative projects and significantly affect the negative results of their implementation.

Thus, consideration of the possibilities of experimental use of the Customer Journey Map and artificial neural networks in the research of improving the quality of interaction of numerous clients based on optimal development programs of start-up projects is an actual challenge.

2. ANALYSIS OF RECENT RESEARCH AND PUBLICATIONS

Application of the project approach to start-up development and effective use of products of these projects at the creation of modern IT were considered in research papers [9–16]. In [9], new models of management of innovations in projects based on the use of Markov processes are offered. This approach can be used for further development

of start-up projects. However, there is no assessment of the value of products for customers of such projects. The [10] contains a description of the competency models that can be applied to Start-up project teams by the customer. However, there are no methods of interaction between users and the customer's team here.

The [11] provides a detailed description of the process approach for managing any kind of projects, and the specifics of innovation project management are considered separately. However, the processes of management of the creation and development of a product are not given. At the same time, [12] describes integrated methodologies for organisational and product management of development projects. However, methods of an estimation of the efficiency of communication interaction are insufficiently described. In [13, 14, 16] approaches to the management of complex IT projects based on proactive (anticipatory, predictive) management with the use of methods of artificial intelligence, in particular, with the use of artificial neural networks, have been proposed. However, studies based on interaction with clients have not been conducted.

The use of methods to assess customer interactions in SaaS and B2B business models based on Customer Journey Map has been highlighted in [4, 17–19]. As noted in these papers, such developments should be based on new tools to assess the productivity (value) of customer interaction with the IT product as feedback. Here we can use well-proven customer journey models though it is not specified, how to process rather big data on numerous indicators analytically.

Then, considering aspects of application of intellectual tools for the analysis and forecasting based on the processing of a big data set from the interaction of clients, it is possible to consider papers [20–27]. In these papers, it is noted that the most acceptable for us, in terms of intellectual forecasting of the composition of programs for the development of complex IT products, is the use of trainable neural networks. Furthermore, it will be useful in our case to build a model of linear regression implementation [23] using the gradient optimisation approach [27].

The goal of the article is to justify and develop a model for predicting the composition of the program for the development of complex products of start-up projects using the analysis of customer interaction data based on the Customer Journey Map and the application of trainable neural networks.

The main objectives of this study are as follows:

1. Identify the interaction elements that form the basis of the Customer Journey Map to create survey datasets within one year.
2. Development of a model for forecasting the composition of development programs, taking into



account maximum customer loyalty and retention in their interaction with the IT system.

3. Development of algorithms for modelling the neural network training processes and setting up the forecasting processes.

3. MODEL DEVELOPMENT AND USE OF MODELLING METHOD

A. Constructing customer interaction models with IT product based on journey maps

As mentioned above, one of the effective tools to assess the quality of user interaction with the IT product is the customer journey map (CJM). It is a tool for visualising of the interaction of the consumer with a product or service. Creating a CJM is both a process of analysis and a method for generating ideas to improve a product or service.

CJM displays a time-bound interaction broken down into small components. The components of interactions refer both to the process (consumers' goals and objectives, their actions, expected results,

problems and barriers preventing the transition to the next step, touchpoints, materials, tools, equipment, KPI from the business point of view, etc.) and to the psycho-emotional state of the consumer (thoughts, feelings, emotions).

This approach to the development of products characterized by multi-channel interaction is particularly useful, that is, for those cases where the customer and the product have many "touchpoints". The customer always has a sum of impressions of the interaction with the product through all available channels, so that even one negative experience is able to vilify the entire product in the eyes of the client and force him to cancel the subscription.

As an example, we can consider the CJM depicted in Fig. 1, which shows the customer's journey from awareness, consideration, purchase and further use of the "instant server" in a particular web application (IaaS distribution model) of the multinational telecommunications company ("Telefónica").

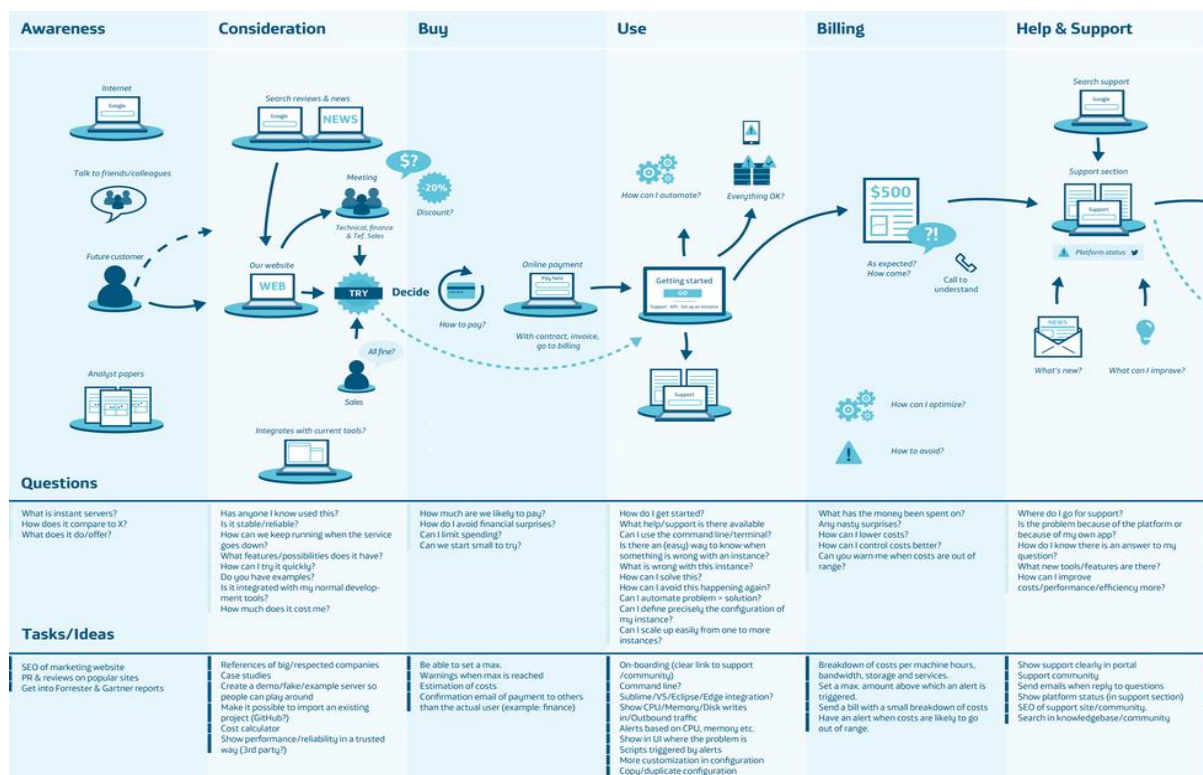


Fig. 1. Example of Customer Journey Map

B. Building predictive models using neural networks

Regression models are used to predict target variables on a continuous scale, which makes them useful for many scientific issues and information industry applications, such as understanding relationships between variables, assessing trends, or making forecasts.

One example of their application can be the prediction of company sales in the coming months:

$$y = w_0 + w_1x, \quad (1)$$

where w_0 – weight, that is the intersection point of the Y-axis; w_1 – the explanatory variable coefficient. The most frequently used is multiple linear regression y :



$$y = w_0x_0 + w_1x_1 + \dots + w_mx_m = \sum_{i=1}^m w_ix_i = w^T x, \quad (2)$$

where w_0 is the intersection point of the Y-axis at m – the number of regression members; w^T – the explanatory variable coefficient.

To determine the number of linear relationships between features, we will now create a correlation matrix.

The correlation coefficients are limited to the range $[-1, 1]$. Two attributes have, respectively, absolutely positive correlation if $r = 1$, no correlation if $r = 0$, and absolutely negative correlation if $r = -1$.

As mentioned earlier, Pearson's correlation coefficient can be calculated simply as the covariance between two attributes x and y – numerator, divided by the product of their standard deviations (denominator):

$$r = \frac{\sum_{i=1}^n [(x^i - \mu_x)(y^i - \mu_y)]}{\sqrt{\sum_{i=1}^n (x^i - \mu_x)^2} \sqrt{\sum_{i=1}^n (y^i - \mu_y)^2}} = \frac{\sigma_{xy}}{\sigma_x \sigma_y}, \quad (3)$$

where n – the number of attributes; μ – the empirical average of these attributes; σ – covariance between attributes.

As can be seen in the resulting figure, the correlation matrix provides us with another final diagram, which can help us select attributes based on their corresponding linear correlations (Fig. 2).

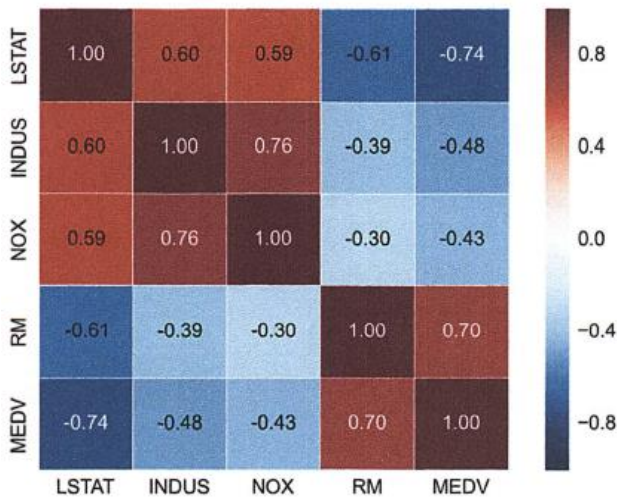


Fig. 2. An example of a correlation matrix

In order to fit a linear regression model, we are interested in those features that have a high correlation with our target variable MEDV. Looking at the above correlation matrix, we can see that our target variable MEDV shows the highest correlation with variable LSTAT (-0.74). The correlation between RM and MEDV is also relatively high (0.70). In the presence of

a linear relationship between these two variables that we have observed in the dot matrix as an explanatory variable, the artificial neuron uses a linear activation function. Moreover, we have defined the $JS(w)$ cost function that we have minimised for weight extraction thanks to optimisation algorithms such as gradient descent (GD) [23] and stochastic gradient descent (SGD) [24].

This cost function in ADALINE is the sum of squared errors (SSE). It is identical to the least squared value function (MLS), which we defined as follows:

$$JS(w) = \frac{1}{2} \sum_{i=1}^n (y^i - \hat{y}^i)^2, \quad (4)$$

where $\hat{y}$ – it is a predicted value $\hat{y} = w^T x$; (note that the coefficient $1/2$ is used simply for the convenience of obtaining an update rule for gradient descent).

In essence, linear regression on MLS can be understood as ADALINE without a single step function. As a result, we get continuous target values instead of class labels – 1 and 1.

4. EXPERIMENTATION

As our target variable, which we want to predict using one or more of these 51 explanatory variables, we will consider the tariffs for interaction with the system (a product of the start-up project) called MEDV.

Before we continue the exploration of this data set, we will put it into *DataFrame* data table of the *pandas* library [27] from the UCI repository (Fig. 3).

```
import pandas as pd

df = pd.read_csv(url, header=None, sep='\s+')
df.columns = ['CRIM', 'ZN', 'INDUS', 'CHAS',
              'NOX', 'RM', 'AGE', 'DIS', 'RAD',
              'TAX', 'PTRATIO', 'B', 'LSTAT', 'MEDV']
df.head()
```

Fig. 3. Downloading the source data

During the training of a linear regression model, it is not required that explanatory or target variables are distributed normally.

First, let us create a matrix of dotted charts that allows visualising pairwise correlations between different attributes in one place in this dataset. To prepare the matrix of point graphs, let's use the *pairplot* function from Python *seaborn* library [30], which is developed based on the *matplotlib* library for building statistical charts (Fig. 4).

Using the *corrcoef* function of the NumPy library on five columns of attributes and the heatmap function of the *seaborn* library, we will get the necessary relationships between the initial indicators



```
import matplotlib.pyplot as plt
import seaborn as sns
sns.set(style='whitegrid', context='notebook')
cols = ['LSTAT', 'INDUS', 'NOX', 'RM', 'MEDV']
sns.pairplot(df[cols], size=2.5)
plt.show()
```

Fig. 4. An example of a matrix of correlation links and point graphs

The following code is used to prepare an array of correlation matrices (Fig. 5).

```
import numpy as np
cm = np.corrcoef(df[cols].values.T)
sns.set(font_scale=1.5)
hm = sns.heatmap(cm,
                  cbar=True,
                  annot=True,
                  square=True,
                  fmt='.2f',
                  annot_kws={'size': 15},
                  yticklabels=cols,
                  xticklabels=cols)
plt.show()
```

Fig. 5. Example of preparing an array of correlation matrices in the form of a heatmap

Linear regression on MLS can be understood as ADALINE without a single step function, resulting in continuous target values instead of class labels 1 and 1 (Fig. 6).

```
class LinearRegressionGD(object):
    def __init__(self, eta=0.001, n_iter=20):
        self.eta = eta
        self.n_iter = n_iter

    def fit(self, X, y):
        self.w_ = np.zeros(1 + X.shape[1])
        self.cost_ = []

        for i in range(self.n_iter):
            output = self.net_input(X)
            errors = (y - output)
            self.w_[1:] += self.eta * X.T.dot(errors)
            self.w_[0] += self.eta * errors.sum()
            cost = (errors**2).sum() / 2.0
            self.cost_.append(cost)
        return self

    def net_input(self, X):
        return np.dot(X, self.w_[1:]) + self.w_[0]

    def predict(self, X):
        return self.net_input(X)
```

Fig. 6. Part of a machine learning program for a classification task without a single step function

To see our *LinearRegressionGD* linear regressor in action, we will use the RM (Profit volume) variable from the IT product development prediction dataset of innovative start-up projects as an explanatory model training variable that can predict MEDV (Interaction Tariffs). Moreover, we standardise the variables for better convergence of

the gradient descent algorithm. The corresponding source code looks like this (Fig. 7).

```
X = df[['RM']].values
y = df['MEDV'].values

from sklearn.preprocessing import StandardScaler
sc_x = StandardScaler()
sc_y = StandardScaler()
X_std = sc_x.fit_transform(X)
y_std = sc_y.fit_transform(y)
lr = LinearRegressionGD()
lr.fit(X_std, y_std)
```

Fig. 7. Standardisation of variables for gradient descent algorithm

Let us plot the value against the number of epochs to check the convergence of the linear regression. As we can see in the chart below, the algorithm of the gradient downturn converged after the fifth epoch (Fig. 8).

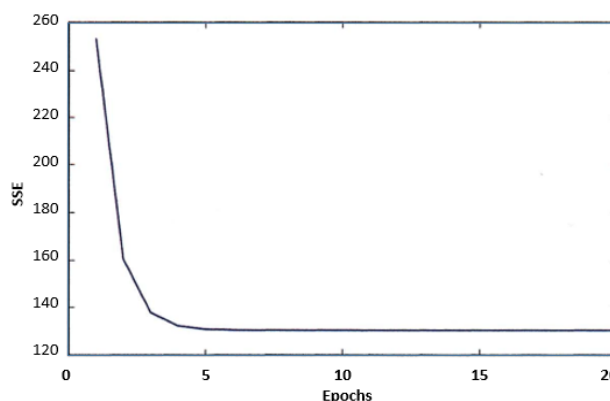


Fig. 8. Graph of cost in relation to the number of epochs

Note that it is more efficient (fewer emissions) to work with non-standardized variables in the linear regression object *LinearRegression* library *scikit-learn*, which uses the dynamic library *LIBLINEAR* and advanced optimisation algorithms. The RANSAC model is used in this case.

As we can see from the execution of the previous source code, the *LinearRegression* model of the *scikit-learn* library, fitted with non-standardized variables RM and MEDV, produced other model coefficients.

Let us compare it with our implementation based on gradient descent by constructing the MEDV chart in relation to RM (Fig. 9).

```
def lin_regplot(X, y, model):
    plt.scatter(X, y, c='blue')
    plt.plot(X, model.predict(X), color='red')
    return None
```

Fig. 9. A code fragment for plotting a point chart of linear regression



Having built a graph of training data and performed the model fitting by executing the above source code, we can now see that the overall result looks identical to our implementation based on the gradient descent (Fig. 10).

Using the RANSAC model, we have reduced the potential impact of emissions in this dataset, but we do not know if this approach has a positive impact on predictive capacity on previously unseen data.

After executing the next source code, we should see a residual graph with a line passing through the beginning of the X-axis, as shown below (Fig. 11).

In case of perfect prediction, the residues would be strictly zeros, which in real and practical applications we will probably never encounter.

However, we expect from a good regression model that the errors are distributed randomly, and the remains are randomly scattered around the midline.

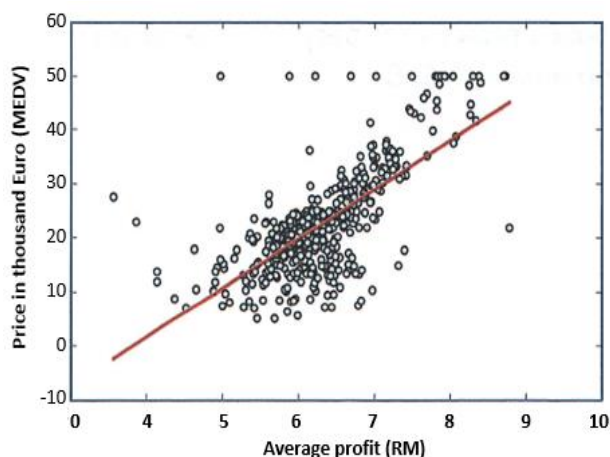


Fig. 10. The trend of IT projects development depending on the level of profit, based on training data

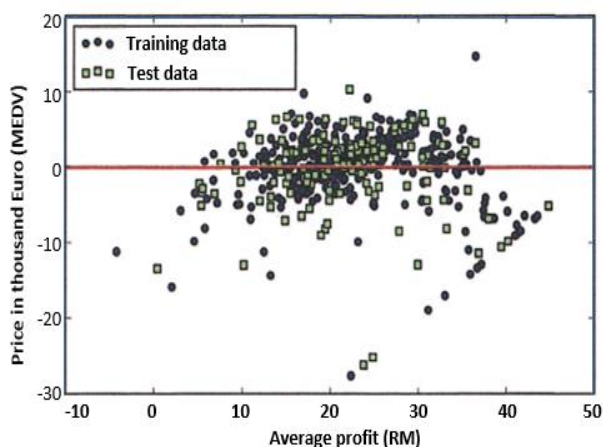


Fig. 11. Errors in fitting training data to the test data

In the case of an ideal prediction, the residuals would be strictly zeros, which we will probably never encounter in real and practical applications.

However, from a good regression model, we expect errors to be randomly distributed, and residuals randomly scattered around the midline

Another useful quantitative measure for assessing the quality of a model is the so-called weighted mean square error (mean squared error, MSE), that is, simply the average value of the SSE cost function, which we minimize to fit the linear regression model. MSE is useful for comparing different regression models or for fine-tuning their parameters by searching the parameter grid and cross-checking:

$$MSE = \frac{1}{N} \sum_{i=1}^n (y^i - \hat{y}^i)^2 \quad (5)$$

We will see that the MSE on the training set is 19.96, and the MSE of the test set is much larger with a value of 27.20.

```
>>>
from sklearn.metrics import mean_squared_error
print('MSE тренировка: %.3f, тестирование: %.3f' % (
    mean_squared_error(y_train, y_train_pred),
    mean_squared_error(y_test, y_test_pred)))
```

Fig. 12. Errors in fitting training data to the test data

5. CONCLUSION

The analysis of the considered approaches to the definition of effective interaction of customers with a mass information system allows defining use of modern mechanisms for the analysis on the base of Customer Journey Maps. In this case, big data sets are generated, which cannot be analysed by traditional methods. To solve this problem, the authors used machine learning algorithms of neural network types.

An exploratory data analysis method was applied to the available data to find repetitive patterns and anomalies, which proved to be effective. It has allowed constructing a model of realisation of linear regression with the use of the approach based on gradient optimisation.

This approach allowed to construct a base for prediction of processes of the development program of the start-up project product.

REFERENCES

- [1] Trends in the development of the global market for information technology. [Online]. Available: <http://eir.pstu.edu/handle/123456789/4299>
- [2] Euripidis Loukis, Marijn Janssen, Ianislav Mintchev, Determinants of software-as-a-service benefits and impact on firm performance, Decision Support Systems, Volume 117, 2019, pp. 38–47.
- [3] B. Brown, K. Swani, Introduction to the special issue: B2B advertising, Industrial Marketing Management, February 2020, DOI: 10.1016/j.indmarman.2020.02.006



- [4] Customer Journey Map: how to understand what the consumer needs. (2019) Available at: <https://www.uplab.ru/blog/customer-journey-map/>
- [5] S.Gordon, M. Linoff, Data mining techniques: for marketing, sales, and customer relationship management. Published by Wiley Publishing Inc., 10475 Crosspoint Boulevard, Indianapolis, 830 p., 2011.
- [6] IBM Institute for Business Value. Analytics: The real-world use of big data in consumer products, 2013.
- [7] Seductive Interaction Design: Creating Playful, Fun, and Effective User Experiences (Voices That Matter) 1st Edition (2019) Available at: <https://www.amazon.com/Seductive-Interaction-Design-Effective-Experiences/dp/0321725522/>
- [8] A. Zudov, Modeling of potential rule interactions in active databases. In journal "Modern problems of science and education" (2015). № 1–1.; URL: <http://www.science-education.ru/ru/article/view?id=17745>.
- [9] V. Gogunskii, O. Kolesnikov, G. Oborska, S. Harelik, D. Lukianov, Representation of project systems using the Markov chain, Eastern-European Journal of Enterprise Technologies, № 1/3 (85), 2017 pp. 25–32.
- [10] International Project Management Association. Individual Competence Baseline Version 4.0. International Project Management Association, 432 p., 2015.
- [11] A Guide to the project management body of knowledge (PMBOK guide). Sixth Edition – USA: PMI Inc., 537 p., 2017.
- [12] R. Turner, Guide to project-based management, tran. from English, Moscow, Grebennikov Publishing House, 552 p., 2007.
- [13] V. Morozov, O. Kalnichenko, S. Bronin, Development Of The Model Of The Proactive Approach in Creation Of Distributed Information Systems. Eastern-European Journal of Enterprise Technologies, № 43/2 (94), pp. 6–15 (2018).
- [14] A. Timinsky, O. Voitenko, I. Achkasov, Competence-based knowledge management in project oriented organisations in bi-adaptive context. Proceedings of the IEEE 14th International Scientific and Technical Conference on Computer Sciences and Information Technologies (CSIT–2019). – Lviv, 2019, pp. 17–20.
- [15] O. Maimon, L. Rokach, Data Mining and Knowledge Discovery Handbook, 2005, [Online]. Available: <http://www.bookmetrix.com/detail/book/ae1ad394-f821-4df2-9cc4-cbf8b93edf40>
- [16] V. Morozov, O. Kalnichenko, M. Proskurin, O. Mezentsseva, Investigation of Forecasting Methods the State of Complex IT-Projects With Using Deep Learning Neural Networks, Published in the book "Lecture notes in computational intelligence and decision making" (series "Advances in intelligent systems and computing"), vol. 1020, 2020, pp. 261–280.
- [17] E. Loukis, M. Janssen, I. Mintchev, Determinants of software-as-a-service benefits and impact on firm performance, Decision Support Systems, Volume 117, 2019, pp. 38–47.
- [18] K. Swani, B. Brown, Su. Mudambi, The untapped potential of B2B advertising: A literature review and future agenda, Industrial Marketing Management, 2019.
- [19] D. Herhausen, K. Kleinlercher, P. Verhoef, T. Rudolph, Loyalty Formation for Different Customer Journey Segments, Journal of Retailing, 2019.
- [20] V. Morozov, O. Kalnichenko, M. Proskurin, Methods of Proactive Management of Complex Projects Based on Neural Networks, Proceedings of the 2019 10th IEEE International Conference on Intelligent Data Acquisition and Advanced Computing Systems: Technology and Applications, 2019, pp. 964–969.
- [21] D. Garaedagi, System thinking. How to manage chaos and complex processes. Platform for modeling business architecture, Grevtsov Buks (Grevtsov Publisher). 480 p., 2011.
- [22] A. Novikov, Ezhov A.A., Rosenblatt's multilayer neural network and its application for solving the problem of signature recognition, Bulletin of TSU. Technical science. No. 2. pp. 188–197, 2016.
- [23] D. Komashinsky, D. Smirnov, Neural networks and their use in control and communication systems, Hotline-Telecom. p. 94, 2003.
- [24] Li, L., Fan, K., Zhang, Z., Xia, Z. Community detection algorithm based on local expansion K-means, Neural Network World, 26(6), 2016, pp.589–605.
- [25] I. Prigogine, G.Nikolis, Knowledge of the complex. Introduction, Per. from English, M.: Lenar, 360 p., 2017.
- [26] G. Carmantini, S. Rodrigues, P. Graben, M. Desroches, A modular architecture for transparent computation in recurrent neural networks, Neural Networks, #85, 2017, pp.85–105.
- [27] A. Hosseini, A non-penalty recurrent neural network for solving a class of constrained optimization problems, Neural Networks, 2016, #73, pp.10–25
- [28] A. Polaine, L. Løvlie, Service Design: From Insight to Implementation, Rosenfeld Media, 2013, 216 p.
- [29] S. Anderson, Seductive Interaction Design: Creating Playful, Fun, and Effective User Experiences, New Riders; 240 p., 2011 <https://www.amazon.com/Seductive-Interaction-Design-Effective-Experiences/dp/0321725522/>
- [30] Installing and getting started, [Online]. Available: <https://seaborn.pydata.org/installing.html>

Стаття надійшла до редколегії

18.10.2020



Використання навчальних нейронних мереж для прогнозування подій у розробленні продукції ІТ

Стан розвитку інновацій в Україні характеризується збільшенням розробок на основі стартап-проектів із використанням як продукту проекту інформаційних систем різної складності. Проведено аналіз слабкої живучості результатів виконання стартап-проектів. Зроблено висновок щодо необхідності прогнозування етапів розвитку продуктів ІТ-проектів на основі аналізу процесів взаємодії користувачів (клієнтів) з інформаційною системою (продуктом). Розглянуто складові моделі прогнозування розвитку ІТ-продуктів інноваційних стартап-проектів, з урахуванням аналізу формуються набори даних взаємодії потенційних клієнтів із такими продуктами. Запропоновано алгоритм формування початкових наборів даних на основі карт подорожей клієнтів (CJM), які є інструментом фіксації подій взаємодії клієнтів із системою. Наведено приклади моделей аналогів карт подорожей клієнтів, які є базою для фіксації та аналізу взаємодій. Цей факт є основою для формування відповідних наборів даних великої розмірності. Як механізм оброблення великих масивів даних і побудови стратегій розвитку ІТ-продуктів запропоновано використання нейронних мереж глибокого навчання. Побудовано математичні моделі для подальшого моделювання й аналізу отриманих результатів. Використано простий лінійний регресійний аналіз для моделювання зв'язку між єдиною пояснювальною змінною і безперервною змінною відгуку (залежною змінною). Для наявних даних застосовано метод розвідувального аналізу даних для пошуку повторюваних образів і аномалій. У ході дослідження побудовано модель реалізації лінійної регресії з використанням підходу на основі градієнтної оптимізації. Також застосовано лінійні моделі бібліотеки *scikit-learn* для завдання регресії і реалізовано стабілізаційний регресійний метод. Проведено моделювання й аналіз отриманих результатів, який показав більшу ефективність щодо збільшеного періоду життєвого циклу продуктів ІТ-проектів.

Ключові слова: запуск; інформаційна взаємодія; карта подорожей клієнта; прогнозування.



Viktor Morozov,
PhD, professor, Head of Dept. of Technology Management of Faculty of Informational Technology of Taras Shevchenko National University of Kyiv.

Віктор Морозов,
кандидат технічних наук, професор, завідувач кафедри менеджменту технологій факультету інформаційних технологій Київського національного університету імені Тараса Шевченка.



Olga Mezentseva,
PhD, Associate professor of Dept. of Technology Management of Faculty of Informational Technology of Taras Shevchenko National University of Kyiv.

Ольга Мезенцева,
кандидат технічних наук, доцент кафедри менеджменту технологій факультету інформаційних технологій Київського національного університету імені Тараса Шевченка.



АНАЛІЗ RTB-ПЛАТФОРМ ДЛЯ АВТОМАТИЗАЦІЇ КУПІВЛІ ТА ПРОДАЖУ МЕДІАКОНТЕНТУ

Проведено аналіз RTB-платформ для автоматизації купівлі та продажу медіаконтенту. За результатами цього аналізу сформовано мету наукового дослідження, яка полягає у проведенні аналізу й обґрунтуванні доцільності використання DSP-, SSP-платформ. Для досягнення поставленої мети проведено аналіз технологій Programmatic і RTB-платформ, переваг і недоліків аукціонних і прямих закупівель, порядку роботи DSP- та RTB-аукціону. З'ясовано, що у RTB-системах взаємодіють між собою рекламні майданчики або біржі, які продають свій інвентар, рекламодавці або агентства, зацікавлені в покупці цього інвентарю, а також відвідувачі сайтів. Визначено, що на відміну від старої моделі медійних закупівель, в RTB на торги виставляють не рекламне місце, а рекламні матеріали для абсолютно конкретного відвідувача, що дає можливість купувати тільки цільову аудиторію. З'ясовано, що інтереси сайтів на аукціоні представляють Sell Side Platform (SSP), де через SSP-майданчики виставляють на аукціон свій рекламний інвентар. Завдяки SSP власники сайтів в автоматичному режимі продають покази реклами найбільшій кількості рекламодавців за максимальною ціною. Визначено, що у SSP зберігається інформація про рекламні майданчики, формати реклами й інформація про відвідувачів сайтів. Зосереджено увагу на тому, що в той самий час інтереси рекламодавців на аукціоні представляють Demand Side Platforms (DSP). З'ясовано, що DSP спрощують процес купівлі реклами на великій кількості рекламних бірж, дозволяють гнучко керувати ціною показів і налаштовувати рекламні кампанії, збирають ставки, настроювання та креативи всіх рекламодавців, які беруть участь в аукціоні. Дійшли висновку, що згідно з eMarketer, програмна реклама постійно набуває популярності на українському медіаринку.

Авторами визначено переваги та недоліки використання DSP. Акцентовано увагу на перевагах, до яких належать: висока ефективність; використання платформ управління даними CRM або DMP; точні можливості з націлювання; підтримка за межами традиційної підтримки клієнтами однієї мережі; високоякісний інвентар. Визначено недоліки, до яких віднесено високу вартість і складність.

Ключові слова: медіаконтент; RTB-платформи; рекламні майданчики; програмна реклама.

1. ВСТУП

Процес закупівлі та продажу реклами або рекламних платформ дуже важкий та дорогий, особливо, якщо цим займається людина. Краще, якщо до цих питань буде залучена машина, що автоматизує вказаний процес і зведе втручання людини до мінімуму. За цей процес і відповідає Demand Side Platform (DSP).

Характерною рисою сучасного методу закупівлі реклами є зростання ролі машини у цьому процесі, яка набуває вирішального значення для економії часу та матеріальних ресурсів. Інформаційна перевага визначається здатністю машини у лічені секунди як купувати медіарекламу,

так і продавати без жодного втручання людини, тобто повністю автоматизувати цей процес [1].

Розглянута у роботі платформа дозволяє:

- спостерігати за ціновою політикою реклами;
- покращувати результати видавців у медіакупівлі;
- забезпечувати користувачів аналітикою їхніх дій;
- здійснювати всі ручні операції в автоматичному режимі;
- підвищувати ефективність купівлі та продажу реклами тощо.

У провідних країнах світу такі системи розглядають як уже невід'ємну частину медіакупів-

© Коротін Д., Поперешняк С. В., 2020



лі, які здатні оперативно забезпечити систематичну купівлю та продаж реклами за вигідними для користувачів тарифами.

Отже, нині застосування автоматизованої купівлі та продажу реклами мають і переваги, і недоліки, що розглядатимуться далі.

2. МЕТА СТАТТІ

Метою статті є проведення аналізу та обґрунтування доцільності використання DSP-, SSP-платформ.

Для досягнення поставленої мети у процесі дослідження розв'язано такі завдання: проведено аналіз технологій Programmatic і RTB, переваг і недоліків аукціонних і прямих закупівель, порядку роботи DSP- та RTB-аукціону.

3. АНАЛІЗ ОСТАННІХ ДОСЛІДЖЕНЬ І ПУБЛІКАЦІЙ

Аналіз останніх досліджень і публікацій показав, що проблематику продажу реклами, а саме швидкості й об'єму продажу, відображено в працях багатьох вітчизняних і зарубіжних учених. Майже всі вони зводяться до того, що обидві платформи (DSP і SSP) є результатом історичної потреби в усуненні людського фактора [2, 3, 4]. У роботі [5] наведено приклад користі систем у контексті – "Процес продажу дорогий і недостовірний, доки DSP і SSP не зробили загальну картину та зробили його сприятливішим та ефективнішим для гаманця". Тобто, у цій праці наголошено на економії грошей за допомогою використання сучасних систем.

Проте є й інший погляд, де результат наукового дослідження, наведеного у роботі [6], показав, що основною складовою ефективності є час і можливість оптимізації рекламних кампаній.

Основними проблемами, які виокремили експерти й учасники ринку в розвитку технології RTB, в Україні вважають такі [7–12]:

- 1) недостатня кількість майданчиків і відповідно рекламного інвентарю;
- 2) відсутність кваліфікованих фахівців;
- 3) відверте шахрайство, пов'язане із ціноутворенням і "закруткою" трафіка з використанням "ботів"; інертність ринку – боязкість рекламодавців ризикувати, експериментувати і пробувати нове;
- 4) брак даних – недостатній обсяг про користувачів, на основі якого створюються таргетинги;
- 5) недостатній ступінь інтеграції учасників RTB-екосистеми;
- 6) опір майданчиків – багато майданчиків бояться використовувати RTB через страх, що це завадить прямим продажам.

Аналіз робіт [13, 14, 15] дав змогу дійти висновків:

- RTB з кожним роком показує все більший приріст, технологія є затребуваною;
- за три роки існування на ринку технологія RTB отримала змішані відгуки, проте знайшла свого споживача;
- розвитку RTB на українському ринку заважають недолік якісних даних про користувачів і недостатня кількість підключених майданчиків, які забезпечували б систему трафіком;
- нова технологія потребує кваліфікованих фахівців, прозорого ціноутворення;
- усі опитані учасники ринку вважають, що український сегмент RTB буде розвиватися і рости.

За три роки на українському ринку технологія RTB пройшла шлях від "революції" до "розчарування" [16].

Учасники українського ринку, натхненні значними успіхами технології на Заході, очікували таких же результатів, але цього не сталося. Цьому перешкоджали недостатня технічна підготовленість і обізнаність про технології, консерватизм великих гравців, недостатній обсяг даних про користувачів на ринку. Проте RTB домігся помітних фінансових успіхів, чим і продемонстрував свою життєздатність на українському ринку [16].

Отже, з метою проведення аналізу реалізації систем DSP, SSP та обґрунтування доцільності їхнього використання, авторами проведено аналіз технологій Programmatic і RTB, розглянуто переваги та недоліки аукціонних і прямих закупівель, порядок роботи DSP- та RTB-аукціону.

4. ВИКЛАД ОСНОВНОГО МАТЕРІАЛУ

Технології Programmatic і RTB здійснили великий прорив у світі інтернет-реклами. З їхньою появою реклама стала розвиватися активніше. На ринок вийшли нові рекламні майданчики і мережі зі своїм інвентарем, а у рекламодавців і агентств з'явилася можливість купувати покази для конкретних аудиторій. Раніше рекламодавцям доводилося самостійно домовлятися з десятками рекламних систем про рекламу. У кожній системі був свій інтерфейс, налаштування, правила, умови роботи тощо.

Вартість реклами диктували власники систем і майданчиків. Причому рекламодавці могли викупити тільки рекламне місце або банер, а не показ реклами цільової аудиторії. Було досить проблематично визначити, на яких рекламних майданчиках і коли потрібна цільова аудиторія [17]. Усе це дуже ускладнювало купівлю реклами. Було потрібно рішення, яке спростило б



процес розміщення реклами, використовуючи одну точку входу на різних біржах і майданчиках. Для розв'язання цих проблем і прийнято рішення розробити єдину систему, яка полегшила б процес купівлі реклами.

4.1. PROGRAMMATIC TA RTB

Programmatic – це технологія автоматизації процесу купівлі та продажу рекламного інвентарю, заснована на знаннях про переваги користувачів. Простіше кажучи, це програма, яка дуже швидко (за частки мілісекунд) приймає рішення про показ реклами конкретного користувача.

Programmatic включає в себе автоматизоване придбання рекламного інвентарю через торги на аукціоні або безпосередньо за домовленістю з рекламними майданчиками [17].

Також часто можна зустріти визначення Programmatic Direct – той же Programmatic, але з прямим придбанням рекламних показів на майданчиках. Programmatic Direct – це процес купівлі рекламного інвентарю за допомогою інструментів Programmatic, але рекламодавець, у цьому випадку безпосередньо спілкується з майданчиком, щоб укласти угоду [17].

RTB (Real Time Bidding) – система, яка призначена для автоматизації купівлі і продажу реклами, в основі якої лежить модель аукціону в реальному часі [17]. На відміну від старої моделі медійних закупівель, у RTB на торги виставляють не рекламне місце, а показ рекламних матеріалів абсолютно конкретному відвідувачеві. Це дає можливість купувати тільки цільову аудиторію.

Крім того, в RTB немає фіксованої ціни за показ реклами. Рекламодавці самі призначають ставки з кроком, наприклад, в один цент. В аукціоні виграє той, хто запропонував найвищу ціну. Причому показ реклами здійснюється за ставкою попереднього конкурента в аукціоні – аукціон другої ціни. Багато фахівців часто плутають Programmatic з аукціоном в реальному часі (RTB). RTB – це одна з моделей програмованих закупівель реклами, тобто частина концепції Programmatic [17]. Виходить, що RTB – це частина Programmatic. Якщо через Programmatic торги за рекламу йдуть через аукціон – це RTB, якщо напряму з майданчиком – це Programmatic Direct.

Щоб розібратися, як влаштовано аукціон у концепції RTB, визначимо учасників аукціону. У RTB-системах взаємодіють між собою рекламні майданчики або біржі, які продають свій інвентар, і рекламодавці або агентства, зацікавлені в покупці цього інвентарю, а також відвідувачі сайтів, за яких і йде "війна". Щоб усім учасникам було простіше знайти спільну мову, взаємодіють вони через спеціалізовані платформи.

Інтереси сайтів на аукціоні представляють SSP-платформи (Sell Side Platform). Через SSP-платформи майданчики виставляють на аукціон свій рекламний інвентар. Завдяки SSP, власники сайтів в автоматичному режимі продають покази реклами найбільшій кількості рекламодавців за максимальною ціною. У SSP зберігається інформація про рекламні майданчики, формати реклами й інформація про відвідувачів сайтів [17]. Інтереси рекламодавців на аукціоні представляють DSP. DSP-системи спрощують процес купівлі реклами на великій кількості рекламних бірж, дозволяють гнучко керувати ціною показів і налаштовувати рекламні кампанії. DSP-системи збирають ставки, настроювання та креативи всіх рекламодавців, які беруть участь в аукціоні, за показ реклами конкретному відвідувачеві [17].

DMP (Data Management Platforms) – це сервіси, які продають SSP- і DSP-систем відсутню інформацію про відвідувачів сайтів. DMP збирають знеособлену інформацію про користувачів, таку як історію відвідин сайтів, соціально-демографічні дані, дані місця розташування, ідентифікатори пристроїв і навіть запити, які вводили в пошукових системах. Як правило, ця інформація зберігається у файлах Cookies [18]. У функції DMP входить збирання й оброблення даних про відвідувачів сайтів, склеювання цієї інформації і зберігання.

4.2. ПОРЯДОК РОБОТИ RTB-АУКЦІОНУ

Припустимо, ми визначили учасників аукціону та програмні продукти, через які відбувається взаємодія. Тепер розглянемо, як улаштований RTB-аукціон. На першому етапі, коли користувач завантажує вебсторінку з рекламним інвентарем, браузер відправляє запит на показ реклами для цього користувача в SSP-платформу.

SSP збирає інформацію про користувача і формує рекламний лот, в якому міститься інформація про майданчик, сторінку показу, формат реклами і знеособлені дані користувача. За необхідності SSP може закупити додаткову інформацію про користувача у DMP-систем. Після SSP відправляє лот DSP-платформам [18]. DSP-платформи повинні за частки мілісекунд визначити, наскільки цінним є цей показ для рекламодавців, і зробити ставки. DSP теж можуть скористатися послугами DMP-систем. Коли ставки зроблено, DSP передає інформацію про ставки SSP-системі. На наступному етапі SSP-система збирає всі ставки рекламодавців, проводить аукціон і вибирає переможця, який запропонував найвищу ставку. Причому переможець платить ставку попереднього рекламодавця. Аукціон завершено, і відвідувач сайту бачить рекламу



переможця аукціону. Весь процес проходить у режимі реального часу і займає всього частки мілісекунд (100–120 мс). Користувачам цей процес абсолютно не помітний [18].

Programmatic Direct – купівля реклами поза аукціоном безпосередньо у майданчиків. Такі види закупівель називають прямими (рис. 1).



Рис. 1. Схема прямих закупівель поза аукціоном

Direct поділяють на два типи закупівель: гарантовані – Programmatic Guaranteed, за яких фіксовані ціни й обсяг показів; і негарантовані – Preferred Deal, коли фіксується ціна за тисячу показів, але не фіксується обсяг показів [18].

У великих рекламних майданчиків є певний обсяг так званого преміального трафіка, який вони не виставляють на аукціон, а реалізують через прямі продажі. Інформація про наявність такого трафіка розсилається лише конкретним рекламодавцям.

Взаємодія майданчиків і рекламодавців у Programmatic Direct також відбувається через SSP- і DSP-системи. Різниця з RTB полягає лише у тому, що преміальний трафік продається до аукціону за попередньою домовленістю. А залишковий трафік, який не був реалізований за моделлю Programmatic Direct, уже виставляють на загальний аукціон [18].

Programmatic Direct ідеально підходить для брендівих рекламних кампаній. За цією моделлю рекламодавці можуть викупити кращі рекламні місця на великих майданчиках. Зрозуміло, ціна такого трафіка набагато вище, ніж через модель RTB.

4.3. ПЕРЕВАГИ ТА НЕДОЛІКИ АУКЦІОННИХ І ПРЯМИХ ЗАКУПІВЕЛЬ

- Після проведеного аналізу, можна визначити такі *переваги моделі RTB*, яка була взята за основу проєкту DSP:

- низький поріг входу;
- щоб запустити рекламу на описаному проєкті не потрібні великі початкові вкладення. Ви можете налаштувати першу рекламну кампанію навіть із бюджетом \$1000;
- оптимізація витрат на рекламу;
- на аукціоні рекламодавці самі встановлюють ціну, яку готові платити за показ. Ставки можна виставляти з кроком в один цент. Це дозволяє мінімізувати вартість залучення кінцевого клієнта;
- точне попадання в цільову аудиторію;
- через нашу платформу можна показувати рекламу цільової аудиторії з конкретними інте-

ресами. Вашу рекламу побачать тільки зацікавлені користувачі;

- велике охоплення трафіка;
- через RTB працює велика кількість рекламних майданчиків різних тематик;
- підходить для представницьких компаній;
- RTB відмінно підходить для компаній, в яких прийнято оцінювати ефективність за ROI або за кількістю реєстрацій чи заявок.

Недоліки представленої платформи:

- остаточний інвентар;
- весь преміальний трафік майданчики намагаються реалізувати безпосередньо, а на аукціон виставляється вже те, що залишилося;
- низька якість майданчиків;
- до RTB-моделі підключено величезну кількість сайтів, зокрема і майданчики другого сорту. Не можна заздалегідь передбачити, на якому саме ресурсі будуть показані ваші рекламні матеріали;
- не підходить для брендівих й іміджевих рекламних кампаній.

Процес купівлі або продажу цифрової реклами протягом останніх 20 років був дорогим і дуже ненадійним. Установлено, що сучасні DSP допомагають зробити цей процес не таким витратним і максимально ефективним, видаляючи на певних етапах фактор людського втручання, а також необхідність безпосереднього узгодження розцінок і пов'язаних із ним "ручних" операцій [19]. RTB-специфікація показує, що DSP влаштовує аукціон, де за рекламну платформу SSP виставляють ціну і DSP вибирає найпривабливіший (оптимальний) для нього лот. Після цього аукціону кожний SSP та DSP можуть подивитися аналітику у своєму особистому кабінеті [19]. Деякою мірою DSP робить багато з того, що раніше пропонували "Ad networks", включаючи доступ до широкого спектра інвентарю і таргетингу. Але його перевагою, на відміну від рекламних мереж, є можливість купити, показати і відстежити оголошення, використовуючи єдиний інструмент. У результаті це дозволяє досить легко оптимізувати рекламні кампанії. Тут усе працює навколо даних. Як правило, рекламні мережі ставлять націнку на медійну рекламу, яку вони продають. DSP, у свою чергу, стягує невелику плату за суттєве спрощення транзакції. DSP тепер працює безпосередньо з клієнтами рекламодавця, ефективно замінюючи агентства, коли справа доходить до закупівель медійної реклами. Клієнти кажуть про те, що вони як і раніше співпрацюють з агентствами для розроблення стратегій або отримання консультацій, але починають активно звертатися і до третіх сторін, включаючи DSP, які допомагають закуповувати рекламу [19].



Supply Side Platform – це платформа, яка торгує рекламним інвентарем або рекламними позиціями інтернет-майданчиків. SSP агрегує пропозиції майданчиків, "збирає" залишковий трафік, а також установлює таку мінімальну вартість, за якою майданчик готовий реалізувати показ. SSP проводить аукціонний торг із DSP, максимально вигідно продаючи інвентар видавцю [19]. DSP не володіє інвентарем, не купує і не продає його, а просто з'єднує видавця з Ad Exchange або SSP, щоб він сам міг продати свій інвентар. SSP – це та ж DSP, але з точки зору видавця або продавця контенту. Це технологічна платформа, яка автоматизує продаж показів для видавців.

Раніше процеси купівлі та продажу оголошень не були такими простими, як сьогодні. Накази щодо вставки вручну та контракти були частиною процесу розгортання оголошення. У минулому це могло тривати кілька днів або тижнів. Тепер за допомогою нашої платформи на стороні попиту і платформ на стороні постачання – це ефективний і майже миттєвий процес.

Платформа зі сторони попиту – це програмне забезпечення, яке рекламодавці використовують для придбання мобільних, пошукових і відеооголошень із ринку, на якому розміщують рекламний ресурс. Ці платформи дозволяють керувати рекламою в багатьох мережах торгів у реальному часі, на відміну від лише однієї, наприклад, Google Ads. Разом із платформами з боку постачальників, DSP дозволяє програмну рекламу. Більшість реклами в інтернеті тепер здійснюється програмним шляхом торгів у режимі реально часу та прямих угод [19].

Ставки в режимі реального часу. Така реклама відбувається в режимі реального часу. Якщо ви вказуєте, з ким бажаєте поширювати об'яви, скільки ви готові витратити, то війна ставок відбувається між вами та всіма іншими рекламодавцями, які намагаються охопити одну й ту ж аудиторію. Перспектива потрапляє на сторінку, і перш ніж сторінка завантажиться повністю, алгоритми визначають, яке оголошення відображатиметься на них. Ці алгоритми враховують такі речі: історія перегляду, час доби, IP-адреса. Той, хто має найвищу ставку для показу, коли всі зібрані, виграє місце розташування. Їхнє оголошення публікується, коли сторінка відвідувача завантажувється повністю.

Пряма реклама. Така реклама більше схожа на традиційну модель, перенесену в інтернет. Це ідеально підходить для підприємств, які хочуть гарантовано розміщувати оголошення в преміальних місцях. Наприклад, домашні сторінки видавців великих імен часто продають свої рекламні місця за допомогою програмних прямих

угод. Видавець надає рекламодавцю інформацію про своїх відвідувачів. Якщо ці відвідувачі є ідеальною аудиторією рекламодавця, рекламодавець може залишити частину місця для публікації для наступної кампанії [19].

Згідно з eMarketer, програмна реклама постійно набуває популярність. До кінця 2019 р. передбачено, що 83,6 % реклами буде куплено та продано програмно [19]. Програма не лише відображає рекламу. Це стосується також продажів оголошень у пошукових мережах, а також будь-якої іншої мережі, купленої програмним забезпеченням. Однак, коли ви купуєте оголошення через ці мережі (наприклад, оголошення Google), то ви не обов'язково можете використовувати платформу на стороні попиту.

4.4. ПОРЯДОК РОБОТИ DSP

Платформи на стороні попиту незалежні від окремих мереж. Якщо ви керуєте об'явами за допомогою менеджера медійної мережі Google, ви купуєте покази лише для видавців Google. Якщо ви використовуєте Менеджер об'яв Facebook, щоб купувати оголошення, ви спеціально купуєте покази на Facebook або Instagram. Це програмне забезпечення третіх сторін, яке дозволяє купувати, аналізувати, керувати оголошеннями в багатьох мережах з одного місця [20].

Зі сторони попиту платформи надають рекламодавцям усю інформацію, необхідну для придбання реклами у видавця. Вони не є власниками або не купують засоби масової інформації безпосередньо у видавців, а замість цього обмінюються інформацією з платформою постачальника.

Платформи на стороні постачання дозволяють видавцям перераховувати ці запаси на обмін об'яв і вони спілкуються з DSP про деталі показу [20, 21]. Якщо це відображення є маркетинговим менеджером, який раніше відвідував вашу демонстраційну сторінку, вона є більш цінною для вас, ніж людина, яка ніколи раніше не відвідувала ваш вебсайт. У такому випадку показник DSP, імовірно, буде вищим для показу [20].

Використання платформи DSP з боку попиту має свої переваги та недоліки. Необхідно добре знати декілька з них, перш ніж ви вкложите значні кошти у програмне забезпечення.

Переваги використання DSP.

1) *Ефективність.* Якщо ви керуєте кампаніями в багатьох мережах, DSP має сенс використовувати та регулювати з однієї панелі.

2) *Дані.* Багато постачальників послуг цифрового телебачення співпрацюють з постачальниками даних третіх сторін, щоб запропонувати рекламодавцям якомога більше інформації. Часто це більше, ніж єдина мережа. Крім того, багато DSP до-



звояють клієнтам імпортувати власні дані з CRM або DMP (платформи управління даними) [22].

3) *Націлювання*. З більшою кількістю даних з'являються точні можливості націлювання. Краще націлювання означає більше персоналізованих оголошень і цільових сторінок, що означає більшу ймовірність переходу.

4) *Підтримка*. Платформи на стороні попиту часто забезпечують підтримку за межами традиційної підтримки клієнтами однієї мережі.

5) *Високоякісний інвентар*. DSP мають доступ до основних мереж, а потім і до деяких інших. Якщо вам потрібні додаткові запаси преміумкласу, можливо, ви шукаєте платформу на стороні попиту. Деякі можуть мати більший доступ, ніж інші, тому важливо з'ясувати, перш ніж вибрати один із них [22].

Недоліки використання DSP.

Вартість і складність. Кожен раз, коли збираєте дані, ви ризикуєте надмірно ускладнити платформу. Деякі рекламодавці можуть виявити, що платформи на стороні попиту занадто складні, аби навчитися досить швидко, щоб побачити переваги [22]. Для деяких підприємств платформа з боку попиту буде мати найбільший сенс. Вона пропонує ефективність і ширший доступ до перспектив для кількох обмінів оголошеннями, включаючи більше інвентаризації преміумкласу. Однак для тих, хто рекламує декілька платформ, вартість і складність DSP не варто використовувати у стеку рекламного інструмента [22].

Існує також платформа Ad Exchange, яка організовує процес спілкування DSP- та SSP-платформами та будує модель "один до багатьох". Дані про платформу Ad Exchange буде детально розкрито у подальших наших дослідженнях.

5. ВИСНОВКИ

Проведено аналіз RTB-платформ для автоматизації купівлі та продажу медіаконтенту. Авторами визначено переваги та недоліки використання DSP.

Для досягнення поставленої мети проведено аналіз технологій Programmatic і RTB-платформ, переваг і недоліків аукціонних і прямих закупівель, порядку роботи DSP- та RTB-аукціону. З'ясовано, що в RTB-системах взаємодіють між собою рекламні майданчики або біржі, які продають свій інвентар, рекламодавці або агентства, зацікавлені в покупці цього інвентарю, а також відвідувачі сайтів. Авторами зроблено висновок, що згідно з eMarketer, програмна реклама постійно набуває популярності на українському медіаринку.

СПИСОК ВИКОРИСТАНОЇ ЛІТЕРАТУРИ

- [1] Е.Л. Головлёва Основы рекламы. – М.: Академичний проект, 2018. – 191 с.
- [2] Г. Грачёв, И. Мельник. Манипулирование личностью: организация, способы и технология информационно-психологического воздействия // Под ред. Г.В. Грачева. – М.: Астра-Пресс, 2016. – 397 с.
- [3] What is Real Time Bidding? How Does it work and Why you Need it in your PPC Campaigns? [Електронний ресурс] – Ресурс доступу: <https://www.omnicoreagency.com/what-is-real-time-bidding>.
- [4] Х. Кафтанджиев. Гармонія в рекламних комунікаціях // Х. Кафтанджиев; пер. з болг. С. Кировой – М.: Эксмо, 2015. – 368 с.
- [5] DSP Vs SSP: Let's Nail This Down Once And For All. [Електронний ресурс] – Ресурс доступу: <https://selfadvertiser.com/blog/dsp-vs-ssp/>.
- [6] Перехід від блоків Директа до programmatic-платформі. [Електронний ресурс] – Ресурс доступу: <https://education.ru/perehod-ot-blokov-direkta-k-programmatic-platforme/>.
- [7] Большие данные: как извлечь из них информацию компании / А. Моррисон [и др.] // Технологический прогноз. Ежекварт. журнал. 2016. – Вып. 3. – С. 42.
- [8] DSP Vs SSP: Let's Nail This Down Once And For All [Електронний ресурс]. – Ресурс доступу: <https://selfadvertiser.com/blog/dsp-vs-ssp/>.
- [9] Види шахрайства в programmatic та як із ними боротися [Електронний ресурс]. – Ресурс доступу: <https://detector.media/production/article/141501/2018-10-04-vidi-shakhraystva-v-programmatic-ta-yak-iz-nimi-borotisy/>.
- [10] What is the Difference Between Ad Exchange and Ad Network? [Електронний ресурс]. – Ресурс доступу: <https://medium.com/swipe-mag/what-is-the-difference-between-ad-exchange-and-ad-network-e6714f950132>.
- [11] What is the difference between RTB and Programmatic? [Електронний ресурс]. – Ресурс доступу: <https://www.gamned.com/what-is-the-difference-between-rtb-and-programmatic-2/>.
- [12] Основы RTB. [Електронний ресурс]. – Ресурс доступу: <https://maddata.agency/mad-blog/osnovy-rtb>.
- [13] Optimal Real-Time Bidding for Display Advertising // Weinan Zhang, Shuai Yuan, Jun Wang, 2014. – С. 6.
- [14] Що робить DSP платформа в Ad Buying? [Електронний ресурс]. – Режим доступу: <https://maddata.agency/mad-blog/chto-delaet-dsp-platforma-v-ad-buying>.
- [15] Реклама в мережі Інтернет // [Електронний ресурс]. – Режим доступу: <http://propel.ru/>.
- [16] Н. Романевич, А. Константинов, Г. Тарасевич. Майбутнє реклами /Експерт online// [Електронний ресурс]. – Режим доступу: http://www.dv-reclama.ru/others/articles/detail.php?ELEMENT_ID=9152.
- [17] Н.Н. Калайтанова. Технология RTB в медийном бизнесе / Н.Н. Калайтанова // Медиаальманах. – 2015. – № 6. – С. 105.
- [18] В. Майер-Шенбергер. Большие данные: революция, которая изменит то, как мы работаем, мыслим и живем = Big Data: A Revolution that Will Transform How We Live, Work, and Think / В. Майер-Шенбергер, К. Кукьер; пер. И. Гайдюк. – М.: Манн, Иванов и Фербер, 2017. – 249 с.
- [19] Budget Constrained Bidding by Model-free Reinforcement Learning in Display Advertising. [Електронний ресурс] – Ресурс доступу: <https://arxiv.org/pdf/1802.08365.pdf>.
- [20] Real-Time Bidding with Multi-Agent Reinforcement Learning in Display Advertising. [Електронний ресурс] – Ресурс доступу: <https://arxiv.org/pdf/1802.09756.pdf>.



[21] O. Starkova, K. Herasymenko, S. Korotin, V. Afanasiev, A. Lisnyk. Development of Recommendations for Ensuring Security in a Corporate Network, 2019 IEEE International Conference on Advanced Trends in Information Theory (ATIT), Kyiv, 2019, pp. 111-115.
DOI: 10.1109/ATIT49449.2019.9030470.

[22] What are the advantages of RTB for buying display advertising? [Електронний ресурс] – Ресурс доступу: <https://econsultancy.com/what-are-the-advantages-of-rtb-for-buying-display-advertising/>.

Стаття надійшла до редколегії

10.10.2020

Analysis of RTB platforms for automating the purchase and sale of media content

In the article analyzes RTB platforms for automating the buying and selling of media content. According to the results of the analysis, the purpose of scientific research is formed, which consists of the analysis and substantiation of the expediency of using DSP-SSP platforms. For achieve this goal, an analysis of Programmatic and RTB platform technologies, the advantages and disadvantages of auctions and direct procurement, the order of operation of the DSP and RTB auction was conducted. In the article found that RTBs interact with advertising sites or exchanges that sell their inventory, advertisers or agencies interested in buying this inventory, as well as site visitors. In the article determines that, unlike the old model of media purchases, RTB does not put up for sale an advertising place, but advertising materials for a very specific visitor, which makes it possible to buy only the target audience. It turned out that the interests of sites at the auction represent Sell Side Platform (SSP). Where through the SSP, the sites auction their advertising inventory.

Thanks to SSP, website owners automatically sell ad impressions to a large number of advertisers at the maximum price. In SSP stores information about advertising platforms, ad formats, and information about site visitors. The focus is that at the same time, the interests of advertisers at the auction represent Demand Side

Platforms (DSP). It was found that DSP simplify the process of buying advertising on a large number of advertising exchanges, allow flexible control of the price of impressions and customize advertising campaigns, collect bids, settings and creative's of all advertisers participating in the auction. It is concluded that, according to eMarketer, programmatic advertising is gaining popularity in the Ukrainian media market. The authors identified the advantages and disadvantages of using DSP. The attention is focused on the advantages, which include: high efficiency; Use of CRM or DMP data management platforms precise targeting capabilities; support beyond traditional customer support of the same network; high quality inventory. Deficiencies to which are assigned are identified – high cost and complexity.

Keywords: media content; RTB platforms; advertising platforms; program advertising.



Денис Коротін,
студент четвертого курсу кафедри програмних систем і технологій факультету інформаційних технологій Київського національного університету імені Тараса Шевченка.

Denis Korotin,
fourth-year student of the Department of Software Systems and Technologies, Faculty of Information Technologies, Taras Shevchenko National University of Kyiv.



Світлана Поперешняк,
кандидат фізико-математичних наук, доцент кафедри програмних систем і технологій факультету інформаційних технологій Київського національного університету імені Тараса Шевченка.

Svitlana Poperehnyak,
Candidate of Physical and Mathematical Sciences, Associate Professor of the Department of Software Systems and Technologies, Faculty of Information Technologies, Taras Shevchenko National University of Kyiv.



ТЕНДЕНЦІЇ РОЗВИТКУ ШТУЧНОГО ІНТЕЛЕКТУ В УКРАЇНІ

Протягом останніх років технології штучного інтелекту почали застосовувати у багатьох сферах людського життя: освіті, медицині, бізнесі, фінансах, менеджменті, маркетингу, промисловості тощо, що стало фактично ключовою тенденцією поточного часу. Наведено інформацію та проаналізовано сучасний стан і тенденції розвитку штучного інтелекту в Україні. Проведено аналіз стратегії розвитку штучного інтелекту, яка запропонована рядом міністерств. Розглянуто концепцію розвитку сфери штучного інтелекту, яка запропонована Державним агентством з електронного урядування. Проаналізовано та надано дані про компанії, які працюють у галузі штучного інтелекту. Визначено основні напрями діяльності українських компаній у галузі штучного інтелекту, зокрема і цілий ряд українських стартапів. Наведено інформацію про наукові школи та наукові дослідження, які виконуються в галузі штучного інтелекту в Україні. Установлено, що основні дослідження зосереджено на таких напрямках: нейромережі, розпізнавання образів, математична інформатика, створення інтелектуальних інформаційних систем і баз знань, розроблення інтелектуальних робототехнічних систем, знання-орієнтовані технології, інтелектуальний пошук, розроблення інтелектуальних навчальних систем, аналіз великих даних, нечіткі системи підтримки прийняття рішень, багатоагентні системи та багато інших. У цілому наукові дослідження, які проводяться в Україні в галузі штучного інтелекту, охоплюють базові напрями робіт, які ведуть у світі. Також проаналізовано напрями дисертаційних досліджень за останні роки по інтелектуальних системах і системах штучного інтелекту. Головні напрями дисертаційних досліджень: інтелектуальний аналіз та оброблення даних, оброблення знань і мультиагентні технології, інтелектуальні технології, зокрема й інтелектуальні технології підтримки прийняття рішень, розпізнавання зображень та ін.

Ключові слова: штучний інтелект; стратегія розвитку; наукові дослідження; наукові школи; перспективи.

1. ВСТУП

Цифрові технології відкривають унікальні можливості для розвитку економіки та підвищення якості життя громадян. За системного державного підходу цифрові технології будуть значно стимулювати розвиток відкритого інформаційного суспільства, підвищення продуктивності, економічного зростання, створення робочих місць, а також підвищення якості життя громадян України. Сучасні напрями застосування цифрових технологій передбачають широке впровадження таких технологій, як штучний інтелект, аналітика великих даних, робототехніка тощо.

Поняття штучний інтелект еволюціонує із часом і в певних випадках під штучним інтелектом розуміють те, що ще не зроблено в комп'ютерних системах. У статті під штучним інтелектом розуміють теорію і практику створення комп'ютерних систем, які здатні розв'язувати задачі, які зазвичай вимагають інтелекту людини. Штучний

інтелект – це штучно розроблена система, яка має людські або близькі до людських інтелектуальні здібності й може виконати будь-яке завдання з можливих для *homo sapiens*.

Наукова дисципліна штучний інтелект містить множини напрямів: машинне навчання, глибинне навчання, представлення знань, робототехніка, експертні системи, багатоагентні системи, розпізнавання образів, оброблення природної мови та багато інших. Поняття штучного інтелекту часто вживають у розрізі машинного навчання і штучних нейронних мереж. Розрізняють поняття слабого і сильного штучного інтелекту. Функції слабого штучного інтелекту передбачають, що система може виконувати різні когнітивні дії, що визначені людиною. Сильний штучний інтелект передбачає, що комп'ютер не тільки обробляє інформацію, а й розуміє її сенс. Дослідженню штучного інтелекту присвячено багато робіт світових учених: А. Тьюрінг, Дж. Маккарти, Г. Саймон, А. Ньюелл, М. Мінський, Ф. Розенблатт та ін. [1].



Основний фактор, що визначає сьогодні розвиток технологій штучного інтелекту – це зростання продуктивності сучасних комп'ютерів у поєднанні з підвищенням якості алгоритмів, які дозволяють розробляти системи, що мають властивості людського інтелекту. Розробки в галузі штучного інтелекту вважають одним із ключових напрямів науково-технічного прогресу XXI ст.

2. ПОСТАНОВКА ПРОБЛЕМИ

Насичення ринку відповідним технічним і програмним забезпеченням дозволяє зробити технології штучного інтелекту більш доступними, частіше використовувати їх промисловістю, бізнесом, університетами і науковими закладами.

Зараз у всьому світі формуються певні напрями застосування штучного інтелекту. Штучний інтелект може поглибити розрив між країнами, посилюючи існуючий цифровий розрив. Країнам можуть знадобитися різні стратегії та відповіді, оскільки рівень сприйняття технологій штучного інтелекту різниться. У відповідь на цей виклик уряди провідних країн затвердили національні програми в галузі штучного інтелекту. Роботи у сфері штучного інтелекту нині ведуть практично всі великі зарубіжні компанії, університети і наукові агентства. Сьогодні в дослідженнях штучного інтелекту беруть участь такі великі компанії: американські компанії – Google, Facebook, Amazon, Microsoft, китайські компанії – Baidu, Alibaba, Tencent. Очікуване зростання світового валового внутрішнього продукту 2030 р. завдяки штучному інтелекту сягне майже 16 трлн доларів [2]. Зростає попит на фахівців у галузі штучного інтелекту. Кількість зарахованих до університетів студентів, націлених на спеціальності в галузі штучного інтелекту, щороку зростає в кілька разів.

У зв'язку з новими викликами постають задачі розвитку штучного інтелекту в Україні. Метою статті є аналіз сучасного стану і напрямів розвитку штучного інтелекту в Україні. Пропонується розглядати проблематику розвитку штучного інтелекту шляхом аналізу результатів досліджень штучного інтелекту в бізнесі і в академічних організаціях та університетах. Розглянуто основи стратегії розвитку штучного інтелекту. У роботі проведено аналіз досліджень у галузі штучного інтелекту, розглянуто задачі, що розв'язуються в наукових дослідженнях із штучного інтелекту в Україні.

3. ОСНОВНА ЧАСТИНА

Роботи в галузі штучного інтелекту в Україні. Аналітичне агентство Deep Knowledge Analytics лондонського інвестиційного фонду Deep

Knowledge Ventures, яке спеціалізується на штучному інтелекті, після дослідження ринку штучного інтелекту у Східній Європі, визначило, що Україна перебуває на лідируючих позиціях. За даними Clutch 28 українських компаній постають рішення для штучного інтелекту порівняно з 226 постачальниками по всьому світу [3]. Згідно LinkedIn, в Україні більш ніж 2000 розробників програмного забезпечення, які спеціалізуються на штучному інтелекті.

Основними напрямками діяльності українських компаній у галузі штучного інтелекту є розроблення програмного забезпечення, розроблення чатботів, інформаційні технології та розважальні продукти. На розвиток і застосування штучного інтелекту найбільше впливають: держава, яка формує політики та пріоритети, державні замовлення, залучення інвестицій і фінансування проєктів; бізнес; державно-приватне партнерство, яке експерти вважають найоптимальнішим варіантом формування нової індустрії [4].

Зараз в Україні роботи зі створення Стратегії розвитку штучного інтелекту тільки розпочато – 2019 р. Міністерство економічного розвитку і торгівлі України створило відповідну робочу групу, а на початку 2020 р. Міністерство цифрової трансформації створило експертний комітет із питань розвитку сфери штучного інтелекту в Україні. До складу комітету входять представники бізнесу, українських і зарубіжних ІТ-компаній, сфери охорони здоров'я та медицини тощо. Державне агентство з питань електронного урядування розпочало роботу над розробленням Концепції розвитку сфери штучного інтелекту в Україні. Концепція має визначити напрями, механізм та умови розвитку штучного інтелекту в Україні. Іншими словами – держава має визначити узгоджену політику щодо врегулювання сфери штучного інтелекту й усунення відставання України в цій сфері [5].

Серед пріоритетів роботи експертного комітету такі [6]:

- створення стратегії розвитку штучного інтелекту – сфери в Україні, виконання якої сприятиме розробленню корисних проєктів, ініціатив і програм та інтеграції надбань штучного інтелекту у сферу держуправління;
- збільшення кількості фахівців у галузі штучного інтелекту – інженерів і підприємців в Україні;
- долучення України до міжнародної спільноти штучного інтелекту, зокрема і допомога українським представникам щодо можливості брати участь у міжнародних конференціях і програмах із розвитку сфери штучного інтелекту;



- стимулювання українського бізнесу використовувати надбання штучного інтелекту.

Вважається, що штучний інтелект може підвищити рівень національної безпеки, освіти й економіки. Пропонується формувати державну політику у сфері розвитку штучного інтелекту в таких напрямках: ринок праці, освіта та наука, бізнес та економіка, етика і законодавство, державне управління та доступні дані. Збільшення кількості в Україні кваліфікованих інженерів у галузях Computer vision, Machine learning, Big data, Natural language processing є критично необхідним. Інакше Україна ще довго не зможе наздогнати розвинені країни.

Як зазначає Ю. Чубатюк, президент групи компаній Everest: "Нам потрібна Національна стратегія розвитку штучного інтелекту, яка дозволить сформувати ключові кейси взаємодії влади, бізнесу, науково-дослідних установ і громадськості, а також розкриє вже наявний у нас потенціал і дасть розуміння того, які рішення ми можемо успішно запозичити та впровадити в наших умовах" [7].

Стартапи і проекти штучного інтелекту в Україні. До українських проектів, які працюють зі штучним інтелектом за даними, складеному аналітичним агентством Deep Knowledge Analytics, належить більше 50 проектів [8]: Let's Enhance (фото на мільйон – це онлайн-сервіс оброблення та масштабування зображення), Captain Growth (онлайн-маркетолог – надання послуги аналізу маркетингових даних), Skillroads (ідеальне резюме – створення програми з написання резюме), Altris.ai (гострий зір – медична платформа, що використовує штучний інтелект для діагностування проблем із зором), 3DLOOK (домашня примірка) та інші проекти.

Рішення, засновані на штучному інтелекті, досить цікаві для інвестицій, коли мають такі характеристики [9]:

- стартап розробив повністю інтегроване рішення або ж full-stack-продукт, приділивши увагу кожному з етапів впровадження;

- автори ідеї мають досвід машинного навчання, а в команді є експерти галузей, де його застосовують;

- стартап має власний актив даних, унікальний і неповторний;

- штучний інтелект і машинне навчання – не просто елементи оптимізації, а основні технології рішення, яке стартап пропонує клієнтам.

До успішних стартапів, які розроблені в Україні в області штучного інтелекту, можна віднести [10]: Minect.ai (система розпізнавання мін і снарядів), Digital Life Lab (новий вид штучного інтелекту – "цифрова форма життя"), Lookserg (застосування штучного інтелекту у соціальних мережах), UniExo (робототехніка), "відкрита влада" (об'єднання ініціатив із дослідження відкритих даних для виявлення прихованих інтересів чиновників) та інші стартапи.

Наукові розробки у сфері штучного інтелекту в Україні. Дослідженню штучного інтелекту присвячено багато робіт учених України: В. М. Глушков, О. І. Кухтенко, В. І. Скурихін, М. М. Амосов, О. Г. Івахненко та ін. В Україні сформовані і діють наукові школи у сфері штучного інтелекту. Провідні позиції в дослідженнях у цій галузі займають такі наукові заклади: Інститут кібернетики імені В. М. Глушкова, Київський національний університет імені Тараса Шевченка, Інститут програмних систем, Харківський національний університет радіоелектроніки, Інститут проблем штучного інтелекту Міністерства освіти і науки України та Національної академії наук України та багато інших. Національні та дослідницькі університети мають кадри і певну матеріально-технічну базу для роботи з технологіями штучного інтелекту. У таблиці наведено наукові центри в галузі штучного інтелекту, які діють в Україні [11].

Таблиця

Наукові дослідження в галузі штучного інтелекту в Україні

№ з/п	Наукова школа, науковий напрям, науковий відділ, науково-дослідна робота	Науковий керівник, засновник
<i>Харківський національний університет радіоелектроніки [12]</i>		
1	Наукова школа "Методи нормалізації, розпізнавання, аналізу й оброблення зображень у системах комп'ютерного зору"	Путятін Є. П.
2	Наукова школа "Ноосферні методології та технології розв'язання проблем управління знаннями та конкурентної розвідки"	Соловйова К. О.
3	Наукова школа "Гібридні системи обчислювального інтелекту для аналізу даних, оброблення інформації та керування"	Бодянський Є. В.
4	Наукова школа "Інтелектуальна обробка інформації"	Руденко О. Г.



Продовження табл.

№ з/п	Наукова школа, науковий напрям, науковий відділ, науково-дослідна робота	Науковий керівник, засновник
Інститут програмних систем НАН України [13]		
1	Науковий напрям "Теоретичні та прикладні питання розроблення систем і технологій програмування, засоби програмної інженерії, проблеми оцінювання і забезпечення якості, стандартизації та сертифікації програмних систем". Методи та засоби створення інтелектуальних сервіс-орієнтованих інформаційно-забезпечуючих систем у середовищі семантичного вебінтерфейсу.	Андон П. І.
2	Науковий напрям "Формально-логічні основи, методи і засоби створення інтелектуальних інформаційних систем, банків даних і знань". Розроблення методів та інтелектуальних інформаційних технологій автоматизації композиції вебсервісів у семантичному вебсередовищі	Андон П. І.
3	Цільова програма "Фундаментальні засади створення перспективних інформаційно-комунікаційних технологій". Розроблення методів кооперації, координації та планування дій агентів у динамічних системах в умовах невизначеності	Яловець А. Л.
4	Цільова програма "Фундаментальні засади створення перспективних інформаційно-комунікаційних технологій та обчислювальних систем". Інтелектуальні інформаційні системи підтримки аналізу великих даних (Big Data)	Андон П. І.
Київський національний університет імені Тараса Шевченка, факультет комп'ютерних наук та кібернетики [14]		
1	Наукова школа "Математична інформатика"	Анісімов А. В.
2	Наукова школа "Алгебро-автоматні методи побудови інтелектуальних інформаційних систем"	Провотар О. І., Кривий С. Л.
Інститут проблем математичних машин і систем НАН України [15]		
1	Напрямок діяльності "Розпізнавання образів"	Вишневський В. В., Калмиков В. Г., Різник О. М.
2	Напрямок діяльності "Нейрокомп'ютери, нейромережі та нейроподібні технології"	Морозов А. О., Різник О. М., Яценко В. О.
Міжнародний науково-навчальний центр інформаційних технологій і систем НАН та МОН України [16]		
1	Науковий напрям "Розроблення загальної теорії інтелектуальних комп'ютерних технологій і систем"	Гриценко В. І., Шлезінгер М. І., Куссульт Е. М., Лебедев Д. В., Файнзільберг Л. С.
2	Науковий напрям "Розроблення методів і засобів інтелектуалізації інформаційних технологій і систем"	
3	Науковий напрям "Розроблення теорії інтелектуальних навчальних систем, комп'ютерної технології навчання та засобів інтелектуалізації діалогу в комп'ютеризованих середовищах"	
4	Науковий напрям "Розроблення нових інформаційних технологій для біології, медицини та екосистем"	
Інститут проблем штучного інтелекту НАН та МОН України [17]		
1	Відділ теоретичних проблем у галузі штучного інтелекту	Шевченко А. І.
2	Відділ інтелектуальних робототехнічних систем	
3	Відділ розпізнавання зорових і мовних образів	
Інститут кібернетики імені В. М. Глушкова НАН України [18]		
1	Відділ інтелектуальних інформаційних технологій	Лаптін Ю. П.
2	Відділ методів і технологічних засобів побудови інтелектуальних програмних систем	Єршов С. В.
3	Відділ методів комбінаторної оптимізації та інтелектуальних інформаційних технологій	Гуляницький Л. Ф.
4	Відділ методів індуктивного моделювання та керування	Гупал А. М.
Інститут телекомунікацій і глобального інформаційного простору НАН України [19]		
1	Науковий напрям "Інформаційно-комунікаційні та знання-орієнтовані технології"	Трофімчук О. М., Устименко В. О., Гуляєв К. Л.
2	Науковий напрям "Математичне моделювання та обчислювальні технології"	



Закінчення табл.

№ з/п	Наукова школа, науковий напрям, науковий відділ, науково-дослідна робота	Науковий керівник, засновник
<i>Інститут проблем реєстрації інформації НАН України [20]</i>		
1	НДР "Розробити та дослідити моделі предметних областей при формуванні баз знань і забезпеченні семантичного пошуку"	Ланде Д. В.
2	НДР "Розробити методи отримання, оброблення та застосування знань різної природи при підтримці прийняття рішень у соціотехнічних системах"	Циганок В. В.
<i>Національний технічний університет України "Київський політехнічний інститут імені Ігоря Сікорського" [21]. Напрями фундаментальних і прикладних наукових досліджень</i>		
1	Апаратне, математичне та програмне забезпечення цифрових систем у сучасних інформаційних технологіях. Методологія та методи побудови інтелектуалізованих інформацій та мережних технологій, баз даних і знань	
2	Програмно-апаратні комплекси розпізнавання образів аудіо- та відеосигналів	

Наукові дослідження в галузі штучного інтелекту виконуються також і в багатьох інших наукових і навчальних закладах:

- Національний університет Львівська політехніка (важковирішувані комбінаторні задачі штучного інтелекту, алгоритми та програмні засоби декомпозиції, апроксимації візуальних образів для їхнього збереження та розпізнавання на основі кластерного аналізу);

- Національний університет Запорізька політехніка (інтелектуальні інформаційні технології побудови автоматизованих систем технічного діагностування);

- Вінницький національний технічний університет (оброблення зображень, застосування технологій софт-комп'ютеринга – нечітка логіка, генетичні і мурашині алгоритми, методи розв'язання прикладних задач для оцінювання і забезпечення надійності людиномашинних систем; технічна і медична діагностика; експертні системи і системи підтримки прийняття рішень; інтелектуальне оброблення даних);

- Одеський національний політехнічний університет (інтелектуальний аналіз даних і системи штучного інтелекту, інтелектуальні технології та системи підтримки прийняття рішень, системи управління базами даних, інтелектуальні системи оброблення та розпізнавання сигналів і зображень, засоби моніторингу динамічних об'єктів, засновані на застосуванні моделей у вигляді нечітких множин, розроблення компонентів інтелектуальних інформаційних систем для моделювання, прогнозування та діагностики складних об'єктів і пристроїв електронної техніки) та в цілому ряді інших закладах освіти і наукових організацій.

Наукові дослідження в галузі штучного інтелекту, які представлено в дисертаціях, в Україні за останні декілька років в першу чергу розвивалися в таких напрямках:

1. Інтелектуальний аналіз та оброблення даних: Еволюційні нейро-фаззі системи в задачах інтелектуального аналізу даних (Бойко О. О.); Методи нечіткої кластеризації на основі ядерних функцій у задачах інтелектуального аналізу даних (Хаустова Я. В.); Нечітка кластеризація часових рядів в інтелектуальному аналізі потоків даних (Кобилін І. О.); Інформаційна технологія ідентифікації знань у наукометричних системах на основі інтелектуального аналізу слабоформалізованих даних (Петрасова С. В.); Інформаційна технологія інтелектуального аналізу фактографічних текстових ресурсів (Дорошенко А. Ю.); Моделі, методи й інтелектуальна інформаційна технологія аналізу неоднорідних послідовностей (Бабій А. С.); Алгоритмічні засоби та програмні компоненти комп'ютерних систем інтелектуального оброблення даних в організаційно-технічних комплексах (Сагайда П. І.); Методи та моделі розподіленого інтелектуального оброблення великих даних у спеціалізованих комп'ютерних системах (Аксак Н. Г.); Методологія інтелектуального аналізу геопросторових даних для задач сталого розвитку (Путренко В. В.);

2. Оброблення знань, мультиагентні технології: Інтелектуальні мультиагентні технології в епідемічних процесах систем популяційної динаміки (Чумаченко Д. І.); Методи та засоби планування дій спеціалізованих інтелектуальних агентів на основі онтологічного підходу (Вовнянка Р. В.); Об'єктно-орієнтована динамічна модель подання знань в інтелектуальних програмних системах (Терлецький Д. О.); Моделі та засоби інтелектуального оброблення даних корпоративних вебресурсів (Зосімов В. В.); Методологія моделювання інноваційних інтелектуальних систем прийняття рішень в управлінні економічними об'єктами (Мисник Б. В.); Інформаційні технології ідентифікації знань в інтелектуальних



системах на основі асоціативно-логічних перетворень (Булкін В. І.);

3. Інтелектуальні технології: Еволюційні штучні нейронні мережі прямого розповсюдження: архітектури, навчання, застосування (Безсонов О. О.); Теоретичні і методологічні основи розроблення інтелектуальних інформаційних технологій аналітичного оцінювання професійної діяльності (Заріцький О. В.); Формування масиву чисельних ознак для класифікації україномовних текстів в інформаційній технології інтелектуального моніторингу (Голуб М. С.); Методи та засоби побудови контекстно-залежних інтелектуальних систем у сфері працевлаштування (Завущак І. І.); Методологічні основи й інформаційна технологія створення розподілених інтелектуальних систем (Прохоров О. В.);

4. Інтелектуальні системи підтримки рішень: Теоретичні основи створення інтелектуальних систем комп'ютерної підтримки рішень під час управління енергозбереженням організацій (Доценко С. І.); Система інтелектуальної підтримки прийняття рішень при керуванні процесом збирання енергетичних культур для біогазових установок (Чирченко Д. В.); Інформаційна технологія підтримки прийняття рішень для забезпечення кібербезпеки розподілених інтелектуальних електроенергетичних систем (Коцюба І. В.); Методологія моделювання інноваційних інтелектуальних систем прийняття рішень в управлінні економічними об'єктами (Мінц О. Ю.);

5. Розпізнавання зображень: Методи та засоби розпізнавання змін властивостей об'єкта за зображенням на основі штучних нейронних мереж (Новосельцев І. В.); Методи сегментації зображень в інтелектуальних системах оброблення відсканованих документів (Іщенко О. В.).

Представлено наукові роботи в галузі штучного інтелекту, а також і в інших напрямках наукових досліджень цієї галузі.

В Україні проводиться багато конференцій, які присвячені штучному інтелекту і машинному навчанню, наприклад:

- Міжнародна науково-технічна конференція "Штучний інтелект та інтелектуальні системи" (AIIS);
- Міжнародна наукова конференція "Інтелектуальні системи прийняття рішень і проблеми обчислювального інтелекту" (ISDMCI);
- Міжнародна науково-технічна конференція "Інформатика, управління і штучний інтелект" (ІУШІ);
- Міжнародна науково-практична конференція "Інтелектуальні системи та інформаційні технології" (ISIT) та ін.

4. ВИСНОВКИ

Практично кожна відома компанія ІТ-ринку (Google, Facebook, Microsoft, IBM) тією або іншою мірою заявила про себе в контексті дослідження штучного інтелекту. Для України надзвичайно важливо продовжувати фундаментальні і прикладні дослідження в галузі штучного інтелекту, адже в майбутньому досягнення із цього напрямку будуть однією з невід'ємних складових економічного процвітання будь-якої держави та її успіху на міжнародному ринку новітніх технологій. В Україні почалися роботи в напрямку розроблення стратегії розвитку штучного інтелекту. Важливим є продовження розвитку на державному рівні.

У роботі проаналізовано напрямки досліджень і розробок у галузі штучного інтелекту в Україні. Нині у багатьох наукових організаціях, університетах виконуються дослідження і розроблення інтелектуальних систем. У статті виконано спробу узагальнити організації, роботи, які вони виконують, наукові школи, що займаються питаннями штучного інтелекту. Також проаналізовано результати дисертаційних досліджень останніх років у цій сфері.

СПИСОК ВИКОРИСТАНОЇ ЛІТЕРАТУРИ

- [1] Рассел С., Норвіг П. Искусственный интеллект: современный подход. М.: Издательский дом "Вильямс", 2016. 1408 с.
- [2] Сабініч А. Все що потрібно знати про штучний інтелект сьогодні, 2019, [Онлайн]. Режим доступу: <https://tokar.ua/read/34132/>
- [3] Україна стала одним із лідерів у сфері штучного інтелекту в Східній Європі. Telekritika. 2019, [Онлайн]. Режим доступу: <https://telekritika.ua/uk/iskusstvennyj-intellekt-ukraina-lidiruet/>.
- [4] Кравцова О. Вперед в будущее: почему правительство должно обратить внимание на искусственный интеллект. Экономическая правда. 2019, [Онлайн]. Режим доступу: <https://www.epravda.com.ua/rus/columns/2019/09/4/651230/>.
- [5] В Україні почали активно працювати над умовами розвитку штучного інтелекту. Юридична газета online. 2019, [Онлайн]. Режим доступу: <https://legalhub.online/investytsiyisfery-biznesu/v-ukrayini-pochaly-aktyvno-pratsyuvaty-nad-umovamy-rozvytku-shtuchnogo-intelektu/>.
- [6] Мінцифра сформувала експертний комітет з питань розвитку сфери штучного інтелекту в Україні. 2020, [Онлайн]. Режим доступу: <https://yur-gazeta.com/golovna/mincifra-sformuvala-ekspertnyy-komitetz-pitan-rozvytku-sferi-shtuchnogo-intelektu-v-ukrayini.html>.
- [7] Чубатюк Ю. Украине нужна Национальная стратегия развития искусственного интеллекта. 2018, [Онлайн]. Режим доступу: https://rus.lb.ua/society/2018/12/11/414650_ukraine_nuzhna_natsionalnaya.html.
- [8] Україна – у трійці лідерів за кількістю компаній в сфері штучного інтелекту у Європі. 2019, [Онлайн]. Режим доступу: <https://www.the-village.com.ua/village/business/news/>



281885-ukrayina-u-spisku-lideriv-v-evropi-za-kilkistyu-kompaniy-v-sferi-shtuchnogo-intelektu.

[9] Старт сфери штучного інтелекту: інформація для інвесторів. 2019, [Онлайн]. Режим доступу: <https://aicongference.com.ua/uk/news/startap-v-sfere-iskusstvennogo-intellekta-informatsiya-dlya-investorov-93160>.

[10] 5 українських стартапів, які працюють зі штучним інтелектом. 2019, [Онлайн]. Режим доступу: <https://www.kyivsmartcity.com/news/5-ukrainskix-startapiv>.

[11] Наукові установи України. [Онлайн]. Режим доступу: <http://www.urau.net.ua/refer/ndu/ndu.htm>.

[12] Харківський національний університет радіоелектроніки. [Онлайн]. Режим доступу: <https://nure.ua/naukovi-shkoli/>

[13] Інститут програмних систем НАН України. [Онлайн]. Режим доступу: <http://www.isofts.kiev.ua/topics-of-research/>

[14] Київський національний університет імені Тараса Шевченка, факультет комп'ютерних наук та кібернетики. [Онлайн]. Режим доступу: <http://www.cyb.univ.kiev.ua/uk/about.html>

[15] Інститут проблем математичних машин і систем НАН України. [Онлайн]. Режим доступу: <http://www.immsp.kiev.ua/activity/>

[16] Міжнародний науково-навчальний центр інформаційних технологій та систем НАН та МОН України. [Онлайн]. Режим доступу: <http://www.nas.gov.ua/UA/Org/ScientificDirections/Pages/Default.aspx?OrgID=0000447>

[17] Інститут проблем штучного інтелекту НАН та МОН України. [Онлайн]. Режим доступу: <http://www.ipai.net.ua/>

[18] Інститут кібернетики імені В.М.Глушкова НАН України. [Онлайн]. Режим доступу: <http://www.nas.gov.ua/>

[19] Інститут телекомунікацій і глобального інформаційного простору НАН України. [Онлайн]. Режим доступу: <http://www.nas.gov.ua/UA/Org/Pages/default.aspx?OrgID=0000315>

[20] Інститут проблем реєстрації інформації НАН України. [Онлайн]. Режим доступу: <http://ipri.kiev.ua/>

[21] Національний технічний університет України "Київський політехнічний інститут імені Ігоря Сікорського". [Онлайн]. Режим доступу: <https://kpi.ua/research>

Стаття надійшла до редколегії

10.11.2020

Trends in the development of artificial intelligence in Ukraine

In recent years, artificial intelligence technologies have begun to be used in many types of human life: education, medicine, business, finance, management, marketing, industry, etc., and have become a key trend of the time. The article provides information and analyzes the current state and trends in the development of artificial intelligence in Ukraine. An analysis of the strategy for the development of artificial intelligence, which is proposed by a number of ministries, considered the concept of development of artificial intelligence, which is proposed by the State Agency for e-Government. Data on companies working in the field of artificial intelligence are analyzed and provided. The main activities of Ukrainian companies in the field of artificial intelligence, including a number of Ukrainian startups, have been identified. Provides information on scientific schools and research conducted in the field of artificial intelligence in Ukraine. It is established that the main research focuses on the following areas: neural networks, pattern recognition, mathematical informatics, creation of intelligent information systems and knowledge bases, development of intelligent robotic systems, knowledge-oriented technologies, intelligent search, development of intelligent learning systems, big data analysis, fuzzy systems decision support, multi-agent systems and much more. In general, scientific research conducted in Ukraine in the field of artificial intelligence covers the basis of areas of work that correspond to research in the world. Also in the article the directions of dissertation researches for the last years on intellectual systems and systems of artificial intelligence are analyzed. The main areas in which the dissertation research was performed: intellectual analysis and data processing, knowledge processing and multi-agent technologies, intelligent technologies, including intelligent technologies to support acceptance, image recognition and a number of others.

Keywords: artificial intelligence; development strategy; scientific research; scientific schools; perspectives.



Георгій Гайна,
кандидат технічних наук, професор,
професор кафедри інтелектуальних
технологій Київського національного
університету імені Тараса Шевченка.

George Gaina,
Candidate of Technical Sciences, Professor,
Professor of the Department of Intellectual
Technologies of Taras Shevchenko National
University of Kyiv.



Г. М. Гнатієнко, orcid.org/0000-0002-0465-50181,
g.gna54@gmail.com

О. Круглов, orcid.org/0000-0003-4347-85842,
awahrhaft@gmail.com

Н. П. Тмєнова, orcid.org/0000-0003-1088-95473,
tmyenovox@gmail.com

Київський національний університет імені Тараса Шевченка, Київ, Україна

ОБЧИСЛЕННЯ РЕЗУЛЬТУЮЧОГО РАНЖУВАННЯ АЛЬТЕРНАТИВ НА ОСНОВІ ВИКОРИСТАННЯ НЕПОВНИХ ЕКСПЕРТНИХ РАНЖУВАНЬ

Неповнота інформації є характерною рисою організаційних систем. Неповнота даних супроводжує особу, що приймає рішення у всіх складових об'єктів корпоративної безпеки: керівництві організації, діяльності персоналу, активах компанії, упровадженіх бізнес-процесах, інформаційних та інших ресурсах, фінансових засобах, використовуваних технологіях, репутації компанії тощо. Незважаючи на це, слід приймати обґрунтоване рішення. Зокрема, поширеною практичною задачею є ранжування альтернатив різної природи. Це здійснюється експертами високої компетентності в межах зон відповідальності. Природним чином виникає ситуація прийняття рішення з неповними даними, на основі якої слід знайти повне результуюче ранжування альтернатив, яке найкращим чином апроксимує інформацію, одержану від експертів, тобто є в якомусь сенсі найближчою до заданих неповних експертних ранжувань. З метою порівняння різних способів досягнення результуючого ранжування альтернатив, розглядається формалізація задачі у класах однокритеріальних і багатокритеріальних моделей для метрик Кука, Хемінга, Евкліда та Литвака. Для розв'язання проблем, які виникають у ситуації неповноти інформації, вводять евристичні – емпіричні методологічні правила, які допомагають знаходити рішення та сприяють до визначеності математично некоректно поставлених задач. Вводиться поняття модифікованої медіани Литвака та компромісної медіани Литвака, яку знаходять із використанням мінімаксного критерію. Описано розроблені авторами алгоритми визначення медіан експертних ранжувань альтернатив: генетичний алгоритм та евристичний алгоритм. Для ілюстрації наведено схеми роботи генетичного алгоритму. Подано основні результати застосування описаних алгоритмів, які ілюструють ефективність їхнього застосування до задач ранжування, які характеризуються неповнотою інформації.

Ключові слова: генетичний алгоритм; організаційна система; евристичний алгоритм; інформаційна безпека; метрика; відстань; медіана; групове упорядкування об'єктів; неповне експертне ранжування.

1. ВСТУП

Неповнота експертної інформації є природним явищем й атрибутом багатьох ситуацій прийняття рішень, нерідко виникає на практиці та є одним із видів невизначеності – разом із нечіткістю, неточністю, ненадійністю, невизначеністю, некоректністю, неадекватністю, недостатністю тощо. Неповнота даних і неможливість їх доповнити закономірно супроводжує експертів та осіб, що приймають рішення, у їхній діяльності [1].

До об'єктів забезпечення корпоративної безпеки або безпеки організації у цілому, як прави-

ло, відносять різноманітні ресурси, складові організації та підсистеми:

- керівництво організації;
- персонал організації;
- активи компанії;
- бізнес-процеси, які застосовують в організації – формалізовані та неформалізовані;
- інформаційні ресурси організації;
- фінансові засоби, що забезпечують функціонування організації;
- технології, які використовує організація;

© Гнатієнко Г. М., Круглов О., Тмєнова Н. П., 2020



- бренд компанії;
- гудвіл – репутацію компанії тощо.

У зв'язку із цим для діяльності організацій є характерним чітке розмежування зон відповідальності менеджерів, і компетентність керівників організації у різних напрямках діяльності не є сталою величиною. Тому інформація, що стосується організації, як правило, є неповною, нечіткою, неоднорідною, і її агрегування вимагає розроблення, обґрунтування та належного застосування нових методів.

2. АНАЛІЗ ОСТАННІХ ДОСЛІДЖЕНЬ І ПУБЛІКАЦІЙ

Загальноприйнятою методологією, що використовується під час створення та дослідження складних соціально-економічних і технічних систем, є системний аналіз. Більшість його етапів базується на проведенні експертиз і застосуванні експертних оцінок, зокрема і в разі вибору структури системи, її оптимізації та розв'язання задач діагностики, класифікації, прогнозування тощо. Відомо також, що у розв'язанні багатьох практичних задач недооцінюються значення коректного одержання та оброблення експертної інформації [2]. Але адекватне використання експертних знань вимагає попереднього аналізу особливостей прийняття рішень людиною, адже дослідження складних, багатоаспектних і суперечливих об'єктів передбачає застосування значних аналітичних зусиль [3, 19].

За вибраними підходами, характером постановки проблем і формою зворотного зв'язку під час проведення експертних оцінок виокремлюють такі основні напрямки [1, 2]:

- абсолютні оцінки;
- бальні оцінки;
- метризовані оцінки відносної ваги альтернатив, їхніх параметрів, критеріїв чи відносної компетентності експертів;
- бінарні відношення, які відображують результати попарних порівнянь альтернатив;
- ранжування альтернатив.

Ця стаття стосується останніх двох підходів і в ній будуть описані задачі, пов'язані саме із цими напрямками.

Зазначимо, що відношення є одним з основних понять сучасної математики, а мову відношень успішно використовують для опису взаємозв'язків між альтернативами. Відношення часто застосовують також для представлення складних структур даних. Зокрема, бінарні відношення є теоретичним підґрунтям теорії прийняття рішень, оскільки властивості бінарних відношень використовуються для оцінювання переваг альтернатив у попарних порівняннях та адек-

ватно інтерпретуються в термінах системи переваг експертів.

Важливим способом представлення експертної інформації, який тісно пов'язаний із бінарними відношеннями, є ранжування альтернатив [2, 5]. Часто послідовність виконання кроків у прийнятті рішення є найсуттєвішим чинником, як захисту, так і забезпечення функціональної стійкості системи й усунення загроз.

Неповні ранжування досліджувалися, наприклад у роботах [3, 6]. Причому в монографії [3] до задач визначення узагальненого ранжування альтернатив застосовувалися класичні методи теорії голосувань. Однак у таких випадках перспективним є саме використання алгебраїчних методів обчислення медіани, яка відображує колективну думку експертів. У статті [6] розглянуто лише загальну схему побудови медіани на основі заданих експертами неповних ранжувань альтернатив. Опишемо докладніше цю процедуру.

3. МЕТА СТАТТІ

Метою статті є розроблення напрямів та алгоритмів до розв'язання задачі результуючого ранжування на основі заданих експертами неповних ранжувань альтернатив та експериментального дослідження зазначених підходів.

На основі запропонованих підходів необхідно розробити алгоритми обчислення результуючого ранжування альтернатив, узгодженого із заданими неповними ранжуваннями експертів з урахуванням різних метрик і різних критеріїв [7].

Для створення алгоритмів порівняння неповних ранжувань альтернатив слід розробити математичні моделі, які міститимуть інструменти, що дозволяють визначати відстані між неповними ранжуваннями. Таким чином виникає можливість застосування алгебраїчного підходу, який позбавлений багатьох недоліків, характерних для інших підходів [5, 8].

4. МЕТОДИ, ЩО ЗАСТОСОВУЮТЬСЯ

Основними методами, які використовуються у цій роботі, є методи теорії прийняття рішень [5], експертні технології [2], генетичний алгоритм [2] та евристичний алгоритм [2]. Методи теорії прийняття рішень [5] беруть початок від теорії дослідження операцій. Експертні технології широко застосовуються у різних галузях людської життєдіяльності й активно розвиваються в останні десятиліття. Одним із ключових понять експертних технологій є евристики [2], якими можуть виступати аксіоми, постулати, припущення, презумпції, парадигми, гіпотези, допущення, пропозиції тощо. Евристики є емпіричними методологічними правилами, які можуть допомогти знайти рішення та сприяти визначе-



ності некоректно поставлених задач. За допомогою евристичних алгоритмів генеруються варіанти рішень, близькі до оптимальних, але не гарантують знаходження оптимального розв'язку задачі. Генетичні алгоритми є еволюційними алгоритмами, які використовуються, зокрема і для розв'язання задач оптимізації, та теоретично можуть забезпечити знаходження глобального оптимуму задачі [9].

5. ВИКЛАД ОСНОВНОГО МАТЕРІАЛУ

Нехай розв'язується задача визначення важливості множини альтернатив деякої складної системи. Для цього здійснюється підготовка й обґрунтування рішення щодо послідовності розташування альтернатив великої та багато-профільної організації. Представники різних підрозділів і служб, які входять до експертної групи, у межах своєї компетенції та уявлення про важливість альтернатив надають індивідуальні пропозиції щодо послідовності реалізації альтернатив, які входять до сфери їхньої відповідальності, або пріоритетності розгляду альтернатив. Серед членів експертної групи часто не збігаються вподобання, що виражається в упорядкуванні ними різних підмножин альтернатив. Це обґрунтовано багатьма чинниками: компетентністю експертів, їхніми інтересами, пріоритетами, акцентами, аспектами, що розглядаються, уявленнями про ідеалізацію моделі тощо. Тому на попередньому етапі агрегування має бути здійснення етапу об'єднання підмножин альтернатив, проранжованих експертами, в єдину множину альтернатив.

5.1. ПОСТАНОВКА ЗАДАЧІ

У роботах [6, 7] уведено поняття неповного ранжування: це бінарне відношення, задане на підмножині альтернатив A' , $A' \subset A$, яке задовольняє властивості повноти, антисиметричності, транзитивності: але тільки на підмножині A' , $A' \subset A$, а не на всій множині A .

Нехай група експертів задає k неповних ранжувань альтернатив R^{iH} , $i = 1, \dots, k$. Необхідно знайти деяке групове (результуюче, агреговане, колективне, консенсусне, інтегративне) ранжування n альтернатив $R^* = (a_{i_1}, \dots, a_{i_n})$, $i_j \in I = \{1, \dots, n\}$, $j \in I$, яке побудоване за логікою, що характеризує процеси функціонування деякої організаційної системи. Тобто ранжування R^* , має бути побудоване на основі індивідуальних упорядкувань задач, які виконуються k елементами системи (експертами) $R^{iH} = (a_{i_1}, \dots, a_{i_{n_i}})$, $i \in J = \{1, \dots, k\}$, де n_i – кількість задач в індивідуальному експертному ранжуванні $i \in J$.

5.2. МЕТРИКИ ТА КРИТЕРІЇ

Відстані між ранжуваннями альтернатив визначаються з використанням метрик [7]:

- метрики Кука незбігання рангів (місць, позицій) альтернатив,

$$d(R^j, R^l) = \sum_{i \in I} |r_i^j - r_i^l|, \quad (1)$$

де r_i^l – ранг i -ї альтернативи у ранжуванні l -го експерта, R^l , $l \in L$, $1 \leq r_i^l \leq n$,

- метрики Хемінга [8];
- метрики Евкліда [8, 10–12];
- вектора переваг, елементами якого є кількість альтернатив, які передують кожній альтернативі у ранжуванні.

Критерії, які найчастіше застосовуються у таких випадках:

- адитивний;
- мінімаксний.

5.3. ФОРМАЛІЗАЦІЯ ЗАДАЧІ ВИЗНАЧЕННЯ РЕЗУЛЬТУЮЧОГО РАНЖУВАННЯ

Найпоширенішим методом знаходження результуючого ранжування альтернатив є обчислення медіани заданих ранжувань. Ця група методів узагальнення експертної інформації є найбільш достовірною та математично обґрунтованою. Розв'язки задачі, які визначаються при застосуванні різних метрик та різних критеріїв є медіанами заданих експертами лінійних порядків.

Позначимо множину усіх можливих ранжувань n альтернатив через Ω^R , множину матриць попарних порівнянь (МПП), які відповідають усім можливим ранжуванням n об'єктів – через Ω^B , множину векторів переваг, що відображують кількість альтернатив, які передують кожній альтернативі у ранжуванні (від 0 до $n-1$) – через Ω^P . Множину заданих експертами ранжувань та бінарних відношень, що їм відповідають, будемо позначати через R^A ; множину відповідних ранжуванням МПП та бінарних відношень між альтернативами – через R^B , а множину векторів переваг та відповідних відношень – через P^A .

Для випадку, який розглядається у цій роботі, потужність множин R^A , R^B та P^A однакова: $|R^A| = |R^B| = |P^A| = n$, $R^l \in R^A$, $B^l \in R^B$, $P^l \in P^A$, $l \in L$. Зрозуміло, що $R^A \subset \Omega^R$, $R^B \subset \Omega^B$, $P^A \subset \Omega^P$. У загальному випадку потужність множини $|\Omega^B| = |\Omega^P| = 2^{n(n-1)/2}$. Але для методу, який описується у цій роботі, будемо розглядати лише



їхні підмножини, позначаючи ці підмножини таким же чином $|\Omega^R| = |\Omega^B| = |\Omega^P| = n!$, оскільки нас не цікавлять нетранзитивні елементи простору розв'язків Ω^B та Ω^P .

5.4. ВИКОРИСТАННЯ МЕТРИКИ КУКА ДЛЯ НЕПОВНИХ РАНЖУВАНЬ АЛЬТЕРНАТИВ

Для метрики Кука (1) у випадку використання адитивного критерію обчислюються:

- медіана Кука – Сейфорда [5]:

$$R^{CS} \in \Omega^{CS} = \text{Arg} \min_{R \in \Omega^R} \sum_{l \in L} d^r(R, R^l); \quad (2)$$

- модифікована медіана Кука – Сейфорда [13]:

$$R^{MCS} \in \Omega^{MCS} = \text{Arg} \min_{R \in R^A} \sum_{l \in L} d^r(R, R^l) \quad (3)$$

При використанні мінімаксного критерію обчислюються:

- ГВ-медіана (компроміс) [5]:

$$R^{GB} \in \Omega^{GB} = \text{Arg} \min_{R \in \Omega^R} \max_{l \in L} d^r(R, R^l); \quad (4)$$

- модифікована ГВ-медіана:

$$R^{MGB} \in \Omega^{MGB} = \text{Arg} \min_{R \in R^A} \max_{l \in L} d^r(R, R^l) \quad (5)$$

Метрика Кука є популярною в задачах ранжування альтернатив [5, 14]. Для її застосування під час розв'язання поставленої задачі аналізу неповних ранжувань уведемо евристики та визначимо на їхній основі відстані від заданих експертами ранжувань до опорного ранжування.

У зв'язку з особливостями обчислення узагальненого ранжування при неповній початковій інформації, у [7] запропоновано застосовувати низку евристик. Зокрема, складові відстаней у разі неповних ранжувань об'єктів описуються таким чином.

Евристика Е1. Відстань від заданих експертами неповних ранжувань $R^{ih}, i=1, \dots, k$, до будь-якого ранжування $R^{*(0)}$ складається з двох складових: визначеної частини відстані та ймовірнісної.

Евристика Е2. Не задана експертом альтернатива породжує невідомі відношення між усіма іншими альтернативами й участі у ранжуванні не бере, тобто ця альтернатива не представлена у неповному ранжуванні. Таким чином, у ході задання неповних ранжувань для кожного експерта маємо кількість альтернатив:

- n_i – заданих ним альтернатив у ранжуванні $R^{ih}, i=1, \dots, k$, які будуть складати визначену частину відстаней;
- $(n - n_i) = v_i$ – не заданих експертом альтернатив у ранжуванні $R^{ih}, i=1, \dots, k$, які складають імовірнісну частину відстаней.

Евристика Е3. Імовірнісна частина відстані від заданого експертом ранжування $R^{ih}, i=1, \dots, k$, до будь-якого опорного ранжування завжди дорівнює $v_i, i=1, \dots, k$, для метрики Кука.

Визначену частину відстаней обчислюють за формулою (1).

5.5. ЕВКЛІДОВА МІРА БЛИЗЬКОСТІ Й ОБЧИСЛЕННЯ СЕРЕДНЬОГО

Відповідно до [5, 11, 13], результуючим ранжуванням може бути вибране середнє $R^S \in \Omega^S = \text{Arg} \min_{R \in R^B} \sum_{l \in L} d^2(R, R^l)$.

Таке результуюче ранжування доцільно використовувати [5, 11–13], якщо відстань між експертними ранжуваннями визначається за допомогою евклідової міри близькості:

$$d^E(B^j, B^l) = \left(\sum_{t \in H} (r_t^j - r_t^l)^2 \right)^{1/2},$$

$$t \in H, \quad j, l \in L.$$

5.6. ВИКОРИСТАННЯ МЕТРИКИ ХЕМІНГА ДЛЯ НЕПОВНИХ РАНЖУВАНЬ АЛЬТЕРНАТИВ

Імовірнісна частина від заданого експертом ранжування $R^{ih}, i=1, \dots, k$, до будь-якого ішого ранжування для метрики Хемінга завжди дорівнює $v_i \cdot (v_i - 1) / 2, i=1, \dots, k$.

Для переходу від простору рангів до простору парних порівнянь альтернатив [7] індивідуальні неповні переваги, задані кожним експертом на підмножинах альтернатив $R^{ih}, l \in L$, представляють у вигляді неповної матриці парних порівнянь:

$$B^{ih} = (b_{ij}^{ih}), \quad j \in I, \quad l \in L, \quad (6)$$

де $b_{ij}^l = 1, i, j \in I, l \in L$ тоді і тільки тоді, коли на думку l -го експерта i -та альтернатива переважає j -ту альтернативу. Причому $b_{ij}^l = -b_{ji}^l, i, j \in I, l \in L$. Коли l -й експерт не задав відношення переваги між альтернативами a_i та a_j , то цей факт позначається $b_{ij}^{ih} = "**", j \in I, l \in L$.

Для визначення відстаней між неповними відношеннями (6) застосовується метрика Хемінга

$$d^h(B^j, B^l) = 0,5 \sum_{i \in I} \sum_{s \in I} |b_{is}^j - b_{is}^l|.$$

Евристика Е4. Математичне сподівання невизначених відстаней між альтернативами у ранжуванні дорівнює одиниці. Тобто відстань між елементами МПП, хоча б один з яких невизначений, має бути рівною 1 – з урахуванням припущення, що рівність його значення "–1" або "1" є однаково ймовірними.



Оскільки матриці відношень B^H та R^H вигляду (1) є кососиметричними, без втрати загальності, будемо використовувати побудовані на їхній основі вектори:

$$c_t = b_{ij}, \quad x_t = r_{ij}, \quad t = (i-1)n + j - (i+1)i/2, \\ 1 \leq i < j \leq n.$$

Позначимо через $N = n(n-1)/2$ кількість елементів векторів c , а через $H = \{1, \dots, N\}$ – множину індексів елементів цих векторів. Тоді відстань між відношеннями B та R запишеться у вигляді

$$d^h(B^j, B^l) = \sum_{t \in H} |c_t^j - c_t^l|, \quad t \in H. \quad (7)$$

Для метрики Хемінга (7) у разі використання аддитивного критерію обчислюються:

- медіана Кемені – Снелла:

$$R^{KC} \in \Omega^{KC} = \text{Arg min}_{B \in \Omega^B} \sum_{l \in L} d^h(B, B^l); \quad (8)$$

- модифікована медіана [13] Кемені – Снелла:

$$R^{MKC} \in \Omega^{MKC} = \text{Arg min}_{B \in R^B} \sum_{l \in L} d^h(B, B^l) \quad (9)$$

Використовуючи мінімаксний критерій, обчислюють:

- ВГ-медіану (компроміс) [2, 5]:

$$R^{BG} \in \Omega^{BG} = \text{Arg min}_{B \in \Omega^B} \max_{l \in L} d^h(B, B^l); \quad (10)$$

- модифікована ВГ-медіана:

$$R^{MBG} \in \Omega^{MBG} = \text{Arg min}_{B \in R^B} \max_{l \in L} d^h(B, B^l) \quad (11)$$

Методи визначення медіан вигляду (2), (4), (8), (10) та їхні особливості розглядаються у монографії [2]. Модифіковані медіани (3), (9), запропоновано застосовувати, зокрема, у роботі [13], але їхнє використання у багатьох практичних задачах є недоцільним, а інколи – необґрунтованим і неправомірним. Пов'язано це з тим, що модифікована медіана суттєво обмежує простір вибору, тому у разі застосування таких критеріїв, як правило, знаходимо неефективні розв'язки: вони домінуються, зокрема власне і медіанами вигляду (3), (5), (9), (11).

5.7. ВИКОРИСТАННЯ ВІДСТАНЕЙ МІЖ ВЕКТОРАМИ ПЕРЕВАГ ДЛЯ НЕПОВНИХ РАНЖУВАНЬ АЛЬТЕРНАТИВ

Для застосування алгебраїчного підходу на множині ранжувань можна також використовувати відстань, яка ґрунтується на векторах переваг [15, 16]: $\pi^l = (\pi_1^l, \dots, \pi_n^l)$, де π_i^l – кількість альтернатив, які передують i -й альтернативі у l -му ранжуванні. У монографії [15, 16] Б. Г. Литвак запропонував для векторів переваг π^1 та π^2 , сформованих на основі ранжувань R^1 та R^2 , визначати відстань за формулою

$$d^\pi(R^1, R^2) = \sum_{i \in I} |\pi_i^1 - \pi_i^2|. \quad (12)$$

Для векторів переваг вигляду (12) у разі використання аддитивного критерію обчислюється медіана Литвака [15, 16]:

$$R^L \in \Omega^L = \text{Arg min}_{R \in \Omega^R} \sum_{l \in L} d^\pi(R, R^l). \quad (13)$$

Уведемо також поняття модифікованої медіани Литвака:

$$R^{ML} \in \Omega^{ML} = \text{Arg min}_{R \in \Omega^A} \sum_{l \in L} d^\pi(R, R^l). \quad (14)$$

Під час використання мінімаксного критерію введемо такі медіани:

- LK-медіана або компромісна медіана Литвака:

$$R^{LK} \in \Omega^{LK} = \text{Arg min}_{R \in \Omega^R} \max_{l \in L} d^\pi(R, R^l); \quad (15)$$

- модифікована LK-медіана:

$$R^{MLK} \in \Omega^{MLK} = \text{Arg min}_{R \in \Omega^A} \max_{l \in L} d^\pi(R, R^l). \quad (16)$$

Для неповних ранжувань і використанні відстані (12) між векторами переваг будемо застосовувати також евристики, подібні до тих, які введено для медіан, що ґрунтуються на застосуванні метрики Кука (1). Подібним чином застосовуються та модифікуються відстані між векторами переваг для визначення медіани Литвака.

5.8. СПОСОБИ ОДЕРЖАННЯ ОПОРНИХ РІШЕНЬ ДЛЯ ЗАСТОСУВАННЯ ГЕНЕТИЧНОГО АЛГОРИТМУ ВИЗНАЧЕННЯ МЕДІАН

Одним із способів визначення медіани експертних ранжувань є застосування розробленого авторами генетичного алгоритму, описаного в роботі [17]. Важливим елементом цього алгоритму є вибір опорного рішення. Можна запропонувати кілька варіантів вибору такого рішення:

- генерувати опорне ранжування, у якого перші елементи повторюють ранжування першого експерта, наступні альтернативи з'являються у ранжуванні мірою доповнення множини альтернатив, одержаних під час вводу неповних ранжувань від наступних експертів;

- як опорні рішення можуть бути застосовані модифіковані медіани Кука – Сейфорда, ГВ-медіана, медіана Кемені – Снелла, ВГ-медіана [2], медіани Литвака та LK-медіани, обчислювальна складність яких невелика;

- вибирати опорні ранжування, одержані за правилами голосування [5]: по Кондорсе, Борда, Сімпсону, Нансону, Копленду, Кемені – Янгу, Тайдеману, Шульце, Болдуїну, альтернативних голосів, відносної більшості тощо.

На наступному етапі серед згенерованих опорних рішень вибираємо для продовження



роботи алгоритму таке, яке має найкраще значення за критерієм, з урахуванням якого розв'язується поточна задача.

5.9. ГЕНЕТИЧНИЙ АЛГОРИТМ ВИЗНАЧЕННЯ МЕДІАН ЕКСПЕРТНИХ РАНЖУВАНЬ АЛЬТЕРНАТИВ

У генетичних алгоритмах застосовують мутації та кросовери для генерування нових поколінь. Але для ранжувань класичні прийоми кросовера не працюють через те, що строге ранжування $R \in \Omega^R$ має складатися з n елементів, які не повинні повторюватися.

У випадку одиначної мутації будемо представляти два випадкові елементи $r_i, r_j \in R, i \neq j, i, j \in \{1, \dots, n\}$, місцями у результуючому ранжуванні R^* відносно їхнього початкового положення. Функція мутації $f(R)$ буде виглядати таким чином [17]:

$$f(R) = \begin{cases} \exists r_i, r_j \in R, i \neq j: \\ (r_{i-1} > r_i > r_{i+1}) \vee (r_{j-1} > r_j > r_{j+1}), \\ \exists r_i^*, r_j^* \in R^*: \\ (r_{i-1} > r_i^* > r_{i+1}) \vee (r_{j-1} > r_j^* > r_{j+1}), \\ r_i = r_i^*, r_j = r_j^*, \end{cases}$$

де R^* – ранжування, одержане в результаті мутації. Таке перетворення повторюється m разів, тобто маємо $f^m(R), m \in \{0, 1, 2, 3, 4\}$.

Кросовером пари R^1, R^2 буде ранжування R^* . Визначимо функцію кросовера $g(R^1, R^2, i, j)$:

$$R^* \cap R^1 \forall r_k \in R^1, i \geq k \geq j,$$

$$R^* \cap R^2 \forall r_k \notin R^* \cap R^1.$$

Таким чином, частина елементів результуючого ранжування R^* буде упорядкована так, як в R^1 , а всі інші елементи R^* будуть упорядковані як в R^2 .

Кожна відстань від поточної медіани до експертних ранжувань $R^{iH} = (a_1, \dots, a_{n_i}), i \in J = \{1, \dots, k\}$, обчислюється за правилами для неповних ранжувань, описаними у пп. 5.4 для простору рангів альтернатив, у пп. 5.6 для простору попарних порівнянь між альтернативами або у пп. 5.7 для простору векторів переваг.

Опишемо покроковий генетичний алгоритм з урахуванням того, що ранжування можна вважати фенотипом – послідовністю генів, де кожен ген буде відповідати порядку конкретної альтернативи у ранжуванні. Популяція – набір експертних ранжувань, доповнених до повних ранжувань шляхом застосування евристик, візьмемо за стартову популяцію (R^0) – ранжирування, що дані нам експертами. Алгоритм пошуку компромісного ранжування виглядатиме таким чином.

Крок 1. Створення дочірніх ранжувань $(R^{(i)*})$ на основі наявних і додавання їх до нової популяції $(R^{(i)} + R^{(i)*})$.

Крок 2. Обчислення фітнес-функції для кожного ранжування.

Крок 3. Сортування ранжувань альтернатив за значеннями їхніх фітнес-функцій.

Крок 4. Відсіювання оптимальних ранжувань у новому поколінні $(R^{(i+1)})$.

Крок 5. Повторення циклу.

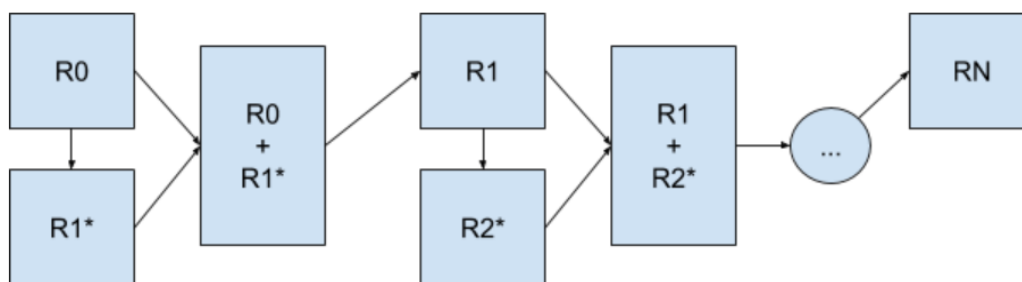


Рис. 1. Схема генерації ранжування

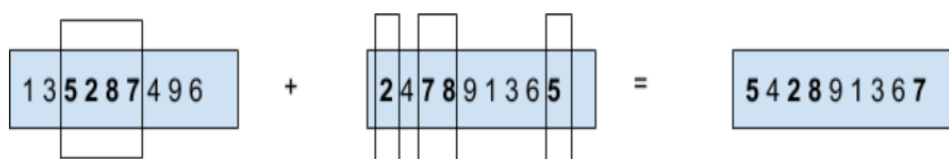


Рис. 2. Схема перетворення елементів у новому ранжуванні



Для генерації нових індивідів застосовують такий підхід (рис. 1, рис. 2).

Крок 1. З вихідної популяції вибирають два випадкових ранжування.

Крок 2. Друге ранжування альтернатив копіюють у результуюче.

Крок 3. Вибирають підпоследовність із першого ранжування.

Крок 4. Переупорядковують елементи нового ранжування, які входять у підпоследовність, у порядку, який відповідає поточній последовності альтернатив.

Крок 5. Мутація: міняють місцями пари елементів у новому ранжуванні альтернатив.

Далі необхідно визначити фітнес-функцію для всіх нових ранжувань. Застосовуючи цільову метрику, ми можемо оцінити, наскільки дане ранжування віддалене від усіх інших. Сума відокремлення від експертних ранжувань і буде фітнес-функцією. Оскільки, у цьому випадку ми вирішуємо завдання мінімізації даного параметра – сортуємо нову популяцію за зростанням. Після чого, відсіюємо "найгірших" індивідів.

Так, ранжування з мінімальною фітнес-функцією в останньому поколінні (RN) і буде нашим оптимальним рішенням.

5.10. РЕЗУЛЬТАТИ ОБЧИСЛЮВАЛЬНОГО ЕКСПЕРИМЕНТУ ІЗ ЗАСТОСУВАННЯМ ГЕНЕТИЧНОГО АЛГОРИТМУ

З метою дослідження описаного алгоритму авторами проведено обчислювальні експерименти з різною кількістю експертів та альтернатив.

Численні обчислювальні експерименти, проведені з використанням генетичного алгоритму, показують перспективність застосування такого підходу. Для випадково згенерованих ранжувань 100–200 альтернатив програма обчислює у просторі всіх можливих ранжувань Ω^R медіани $R^{KS1}, R^{G1}, R^{KS2}, R^{G2}$, які за критеріями (5)–(6) є приблизно на 10 % ближчими до медіан, заданих експертами $R^l, l \in L$, ніж опорні ранжування. Застосовуючи генетичний алгоритм, покращуємо їх у кожному з вибраних восьми напрямків.

5.11. ЕВРИСТИЧНИЙ АЛГОРИТМ ВИЗНАЧЕННЯ МЕДІАН ІЗ ВИКОРИСТАННЯМ МЕТРИКИ ХЕМІНГА

У роботі [18] наведено евристичний алгоритм визначення медіани R^* заданої множини неповних ранжувань $R^{ih}, i \in I$, у вигляді медіани Кемені – Снелла, що розглянуто у роботі [18].

Крок 1. Запишемо верхні наддіагональні трикутні частини матриць $B^{ih}, i \in I$, у вигляді ряд-

ків нової матриці: $C = (c_{ij}), i \in I, j \in J = \{1, \dots, N\}$, $N = n * (n - 1) / 2$. Елементи матриці C визначають таким чином:

$$c_{ij} = b_{lt}^i, \quad j = (l - 1) * n + t - l * (l + 1) / 2, \\ 1 \leq l \leq t \leq n, \quad i \in I.$$

Крок 2. Визначимо метризовану матрицю пріоритетів $M = (m_{lt}), l, t \in I$, елементи якої обчислюють як

$$m_{lt} = \sum_{\substack{i \in I, \\ c_{ij} \geq 0}} c_{ij} / \sum_{\substack{i \in I, \\ c_{ij} < 0}} |c_{ij}|, \quad 1 \leq l \leq t \leq n, \quad j \in J, \quad (17)$$

де $|x|$ – абсолютне значення числа x , $l = [j / (n - [j / n])] + 1$, $t = j - (l - 1) * n + l * (l + 1) / 2$, де $[x]$ – ціла частина числа x . Причому значення симетричних елементів матриці перебувають у такому співвідношенні: $m_{tl} = 1 / m_{lt}$, $m_{tt} = 1, \forall t, l \in I$. Зрозуміло, що елементи $c_{ij}, i \in I, j \in J$, значення яких не визначено експертами, тобто $c_{ij} = '*'$, не беруть участі у формуванні значень m_{lt} , $1 \leq l \leq t \leq n, j \in J$, вигляду (1).

Крок 3. Побудова матриці нев'язок $P = (p_{ij}), i = 1, \dots, n, j = 1, \dots, N$, на основі аналізу всіх можливих відношень між трійками альтернатив.

$$p_{1j} = m_{lt}, \quad j = (l - 1) * n + t - l * (l + 1) / 2, \\ 1 \leq l \leq t \leq n,$$

$$p_{ij} = \begin{cases} m_{ls} / m_{ts}, & l < s, \quad s > t, \\ m_{ls} * m_{ts}, & l < s, \quad s < t, \\ m_{st} / m_{sl}, & l > s, \quad s > t, \end{cases} \quad (18)$$

$$s = 1, \dots, n, \quad s \neq l, \quad s \neq t, \quad i = l + 1, \quad l = 1, \dots, n - 1, \\ t = l + 1, \dots, n, \quad j = (l - 1) * n + t - l * (l + 1) / 2.$$

Крок 4. Здійснення для кожного індексу $i = 2, \dots, n$, заміни одного з елементів матриці $M = (m_{lt}), l, t \in I$, елементом матриці $P = (p_{ij}), i = 2, \dots, n, j = 1, \dots, N$, якщо не збігаються елементи: $m_{lt} \neq p_{ij}$, для $j = (l - 1) * n + t - l * (l + 1) / 2$.

Крок 5. Для кожної чергової згенерованої у пункті 4 матриці визначаємо вагові коефіцієнти кожного з n альтернатив методом рядкових або стовпчикових сум, які є еталоном обчислювальної простоти визначення "ваги".

Крок 6. Упорядкування елементів вектора, одержаного у пункті 5, за спаданням значень (для рядкових сум) або зростанням значень (для рядкових сум). Індеси упорядкованого таким чином вектора вважатимемо індексами альтернатив у результуючому ранжуванні.



Крок 7. Визначення суми відстаней від одержаного в пункті 6 ранжування до ранжувальних заданих експертами, за правилами, описаними для неповних матриць попарних порівнянь.

Крок 8. Продовження процедур, описаних у пунктах 4–7, поки будуть обчислені всі модифіковані матриці $M = (m_{lt})$, $l, t \in I$, шляхом заміни в них чергового елемента матриці $P = (p_{ij})$, $i = 2, \dots, n$, $j = 1, \dots, N$.

Одержані таким чином найкращі матриці, визначені в пунктах 4–8, і складатимуть множину медіан Кемені – Снелла. ВГ-медіана також може бути визначена за допомогою описаного алгоритму, оскільки з його допомогою генеруються варіанти, близькі до оптимальних розв'язків.

5.12. РЕЗУЛЬТАТИ ОБЧИСЛЮВАЛЬНОГО ЕКСПЕРИМЕНТУ ІЗ ЗАСТОСУВАННЯМ ЕВРИСТИЧНОГО АЛГОРИТМУ

Авторами проведено обчислювальний експеримент із застосуванням евристичного алгоритму для матриць розмірностей від 6×6 до 10×10 , тобто коли задаються кілька десятків випадково упорядкованих 6–10 альтернатив. Експерименти із застосуванням описаного методу, перевірені точними методами, дали такі результати:

- за слабкої узгодженості заданих експертами матриць множина ефективних розв'язків може бути дуже великою – кількість медіан Кемені – Снелла для деяких профілів переваг складає до 15 % загальної кількості ранжувальних на множині альтернатив, тобто – до $0,15 \cdot n!$;
- серед знайдених за допомогою описаного алгоритму розв'язків при застосуванні методу строкових сум до 50 % ранжувальних альтернатив є медіанами Кемені – Снелла, а при застосуванні методу рядкових сум – до 83 %;
- застосування описаного алгоритму у деяких випадках дозволяє знайти до 39 % ранжувальних, які належать до множини медіан Кемені – Снелла;
- серед знайдених за допомогою описаного алгоритму компромісних ранжувальних вигляду (18) до 22 % є одночасно ГВ-медіанами.

Таким чином, описаний алгоритм є зручним евристичним алгоритмом для визначення множин медіан Кемені – Снелла та ГВ-медіан. Хоча за допомогою цього алгоритму всю множину ефективних розв'язків визначити неможливо, але деякі з медіан гарантовано обчислюються. Проведене дослідження підтверджує взаємозв'язок між ординальними та кардинальними моделями задання експертних переваг, а також перспек-

тивність застосування евристичних алгоритмів у задачах експертного оцінювання.

6. ВАРІАНТИ ЗАСТОСУВАННЯ РОЗРОБЛЕНИХ АЛГОРИТМІВ

Описані у цій роботі алгоритми обчислення результуючого ранжування заданих неповних експертних ранжувальних альтернатив можуть бути застосовані для розв'язання різноманітних задач у різних предметних сферах. Зокрема, з допомогою описаних підходів і наведених алгоритмів можуть бути вирішені такі проблеми:

- розроблення системи комплексних заходів кібербезпеки [6];
- розроблення та впровадження процедур швидкого реагування на зовнішні загрози для організаційної системи;
- пошук ранжування найпопулярніших криптовалют на торговельних майданчиках і визначення інтегрального упорядкування популярності криптовалют;
- вибір підручників для формування переліку літератури та визначення послідовності викладення навчальної дисципліни;
- визначення послідовності підготовки добірок книг у бібліотеках для гармонійного формування особистості;
- побудова послідовності лекційних курсів у задачах розроблення навчальних планів освітніх програм тощо.

7. ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

У багатьох задачах події, які слід упорядкувати шляхом використання неповних експертних ранжувальних, мають виконуватися паралельно або навіть відбуватися одночасно. Тому логічно формалізувати поставлену задачу у класі обчислення колективного квазіпорядку [19].

Перспективним є також розроблення паралельних алгоритмів із використанням методів штучного інтелекту, застосування яких за описаних підходів може сприяти одержанню синергетичного ефекту у випадку:

- формалізації та подальшої оптимізації бізнес-процесів;
- розв'язання задач відновлення інформації у відношеннях переваги експертів на основі визначення групового ранжування;
- застосування формалізмів задачі визначення колективного ранжування до широкого класу класичних комбінаторних задач в описах відповідних постановок для адаптації та інтерпретації постановки.



8. ВИСНОВКИ

Таким чином у цій роботі досліджено підходи до визначення результуючого ранжування альтернатив на основі неповних експертних ранжувань та одержано такі основні результати:

- запропоновано постановки задач визначення групового ранжування альтернатив на основі неповних експертних ранжувань;
- уведено поняття компромісної медіани Литвака;
- розглянуто підходи до агрегування експертних даних з урахуванням особливостей неповної інформації, одержаної від експертів;
- розроблено алгоритми для розв'язання задач обчислення колективного ранжування великого розміру;
- проведено обчислювальні експерименти з дослідження описаних алгоритмів та особливостей задач ранжування;
- встановлено, що модифіковані медіани заданої множини експертних ранжувань завжди закономірно й обґрунтовано домінуються медіанами, обчисленими у повному просторі всіх можливих ранжувань альтернатив.

СПИСОК ВИКОРИСТАНОЇ ЛІТЕРАТУРИ

- [1] О.Ф. Волошин, Г.М. Гнатієнко, О.В. Дробот, "Метод косвенного определения интервалов весовых коэффициентов параметров для метризованных отношений между объектами", Проблемы управления и информатики, №2, С.34–41, 2003.
- [1] Г.М. Гнатієнко, В.Є. Снитюк, Экспертні технології прийняття рішень, ТОВ "Маклаут", Київ, 2008. – 444 с.
- [2] A. Dodonov, D. Lande, V. Tsyganok, O. Andriichuk, S. Kadenko, A. Graivoronskaya (2019). Information Operations Recognition. From Nonlinear Analysis to Decision-Making, Lambert Academic Publishing, 2019, 275 p.
- [3] С.В. Каденко, В.В. Циганок, О.В. Андрійчук, О.В. Карабчук, "Аналіз інструментарію підтримки прийняття рішень у контексті вирішення задач стратегічного планування", Реєстрація, зберігання і обробка даних, Т. 22, № 2, С.77–91, 2020.
- [4] О.Ф. Волошин, С.О. Машенко, Теорія прийняття рішень: Навчальний посібник, Видавничо-поліграфічний центр "Київський університет", Київ, 2006, 304 с.
- [5] Г.М. Гнатієнко, Н.П. Тмєнова, "Визначення пріоритетності заходів кібербезпеки при неповних експертних ранжуваннях", Безпека інформаційних систем і технологій, №1(2), с. 9–15, 2020.
- [6] Н. Hnatiienko, N. Tmienov, A. Kruglov, "Methods for Determining the Group Ranking of Alternatives for Incomplete Expert Rankings". in: Shkarlet S., Morozov A., Palagin A. (eds) Mathematical Modeling and Simulation of Systems (MODS'2020), Advances in Intelligent Systems and Computing, vol 1265. Springer, Cham. https://doi.org/10.1007/978-3-030-58124-4_21, 2021, 1265 AISC, pp. 217–226.
- [7] И.М. Макаров, Т.М. Виноградская, А.А. Рубчинский, Теория выбора и принятия решений, Наука, Москва, 1982, 328 с.

- [8] В.Є. Снитюк, Прогнозування. Моделі. Методи. Алгоритми: Навчальний посібник, "Маклаут", Київ, 2008, 364 с.
- [9] Ю.Я. Самохвалов, Е.М. Науменко, Экспертное оценивание. Методический аспект. Монография, ДУИКТ, Киев, 2007, 262 с.
- [10] П.М. Зінько, Методичні рекомендації щодо забезпечення самостійної роботи студентів з дисципліни "Системи та методи прийняття рішень" (для бакалаврів), ДП "Видім "Персонал", Київ, 2014, 36 с.
- [11] Методика оцінювання наукових напрямів закладів вищої освіти під час проведення державної атестації закладів вищої освіти в частині провадження ними наукової (науково-технічної) діяльності. Затверджена наказом Міністерства освіти і науки України 12 березня 2019 р. №338, Зареєстроване в Міністерстві юстиції України 27 червня 2019 р. за №688/33659.
- [12] А.И. Орлов, Методы принятия управленческих решений: учебник, КНОРУС, Москва, 2018, 286 с.
- [13] Г.М. Гнатієнко, В.Є. Снитюк, "Апостеріорне визначення компетентності експертів в умовах невизначеності", Інформаційні технології та безпека. Матеріали XIX Міжнародної науково-практичної конференції ІТБ-2019, ООО "Инжиниринг", Киев, 2019, С. 184–187.
- [14] Б.Г. Литвак, Экспертная информация: Методы получения и анализа, Радио и связь, Москва, 1982, 184 с.
- [15] В.А. Болтенков, В. И. Куваева, А. В. Позняк, "Анализ медианных методов консенсусного агрегирования ранговых предпочтений", Информатика та математичні методи в моделюванні, Т. 7, № 4, 2017, С. 307–317.
- [16] Н.М. Hnatiienko, A.I. Kruglov, "Definition of a compromise ranking on the set of individual rankings using the genetic algorithm", Міжнародний науковий симпозиум "ІНТЕЛЕКТУАЛЬНІ РІШЕННЯ". Обчислювальний інтелект (результати, проблеми, перспективи): праці міжнар. наук.-практ. конф., 15-20 квітня 2019 р., Ужгород, ДВНЗ "Ужгородський національний університет", та [ін.]; наук. Ред. В.Є. Снитюк. С.63–64.
- [17] Н.М. Hnatiienko, V.H. Hnatiienko, "Heuristic algorithm for determining compromise rankings on a set of individual expert rankings", INTELLIGENT SOLUTIONS. Decision Making Theory: Proceedings of the International School-Seminar, April 15–20, 2019, Uzhhorod, Ukraine, Uzhhorod national university and [etc]; Leonid F. Hulianyskyi (Editor). – pp.53–54.
- [18] О.Ф. Волошин, Г.М. Гнатієнко, В.І. Кудін, Послідовний аналіз варіантів. Технології та застосування, Стило, Київ, 2013, 304с.

Стаття надійшла до редколегії

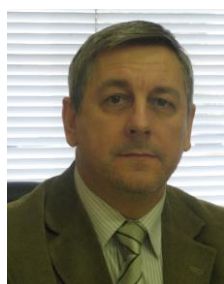
10.11.2020



Calculation of the resulting ranking of alternatives based on the use of incomplete expert rankings

Incomplete information is a characteristic feature of organizational systems. Incomplete data accompanies the decision-maker in all components of corporate security, namely the management of the organization, staff activities, company assets, implemented business processes, information and other resources, financial resources, used technologies, the company's reputation, etc. Nevertheless, a reasonable decision should be made. In particular, a common practical task is to rank alternatives of different nature. This is done by experts of high competence within the areas of responsibility. Naturally, there is a situation of decision-making with incomplete data, on the basis of which it is necessary to find a complete resulting ranking of alternatives, which best approximates the information obtained from experts, ie is in some sense closest to the given incomplete expert rankings. In order to compare different ways to achieve the resulting ranking of alternatives, the formalization of the problem in the classes of single-criteria and multicriteria models for the metrics of Cook, Heming, Euclid and Litvak is considered. To solve the problems that arise in a situation of incomplete information, a number of heuristics that are empirical methodological rules that help to find solutions and contribute to the definition of mathematically incorrect problems are introduced. The notion of the modified Litvak median and the Litvak compromise median, which is used using the minimax criterion, is introduced. The algorithms developed by the authors for determining the medians of expert rankings of alternatives, namely the genetic algorithm and the heuristic algorithm are described. To illustrate the results the schemes of the genetic algorithm are given. The main results of the application of the described algorithms, which illustrate the efficiency of their application to ranking problems, that are characterized by incomplete information are given.

Keywords: genetic algorithm; organizational system; heuristic algorithm; informational security; metrics; distance; median; group arrangement of objects; incomplete expert ranking.



Григорій Гнатієнко,
заступник декана факультету інформаційних технологій Київського національного університету імені Тараса Шевченка.

Hryhorii Hnatienko,
Deputy Dean of the Faculty of Information Technology, Taras Shevchenko National University of Kyiv.



Наталія Тмєнова,
доцент кафедри інтелектуальних технологій факультету інформаційних технологій Київського національного університету імені Тараса Шевченка.

Natalia Tmenova,
Associate Professor, Department of Intellectual Technologies, Faculty of Information Technologies, Taras Shevchenko National University of Kyiv.



Олександр Круглов,
асистент кафедри інтелектуальних технологій факультету інформаційних технологій Київського національного університету імені Тараса Шевченка.

Alexander Kruglov,
Assistant of the Department of Intellectual Technologies, Faculty of Information Technologies, Taras Shevchenko National University of Kyiv.



ІНФОРМАЦІЙНА ТА КІБЕРНЕТИЧНА БЕЗПЕКА

УДК 004.415.056.5

DOI <https://doi.org/10.17721/ISTS.2020.4.38-47>С. В. Толупа, orcid.org/0000-0002-1919-9174, tolupa@i.uaН. В. Лукова-Чуйко, orcid.org/0000-0003-3224-4061, lukova@ukr.netВ. С. Наконечний, orcid.org/0000-0002-0247-5400, nvc2006@i.uaВ. Г. Сайко, orcid.org/0000-0002-3059-6787, vgsaiko@gmail.com

Київський національний університет імені Тараса Шевченка, Київ, Україна

О. Кулінич, orcid.org/0000-0002-0643-6898,

oleh_nicol@gmail.com

Національний технічний університет України імені Ігоря Сікорського

ФОРМУВАННЯ СТРАТЕГІЇ УПРАВЛІННЯ РЕЖИМАМИ РОБОТИ СИСТЕМ ЗАХИСТУ НА ОСНОВІ МОДЕЛІ ІГРОВОГО УПРАВЛІННЯ

Основними сферами застосування теорії ігор є економіка, політологія, тактичні й воєнно-стратегічні задачі, еволюційна біологія і, останнім часом, інформаційні технології, безпека та штучний інтелект. Теорія ігор вивчає задачі прийняття рішень декількох осіб (гравців). Вона стосується поведінки гравців, рішення яких впливають один на одного. Теорія ігор призначена для вирішення ситуацій, в яких результат рішення гравців залежить не лише від того, як вони їх вибирають, а й від вибору рішень інших гравців, з якими вони взаємодіють. Якщо розглядати сферу інформаційної безпеки, то особливістю інформаційного конфлікту системи оперативного управління захистом інформації та порушника, який намагається здійснити несанкціонований доступ, є те, що протидіючі сторони, які мають декілька способів дій, можуть застосовувати їх багаторазово, вибираючи найкращий спосіб з урахуванням інформації про дії протилежної сторони. Причому кожен крок розв'язання конфлікту характеризується не фінальним станом, а деякою платіжною функцією. У багатьох ситуаціях під час проектування систем захисту інформації виникає необхідність розроблення та прийняття рішень в умовах невизначеності. Невизначеність може мати різний характер. Невизначеними є сплановані дії хакерів, які скеровані на зменшення ефективності систем захисту. Невизначеність може стосуватися ситуації ризику, в якій система управління інформаційної мережі, що приймає рішення щодо застосування системи захисту, здатна встановлювати не тільки всі можливі результати рішень, але й вірогідність можливих умов їхньої появи. Умови проектування впливають на прийняття рішень підсвідомо, незалежно від дій суб'єкта, що приймає рішення. Коли відомі всі наслідки можливих рішень, але невідомі їхня вірогідність, очевидно, що рішення приймають в умовах повної невизначеності. Основною перспективною теорією аналізу процесів прийняття рішень на етапі проектування систем захисту інформації є теорія ігор. Застосування теорії ігор в області моделювання процесів прийняття рішень у сфері інформаційної безпеки має різні підходи, які нині не систематизовані, а інколи і вступають у протиріччя між собою. Тому виникає необхідність розроблення методів оперативного (адаптивного) управління захистом інформації залежно від наявності апріорної інформації про можливість атак порушника й реалізовані ним стратегії створення несанкціонованого доступу до інформаційного ресурсу. Теорія ігор дозволяє запропонувати рекомендації із формування стратегії управління режимами роботи систем захисту.

Ключові слова: захист інформації; система безпеки; теорія ігор; оптимальна стратегія; система порушника; прийняття рішення.

1. ВСТУП

У дослідженнях конфліктів існує значна кількість наукових проблем, для розв'язання яких теорія ігор може бути екстремально корисною. Наприклад, важливим, але недостатньо дослідженим напрямком є аналіз дій супротивників у

кількох періодах, стратегічної поведінки учасників в умовах неповної інформації. Іншою сферою, яка практично не береться до уваги як у класичних, так і в сучасних дослідженнях та іграх – це моральні норми, цінності суспільства, моральна сторона кожної стратегії.

© Толупа С. В., Лукова-Чуйко Н. В., Наконечний В. С., Сайко В. Г., Кулінич О., 2020



У багатьох ситуаціях проектування систем захисту інформації виникає необхідність розроблення та прийняття рішень в умовах невизначеності. Невизначеність може мати різний характер. Останнім часом досить гостро постало питання протидії атакам в інформаційні системи [1, 2]. Так, невизначеними є сплановані дії хакерів, які скеровані на зменшення ефективності систем захисту. Невизначеність може стосуватися ситуації ризику, в якій система управління інформаційної мережі, що приймає рішення із застосування системи захисту, здатна встановлювати не тільки всі можливі результати рішень, але й вірогідність можливих умов їхньої появи. Умови проектування впливають на прийняття рішень підсвідомо, незалежно від дій суб'єкта, що приймає рішення. Коли відомі всі наслідки можливих рішень, але невідома їхня вірогідність, очевидно, що рішення приймають в умовах повної невизначеності. Основною перспективною теорією аналізу процесів прийняття рішень на етапі проектування систем захисту інформації є теорія ігор [3–6]. Застосування теорії ігор в області моделювання процесів прийняття рішень має різні підходи, які нині не систематизовані, а інколи і вступають у протиріччя між собою. Тому дослідження вказаної тематики є актуальними науковими завданнями.

2. ПОСТАНОВКА ПРОБЛЕМИ

Теорія ігор призначена для врегулювання ситуацій, в яких результат рішення гравців залежить не тільки від того, як вони їх вибирають, а й від вибору рішень інших гравців, з якими вони взаємодіють. Ігри залежать від кількості гравців; до них належать ігри з нульовою сумою (антагоністичні) і з ненульовою сумою. Множини стратегій можуть бути скінченними або нескінченними (матричні ігри й ігри на компактах відповідно). Також гравці можуть діяти незалежно один від одного або утворювати коаліції. Відповідними моделями є некооперативні та кооперативні ігри. Існують також ігри з повною або частковою інформацією [7, 8].

Якщо розглядати сферу інформаційної безпеки (ІБ), то особливістю інформаційного конфлікту системи оперативного управління захистом інформації та порушника, який намагається здійснити несанкціонований доступ (НСД) є те, що протидіючі сторони, які мають декілька способів дій, можуть застосовувати їх багаторазово, вибираючи найкращий спосіб з урахуванням інформації про дії протилежної сторони. Причому кожний крок вирішення конфлікту характеризується не фінальним станом, а деякою платіжною функцією. Традиційний ігровий підхід до аналізу дій порушника не дозволяє врахувати багатокро-

ковість конфлікту й не відображує залежність способів дій сторін від інформації про протилежну сторону, а відомий конфліктний підхід на основі розрахунку фінальних імовірностей перебування систем у стані виграшу до заданого моменту часу не відображує багатоваріантність дій сторін і неповну завершеність конфлікту на кожному кроці [9, 10].

Якщо розглядати конфліктність між системою оперативного управління захистом інформації (ЗІ) та порушником, то можна розглядати різні стратегії ведення ігор.

Так модель чистої стратегії ведення ігор – це пара стратегій (одна – для загрози, а друга – для механізму захисту), які перехресжуються в сідловій точці, яка в цьому випадку і визначає ціну гри.

Нехай загроза A обрала стратегію A_i , тоді у найгіршому разі вона отримає виграш, що дорівнює $\min_j a_{ij}$, тобто навіть тоді, якщо механізм захисту B знав би стратегію загрози A . Передбачаючи таку можливість, загроза A має вибрати таку стратегію, щоб максимізувати свій мінімальний виграш, тобто $\alpha = \max_i \min_j a_{ij}$.

Така стратегія загрози A позначається A_{i_0} і має назву максимінної, а величина гарантованого виграшу цього гравця називається нижньою ціною гри.

Механізм захисту B , який програє суми у розмірі елементів платіжної матриці, навпаки має вибрати стратегію, що мінімізує його максимально можливий програш за всіма варіантами дій загрози A . Стратегія механізму захисту B позначається через B_{j_0} і називається мінімаксною, а величина його програшу – верхньою ціною гри, тобто $\beta = \min_j \max_i a_{ij}$.

Оптимальний розв'язок цієї задачі досягається тоді, коли жодній стороні не вигідно змінювати вибрану стратегію, оскільки її противник може у відповідь вибрати іншу стратегію, яка забезпечить йому кращий результат.

Якщо $\max_i \min_j a_{ij} = \min_j \max_i a_{ij} = \nu$, тобто, якщо $\alpha = \beta = \nu$, то гра називається цілком визначеною. У такому разі виграш загрози A (програш механізму захисту B) називається значенням гри і дорівнює елементу матриці $a_{i_0 j_0}$.

Цілком визначені ігри називають іграми із сідловою точкою, а елемент платіжної матриці, значення якого дорівнює виграшу загрози A (програшу механізму захисту B) і є сідловою точкою. У цій ситуації оптимальним рішенням гри для обох сторін є вибір лише однієї з можливих, так званих чистих стратегій – максимінної



для загрози A та мінімаксною для механізму захисту B , тобто якщо один із гравців притримується оптимальної стратегії, то для другого відхилення від його оптимальної стратегії не може бути вигідним. Такий приклад можемо розглядати як частковий випадок аналізу захищеності системи захисту інформації.

Також широке застосування отримала модель змішаної стратегії ведення гри – модель стратегії ведення гри, яка полягає у тому, що гравець застосовує одну із своїх чистих стратегій, обрану в кожній грі за випадковим законом.

Змішану стратегію можна ототожнити з імовірнісною мірою на множині можливих для гравця дій, тобто його чистих стратегій.

Уведенням змішаної стратегії розширюють клас допустимих дій гравця для того, щоб допомогти існування розв'язків гри, що вимагається принципом здійснення мети [11, 12].

Імовірності (або частоти) вибору кожної стратегії задають відповідними векторами: для загрози A – вектор $X = (x_1, x_2, \dots, x_m)$, де $\sum_{i=1}^m x_i = 1$; для механізму захисту B – вектор $Y = (y_1, y_2, \dots, y_n)$, де $\sum_{j=1}^n y_j = 1$. Очевидно, що $x_i \geq 0 (i = \overline{1, m})$; $y_j \geq 0 (j = \overline{1, n})$.

Виявляється, що коли використовуються змішані стратегії, то для кожної скінченної гри можна знайти пару стійких оптимальних стратегій. Існування такого розв'язку визначає основна теорема теорії ігор.

Теорема (основна теорема теорії ігор). Кожна скінченна гра має, принаймні, один розв'язок, можливий в області змішаних стратегій.

Нехай маємо скінченну матричну гру з платіжною матрицею:

$$A = \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \dots & \dots & \dots & \dots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{pmatrix}.$$

Оптимальні змішані стратегії гравців A і B за теоремою визначають вектори $X^* = (x_1^*, x_2^*, \dots, x_m^*)$ і $Y^* = (y_1^*, y_2^*, \dots, y_n^*)$, що дають змогу отримати виграш: $\alpha \leq v \leq \beta$.

Використання оптимальної змішаної стратегії загрозою A має забезпечувати виграш на рівні, не меншому, ніж ціна гри за умови вибору механізмом захисту B будь-яких стратегій. Математично ця умова записується так:

$$\sum_{i=1}^m a_{ij} x_i^* \geq v \quad (j = \overline{1, n}).$$

З другого боку, використання оптимальної змішаної стратегії механізмом захисту B має забезпечувати за будь-яких стратегій загрози A програш, що не перевищує ціну гри v , тобто

$$\sum_{j=1}^n a_{ij} y_j^* \leq v \quad (i = \overline{1, m}).$$

Ці співвідношення використовують для знаходження розв'язку гри [13, 14].

Зауважимо, що в цьому разі розраховані оптимальні стратегії завжди є стійкими, тобто якщо один із гравців притримується своєї оптимальної змішаної стратегії, то його виграш залишається незмінним і дорівнює ціні гри v незалежно від того, яку з можливих змішаних стратегій вибрав інший гравець.

Оптимальне рішення гри є його виграшем. Тому існують моделі оптимальних змішаних стратегій гри, основою яких є змішані моделі стратегій. Оптимальна модель стратегії гри у змішаних моделях стратегій має такі властивості: кожен гравець не зацікавлений у відході від своєї оптимальної змішаної стратегії, якщо його противник застосовує оптимальну змішану стратегію так, як йому не вигідно. Чисті стратегії гравців у їхніх оптимальних змішаних стратегіях мають назву активних.

Застосування оптимальної змішаної стратегії забезпечує гравцю максимальний середній виграш (або мінімальний середній програш), який дорівнює ціні гри, незалежно від того, які дії застосовує інший гравець, якщо тільки він не виходить за межі своїх активних стратегій. Такі оптимальні рішення більш імовірні для часткових прикладів під час аналізу проєктів окремих складових системи ЗІ.

Тому виникає необхідність розроблення методів оперативного (адаптивного) управління захистом інформації залежно від наявності апріорної інформації про можливість атак порушника й реалізовані ним стратегії створення НСД [15, 16].

3. ВИКЛАД ОСНОВНОГО МАТЕРІАЛУ

Для характеристики поточного стану конфлікту будемо використовувати показник захищеності системи $a_{ij} = P_{зах}$ при реалізації в ній i -ї, $i \in I = \{1, 2, \dots, n\}$, стратегії (способу) захисту й застосуванні j -ї, $j \in J = \{1, 2, \dots, m\}$, стратегії (способу) створення контуру безпеки, m і n – кількість стратегій захисту і створення контурів безпеки, реалізованих у системі безпеки (СБ) і в системі порушника відповідно.

Стороною A назвемо підсистему оперативного управління ЗІ, стороною B – систему протидії цьому захисту, а величину a_{ij} – виграшем сторо-



ни A (програшем сторони B) у ситуації (i, j) . За традиційного ігрового підходу до аналізу систем безпеки передбачається, що сторонам відома матриця гри і скінченна множина стратегій порушника, але невідомо, яка стратегія реалізується в конкретній ситуації. У цьому випадку в матричній грі формалізується ситуація вибору стратегій захисту в умовах невизначеності. Однак такий підхід не відображує динаміку конфлікту, а також можливість цілеспрямованого вибору стратегій захисту на кожному кроці залежно від інформації про дії системи порушника. Тому пропонується для опису розглянутого конфлікту використати модель крокової матричної гри із запізненням і помилками в інформованості сторін про дії порушника (матрично-ігрового процесу). Позначимо: $T_{C3}(T_{CП})$ – час однократної реалізації стороною A (B) своєї чистої стратегії; $t_{C3}(t_{CП})$ – час реакції сторони A (B), який дорівнює інтервалу часу від моменту початку реалізації стратегії стороною B (A) до моменту початку реалізації відповідної стратегії стороною A (B).

Будемо вважати, що сторонам відома: матриця гри $A = (a_{ij})_m^n$, множини активних стратегій I, J і оцінки величин $T_{C3}(T_{CП})$ і $t_{C3}(t_{CП})$; матриця гри A є невід’язною і має рішення у вигляді цін гри v і векторів оптимальних змішаних стратегій сторони $A - P^* = (P_1^*, P_2^*, \dots, P_n^*, \dots, P_n^*)$ і сторони $B - Q^* = (Q_1^*, Q_2^*, \dots, Q_j^*, \dots, Q_m^*)$; протягом часу гри T відсутня післядія, а множини I і J є незмінними.

Методика призначена для адаптивної зміни параметрів і режимів роботи СБ за ігровим алгоритмом залежно від наявності апріорної інформації про параметри системи порушника і стратегії створення нею атак на ІС.

Сутність ігрового алгоритму управління полягає в порівнянні великої кількості можливих у цих умовах якісно різних рішень, визначенні оптимального або найкращого з урахуванням всіх обмежень рішення та формування відповідної команди.

Для підвищення ефективності у разі вирішення динамічних ігор використовують метод прогнозування.

Одним із можливих рішень ігор у змішаних стратегіях є, як зазначалося вище, збільшення швидкості реакції (зниження часу адаптації) однієї зі сторін, що дозволяє підвищити результативність використання стратегій.

Розповсюджений спосіб рішення матричної гри у змішаних стратегіях, наприклад, методами лінійного програмування, істотно ускладнюється для матриць великої розмірності.

Застосування декомпозиційних методів не завжди можливе, а ітеративні методи рішення,

такі, наприклад, як метод Брауна – Робінсона, мають найчастіше недостатньо високу швидкість збіжності. Як альтернативний може бути застосований метод динамічного програмування, що використовує результати короткочасного та довгострокового прогнозування [9].

Розглянемо алгоритм рішення матричних ігор, що використовує метод динамічного програмування. Стосовно до розглянутих випадків довгострокове прогнозування дозволяє з досить великим ступенем надійності обмежити кількість імовірних стратегій системи порушника й редукувати ігрову матрицю.

Рішенням матричної гри з урахуванням принципу прогнозування на основі марковського підходу є оптимізація умовної стратегії СБ на N циклів уперед по прогнозованій стратегії системи порушника. Очевидно, що зі зростанням N точність прогнозування знижується. У зв’язку із цим розглянемо випадок, коли $N = 1$. Можна виокремити три етапи алгоритму формування оптимальної стратегії СБ.

Методика управління системою безпеки на базі методів теорії ігор, блок-схема алгоритму реалізації, складається з таких етапів (рис. 1).

Уведення вихідних даних. Уводяться параметри засобів безпеки і каналу прийняття рішення $\Psi = \{\psi_i\}$, а також значення допустимої величини ймовірності помилкового прийняття рішення.

Отримання інформації про дії системи порушника. За допомогою одного з методів контролю стану системи безпеки визначає стратегію або розпізнає факт дії на ній системи порушника.

Визначення оптимальної стратегії СБ. Задача оптимізації алгоритмів функціонування СБ полягає у визначенні такої оптимальної стратегії $\hat{a}^* \in \hat{A}^*$, при якій забезпечується максимальна ефективність функціонування СБ протягом необхідного часу функціонування. Для підвищення ефективності під час вирішення динамічних ігор використовується метод прогнозування.

Одним із можливих рішень ігор у змішаних стратегіях є, як зазначено вище, збільшення швидкості реакції (зниження часу адаптації) однієї зі сторін, що дозволяє підвищити результативність використання стратегій.

Коефіцієнт адаптації СБ залежить від співвідношення $T_{CБ}/T_{CП}$, а значення $T_{CБ}, T_{CП}$ – від тривалостей часу регулювання і зміни режимів роботи засобів захисту, які залежать від їхнього стану в попередньому циклі. Тривалість переходу СБ зі стану H_n у стан H_m на етапах регулювання (H_n і H_m – вектори станів засобів захисту) задається заздалегідь відомою квадратною матрицею часу переходу з будь-якого можливого (узятого з області визначення) стану в будь-який можливий.

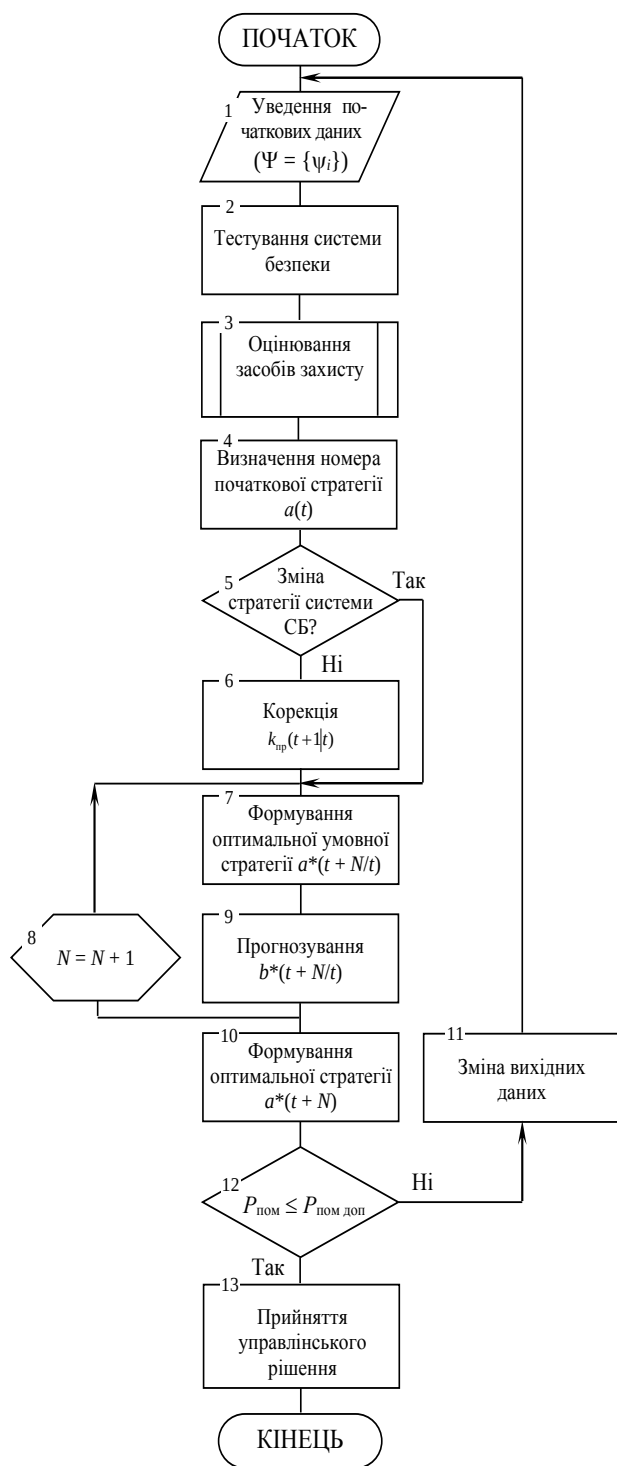


Рис. 1. Блок-схема алгоритму реалізації методики управління системою безпеки на основі моделі ігрового управління

Елементами матриці будуть T_{CB}^{reg} , що входять до складу T_{CB}^H на етапі регулювання параметрів і зміни режимів роботи системи. Тоді процес переходу з H_n в H_m з урахуванням можливих проміжних станів може бути описаний апаратом однорідних марковських ланцюгів із дискретними станами в дискретному часі. Порядок переходу від $H_n(t)$ до $H_m(t+1)$ задається відповідною

стратегією $a_i \in S_{CB}$. Аналогічний стан системи порушника при переході визначається як $H_{CPln}(t)$ і $H_{CPlm}(t+1)$.

Тому задача умовної оптимізації часу адаптації на етапі регулювання полягає у виборі такої стратегії a^* на циклі $(t+1)$, при якій

$$\begin{cases} T_{CB}((t+1), a^*) = \min_{a_i \in S_{CB}} T_{CB}(H_n(t), H_m(t+1), a_i); \\ T_{CPl}((t+2), a^*) = \max_{a_i \in S_{CPl}} T_{CPl}(H_{CPln}(t+1), H_{CPlm}(t+2), a_i) \end{cases} \quad (1)$$

за умови $T_{CB} < T_{CPl}$, що й відбувається у процесі рішення гри за рахунок уведення k_a під час розрахунку елементів матриці.

Чисельно точність передбачуваного значення функції виграшу задається деяким коефіцієнтом помилки прогнозування:

$$|R_{(per)nm}| N \times M, n = \overline{1, N}, m = \overline{1, M}.$$

$$k_{пр}(t+1|t) = \frac{\Phi(t+1)}{\Phi_{RL}(t+1)},$$

де $\Phi_{RL}(t+1)$ обчислено в разі досягнення $(t+1)$ у результаті моніторингу. Так, у випадку незмінної протягом ряду циклів стратегії системи СП при $\Phi(t)$, відбувається корекція $\Phi(t+1)$ за рахунок $k_{пр}(t+1|t)$, що входить до складу коефіцієнта $\beta_m(t+1)$. Це усуває систематичну помилку в розрахунках значень $\Phi(t)$ і деякою мірою впливає на вибір $a^*(t+1)$ у ході рішення матричної гри.

Оскільки коефіцієнт помилки прогнозування перебуває у зворотній залежності від коефіцієнта інформованості k_{inf}^a , функція $f(k_{inf}^a) = |k_{пр}(t+1|t) - 1|$. У цьому випадку має місце постановка наступної задачі умовної оптимізації:

$$k_{inf}^a \rightarrow \max,$$

де $\max k_{inf}^a = k_{inf}^S$ при обмеженні

$$|k_{пр}(t+1|t) - 1| \leq \delta_{пом},$$

де $\delta_{пом}$ – деяка центрована випадкова величина з нульовим математичним очікуванням і дисперсією δ^2 , яка визначає деяке граничне значення помилки.

4. РЕЗУЛЬТАТИ

Розглянемо алгоритм рішення матричних ігор, що використовує метод динамічного програмування. Стосовно до розглянутих випадків довгострокове прогнозування дозволяє з досить великим ступенем надійності обмежити кількість імовірних стратегій системи порушника в наступних циклах управління до 2...4 і редукувати ігрову матрицю. Рішенням матричної гри з урахуванням принципу прогнозування на основі марковського підходу є оптимізація умовної стратегії СБ на N циклів уперед по прогнозований



стратегії системи порушника. Очевидно, що зі зростанням N точність прогнозування знижується. У зв'язку із цим розглянемо випадок, коли $N = 1$. Можна виокремити три етапи алгоритму формування оптимальної стратегії СБ.

На першому етапі на підставі інформації про поточний стан засобів захисту, передбачуваного значення перехідних ймовірностей СБ, яка застосовує на циклі управління t стратегії системи порушника $b(t)$, а також з урахуванням попередніх стратегій СП формується оптимальна умовна стратегія $a^*(t+1|t)$:

$$a^*(t+1|t) = \arg \left[\max_{a \in S_{CS}} P(a(t), b(t), H(t)) \right]. \quad (2)$$

На другому етапі розв'язується завдання прогнозування стратегії, яка використовується на циклі управління $t+1$ системою порушника і яка забезпечить мінімізацію функціонала

$$b^*(t+1|t) = \arg \left[\min_{b \in S_{PEH}} P(a^*(t+1), b(t+1)) \right]. \quad (3)$$

На третьому етапі формується оптимальна стратегія управління СБ з урахуванням прогнозованої стратегії системи порушника і поточного стану засобів захисту:

$$a^*(t+1|t) = \arg \left[\begin{array}{l} \max_{a \in S_{SS}} P(a(t+1), \\ b^*(t+1|t), \\ H(t+1), \\ \beta_m(t+1), \\ k_a(H(t+1)|H(t))) \end{array} \right]. \quad (4)$$

Для підвищення надійності результату алгоритм може повторюватися обмежену кількість разів за наявності певної дисперсії розподілу ймовірностей застосування стратегії $a_1^*(t+1|t) \dots$

$a_n^*(t+1|t)$ з їх наступним оцінюванням на основі критеріїв переваги, що вводяться. У випадку неможливості визначення такої стратегії $a^*(t+1)$, за якої втрати не перевищують припустимого значення, розв'язується завдання розширення множини допустимих стратегій СБ, після чого знову визначається $a^*(t+1)$. Аналогічно формують стратегії СБ шляхом оптимізації умовної стратегії управління СБ по прогнозованій на N кроків стратегії системи порушника. Третій етап алгоритму в цьому випадку матиме вигляд

$$\begin{aligned} a^{*(N)}(t+N) &= \arg \left[\max_{a \in S_{SS}} P(a(t+N), \right. \\ &\quad b^*(t+N|t+N-1), \\ &\quad H(t+N), \\ &\quad \beta_m(t+N), \\ &\quad \left. k_a(H(t+N)|H(t+N-1)) \right], \quad N = 2, 3, \dots \end{aligned} \quad (5)$$

Апарат марковських кіл використовується на другому і третьому етапах, що дозволяє розраховувати ймовірності застосування тієї або іншої стратегії на черговому циклі управління і вибору оптимальної стратегії.

Припустимо, що в циклі управління t використовується стратегія системи порушника: b_2 . Нехай, за критерієм $\max_{a_1} P_{i,2}(t) = P_{3,2}$ вибирається стратегія a_3 (табл. 1).

Таблиця 1

Алгоритм прогнозованого переходу станів СБ

t	Процес прогнозування				$t+1$
1-й етап	2-й етап		3-й етап		
Рішення на циклі	Ймовірність переходу $P_{3j} = 1 - k_N \Phi_{3j}$	$\min_{b_j} (1 - P_{3j}(t+1))$	Ймовірність переходу $P_{i4} = k_N \Phi_{i4}$	$\max_{a_1} P_{i2}(t+1)$	
a_3	$P_{31} \Rightarrow$	b_1	b_4	$P_{14} \Rightarrow$	a_1
	$P_{32} \Rightarrow$	b_2		$P_{24} \Rightarrow$	a_2
	$P_{33} \Rightarrow$	b_3		$P_{34} \Rightarrow$	a_3
	$P_{34} \Rightarrow$	b_4		$P_{44} \Rightarrow$	a_4
	$\dots$	$\dots$		$\dots$	$\dots$
$a_3; b_2$	$P_{3j} \Rightarrow$	a_j		$P_{i4} \Rightarrow$	a_{i4}
					$a_2; b_4$

Одночасно із цим реалізується алгоритм прогнозування. У табл. 1 представлено спрощений приклад прогнозованого переходу системи зі стану в циклі t у стан $(t+1)$ на основі неоднорідних марковських кіл. Відповідно до принципу оптималь-

ності, у разі пошуку оптимального рішення в багатокроковому завданні оптимізації вибір стратегії управління $a(t)$ на кожному кроці незалежно від початкового стану має бути спрямований на оптимізацію не тільки даного, але й усіх наступних



кроків. З урахуванням прогнозування на $(t + N)$ кроків уперед (у цьому випадку, не більше трьох кроків) механізм вибору оптимальної стратегії $a^*(t)$ у циклі t буде також визначатися обчисленням зворотної функції Беллмана останніх прогнозованих $N - t + 1$ циклів управління. Так, для $t = N$:

$$B_N(H(N-1)) = \max_{a(N) \in \Lambda_N(H(N-1))} P_N(H(N-1), a(N)),$$

де $H(N-1)$ – стан СБ на $(N-1)$ -му циклі управління; $a(N)$ – стратегія управління на циклі N ; $\Lambda_N(H(N-1))$ – скінченна множина допустимих стратегій на циклі $(N-1)$.

Метод Беллмана використовують для підвищення точності прогнозу, обґрунтованості вибору поточних стратегій і підтримки прийняття рішень пристроєм управління системи безпеки.

Основні проблеми за використання теорії ігор виникають під час визначення функції виграшу для конкретної ситуації. Для завдань, які розв'язує СБ-функція виграшу, у першу чергу, повинна відображати зміну якості системи безпеки.

Розглянемо приклад, що відображає реальну ситуацію, що описана матрицею виграшу 2×2 . Розмірність указує на кількість чистих стратегій у гравців. Як гравці виступають СБ і СП. Вибір цих прикладів обумовлений такими міркуваннями. Гра 2×2 дозволяє яскраво продемонструвати можливості й обмеження теорії ігор.

СБ і СП можуть вибрати різні змішані стратегії: для СБ маємо $S_{CB}^* = (P_1, 1-P_1)$, для СП – $S_{СП}^* = (P_2, 1-P_2)$. У результаті рішення гри мають бути отримані оптимальні стратегії гравців і значення гри, тобто трійка $(S_{CB}^*, S_{СП}^*, V')$. У нашому випадку потрібно визначати P_1' і P_2' . Значенням гри буде математичне очікування ймовірності правильного прийняття рішення, обчислене з урахуванням усіх можливих ситуацій.

У табл. 2 і на рис. 2 наведено результати вирішення гри.

Оптимальними стратегіями з позиції теорії ігор є: $P_1' = 0,466$ і $P_2' = 0,534$. Отже, СБ повинна вибирати перший режим роботи з ймовірністю 0,466, а другий з ймовірністю $1 - 0,466 = 0,534$. Система СП має вибирати свої режими з ймовірностями 0,534 і 0,466. Видно, що, вибравши оптимальну стратегію, СБ гарантовано забезпечить собі математичне очікування виграшу 0,466 незалежно від того, як діятиме супротивник.

Виграш невеликий, але зрозуміло, що математичне очікування величини не може бути більше 0,5, якщо в половині випадків вона набуває нульових значень. Причому, якщо СП дотримується оптимальної стратегії, то підвищити матема-

тичне очікування ймовірності правильного прийняття рішення при забезпеченні контуру безпеки не вдасться.

Якщо вказана ситуація не влаштовує СБ, то слід вживати заходів щодо підвищення виграшу за певних поєднань методів і засобів безпеки. Проте видно, що якщо СП відхиляється від своєї оптимальної стратегії, то СБ має можливість підвищити свій виграш, також відхилившись від оптимальної стратегії.

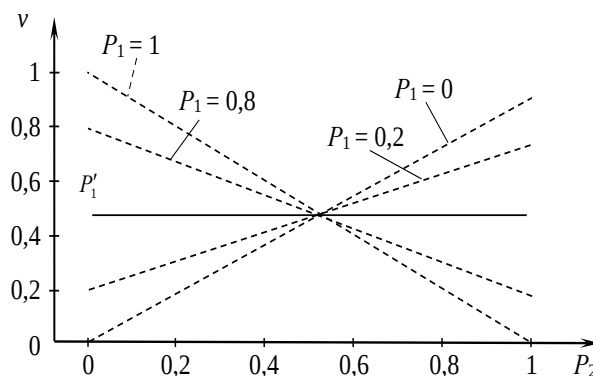


Рис. 2. Залежності виграшу від вибору стратегій СБ і СП

Таблиця 2
Виграш у разі різних стратегій СБ і систем порушника

$P_1 \backslash P_2$	0	0,2	0,4	0,6	0,8	1	$P_1' = 0,534$
0	0	0,175	0,349	0,524	0,699	0,874	0,466
0,2	0,200	0,300	0,399	0,499	0,599	0,699	0,466
0,4	0,400	0,425	0,450	0,474	0,499	0,524	0,466
0,6	0,600	0,550	0,500	0,450	0,399	0,349	0,466
0,8	0,800	0,675	0,550	0,425	0,300	0,175	0,466
1	0,999	0,800	0,600	0,400	0,200	0	0,466
$P_2' = 0,466$	0,466	0,466	0,466	0,466	0,466	0,466	0,466

5. ВИСНОВКИ

Таким чином, у процесі умовної оптимізації з урахуванням поточної ігрової матриці будуть формуватися умовно оптимальні стратегії, що визначають фазову траєкторію СБ, починаючи із заключного циклу прогнозування $t = N$ до поточного значення t .

Основні проблеми під час використання теорії ігор виникають при визначенні функції виграшу для конкретної ситуації. Для завдань, які розв'язує система безпеки, функція виграшу, у першу чергу, має відображати зміну якості системи безпеки.

У разі, якщо вказана ситуація не влаштовує СБ, слід вживати заходів щодо підвищення виграшу за певних поєднань методів і засобів захисту.

Якщо злоумисник відхиляється від своєї оптимальної стратегії, то СБ має можливість підвищити свій виграш, також відхилившись від оптимальної стратегії.



Результати імітаційного моделювання процесу функціонування СБ за запропонованим ігровим алгоритмом показав, що додаткове застосування прогнозування стратегії на N циклів уперед дозволяє підвищити ефективність функціонування системи на 5–8 %.

Отже, теорія ігор дозволяє запропонувати рекомендації з формування стратегії управління режимами роботи систем захисту. Причому, принаймні, для певних типів конфліктів і матриць виграшів ці рекомендації дозволяють СБ отримати виграш і досягнути поліпшення своїх технічних характеристик.

Аналіз виграшу, який отримує СБ в різних ситуаціях, показав, що теорія ігор не тільки дозволяє сформувати оптимальну стратегію, що забезпечує гарантії певного виграшу, але й дає змогу видати рекомендації з її зміни з метою збільшення виграшу, якщо система порушника відступає від своєї оптимальної стратегії. Коли система порушника слідує своїй оптимальній стратегії, теорія ігор дозволяє оцінити цю ситуацію. Якщо результати оцінювання не влаштовують, то необхідно вживати заходів зі зміни ситуації.

Подальші напрямки досліджень на основі теорії ігор. Застосування ігрового підходу до розв'язання проблем безпеки (що включає моделювання, постановку проблеми і її розв'язання) усе ще сильно залежить від вибраної схеми ігрової взаємодії користувачів, яка, як правило, є спрощенням реального світу. Наприклад, якщо розглядаються ігри між одним нападником і системою захисту, то використовують ігри двох учасників. Якщо ситуація розгортається у часі, використовують динамічну постановку. Важливим є також припущення про інформованість учасників, залежно від чого використовуються ігри з досконалою або недосконалою та повною або неповною інформованістю.

Найбільш перспективні можливі напрямки досліджень, у яких буде розвиватись теоретико-ігровий підхід у області кібербезпеки, – це хмарні технології. Сучасна хмарна система існує в середовищі постійних змін. Змінюються технології, протоколи та програмне забезпечення. Змінюються користувачі, їхні пріоритети, задачі й поведінка. Усі ці зміни є непередбачуваними. Тому теорія ігор, яка вже має історію успішного застосування до розв'язання задач маршрутизації, планування, керування потоками даних і перевантаженнями, є головним напрямком аналітичного моделювання хмарних систем. Застосування теорії ігор використовують у технологіях інтернет-речей, що об'єднує в одну мережу всі розумні прилади, що вже в недалекому майбутньому забезпечуватимуть людське життя. Проникнення розумних речей у

наше життя стає тотальним, а тому і загрози, які при цьому виникають, також є тотальними. Слід виокремити технологію блокчейн, на основі якої будуються криптовалюти і деякі нові сервіси. Однією з ключових проблем є забезпечення ефективного та безпечного "майнінгу" (тобто обчислення підтвердження нових операцій або одиниць валюти). Теорія ігор моделює процес майнінгу як взаємодію автономних агентів, що конкурують за ресурси. Враховуючи сучасну вартість криптовалюти біткоїн, атака на механізми обчислення може бути надзвичайно прибутковою. Таким чином можна сказати, що застосування теоретико-ігрового підходу до проблем безпеки, хоч і продовжується два десятиліття, усе ще є новим підходом із багатьма нерозв'язаними проблемами. Складність ігор із багатьма учасниками за умов конфлікту і невизначеності стимулює дослідників до створення нових моделей і методів, які допоможуть створити новий безпечний і ефективний кіберпростір.

СПИСОК ВИКОРИСТАНОЇ ЛІТЕРАТУРИ

- [1] Fei He, Jun Zhuang, and United States. Game-theoretic analysis of attack and defense in cyber-physical network infrastructures. In Proceedings of the Industrial and Systems Engineering Research Conference. 2012.
- [2] Johnson, Benjamin, et al. "Game-theoretic analysis of DDoS attacks against Bitcoin mining pools." International Conference on Financial Cryptography and Data Security. Springer, Berlin, Heidelberg, 2014.
- [3] Бурячок В. Теорія ігор, як метод управління інформаційною безпекою / Володимир Бурячок, Анатолій Шиян // Правове, нормативне та метрологічне забезпечення системи захисту інформації в Україні : науково-технічний збірник. – 2013. – Вип. 2(26). – С. 21–28.
- [4] Петросян Л.А. Теория игр / Л.А. Петросян, Н.А. Зенкевич, Е.А. Семина. – М.: 1998. – Вища школа. – 304 с.
- [5] Дюбин Г.Н. Введение в прикладную теорию игр / Г.Н. Дюбин, В.Г. Суздаль. – М.: 1981. – Наука. – 336 с.
- [6] Воробьев Н.Н. Бесконечные антагонистические игры / под ред. Н.Н. Воробьева. – М.: 1993. – Вища школа. – 505 с.
- [7] Гришук Р.В. Теоретичні основи моделювання процесів нападу на інформацію методами теорії диференціальних ігор та диференціальних перетворень / Р.В. Гришук. – Монографія. – Житомир. – 2010. – 280 с.
- [8] Дослідження операцій. Ч. 3. Ухвалення рішень і теорія ігор / М. Я. Бартіш, І. М. Дудзяний. – Львів: Видавничий центр Львівського національного університету ім. І. Франка, 2009. – 277 с. : іл. – Бібліогр.: с.271–272 (36 назв). – ISBN 966-613-496-9
- [9] Baranovska L. V. Mixed strategy Nash equilibrium in one game and rationality / L. V. Baranovska, O. M. Bukovskiy // International Scientific and Practical Conference "WORLD SCIENCE". Proceedings of the III International Scientific and Practical Conference "Scientific Issues of the Modernity" (April 27, 2017, Dubai, UAE). – 2017. – No 5(21), Vol. 1, May. – Pp. 4–8.
- [10] Alpcan T., Başar T. Network security: A decision and game-theoretic approach. Cambridge University Press, 2010.
- [11] Толіпа С.В., Павлов І.М. Аналіз підходів моделювання процесів прийняття рішень при проектуванні систем захисту інформації. // Науково-технічний журнал "Сучасний захист інформації". – 2014. – №2. – С. 96–104.



[12] Толіюпа С.В., Павлов І.М. Аналіз підходів оцінки ефективності математичних моделей при проектуванні систем захисту інформації. // Науково-технічний журнал "Сучасний захист інформації". – 2014. – №3. – С. 36–44.

[13] Alpcan T., Başar T. A game theoretic approach to decision and analysis in network intrusion detection. Decision and Control, 2003. Proceedings. 42nd IEEE Conference on. Vol.

[14] Roy, Sankardas, et al. A survey of game theory as applied to network security. System Sciences (HICSS), 2010 43rd Hawaii International Conference on. IEEE, 2010.

[15] Liang, Xiannuan, and Yang Xiao. Game theory for network security. IEEE Communications Surveys & Tutorials 15.1. 2013. P. 472–486.

[16] Do, Cuong T., et al. Game Theory for Cyber Security and Privacy. ACM Computing Surveys (CSUR) 50.2. 2017. 30 p.

Стаття надійшла до редколегії

14.10.2020

Formation of a strategy for managing the modes of operation of protection systems based on the game management model

The main areas of application of game theory are economics, political science, tactical and military-strategic tasks, evolutionary biology and, more recently, information technology, security and artificial intelligence. Game theory studies the problems of decision-making of several people (players). It concerns the behavior of players whose decisions affect each other. The application of game theory in the field of modeling decision-making processes has different approaches, which are not systematized in the future, and sometimes contradict each other. Game theory is designed to solve situations in which the outcome of players' decisions depends not only on how they choose them, but also on the choices of other players with whom they interact. If we consider the field of information security, the peculiarity of the information conflict between the operational management system of information protection and the infringer who tries to gain unauthorized access is that opposing parties who have several ways of action can apply them repeatedly, choosing the best way based on information about the opposite parties. In this case, each step of resolving the conflict is characterized not by the final state, but by some payment function. In many situations, when designing information security systems, there is a need to develop and make decisions in conditions of uncertainty. Uncertainty can be of different nature. The planned actions of hackers, which are aimed at reducing the effectiveness of security systems, are uncertain. Uncertainty may relate to a risk situation in which the management system of the information network that decides on the application of the protection system is able to establish not only all possible results of decisions, but also the probability of possible conditions for their occurrence. Design conditions affect decision-making subconsciously, regardless of the actions of the decision-maker. When all the consequences of possible decisions are known, but their probability is unknown, it is obvious that decisions are made in conditions of complete uncertainty. The main promising theory of analysis of decision-making processes at the stage of designing information security systems is game theory. Therefore, there is a need to develop methods of operational (adaptive) management of information protection, depending on the availability of a priori information about the possibility of attacks by the infringer and his strategy to create unauthorized access to information resources. Game theory allows us to offer recommendations for the formation of management strategies for protection systems.

Keywords: information protection; security system; game theory; optimal strategy; system of the violator; making a decision.



Сергій Толіюпа,

доктор технічних наук, професор, професор кафедри кібербезпеки та захисту інформації Київського національного університету імені Тараса Шевченка.

Serhii Toliupa,

Doctor of Technical Sciences, Professor, Professor of the Department of Cybersecurity and Information Protection, Taras Shevchenko National University of Kyiv.



Наталія Лукова-Чуйко,

доктор технічних наук, професор завідувач кафедри кібербезпеки та захисту інформації Київського національного університету імені Тараса Шевченка.

Natalia Lukova-Chuiko,

Doctor of Technical Sciences, Professor, Head of the Department of Cybersecurity and Information Protection of the Kyiv National Taras Shevchenko University of Kyiv.



Володимир Наконечний,
доктор технічних наук, старший науковий співробітник, професор кафедри кібербезпеки та захисту інформації Київського національного університету імені Тараса Шевченка.

Volodymyr Nakonechnyi,
Doctor of Technical Sciences, Senior Research Fellow, Professor of the Department of Cybersecurity and Information Protection, Taras Shevchenko National University of Kyiv.



Володимир Сайко,
доктор технічних наук, професор, професор кафедри прикладних інформаційних систем Київського національного університету імені Тараса Шевченка.

Volodymyr Saiko,
Doctor of Technical Sciences, Professor, Professor of the Department of Applied Information Systems, Taras Shevchenko National University of Kyiv.



Олег Кулінич,
кандидат технічних наук, доцент, доцент спеціальної кафедри Інституту спеціального зв'язку та захисту інформації Національного технічного університету України "Київський політехнічний інститут імені Ігоря Сікорського".

Oleg Kulynich,
Candidate of Technical Sciences, Associate Professor, Associate Professor of the Special Department of the Institute of Special Communications and Information Protection of the National Technical University of Ukraine "Kyiv Polytechnic Institute named after Igor Sikorsky".



О. І. Махович, orcid.org/0000-0002-4684-9881,

makhovych.oleksandr@univ.kiev.ua

Р. В. Миколайчук, orcid.org/0000-0001-5349-4487,

mykroman@ukr.net

Київський національний університет імені Тараса Шевченка, Київ, Україна

В. Г. Миколайчук, orcid.org/0000-0002-2532-5771,

mukwira@ukr.net

Державний університет телекомунікацій, Київ, Україна

РЕКОМЕНДАЦІЇ ЩОДО ВИБОРУ СПОСОБУ БЕЗПЕЧНОГО ЗБЕРІГАННЯ ПАРОЛІВ

Використання паролів залишається найпоширенішим способом автентифікації користувачів для різного роду інформаційних систем. У зв'язку із цим виникає задача забезпечення безпеки зберігання інформації, що стосується даних автентифікації користувачів, та її захисту від несанкціонованого доступу. На практиці набули широкого розповсюдження різноманітні алгоритми безпечного зберігання паролів. Взаємосуперечні вимоги до таких алгоритмів безпечного зберігання паролів, які з одного боку мають бути достатньо складними для протидії різноманітним атакам, а з іншого – простими для забезпечення швидкодії інформаційної системи – перебору, особливо враховуючи те, що обчислювальна потужність центральних і графічних процесорів постійно зростає. Тому виникає потреба мати можливість змінювати складність обчислення хеш-коду, а отже й обсяг обчислень і час так, щоб значно ускладнити здійснення атаки, але не спричиняти дискомфорту кінцевому користувачу через затримку перевірки достовірності пароля. Серед відомих способів безпечного зберігання паролів розглянуто шифрування паролів, використання хеш-функцій у класичному варіанті, а також із додаванням солі та застосуванням ітерацій для обчислення хеш-коду. У роботі проведено порівняльний аналіз наведених способів, встановлено їхні переваги й недоліки, окреслено доцільні галузі застосування кожного способу, розроблено відповідні рекомендації. Для проведення обчислювального експерименту використовувалися засоби платформи Microsoft .NET Core 3.1, що дало змогу встановити часові показники роботи алгоритму отримання хеш-коду залежно від установлених параметрів алгоритму. Отримані за результатами експерименту дані можуть бути використані для вибору способу безпечного зберігання паролів.

Ключові слова: захист інформації; автентифікація користувача; шифрування; алгоритми обчислення хеш-коду; обчислювальний експеримент.

1. ВСТУП

Використання паролів залишається найпоширенішим способом автентифікації користувачів. Для забезпечення безпеки зберігання інформації та її захисту від несанкціонованого доступу набули широкого розповсюдження алгоритми безпечного зберігання паролів [1–3]. Водночас для уникнення можливості компрометації інформаційної системи під впливом зовнішніх загроз, необхідно підвищувати складність зазначеного алгоритму. З іншого боку, надмірне ускладнення процедури автентифікації призводить до зниження швидкості роботи системи. У сукупності наведені взаємосуперечні вимоги до алгоритмів безпечного зберігання паролів складають актуальну наукову проблему, одному з елементів якої,

а саме дослідженню способів безпечного зберігання паролів, присвячена ця стаття. Метою статті є обґрунтування рекомендацій щодо вибору способу безпечного зберігання паролів.

2. СПОСОБИ ЗБЕРІГАННЯ ПАРОЛІВ

Як зазначено раніше, важливим аспектом застосування парольної автентифікації є безпечне зберігання паролів. Для забезпечення зазначеної функції інформаційної системи використовують способи, що будуть розглянуті далі.

2.1. ШИФРУВАННЯ ПАРОЛІВ

Шифрування є звичним засобом забезпечення конфіденційності інформації. Цей процес передбачає застосування двох процесів:



зашифровування та розшифровування даних із використанням секретного ключа. Шифрування може бути застосоване для зберігання паролів, але при цьому можуть виникнути певні труднощі й задачі:

- організація керування ключами;
- забезпечення надійного зберігання секретного ключа;
- визначення обмеження доступу клієнтського додатку до секретного ключа для шифрування пароля перед його відправкою на сервер;
- організація безпечної передачі ключів.

Таблиця 1
Хеш-функції, доступні в .NET Core

Алгоритм	Розмір хешу (байт)	Опис
MD5 (<i>Message Digest Algorithm</i>)	16	Не рекомендується для нових проєктів, оскільки нестійкий до колізій, але швидкий у роботі
SHA1 (<i>Secure Hashing Algorithm</i>)	20	Використання в інтернеті не рекомендується з 2011 р.
SHA256 SHA384 SHA512 (<i>Secure Hashing Algorithm</i>)	32 48 64	Алгоритми із сім'ї "Безпечний алгоритм хешування другого покоління" із різними розмірами хеш-значень

Відсутність необхідності використовувати відкритий пароль та складність управління ключами шифрування ставить під сумнів оптимальність застосування шифрування для організації безпечного зберігання паролів та може породити загрози інформаційній безпеці. У випадку компрометації секретного ключа та таблиці паролів, зловмисник може отримати повний доступ до всієї інформації.

2.2. ХЕШУВАННЯ ПАРОЛІВ

Криптографічна хеш-функція H є однобічним перетворенням бітового рядка M довільної довжини на бітовий рядок (блок) фіксованої довжини h [1]. Вона має такі властивості [2]:

- 1) хеш-код повинен легко обчислюватися для довільного вхідного повідомлення;
- 2) результат роботи хеш-функції має залежати від усіх двійкових символів вихідного повідомлення, а також від їхнього взаємного розташування (незначна зміна вхідного повідомлення має призводити до тотальної зміни значення хеш-функції – так званий "лавинний ефект"). Значення хеш-коду не повинно давати витоку інформації навіть про окремі біти аргументу;

3) хеш-функція має бути стійкою до відновлення прообразу (односторонність) – відсутність практичної можливості за прийнятний час відтворити повідомлення M за його хеш-кодом h так, щоб $H(M) = h$;

4) хеш-функція має бути стійкою до виникнення колізій:

- *стійкість до колізій першого роду*: для заданого повідомлення M має бути неможливо за допомогою обчислень за прийнятний час підібрати інше повідомлення N , для якого $H(N) = H(M)$;

- *стійкість до колізій другого роду*: має бути неможливо за допомогою обчислень підібрати пару повідомлень (M, N) і $M \neq N$ таких, що мають однаковий хеш $H(M) = H(N)$.

На платформі Microsoft .NET Core є доступними для використання [3] реалізації декількох алгоритмів хешування (табл. 1). Обираючи алгоритм хешування, потрібно враховувати два важливих фактори:

- 1) стійкість до колізій;
- 2) стійкість до знаходження прообразу.

Використання хеш-функцій є поширеним для підтримки цілісності даних, проте їхнє застосування для зберігання паролів може породжувати певні загрози інформаційній безпеці.

Оскільки хеш-функція для однакових вхідних даних повертає ідентичний хеш-код, то зловмисник може використати словник найбільш уживаніших паролів для атаки або генерувати їх програмно. При цьому для кожного ймовірного пароля послідовно обчислюється хеш-код і порівнюється з атакованим хеш-кодом.

Такий підхід вимагає значних обчислювальних ресурсів, тому інший тип атаки передбачає використання величезних масивів наперед обчислених хеш-кодів (так званих *райдужних таблиць*) для різноманітних паролів. Пошук аргументу для хеш-функції в цьому випадку може відбуватися надзвичайно швидко [4].

2.3. ДОДАВАННЯ "СОЛІ"

Атаки на хеш можливі лише тому, що обчислення хеш-кодів здійснюється кожного разу ідентичним способом. Якщо для кожного пароля спосіб обчислення хешу модифікувати випадковим чином, то можна уникнути вразливості до атак із використанням райдужних таблиць. У цьому випадку доцільно додавати до пароля додаткову ентропію у вигляді криптографічно стійкої послідовності випадкових значень, яку називають *сіллю* [5].

Сіль додається до пароля перед обчисленням хеш-коду. Тому навіть якщо двоє різних користувачів використовують однаковий пароль, хеш-коди у



них будуть різні. У випадку, коли сіль для кожного пароля генерується окремо й ніколи не використовується повторно, зломисник не зможе наперед обчислити райдужні таблиці та скористатися ними для атаки. Оскільки сіль потрібно використовувати при обчисленні хешу для введеного пароля користувача, то її зберігають, як правило, у базі даних поряд з хеш-кодом, або як частину хеш-коду.

2.4. СТАНДАРТ PBKDF2

Використання солі при обчисленні хеш-коду пароля дозволяє уникнути вразливості з використанням райдужних таблиць. Але загроза прямого перебору залишається, особливо враховуючи те, що обчислювальна потужність центральних і графічних процесорів постійно зростає. Тому виникає потреба мати можливість змінювати складність обчислення хеш-коду, а отже й обсяг обчислень та час так, щоб значно ускладнити здійснення атаки, але не спричиняти дискомфорту кінцевому користувачу через затримку перевірки достовірності пароля.

Таке ускладнення обчислень досягається використанням спеціальних повільних алгоритмів обчислення хеш-коду, наприклад, PBKDF2 (Password-Based Key Derivation Function) [6]. PBKDF2 генерує похідний ключ, базуючись на основному ключі та додаткових параметрах. У ролі головного ключа виступає пароль, а додатковими параметрами є сіль та кількість ітерацій обчислення хеш-коду функцією, що лежить в основі алгоритму:

$$DK = \text{PBKDF2}(P, S, c, \text{dkLen}),$$

де P – пароль, S – "сіль", c – кількість ітерацій, dkLen – довжина похідного ключа, DK – похідний ключ.

Для дослідження залежності часу, необхідно для обчислення хеш-коду пароля, від кількості ітерацій проведено серію обчислювальних експериментів.

Для цього створено консольний додаток на платформі Microsoft .NET Core 3.1, у якому для кожного базового алгоритму хешування та кожної кількості ітерацій проводилося 10 послідовних обчислень зі знаходження хеш-коду пароля. При цьому фіксувався час у мілісекундах, необхідний для виконання обчислень. Для фіксування часу обчислень використовувалися функції `Start()`, `Stop()` і властивість `ElapsedMilliseconds` класу `Stopwatch` із простору імен `System.Diagnostics`. Як результат обирали середнє арифметичне отриманих значень затраченого часу. Експерименти повторювали для кількості ітерацій від 10000 до 1280000, де кожне наступне значення обиралося вдвічі більше за попереднє.

Для генерації солі використовувався клас `RNGCryptoServiceProvider` із простору імен `System.Security.Cryptography`. Довжина солі становила 32 байти.

Обчислення хеш-коду за алгоритмом PBKDF2 виконувалося за допомогою функції `GetBytes()` класу `Rfc2898DeriveBytes` простору імен `System.Security.Cryptography`.

Для базових алгоритмів SHA1, SHA256, SHA384 та SHA512 були обрані ефективні довжини згенерованого хеш-коду 20, 32, 48 та 64 байти відповідно (табл. 2).

Таблиця 2

Результати обчислювальних експериментів

Базова функція хешування	Довжина хеш-коду (байт)	Кількість ітерацій	Середнє значення витраченого часу (мс)
SHA1	20	1,00E+04	22,20
		2,00E+04	25,40
		4,00E+04	38,40
		8,00E+04	77,50
		1,60E+05	132,00
		3,20E+05	244,70
		6,40E+05	473,40
		1,28E+06	793,80
SHA256	32	1,00E+04	27,10
		2,00E+04	27,80
		4,00E+04	60,40
		8,00E+04	144,50
		1,60E+05	246,10
		3,20E+05	504,90
		6,40E+05	854,30
		1,28E+06	1262,80
SHA384	48	1,00E+04	35,50
		2,00E+04	42,00
		4,00E+04	76,90
		8,00E+04	155,20
		1,60E+05	276,60
		3,20E+05	544,10
		6,40E+05	933,50
		1,28E+06	1497,30
SHA512	64	1,00E+04	28,70
		2,00E+04	29,90
		4,00E+04	80,90
		8,00E+04	157,70
		1,60E+05	289,70
		3,20E+05	556,50
		6,40E+05	887,70
		1,28E+06	1480,70



Для кожного із зазначених алгоритмів за допомогою комп'ютерного моделювання визначено числові характеристики часу роботи алгоритму, залежно від установленної кількості ітерацій обчислення хеш-коду.

Обираючи кількість ітерацій для обчислення хеш-коду пароля за алгоритмом PBKDF2, потрібно враховувати, що при автентифікації користувачів буде зростати час відгуку сервера, що може спричинити негативне враження на користувача. Справді, для кожного випадку автентифікації потрібно буде обчислювати хеш-код для введенного пароля і порівнювати його з хеш-кодом, що зберігається у базі даних. Тому для інформаційних систем із великою інтенсивністю взаємодій із користувачами потрібно дотримуватися балансу між надійністю зберігання паролів і продуктивністю всієї системи (рисунок).

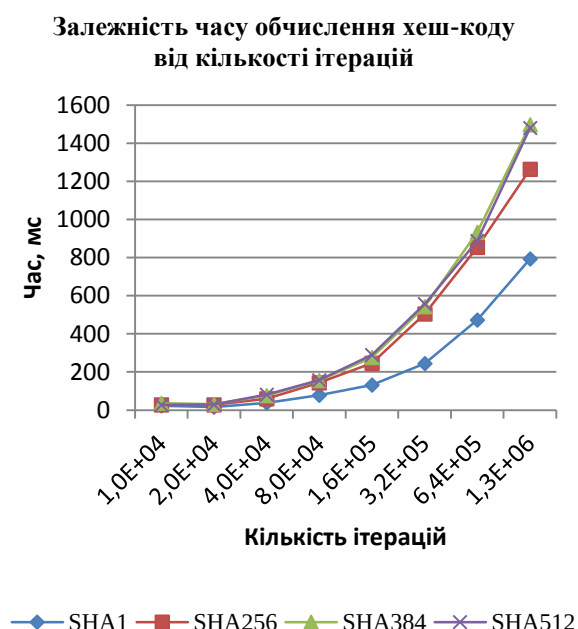


Рисунок. Графік залежності часу, необхідного для знаходження хеш-коду за алгоритмом PBKDF2, від кількості ітерацій

3. РЕКОМЕНДАЦІЇ ЩОДО ВИБОРУ СПОСОБІВ ЗБЕРІГАННЯ ПАРОЛІВ

На основі наведених у попередньому розділі статті результатів, можливо запропонувати такі рекомендації щодо вибору способу безпечного зберігання паролів:

1. Шифрування паролів слід використовувати лише в інформаційних системах, для яких критично важливим є забезпечення можливості відновлення оригінального пароля.

2. Використання “класичної” хеш-функції є доцільним в успадкованих інформаційних системах для забезпечення сумісності, або у системах,

де критичною є швидкість обчислень. Причому рекомендується застосовувати алгоритми SHA256, SHA384, SHA512.

3. Розмір солі для хешування має бути не менше ефективної довжини згенерованого хеш-коду, для кожного випадку вона має генеруватися окремо у вигляді криптографічно стійкої послідовності випадкових значень, одна й та ж сіль ніколи не повинна використовуватися повторно.

4. Для найбільш критичних до стійкості від зовнішніх атак додатків доцільно обирати алгоритм PBKDF2 хешування паролів, як досить гнучкий у налаштуваннях і порівняно стійкий до різного роду атак на хеш-код, у цьому разі кількість ітерацій має забезпечувати достатню обчислювальну складність для уникнення атаки прямим перебором, але разом із тим не перевантажувати власний сервер.

4. ВИСНОВКИ

Таким чином у результаті дослідження проведено порівняльний аналіз відомих способів безпечного зберігання паролів, установлено їхні переваги й недоліки, експериментально визначено часові показники застосування алгоритму стандарту PBKDF2, окреслено доцільні галузі застосування кожного способу.

Це дало змогу запропонувати рекомендації, які можливо використовувати для вибору способу зберігання паролів, залежно від типу інформаційної системи та важливості даних, що в ній зберігаються.

Отримані за результатами обчислювального експерименту дані можуть бути використані для вибору способу безпечного зберігання паролів.

Подальші дослідження доцільно проводити у напрямку створення моделей оцінювання ефективності способів зберігання паролів в інформаційних системах.

СПИСОК ВИКОРИСТАНОЇ ЛІТЕРАТУРИ

- [1] Н. Смарт. *Криптографія*, Техносфера, Москва, 2005, 528 с.
- [2] Stephen Haunts, *Applied Cryptography in .NET and Azure Key Vault. A Practical Guide to Encryption in .NET and .NET Core*, Apress, Berkeley, CA, 2019, 228 p.
- [3] *Microsoft documentation*, [Online]. Available: <https://docs.microsoft.com/en-us/dotnet/api/system.security.cryptography.hashalgorithmname?view=netcore-3.1>
- [4] *Salted Password Hashing – Doing it Right*, [Online]. Available: <https://crackstation.net/hashing-security.htm>
- [5] Yasser M. Alginahi, Muhammad Nomani Kabir, *Authentication Technologies for Cloud Computing, IoT and Big Data*, The Institution of Engineering and Technology, London, UK, 2019, 354 p.
- [6] *RFC2898 PKCS #5: Password-Based Cryptography Specification, Version 2.0*, [Online]. Available: <https://tools.ietf.org/html/rfc2898>

Стаття надійшла до редколегії

17.11.2020



Recommendations for choosing a method of safe storage of passwords

Using passwords remains the most common way to authenticate users for various types of information systems. This poses the challenge of securing the storage of user authentication information and protecting it from unauthorized access. In practice, various algorithms for secure password storage have become widespread. Mutually contradictory requirements for such algorithms for secure password storage, which on the one hand must be complex enough to counter various attacks, and on the other – simple to ensure the speed of the information system – determine the relevance of the study. There is a significant threat of direct search, especially given the fact that the computing power of CPUs and GPUs is constantly growing. Therefore, there is a need to be able to change the complexity of the hash code calculation, and therefore the amount of computation and time so as to significantly complicate the attack, but not cause discomfort to the end user due to the delay in password verification. Among the known methods of secure password storage are password encryption, the use of the hash function in the classical version, as well as the addition of salt and the use of iterations to calculate the hash code. The comparative analysis of the given methods is carried out in the work, their advantages and disadvantages are established, expedient areas of application of each method are outlined, the corresponding recommendations are developed. For the computational experiment, the tools of the Microsoft .NET Core 3.1 platform were used, which made it possible to set the time indicators of the hash code generation algorithm depending on the set parameters of the algorithm. The data obtained from the experiment can be used to select a method of securely storing passwords.

Keywords: information protection; user authentication; encryption; hash code calculation algorithms; computational experiment.



Олександр Махович,
кандидат технічних наук, асистент кафедри мережних та інтернет-технологій факультету інформаційних технологій Київського національного університету імені Тараса Шевченка.

Olexander Mahovich,
Candidate of Technical Sciences, Assistant of the Department of Network and Internet Technologies, Faculty of Information Technologies, Taras Shevchenko National University of Kyiv.



Роман Миколайчук,
доктор технічних наук, доцент, доцент кафедри мережних та інтернет-технологій факультету інформаційних технологій Київського національного університету імені Тараса Шевченка.

Roman Mykolaichuk,
Dr. Tech. Sciences, Associate Professor, Associate Professor of the Department of Network and Internet Technologies, Faculty of Information Technology, Taras Shevchenko National University of Kyiv.



Віра Миколайчук,
асистент кафедри інформаційних систем та технологій інституту інформаційних технологій Державного університету телекомунікацій.

Vira Mykolaichuk,
Assistant of the Department of Information Systems and Technologies of the Institute of Information Technologies of the State University of Telecommunications.



УДК 004.056

DOI <https://doi.org/10.17721/ISTS.2020.4.53-57>

V. I. Ignisca, orcid.org/ 0000-0001-9295-0983,
veravialkova@gmail.com

D. Vdovenko, orcid.org/0000-0002-3263-867X,
vdovenko.danylo@gmail.com

Taras Shevchenko National University of Kyiv, Kyiv, Ukraine

ANALYSIS OF METHODS DATA SECURITY

The article analyzes the main methods of information protection, from which it is possible to conclude that no method of data protection is ideal for all situations. It is important to choose an enterprise solution that provides comprehensive functionality, a flexible range of data protection options, broad support for platform and data types, and proven success in production implementations. The choice of method of information protection should take into account many circumstances that may arise during the implementation of a particular method. Due to the variety of data generated today, in addition to increasing the number of new platforms, flexibility can be a critical aspect of the data protection solution. A careful review of the requirements should make it easy to compare them with the relevant data protection methods, and it is necessary to make sure that the solution includes everything necessary to meet these requirements. Choosing the right method of information protection becomes much more difficult when more complex environments with many conflicting variables are involved, as it must support several options to provide flexibility to protect and meet data confidentiality, integrity and availability requirements. Only the integrated use of different measures can ensure reliable protection of information, because each method or measure has weaknesses and strengths. In some situations, internal security policies or regulations may forcibly change one method of data protection to another. Today, most standards, such as PCI DSS and HIPAA, allow a combination of the aforementioned methods, but these standards usually lag behind available or new data protection technologies. The set of methods and means of information protection includes software and hardware, protective transformations and organizational measures. A set of such methods, which are focused on protecting information, should protect them depending on whether the information is stored, moved or copied, accessed or used.

Keywords: start-up, information interaction, customer journey map, forecasting.

1. INTRODUCTION

Modern corporate solutions and solutions for enterprises require flexibility, versatility, stability and a high level of security to protect information. Information security solutions should ensure the integrity, confidentiality and availability of information. Compliance with information properties such as privacy, integrity, and accessibility is critical.

Information security can be presented in certain forms, such as, common forms and methods include: data management, network firewalls, intrusion prevention systems (IPS), semantic security at the row and column level (RLS / CLS), identity management (IDM), role-based access control (RBAC), activity monitoring, encryption, tokenization, encryption, data masking and more. The set of methods and means of information protection includes software and hardware, protective transformations and organizational measures. A set of such methods that focus on protecting information should protect them

depending on whether the information is stored, moved or copied, accessed or used [1].

Organizational measures for information protection include a set of actions for the selection and verification of personnel involved in the preparation and operation of programs and information, clear regulation of the development and operation of the information system. Only the integrated use of different measures can ensure reliable protection of information, because each method or measure has weaknesses and strengths. Technical (hardware) means. These are different types of devices (mechanical, electromechanical, electronic, etc.), which hardware solve the problem of information protection [2].

Software includes programs for user identification, access control, encryption of information, removal of residual information such as temporary files, test control of the security system.

Mixed hardware and software implement the same functions as hardware and software separately



and have common properties. Organizational means consist of organizational and technical (preparation of premises with computers, laying of cable system taking into account requirements of restriction of access) and organizational and legal.

2. ANALYSIS OF INFORMATION PROTECTION METHODS

Dynamic data masking. Data is created that is completely or partially disguised when users or services without access to view it receive it. This technology does not change the data of the original text during storage on media (Figure 1). In most cases, this method is used as additional protection, most databases use this technology for viewing mode [3].

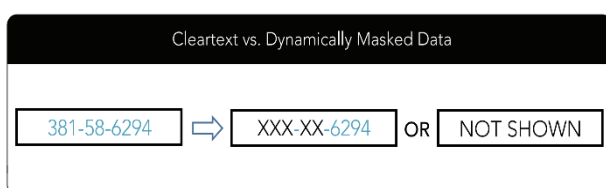


Fig. 1. An example of dynamic data masking

Static data masking. When using this method, the data cannot be reset. This method uses one-way hash algorithms along with encryption technology to achieve a result when the hash value is represented in binary form and cannot be stored using the original data type (Figure 2). Most static data masking tools generate data that is similar to the original and can be stored using the same data type and set of characters.

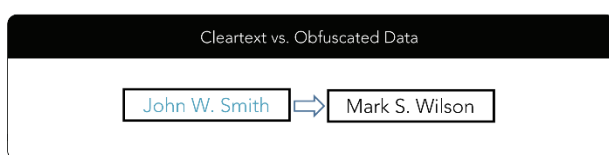


Fig. 2. An example of static data masking

Encryption technology uses mathematical algorithms and cryptographic keys to convert data into binary ciphertext (Figure 3). Resetting the data is possible only with the correct key with the algorithm. There are many forms of data encryption, various key benefits and other parameters.

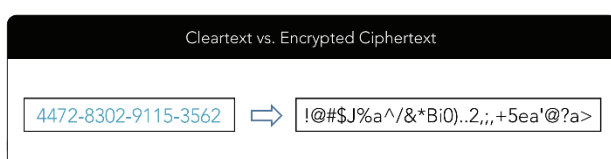


Fig. 3. An example of encryption technology

Tokenization. This method replaces a randomly or in some way generated value (token) with the value of the original text, stores this value in a table that corresponds to the value of the original text to the generated token (Figure 4).

The type and length of the token data remain the same as the original text, the token search table becomes the "key", which allows you to get the value of the original text from the token. With the increase in the number of such tokenized data and tables, along with the complexity of the IT infrastructure, leads to the fact that token search tables make it impossible to quickly and efficiently search.

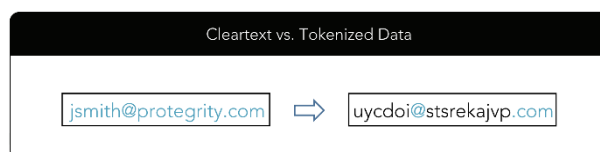


Fig. 4. An example of data tokenization

Encryption while preserving the data format. This method combines the advantages of both encryption (using a mathematical algorithm for encryption) and tokenization (the same type of data is stored).

Format encryption requires the same CPU cycles for encryption and then additional processing to convert binary ciphertext to the same data type as the original text and avoid possible "collisions" (same output for two different input values) by converting a larger binary fields in a smaller alphanumeric or alphanumeric data type.

Deidentification is a more general term for "anonymizing" data, used, for example, to display user data that is almost impossible to identify using available data.

FIRST	MI	LAST	TYPE	CITY	STATE	SSN
John	W	Smith	Owner	Detroit	MI	248-632-1292
Ueqa	K	Hvapi	Owner	Orlaqnt	MI	248-999-9999

Fig. 5. An example of data deidentification

For sufficient identification or anonymization of field records to be protected in databases, these fields must be defined in the organization's data security policy. The data classification scheme should specify a minimum list of fields to be protected (for example, credit card number, last name, first name, etc.), as well as a list of recommendations for additional fields for less secure or more widely



available systems. These additional fields (for example, city, e-mail address, telephone number, etc.) will also be protected in cases where this is clearly justified by the security policy of the organization or enterprise [4].

3. DATA CLASSIFICATION AND DATA SECURITY POLICY

The data classification scheme should also determine which data should be protected on media, during data transmission and during their use [5]. The protection method should be platform-specific, using a centralized approach focused on maintaining multiple data protection options as needed.

Data protection on media.

The use of a particular method of data protection on media depends on the environment and the criteria that the method must meet. Full encryption is a common choice for portable media such as laptop hard drives or flash cards. Incomplete data encryption of certain fields has disadvantages, such as: changing the data type in VarByte (binary), larger field size, problems with systems or applications that do not support binary data fields.

Tokenization. It's a modern solution for data protection on media. Gained popularity for the following reasons:

- data can be easily moved between systems open to databases or programs where access to the original text values is not required;
- if necessary, it is possible to find the correspondence of the original data using the correspondence table;
- Data type, field size, or supported character sets do not change.

Data protection during transportation.

Network traffic encryption protocols (SFTP, HTTPS, SSL, TLS) are most commonly used when transporting data, but field-level protection such as tokenization or encryption can also be used to add an extra layer of security to data transported between systems.

Data protection during use.

The biggest problem in protecting particularly sensitive data fields is when used by users or the corporate environment [6]. Usually only 1% to 3% of the total data in a large database guarantees the use of an incomplete type of information protection. In addition, 80% to 90% of all information required for business activities can be submitted in a secure form. Therefore, only 10% to 20% of the information required for operational activities will require access to 1%–3% of data that will require additional processing.

Sensitive user data fields, such as credit card numbers and other vulnerable data, require access to the original text, so both encryption and tokenization can be used without compromising the reference integrity of the data.

Incomplete data encryption gives users greater access to data, while improving the overall security framework and adhering to security policies and regulations.

Data security policy sequence.

It is equally important to consider the ability to consistently apply data security policies at all of the above stages and in all environments of the organization. The same method of protection, limited rights of access to information, accountability and audit, protection against unauthorized protection, etc. must be applied consistently throughout the flow of data exchange [7]. Data must be equally secure, processes must be consistent with security policies from start to finish, regardless of the platform used for data processing, analysis and storage.

The choice of method of information protection should take into account many circumstances that may arise during the implementation of a particular method.

The choice of the necessary method of information protection becomes much more difficult when more complex environments with numerous conflicting variables are involved. A solution for a particular enterprise or organization must support several options to provide flexibility to protect and meet the requirements of confidentiality, integrity and availability of data.

When protecting small field data (1 or 2 characters), encryption is the best option, as even small fields such as Boolean logical fields can be encrypted (with an initialization vector). Tokenization is limited by the width of the token search table used and is typically used for fields with three or more characters, but very large fields, such as more than 100 characters and can contain hundreds or thousands of characters, are not suitable for tokenization.

There are several approaches to consider for the protection of sensitive data using external cloud services, in addition to the obvious transition to management using a contractual approach. New, lower data processing costs can be provided by the heads of information security departments, finding a way to achieve the same or even improve the level of information protection through the use of cloud services [7].

Protect data fields of different sizes

When protecting small field data (1 or 2 characters), encryption is the best option, as even small fields



such as Boolean logical fields can be encrypted (with an initialization vector).

If you want to use a protected field often enough, encryption can be faster than tokenization, especially for large fields, because encryption takes the same time to process a 2-character field and a 16-character field. However, for standard structured fields of limited length (3 to 15 characters), tokenization is an ideal solution.

Storage regulations have a significant impact when it is necessary to transport information between several services, such as in some cloud applications. They often require data deidentification, but in some cases they do not allow certain particularly sensitive data to be transported between media at all.

4. CONCLUSION

The paper proposes analyzes of the main methods of information protection, from which it is possible to conclude that no method of data protection is ideal for all situations. It is important to choose an enterprise solution that provides comprehensive functionality, a flexible range of data protection options, broad support for platform and data types, and proven success in production implementations.

Due to the variety of data generated today, in addition to increasing the number of new platforms, flexibility can be a critical aspect of the data protection solution. A careful review of the requirements should make it easy to compare them with the relevant data protection methods, and it is necessary to make sure that the solution includes everything necessary to meet these requirements.

To protect data in web services, the best solution is to use methods to protect information using tokens, because modern enterprise and enterprise solutions require flexibility, versatility, stability and a high level of information protection.

REFERENCES

- [1] Averchenkov VI Information security audit.– M.FLINTA, 2016 – 269 p.c.
- [2] Gvozdeva, T.V. Design of information systems / T.V. Gvozdev, B.A. Ballod. – M.: Fenix, 2009. – 512 p.
- [3] Protegrity. Methods of Data Protection [Electronic resource] / Protegrity Access mode to the resource: <http://sfbay.issa.org/comm/presentations/2015/Methods%20of%20Data%20Protection%20White%20Paper%20-%20Protegrity%20Sept%202015.pdf>.
- [4] Goodson, John A Practical Guide to Data Access (+ DVD-ROM) / John Goodson, Rob Steward. – M.: BHV-Petersburg, 2013. – 304 p. Голицына, О. Л. Основы алгоритмизации и программирования / О. Л. Голицына, И. И. Попов. – М.: Форум, 2010. – 432 с.
- [5] Pollis, Gary Software Development. Based on the Rational Unified Process (RUP) / Gary Pollis et al. – M.: Binom-Press, 2011. – 256 p.

[6] Vsyakikh, Ye.I. Practice and problems of modeling business processes. – M.: Book on Demand, 2008. – 246 p.

[7] Internet Engineering Task Force. JSON Web Token (JWT) [Electronic resource] / Protegrity – Access mode to the resource: <https://tools.ietf.org/html/rfc7519>.

Стаття надійшла до редколегії

07.12.2020



Аналіз методів захисту даних

Проаналізовано основні методи захисту інформації, з яких можна зробити висновок, що жоден метод захисту даних не ідеальний для всіх ситуацій. Важливо вибрати корпоративне рішення, яке надає широкі функціональні можливості, гнучкий спектр варіантів захисту даних, широку підтримку платформ і типів даних, а також доведений успіх у виробничих реалізаціях. У випадку вибору методу захисту інформації слід урахувати багато обставин, які можуть виникнути під час упровадження певного методу. Завдяки різноманітності даних, що генеруються нині, крім збільшення кількості нових платформ, гнучкість може бути критичним аспектом рішення щодо захисту даних. Ретельний огляд вимог має полегшити їхнє порівняння з відповідними методами захисту даних, і необхідно переконатися, що рішення містить усе необхідне для задоволення цих вимог. Вибір правильного методу захисту інформації стає набагато складнішим, коли задіяні більш складні середовища з багатьма конфліктуючими змінними, оскільки він повинен підтримувати кілька варіантів, щоб забезпечити гнучкість захисту та задоволення вимог щодо конфіденційності, цілісності та доступності даних. Тільки комплексне використання різних заходів може забезпечити надійний захист інформації, оскільки кожен метод або захід має слабкі та сильні сторони. У деяких ситуаціях політика чи правила внутрішньої безпеки можуть примусово змінити один спосіб захисту даних на інший. Сьогодні більшість стандартів, таких як PCI DSS та HIPAA, дозволяють поєднувати зазначені вище методи, але ці стандарти зазвичай відстають від наявних або нових технологій захисту даних. Сукупність методів та засобів захисту інформації включає програмно-апаратні засоби, захисні перетворення й організаційні заходи. Набір таких методів, орієнтованих на захист інформації, має захищати їх залежно від того, зберігається, переміщується чи копіюється інформація, доступ до неї чи використання.

Ключові слова: стартап, інформаційна взаємодія, карта шляху клієнта, прогнозування.



Vira Igniska,

PhD, associate professor of Dept. of cyber security and information protection of Faculty of Information Technology of Taras Shevchenko National University of Kyiv.

Віра Ігніска,

кандидат технічних наук, доцент кафедри кібербезпеки та захисту інформації факультету інформаційних технологій Київського національного університету імені Тараса Шевченка.



Danylo Vdovenko,

Student of Faculty of Informational Technology of Taras Shevchenko National University of Kyiv.

Данило Вдовенко,

студент факультету інформаційних технологій Київського національного університету імені Тараса Шевченка.



КОМП'ЮТЕРНА ІНЖЕНЕРІЯ ТА ПРОГРАМНЕ ЗАБЕЗПЕЧЕННЯ

УДК 005.8:316.422

DOI <https://doi.org/10.17721/ISTS.2020.4.58-62>Б. Бондаренко, orcid.org/0000-0001-6431-3434,
daren.24@ukr.netЮ. Я. Самохвалов, orcid.org/0000-0001-5123-1288,
yu131053@gmail.com

Київський національний університет імені Тараса Шевченка, Київ, Україна

ПОШУК МУЛЬТИМЕДІЙНОЇ ІНФОРМАЦІЇ НА ОСНОВІ НЕЙРОННИХ МЕРЕЖ

Розглянуто підходи до використання нейронних мереж для пошуку мультимедійної інформації. Розроблення методів пошуку мультимедійної інформації необхідне через велику кількість такої інформації. Традиційні методи пошуку мультимедійної інформації мають високу швидкість оброблення даних, але низьку точність через відсутність можливості виконання семантичного пошуку. Використання нейронних мереж дозволяє здійснювати семантичний пошук, що збільшує його точність і повноту. Наведено підходи використання нейронних мереж на етапах індексування та пошуку мультимедійної інформації. За допомогою нейронної мережі аналізують мультимедійний файл і виконують його класифікацію. Результат класифікації файла використовують для створення його текстового опису – анотації, яку порівнюють із запитом для визначення релевантності. Існує багато готових мереж для класифікації, використання яких дозволяє пришвидшити процес створення системи пошуку мультимедійної інформації, але неможливо створити нейронну мережу для класифікації всіх об'єктів реального світу, тому потрібно застосовувати декілька нейромереж. Також за допомогою нейронних мереж будують вектори ознак для мультимедійного файла та пошукового запиту. Прості функції подібності, такі як косинус подібності, застосовують до побудованих векторів, для визначення семантичної близькості запиту та мультимедійного файла. Причому пошуковий запит може бути як у текстовій формі, так і у вигляді будь-якого формату пошуку в мультимедійному файлі. Цей підхід дозволяє будувати оптимальну нейронну мережу під конкретну задачу. Нейронні мережі застосовують для порівняння побудованої анотації файла та запиту, що підвищує точність і повноту пошуку, порівняно з традиційними методами, за рахунок здатності нейромереж враховувати семантичне значення тексту.

Ключові слова: нейронні мережі; пошук мультимедійної інформації; семантичний пошук; системи пошуку інформації; анотація даних.

1. ВСТУП

З поширенням інтернету з'явилась можливість передачі та зберігання великих обсягів інформації, зокрема мультимедійної. Це викликало потребу в розробленні методів пошуку інформації мультимедіафайлів. Існуючі підходи до пошуку текстової інформації не можливо застосовувати для пошуку мультимедійних даних. Тому створення ефективних методів пошуку мультимедійної інформації є важливим напрямом досліджень. Наявні методи, такі як аналіз гістограми кольорів зображення або частоти звуку [1], мають високу швидкість оброблення, але не здатні

виконувати семантичний пошук. Використання нейронних мереж дозволяє обробляти мультимедіафайли та виконувати семантичний пошук, що підвищує точність і повноту пошуку.

2. АНАЛІЗ ОСТАННІХ ДОСЛІДЖЕНЬ І ПУБЛІКАЦІЙ

Питання пошуку мультимедійної інформації висвітлено в джерелах [1–4]. Вони розглядалися відомими компаніями, такими як, Google, Apple та Amazon, які займаються збиранням і зберіганням даних. Указаній тематиці присвячена велика кількість міжнародних конференцій, се-

© Бондаренко Б., Самохвалов Ю. Я., 2020



мінарів і круглих столів. Проте питання пошуку мультимедійної інформації на основі використання нейромереж ще не достатньо повно висвітлено у працях вітчизняних і зарубіжних вчених. Тому застосування нейромережного підходу у пошуку мультимедійної інформації є актуальним напрямом досліджень.

3. ОСНОВНИЙ ТЕКСТ

Інформація в базі даних, з якою працює інформаційно-пошукова система, може бути структурованою або неструктурованою. Саме структура даних вимагає відповідного підходу до отримання корисних знань з інформаційних масивів.

Структуровані дані – це інформація, упорядкована й організована певним чином із метою її ефективного оброблення. [3]. Неструктуровані дані – це інформація, яка або не має наперед визначеної структури, або не організована в установленому порядку. Підвидом неструктурованих даних є мультимедійні, до яких належать зображення, аудіо- та відеофайли. Такі дані зберігають у первісному вигляді й аналізують за допомогою спеціальних методів.

Система пошуку інформації повинна мати можливість представляти та зберігати інформацію таким чином, щоб забезпечити її швидкий пошук [2]. Вона має бути здатною працювати з різними типами інформації, зокрема й мультимедійною. Щоб здійснювати пошук мультимедійних даних, система, як правило, повинна визначати їхні особливості, наприклад різноманітність кольорів і текстуру зображення або ритмічність аудіо. На рис. 1 показано загальну структуру системи пошуку інформації.

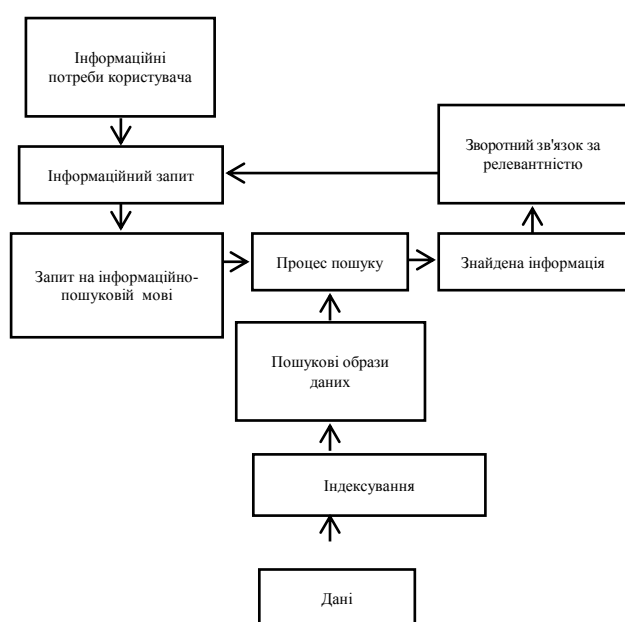


Рис. 1. Загальна структура системи пошуку інформації

Основними етапами роботи системи пошуку інформації є: індексування документів і створення їхнього пошукового образу (ПО); перетворення запиту користувача, сформульованого на природній мові, на запит мовою інформаційно-пошукової системи; пошук інформації і надання її користувачу.

Серед цих етапів індексування й пошук інформації є ключовими, тому що повинні враховувати семантичне значення документів. Тому використання на цих етапах нейронних мереж дозволить підвищити точність і повноту пошуку.

3.1. ІНДЕКСУВАННЯ МУЛЬТИМЕДІЙНОЇ ІНФОРМАЦІЇ

Щоб мати можливість швидкого пошуку інформації, системі потрібно побудувати її пошуковий образ. Для цього використовується індекс. Індексация дозволяє структурувати дані так, щоб була можливість пошуку документів за одним або кількома критеріями подібності [5]. У процесі індексації створюється пошуковий образ файлу, який потім порівнюється із запитом для визначення його релевантності. Для проведення індексації мультимедійної інформації розроблено різні методи. Наприклад, використання гістограм кольорів для оброблення зображень і використання прихованої марковської моделі для аудіофайлів [1, 2]. Одним із методів індексації є створення анотації – текстового опису файлу. Анотування файлу – це процес присвоєння йому ключового слова або списку ключових слів, що описує семантичне значення його вмісту.

Анотування зображень. Під час оброблення зображень і відео (як набору зображень) анотація може виконуватися на двох рівнях: локальному та глобальному. На локальному рівні зображення розглядається як сукупність об'єктів. Така анотація має на меті описати кожен об'єкт на зображенні у вигляді ключового слова або списку ключових слів. Глобальний рівень описує весь образ зображення і створює список ключових слів для опису його загального змісту.

Об'єкти, що мають однакове семантичне значення, можуть мати різний вигляд на фото [4]. Анотування зображень традиційними методами є малоефективним, адже вони використовують прості ознаки зображення (наприклад, форма, гістограми кольорів і текстура), які обчислюються на основі характеристик пікселів, таких як, положення та колір. Під час створення анотації зображень за допомогою нейронних мереж класифікують об'єкти та явища, що зображені на них, та визначають пов'язані з ними ключові слова.

Існують різні архітектури нейронних мереж для класифікації. Майже всі вони використовую-



ють як основу згорткові шари та шари агрегування. За допомогою спеціально натренованих згорткових нейромереж можна виконувати класифікацію та опис зображень. Після класифікації зображення створюється його текстовий опис. У подальшому виконується семантичний пошук по описах зображень і знаходяться не лише конкретні об'єкти із зображень, а близькі до них за значенням. Наприклад, у разі пошуку слова *стілець*, система знаходить усі зображення на яких є *стілці*, а також об'єкти зі схожим значенням, наприклад *стіл*. Прикладом мережі класифікації зображень є Inception-ResNet-v2 – це згорткова нейронна мережа, яка натренована на більш ніж мільйони зображень із бази даних ImageNet [6]. Її архітектуру наведено на рис. 2. Мережа складається зі 164 шарів і може класифікувати зображення за 1000 категоріями різних предметів і тварин. Як результат, мережа отримала широкі можливості для класифікації зображень [6].

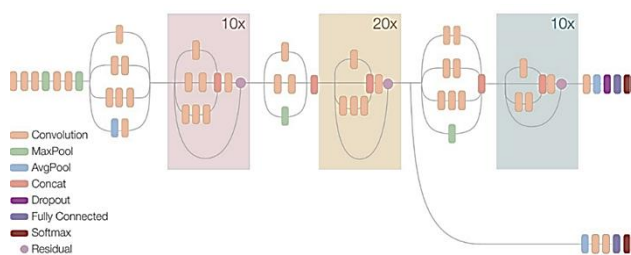


Рис. 2. Архітектура мережі для класифікації зображень

Мережа Inception-ResNet-v2 приймає на вхід зображення розміром 299 на 299 пікселів та визначає, що зображено на ньому. Її глибока архітектура складається з більше 30 згорткових шарів, що дозволяє показувати високу точність класифікації, але потребує значних затрат ресурсів на оброблення одного зображення.

Окрім класифікації, нейронні мережі можуть виконувати сегментацію, тобто визначати розмір і положення об'єкта на зображенні. Результат сегментації також може бути корисним при побудові анотації. Існуючі нейронні мережі не можуть класифікувати будь-які об'єкти, тому доцільно використовувати декілька нейромереж, натренованих для класифікації різних типів об'єктів. Створення універсальної системи класифікації – це ще один напрям для розвитку.

Анотування аудіоінформації. Пошук певного звуку або типу звуку (наприклад, промови певної людини або музики) може бути непростим завданням [7]. Немає стандарту класифікації звуків: два користувачі дадуть різний опис для одного звуку.

Для створення анотації аудіофайлів потрібно використовувати різні типи нейронних мереж,

адже кожен аудіофайл несе своє значення – якщо це запис промови, то потрібно розпізнати слова, якщо це музика, то можливо визначити її жанр. Часто в основі мереж для розпізнавання звуків лежать рекурентні нейронні мережі, адже вони здатні працювати з послідовностями даних, розподіленими в часі. Їх використовують для розпізнавання природної мови та прогнозування даних. Рекурентна нейронна мережа на виході додатково виводить прихований стан, який додається до наступних даних, та знову подається на вхід цієї ж мережі.

Різновидом рекурентних мереж є вентильний рекурентний вузол (Gated Recurrent Unit, GRU), який працює таким чином [8]. На вхід подають значення вхідного вектора x_t та вихідне значення попереднього блоку h_{t-1} . За ними розраховують значення векторів вузлів уточнення z_t та скидання r_t :

$$z_t = \sigma(W_z x_t + U_z h_{t-1} + b_z), \quad (1)$$

$$r_t = \sigma(W_r x_t + U_r h_{t-1} + b_r), \quad (2)$$

де W та U – матриці ваг, а b – вектор параметрів.

Значення вузла уточнення z_t допомагає моделі визначити, яку кількість минулої інформації (з попередніх етапів часу) потрібно передати далі. А вузла скидання r_t – яку кількість минулої інформації забути (не враховувати). Ці два вектори використовують для розрахунку прихованого стану даної ітерації мережі. Вихідні значення даного блоку h_t розраховують за формулою

$$h_t = z_t \circ h_{(t-1)} + (1 - z_t) \circ \tanh(W x_t + r_t \circ U h_{(t-1)} + b_h), \quad (3)$$

де $\circ$ – добуток Адамара.

Якщо значення елемента вектора r близьке до 1, то більшість інформації з попередньої ітерації не буде враховано. Причому, чим ближче значення z_t до 1, тим більше інформації з попереднього виходу мережі буде передано на наступну ітерацію.

Аудіофайл поділяють на частини однакової довжини й аналізують за допомогою рекурентних нейромереж. Це дозволяє виконати розпізнавання мови, записаної на аудіофайлі, або класифікувати його жанр, якщо це музика. Розпізнана інформація файлу використовується для формування його анотації.

Побудова вектора ознак. Для пошуку мультимедійної інформації, за допомогою нейронних мереж, будують вектор ознак медіафайла. Використовуючи мережу зі спеціально побудованою архітектурою, мультимедіафайл перетворюють на n -вимірний вектор, створюючи його пошуковий образ [3]. На рис. 4 зображено загальну архітектуру нейронної мережі перетворення зображення на вектор ознак. Перші шари таких мереж зазвичай є шарами згортки й агрегування, які потрібні для зменшення розмірності вхідного файлу. Після цих

шарів розташовують повнозв'язний шар, вихідне значення елементів якого стає вектором ознак. Для покращення точності пошуку, між шарами згортки та повнозв'язним шаром доцільно розташовувати декілька прихованих шарів.

Побудований вектор ознак для мультимедіа файла зберігає в собі особливості даних, наприклад, для зображення – це різноманіття кольорів, пропорції фігур, текстура тощо.

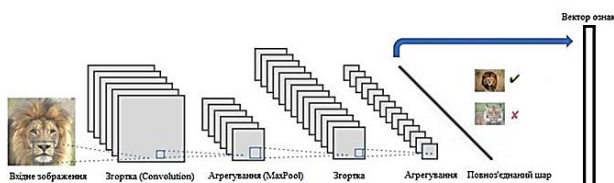


Рис. 3. Загальна архітектура мережі для перетворення зображення на вектор ознак

Отриманий від користувача пошуковий запит також обробляється нейронною мережею та будується вектор його ознак. Пошуковий запит може бути як у текстовій формі, так і у вигляді відповідного формату шуканого медіафайла – при пошуку зображення, порівнювати зі схожим зображенням; шукаючи аудіофайл, виконувати пошук за схожим відрізком звуку або схожою ритмічністю тощо.

Указаний підхід дозволяє зменшити витрати часу на аналіз файлу, але він потребує налаштування архітектури мережі для досягнення найбільшої точності пошуку. Оброблення великих векторів нейронною мережею потребує значних затрат ресурсів, тому зменшення розмірів векторів має значення для підвищення ефективності роботи системи при використанні цього методу.

Семантичний пошук. Після побудови анотації для мультимедіафайла, вона порівнюється з пошуковим запитом і визначається їхня відповідність. Існує два підходи до порівняння анотації та запиту – лексичний і семантичний. За використання лексичного підходу, відповідність оцінюють лише на основі лексичних збігів термінів із запиту в документі, тобто здійснюють пошук однакових слів. Такі класичні алгоритми для оцінювання важливості слова в контексті документа, як TF-IDF та Okapi BM25, є методами лексичного пошуку [9].

Але лексичний метод пошуку неефективний, якщо немає повного збігання слів у запиті та текстовому описі. Цю проблему розв'язує другий метод пошуку – семантичний. Семантичний пошук може віднаходити синоніми та споріднені слова до слів запиту, а також він стійкий до помилок у словах запиту.

Ефективне використання нейронних мереж для семантичного пошуку пов'язано з тим, що їхня глибока ієрархічна структура створює високі рівні абстракцій, аналізуючи необроблені дані, й автоматично визначає семантичні зв'язки. Загальну архітектуру нейромережі для виконання семантичного пошуку показано на рис. 3. Метою роботи даної мережі є представлення тексту запиту та тексту анотації файлу у вигляді таких векторів, скалярний добуток яких тим більший, чим більше релевантний документ до запиту – дані вектори називають векторами прихованих семантичних ознак. Перетворення вхідного тексту на вектори семантичних ознак відбувається за допомогою прихованих шарів. Для оцінювання релевантності документа до запиту, два отримані вектори семантичних ознак порівнюють один з одним, використовуючи деяку функцію подібності [9].

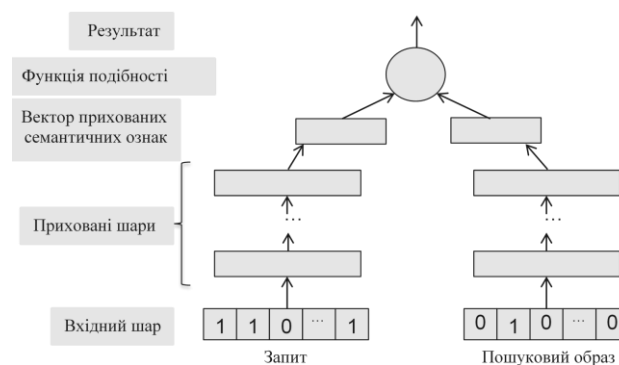


Рис. 4. Архітектура мережі для семантичного пошуку

Прикладом нейромережі для семантичного пошуку інформації є DSSM (Deep Semantic Similarity Model) [10]. Вхідні дані моделі DSSM – це багатовимірний вектор термінів у запиті або документі. Терміни розбиваються на триграми. Елементи вектора, що відповідають триграмам, присутнім у вхідному тексті та документі, набувають значення 1, в іншому випадку – 0. Вхідні вектори обробляються прихованими шарами, які перетворюють вхідний вектор на вектор прихованих семантичних ознак. Для обчислення оцінки релевантності між документами та запитом використовується функція косинусної подібності відповідних їм семантичних векторів понять [8]. Ціль навчання даної нейронної мережі – максимізувати вихідне значення функції подібності.

4. ВИСНОВОК

Великий обсяг мультимедіафайлів, викликав потребу в розробленні методів пошуку інформації серед таких файлів. Традиційні методи пошуку мультимедійної інформації, такі як, аналіз



гістограми кольорів зображення або частоти звуку, мають високу швидкість оброблення даних, але низьку точність через відсутність можливості виконання семантичного пошуку. Використання нейронних мереж дозволяє обробляти та шукати мультимедійну інформацію, а їхня здатність віднаходити семантичні зв'язки дає змогу виконувати семантичний пошук. Створення анотації мультимедіафайлів за допомогою нейронних мереж, підвищує точність і повноту пошуку, порівняно з традиційними методами. Але через неможливість створення нейромережі для класифікації всіх об'єктів реального світу, потрібно використовувати декілька нейромереж. Побудова вектора ознак файла за допомогою нейромереж також враховує його семантичне значення. Використовуючи вказаний підхід, можна будувати оптимальну мережу під конкретну задачу, що зменшить затрати ресурсів у разі побудови вектора, при цьому важливим є правильне налаштування нейронної мережі, адже це впливає на швидкість оброблення файла та точність пошуку.

СПИСОК ВИКОРИСТАНОЇ ЛІТЕРАТУРИ

- [1] Sagarmay Deb. Multimedia Systems and Content-Based Image Retrieval. IGI Global, 2003, 406 p.
- [2] Gerald J. Kowalski, Mark T. Maybury. Information Storage and Retrieval Systems. Kluwer Academic Publishers, 2000, 333 p.
- [3] Athman Bouguettaya, Boualem Benatallah, Ahmed K. Elmagarmid. Interconnecting Heterogeneous Information Systems. Springer Science+Business Media, 1998, 228 p.
- [4] Yu-Jin Zhang. Semantic-Based Visual Information Retrieval. IRM Press, 2006, 385 p.
- [5] Zongmin Ma. Artificial Intelligence for Maximizing Content Based Image Retrieval. Information Science Reference, 2009, 429 p.
- [6] Alex Alemi. Improving Inception and Image Classification in TensorFlow. URL: <https://ai.googleblog.com/2016/08/improving-inception-and-image.html> (дата звернення: 7.02.2021).
- [7] Maha Mahmood, Wijdan Jaber Al-Kubaisy, Belal Al-Khateeb. Using Artificial Neural Network for Multimedia Information Retrieval. Journal of Southwest Jiaotong University. Vol 54, No 3 (2019).
- [8] Kostadinov S. Understanding GRU Networks. URL: <https://towardsdatascience.com/understanding-gru-networks-2ef37df6c9be> (дата звернення 7.02.2021).
- [9] Po-Sen Huang, Xiaodong He, Jianfeng Gao. Learning Deep Structured Semantic Models for Web Search using Clickthrough Data. At ACM International CIKM. October 2013.
- [10] Gia-Hung Nguyen, Lynda Tamine, Laure Soulier, Nathalie Souf. Toward a Deep Neural Approach for Knowledge-Based IR. URL: <https://arxiv.org/pdf/1606.07211.pdf> (дата звернення 7.02.2021).

Стаття надійшла до редколегії

10.10.2020



Searching for multimedia information based on neural networks

The article considers approaches to the use of neural networks in multimedia information retrieval. The development of methods for multimedia information retrieval is necessary due to the large amount of such information. Traditional methods of multimedia information retrieval have a high speed of data processing, but low accuracy due to the inability of semantic search. The use of neural networks allows for semantic search, which increases its accuracy and completeness. Approaches to the use of neural networks at the stages of indexing and retrieval of multimedia information are considered. With the help of a neural network, a multimedia file is analyzed and classified. The result of classifying a file is used to create its textual description - an annotation that is compared to the search query to determine relevance. There are many ready-made classification networks that can be used to speed up the process of creating a multimedia search system, but it is not possible to create a neural network to classify all real-world objects, so multiple neural networks should be used. Neural networks are also used to build feature vectors for a media file and a search query. Similarity functions, such as cosine of similarity, are applied to constructed vectors to determine the semantic similarity of a query and a media file. In this case, the search query can be both in text form and in the form of the appropriate format of the desired media file. This approach allows to build an optimal neural network for a specific task. Neural networks are used to compare the constructed annotation of a file and a query, which increases the accuracy and completeness of the search, compared to traditional methods, due to the ability of neural networks to take into account the semantic meaning of the text.

Keywords: neural networks; search for multimedia information; semantic search; information search systems; data annotation.



Богдан Бондаренко,
магістрант кафедри інтелектуальних технологій Київського національного університету імені Тараса Шевченка.

Bogdan Bondarenko,
Master's student of the Department of Intellectual Technologies of Taras Shevchenko National University of Kyiv.



Юрій Самохвалов,
доктор технічних наук, професор кафедри інтелектуальних технологій. Київський державний університет імені Тараса Шевченка.

Yuri Samokhvalov,
Doctor of Technical Sciences, Professor of the Department of Intellectual Technologies. Taras Shevchenko National University of Kyiv.



УДК 621.391

DOI <https://doi.org/10.17721/ISTS.2020.3-4.63-70>М. Завалій, orcid.org/0000-0003-1088-67873, _zavaley@ukr.netА. Завалій, orcid.org/0000-0002-0465-58221, _zavaley@ukr.netЛ. В. Дакова, orcid.org/0000-0001-6104-8217, dacova@ukr.net

Державний університет телекомунікацій, Київ, Україна,

С. Ю. Даков, orcid.org/0000-0001-9413-3709, dacov@ukr.netІ. І. Пархоменко, orcid.org/0000-0001-6889-9284, parkh08@gmail.com

Київський національний університет імені Тараса Шевченка, Київ, Україна

ДОСЛІДЖЕННЯ БЕЗСЕРВЕРНИХ ОБЧИСЛЕНЬ ТА ЇХНЕ ВИКОРИСТАННЯ В МЕРЕЖАХ МОБІЛЬНОГО ЗВ'ЯЗКУ

Досліджено безсерверні обчислення та їхнє використання в мережах мобільного зв'язку. Нині без розрахунків неможливо уявити цілий світ. Вони оточують нас в усіх сферах життя – від звичайної оплати карткою в супермаркеті до наукових обчислень. Основною проблемою розрахунків є зростання кількості запитів і складність надання послуг. Обчислювальні компанії, яким потрібно постійно збільшувати пропускну здатність серверів, мають працювати над мережною інфраструктурою та наймати фахівців для підтримки таких систем. Усі витрати та складність розроблення серверних систем лягають на компанію-розробника, що є дуже дорогим фактором для нових компаній, діяльність яких пов'язана з обчисленнями й обробленням даних. Для вирішення питань розробляється та застосовується технологія "безсерверних обчислень", яка є моделлю хмарних обчислень. Ця технологія дозволяє використовувати ресурси, які їм фізично не належать. Усі проблеми з налаштуванням та обслуговуванням обладнання лягають на постачальників послуг. Також за допомогою інтелектуального відстеження навантаження система автоматично розподіляє необхідну потужність і ресурси для виконання розрахунків на даний момент. Основна проблема технології "безсерверних обчислень" – це не адаптація до потреб багатьох галузей і проблеми із сертифікацією та стандартизацією, але вони поступово розв'язуються. Дослідження використання таких технологій у телекомунікаційних системах є актуальним питанням розвитку науки та компаній у країні.

Розглянуто структуру мобільного оператора, технологію LTE, можливість прийому й аналізу даних від абонентів, методи відстеження абонентів у мережі за унікальними ідентифікаторами, а також можливість передачі інфраструктури оператора безсерверним технологіям. Наведено питання оброблення даних, існуючі методи та технології оброблення даних Big Data, системи оброблення/аналізу великих даних операторами в Україні та системи відслідковування й аналізу трафіка абонентів. Розглянуто комплексні рішення у вигляді безсерверних (хмарних) технологій, класифікацію таких рішень, огляд існуючих сервісів.

Ключові слова: безсерверні технології; хмарні технології; мережа LTE; великі дані; мобільний оператор.

1. ВСТУП

Нині в Україні стрімко розвиваються мережі стільникового зв'язку LTE. Трафік зростає щомісяця. Основна проблема обчислень – зростаюча кількість запитів і складність надання послуг у компаніях, що пов'язані з обчисленнями, для яких необхідно весь час збільшувати серверні потужності, працювати над інфраструктурою мережі та наймати спеціалістів для підтримки таких систем. Усі витрати та складність розроблення серверних систем припадають на компанію-розробника, що є дуже витратним факто-

ром для нових компаній, діяльність яких пов'язана з обчисленнями та обробленням даних. Для розв'язання таких проблем розробляють і застосовують технологію "без серверних обчислень", технологія являє собою модель хмарних обчислень, для яких платформа динамічно керує виділенням машинних ресурсів.

2. АНАЛІЗ ОСТАННІХ ДОСЛІДЖЕНЬ І ПУБЛІКАЦІЙ

Указана технологія надає змогу використовувати ресурси, що їм фізично не належать, а всі

© Завалій М., Завалій А., Дакова Л. В., Даков С. Ю., Пархоменко І. І., 2020



проблеми з налаштуванням та підтримкою обладнання припадають на постачальників послуг [1]. Також шляхом інтелектуального відстеження навантаження, система автоматично виокремлює необхідну потужність і ресурси для проведення обчислень у даний момент. Користувачі сплачують лише за ті послуги, які вони використовують [2]. Основна проблема технології “без серверних обчислень” – непристосованість до потреб багатьох сфер і проблеми із сертифікацією та стандартизацією, але вони поступово вирішуються. Дослідження використання таких технологій у телекомунікаційних системах є актуальним питанням для розвитку науки та компаній у країні [3].

3. ОСНОВНА ЧАСТИНА

Розглянемо структуру оператора мережі LTE та проведемо аналіз даних, які надходять до системи. Розглянемо технології оброблення даних і систем хмарних обчислень.

Розвитком світових телекомунікаційних технологій у галузі мобільного зв'язку є розроблення й упровадження стандартів четвертого (4G) і п'ятого покоління (5G), що забезпечують прискорення передавання даних і відповідне покращення якості надавання пропонованих послуг, при загальному зниженні витрат в експлуатації телекомунікаційного обладнання. Технологія, що покликана розв'язувати сучасні проблеми телекомунікацій, є LTE-технологія. Технологія LTE є актуальною на території України. Мережа LTE забезпечує підтримку пакетного трафіка з мінімальними затримками доставлення пакетів, без втрати пакетів передачі даних (seamless), мобільністю та високими показниками якості обслуговування (QoS). Основні два типи мобільності, це дискретна мобільність (роумінг) і безперервна мобільність (хендовер). Мережі LTE мають підтримувати процедури роумінгу і хендовера з усіма наявними мережами, для кінцевих споживачів LTE має забезпечуватися повсюдне покриття послуг бездротового широко-смугового доступу.

Пакетна передача даних дозволяє забезпечити всі послуги, включаючи передачу, призначену для кінцевого споживача голосового трафіка. У минулих поколіннях, спостерігалася досить висока розділена мережна відповідальність. Архітектуру мереж LTE можна назвати “плоскою”, оскільки практично вся мережна взаємодія відбувається між двома вузлами: базовою станцією (БС), яка в технічних специфікаціях називається В-вузлом (Node-B, eNB) і блоком управління мобільністю БУМ (ММЕ, Mobility Management Entity). Реалізаційна, як правило, включає і мережний шлюз (Gateway). Мають місце комбіно-

вані блоки (ММЕ / GW). БУМ відповідає лише за службову інформацію. IP – пакети, що містять інформацію про користувача, через нього не проходять. Основною перевагою такого окремого блоку є те, що пропускна здатність може незалежно нарощуватися для трафіка і для передавання службової інформації. Головною функцією БУМ є управління терміналами, що перебувають в режимі очікування, включаючи перенаправлення і виконання викликів, авторизацію і аутентифікацію, роумінг і хендовер, установаження службових і призначених каналів.

У стільникових мережах (3G) в основу побудови LTE закладено два типи: фізична реалізація блоків мережі та формування функціональних взаємозв'язків. У процесі фізичної реалізації можна пройти через поняття предметної області (домену), а функціональні зв'язки розглядатимуться в межах (stratum). Поділом архітектури мережі представлений фізичний рівень, на області обладнання кінцевого споживача (UED) і області мережної інфраструктури (ID). Устаткування користувача – це сукупність різних видів пристроїв з різноманітними функціональними можливостями. У сфері споживачів розроблено протоколи, що забезпечують передавання даних. У сфері управління є протоколи, які в різних аспектах гарантують присутність зв'язку між абонентом і мережею. Крім того, існують протоколи, призначені для простого передавання інформації так, що вона доступна для надання різних послуг. Область радіодоступу розділена на два рівні: рівень радіозв'язку (RNL) і рівень транспортної мережі (TNL).

Розглядаючи призначення функціональних блоків мережі радіодоступу побачимо, що на базові станції LTE покладено виконання багатьох функцій:

- управління радіоресурсами;
 - стиснення шифрування для IP-пакетів;
 - обирання БУМ при підключенні, за умови відсутності інформації про минулі підключення;
 - маршрутизація даних;
 - диспетчеризація повідомлень;
 - PWS – система тривожних сповіщень.
- БУМ забезпечує виконання таких функцій:
- передавання захищеної інформації і даних у БС;
 - управління зонами відстеження користувачів і БС, що перебувають у стані очікування та перенаправлення викликів;
 - обирання обслуговуючого, пакетного шлюзу мережі різних стандартів і нового БУМ під час виконання сесії;



- роумінг;
- аутентифікація;
- управління радіоканалом і підтримка передавання сповіщень PWS, включаючи виділені канали.

Вузол, що обслуговує, відповідає за такі функції:

- вибір локального місця розташування (LMA);
- буферизація даних для терміналів, які перебувають у стані очікування;
- санкціоноване отримання інформації від абонентів;
- маршрутизація даних і маркування пакетів на транспортному рівні;
- створення облікових записів та їхня ідентифікація для впровадження тарифікації.

Шлюз мережі (SG) забезпечує виконання функцій:

- фільтрація пакетів споживачів;
- санкціонований аналіз інформації споживачів;
- розподіл IP-адрес для терміналів;
- маркування транспортних пакетів напрямку;
- тарифікація послуг та їхня селекція.

У LTE визначено два стани рівня контролю радіоресурсу (RRC): з'єднання (RRC CONNECTED) й очікування (RRC IDLE). Функціонуючи, зі стану очікування у стан з'єднання, після успішного встановлення з'єднання, термінал може повернутися у стан очікування, розірвавши з'єднання на підрівні. Перебуваючи у стані очікування, система може моніторити канали виклику, отримувати інформацію, здійснювати вимірювання стільників і за необхідності змінювати їх. Щоб визначити графік проходження пакетів, термінал здійснює моніторинг каналів, для яких забезпечується зворотний зв'язок, що дає інформацію про поточний рівень якості. На відміну від стану очікування, у стані з'єднання мережею проводиться управління маневреністю.

Найважливішою особливістю всіх мереж, є підтримка безшовного зв'язку абонента щодо базових змін. Вимоги до ефективності підвищуються під час використання чутливих до затримок пакетів програм, таких як VoIP. В основі безшовного зв'язку лежать різні процедури хендовера. Тому для підготовки до виконання хендовера використовують сигналізацію з інтерфейсу X2, між різними БС. Ефективність виконання хендовера є одним з найважливіших показників якості роботи мереж. Погано відрегульовані параметри можуть привести до навантаження на службові канали та ведуть до втрат сеансів зв'язку. Опису алгоритмів хендовера присвячена велика кількість специфікацій (процедура хендовера детально описана у специфікації TS 36.413). На відміну від мереж GSM, де аналіз навколишнього оточення здійснюється контролером базових

станцій, у мережах LTE подібні дії довірено самому терміналу, а остаточне рішення про хендовер приймається мережею. Під час підготовки та виконання хендовера можуть бути встановлені тунелі: один – для передавання даних висхідного напрямку, інший — для передавання даних спадного напрямку. Це робиться для того, щоб забезпечити передачу довгих пакетів у разі переповнення буферів. У разі виконання процедури хендовера абоненту присвоюється тимчасовий ідентифікатор C-RNTI. Аналогічні ідентифікатори присвоюють абоненту під час проведення різних інших процедур, пов'язаних із мережею радіодоступу. Так, у механізмах управління потужністю TPC по фізичних каналах PUSCH і PUCCH використовують відповідні ідентифікатори TPC-PUSCH і TPC-PUCCH. Глобальна ідентифікація мереж LTE здійснюється за допомогою ідентифікатора ECGI, який формується додаванням до локального мережного ідентифікатора. Аналогічним чином здійснюється глобальна ідентифікація базових станцій. Для управління мобільністю, що перебуває у стані очікування, вводять поняття зони відстеження як площі, що покриває зону обслуговування декількох базових станцій.

Розглядаючи збір даних від абонента для аналізу ми бачимо, що оператор має можливість збору даних від абонента (терміналу), наприклад під час його підключення до мережі, зміні БС, реєстрації в мережі тощо. Також в оператора є можливість збору даних абонента через його підключення до інтернету, наприклад від різних джерел, таких як координати WI-FI або за допомогою GPS і сервісів. Існують як "легальні сервіси" (тобто з дозволу абонента), так і такі, які існують без відома абонента, але необхідні з технічного боку обладнання. HSS (Home Subscriber Server) – сервер абонентських даних LTE – являє собою велику базу даних і призначений для зберігання даних про абонентів. HSS змінює набір реєстрів (VLR, HLR, AUC, EIR) у мережах 2G і 3G.

HSS служить для зберігання такої інформації:

- інформація для доступу, аутентифікації і авторизації;
- інформація про місцезнаходження абонента на міжмережному рівні, тобто про те, в яку мережу перейшов у разі вхідного дзвінка;
- інформація про профіль користувача.

HSS генерує дані, необхідні для шифрування, аутентифікації та ін. Мережа LTE має один або декілька HSS. Кількість HSS залежить від географічної структури мережі та кількості абонентів. Присутня можливість пошуку абонентів за допомогою IMEI. Методи стеження за переміщенням у мережі, за допомогою:



- унікальних ідентифікаторів у мережі (номером телефона), приблизна точність 150 м;
- IMEI, полягає у визначенні номера та потім його місцезнаходження;
- через GPS-координати, метод дає найбільшу точність, але потребує додаток на смартфон із дозволом використання GPS-локації;
- локації по координатах Wi-Fi.

За оброблення інформації відповідає оператор, це змушує його зберігати, оброблювати великі об'єми даних, що у свою чергу потребує обладнання та кваліфікованого персоналу. Під час розроблення LTE розглядалась можливість застосування хмарних технологій і зменшення участі мобільного оператора в обробленні даних для надання послуг, передбачалось лише здійснення контролю за білінгом і підключення нових абонентів. Завдяки цьому оператор як власник усієї технологічної інфраструктури і контролер абонентів і ЦОД, втрачає власність. Такі зміни потребують перегляду в діючій моделі.

Під час оброблення даних від телефонного апарату на БС, відбувається переключення між базовими станціями, у середньому за день це відбувається 10 разів, а якщо в мережі 1 млн абонентів, то це 10 млн записів на день і це тільки записи, що надходять від абонентів і вони можуть складати десятки ГБ. Для оброблення таких об'ємів використовують методи, алгоритми та підходи, що називаються Big Data, як набір інструментів оброблення структурованих і неструктурованих великих даних. Big Data, використовують найбільш поширене визначення семи "V":

- Volume – об'єм даних;
- Velocity – оброблення з великою швидкістю;
- Variety – різноманітність і структурованість даних;
- Veracity – достовірність даних;
- Variability – мінливість і точність сприйняття контексту;
- Visualization – візуалізація;
- Value – цінність, витяг користі для збільшення результату.

Зараз більшість експертів сходяться на думці, що прискорення зростання обсягу даних є об'єктивною реальністю. Джерелами генерування великих обсягів інформації виступають соціальні мережі, мобільні пристрої, різноманітні цифрові пристрої та бізнес-інформація. За існуючих темпів зростання кількість даних у світі буде щорічно подвоюватися.

Наукова модель Big Data визначає такі завдання:

- зберігання та управління даними;
- обробка та структурування даних;

- аналітика накопичених даних (Data Mining) для їхнього корисного використання;

- візуалізація даних, із використанням інтерактивних можливостей та анімації.

Робота з великими даними повинна підпорядковуватися таким принципам:

- горизонтальна масштабованість, яка має легко розширюватися і бути пропорційною до цього зростання;
- відмовостійкість, що має на увазі методи роботи з великими даними, які повинні враховувати можливість збоїв і переживати їх без наслідків;
- локальність даних. Big Data принцип рішень – оброблення даних відбувається на тому пристрої, де вони зберігаються.

Однією з ключових технологій Big Data є платформа Hadoop. Apache Hadoop – це фреймворк, що дозволяє виробляти розподілене оброблення великих масивів даних на кластерах, що складаються із численних вузлів, використовуючи прості моделі програмування. Надійність забезпечується здатністю виявляти збої на програмному рівні. Система Apache Hadoop (HDFS) – це розподілена файлова система, яка створена для зберігання великого обсягу інформації з високою швидкістю доступу до інформації. У HDFS файли зберігаються в надлишкової формі, де використовується хмарно-структурована файлова система. Структура метаданих в HDFS змінюється великою кількістю одночасних клієнтів. Для збереження цілісності всі звернення обробляється однією машиною NameNode. NameNode зберігає всі метадані файлової системи. Для відкриття та оброблення даних за допомогою Hadoop використовують MapReduce програму від компанії Google. Оброблення даних відбувається у три стадії. Основним недоліком MapReduce є те, що зазвичай при обробленні даних відбуваються кілька послідовних і паралельних операцій, де всі результати записуються на диск. Завдяки чому час роботи з диском у кілька разів перевищує час самих обчислень. Для вирішення цієї проблеми розроблено програмні платформи – Spark та Tez. Spark обробляє дані в оперативній пам'яті, завдяки чому дозволяє отримувати значний вииграш у швидкості роботи. Tez представляє задачу у вигляді ациклічного графа (DAG).

Крім оброблення, важливим є візуальне відображення результатів для аналізу. Kibana – фреймворк для візуалізації великих обсягів даних, розташованих в Elasticsearch кластері. Elasticsearch – це пошуковий движок, який працює з json-документами та добре підходить для роботи з великими обсягами даних. Для відображення даних може використовуватися платформа Mapbox, яка



дозволяє створювати кастомізовані карти і виводити на них необхідну інформацію за допомогою вебінтерфейсу або із застосуванням API. Інформація береться з відкритих джерел, та немає обмежень на кількість звернень до геокодеру протягом доби з одної IP-адреси.

Системи хмарних обчислень та їхні можливості – це технологія, в якій комп'ютерні ресурси і потужності надаються як інтернет-сервіс. Основні чинники доступності технології:

- розвиток процесорів і технології літографії;
- збільшення розмірів дисків і зниження вартості зберігання інформації;
- розвиток технології з розподіленням потоків;
- розвиток технологій віртуалізації;
- збільшення пропускної здатності.

Перераховані фактори створили конкуренцію звичайним серверним обчисленням і поширили хмарні обчислення на маси. Основні недоліки – це постійне з'єднання з мережею, програмне забезпечення та його кастомізація, конфіденційність, надійність, безпека, вартість обладнання для побудови хмари. Але є і значні переваги, такі як послуги та сервіси, залежно від класифікації "хмар".

Нині виокремлюють три категорії "хмар":

Публічна хмара – це IT-інфраструктура, яка використовується одночасно безліччю компаній і сервісів. Користувачі не мають можливості управляти вказаною хмарою та обслуговувати її, уся відповідальність лежить на власнику. Абонентом може стати будь-яка компанія і користувач. Приклади: онлайн сервіси Amazon EC2 і Simple Storage Service (S3), Google Apps/Docs, Salesforce.com, Microsoft Office Web.

Приватна хмара – це IT-інфраструктура, контрольована й експлуатована в інтересах однієї організації. Організація може керувати приватною хмарою самостійно. Інфраструктура може розміщуватися у приміщеннях замовника, у зовнішнього оператора, або частково у замовника і частково в оператора.

Гібридна хмара – це IT-інфраструктура яка використовує якості публічної і приватної хмар. Такий тип хмар застосовують, коли наявні сезонні періоди активності.

Для реалізації системи потрібно таке.

- 1) Створити віртуальну машину генерації трафіка.
- 2) Налаштувати та запустити потік даних.
- 3) Запустити завдання потоку даних.
- 4) Отримати протікання даних типу Dataflow.
- 5) Підключити потік даних до Pub/Sub та BigQuery.
- 6) Перевірити, як автоматичне масштабування Dataflow налаштовує обчислювальні ресурси для оптимального оброблення вхідних даних.

7) Дослідити навантаження на систему в реальному часі.

Необхідно зареєструватись і створити віртуальну машину. Після успішного створення віртуальної машини перейти до запуску генератора даних.

Наступний крок – запуск генератора даних, що отримує дані у форматі CSV і передає їх на сервер через спеціальний потік Dataflow.

Необхідно зайти на сервер і підключити до проєкту, використовуючи команди:

```
$ gcloud auth application-default login
$ sudo pip install google-cloud-pubsub
$ virtualenv cpb104
$ source cpb104/bin/activate
$ pip install google-cloud-pubsub==0.38.0
$ gcloud auth application-default login
```

Після виконання команд, запустити генератор даних:

```
$ ./send_sensor_data.py--speedFactor=60--project=YOUR-PROJECT-ID
```

Необхідно відкрити ще одну консоль та запустити:

```
$ training-data-analyst/courses/streaming/process/run_oncloud.sh yourproject
yourbucket AverageSpeeds
```

Можна побачити навантаження, аналітику потоку даних і використання каналу.

На основі розглянутих технологій, можна дійти висновку, що мобільні оператори будуть використовувати технології мережі LTE та наступних поколінь, які розробляються з урахуванням хмарних технологій та їхніх можливостей. Система, яка дозволяє передавати та зберігати отримані дані від користувачів із можливістю одночасної фільтрації. Результатом є автоматизована система, що виконує з'єднання з хмарними сервісами, створює потік і передавання на хмарний сервер із можливістю перегляду даних у реальному часі.

4. ВИСНОВОК

При описі розглянуто структуру мобільного оператора, технологію LTE (структуру підключення абонентів, хендвера тощо), можливість отримання й аналізу даних від абонентів, методи відслідковування абонентів у мережі по унікальних ідентифікаторах, а також можливість перенесення інфраструктури оператора на безсерверні технології, тобто втрати позиції оператора як обов'язкового власника всієї технологічної інфраструктури.

Розглянуто питання оброблення даних, існуючі методи і технології оброблення даних Big Data, системи оброблення й аналізу великих да-



них операторами в Україні та системи відслідковування й аналізу трафіка абонентів.

Розглянуто комплексні рішення у вигляді безсерверних технологій, класифікація таких рішень, огляд існуючих сервісів. Також обрано найоптимальнішу систему хмарних обчислень.

На підставі проведеної роботи з дослідження безсерверних обчислень і їхнього використання в мережах мобільного зв'язку, можна дійти висновку, що впровадження хмарних технологій у мобільних мережах значно зменшить витрати оператора на обслуговування інфраструктури і допоможе краще зорієнтуватися в наданні якісних послуг.

- [20] Вікіпедія Хмарні обчислення (Електронний ресурс)
- [21] Облачные технологии и моё будущее (Електронний ресурс)
- [22] Вікіпедія Google Cloud Platform (Електронний ресурс)
- [23] Вікіпедія Amazon Web Service (Електронний ресурс)
- [24] Вікіпедія Microsoft Azure (Електронний ресурс)

Стаття надійшла до редколегії

23.06.2019

СПИСОК ВИКОРИСТАНОЇ ЛІТЕРАТУРИ

- [1] Гельгор А.Л. Технология LTE мобильной передачи данных: учеб. пособие / Гельгор А.Л., Попов Е.А. – СПб.: Изд-во Политехн. ун-та, 2011. – 204с.
- [2] Тихвинский В.О., Терентьев С.В., Минаев И.В./ Стандартизация, спецификации, эволюция технологии и архитектура базовой сети LTE // Сети и средства связи, No 2 (10), 2009.
- [3] ETSI TS 136 413 V12.3.0 (2014-09)LTE;Evolved Universal Terrestrial Radio Access Network (E-UTRAN); S1 Application Protocol (S1AP) (3GPP TS 36.413 version 12.3.0 Release 12)
- [4] Sustainable E-commerce: Improving Eco-Efficiency in E-Commerce
- [5] Навч. посібник/ І.М. Пістунів, О.П. Антонюк – / Мвоосвітні науки України; Нац. Гірни. ун-т. – Д.: НГУ, 2017. – 291 с.
- [6] Електроннаекономіка. Том 1. Криптовалюта. Big Data [Електронний ресурс]: Режим доступу: http://pistunovi.inf.ua/EE_KC_BD.pdf
- [7] Как установить и настроить Elasticsearch, Logstash, Kibana (ELK Stack) на Ubuntu/Debian/Centos (Електронний ресурс)/Режим доступу: <https://serveradmin.ru/ustanovka-i-nastroyka-elasticsearch-logstash-kibana-elk-stack/>
- [8] Мобильщики следят за украинцами (Електронний ресурс)/Режим доступу: <https://ubr.ua/market/telecom/kiyevstar-nachinaet-sledit-za-ukraintsami3881682>
- [9] <http://ursa-tm.ru/forum/index.php?topic/261106-avtonomnyy-sposob-obhoda-dpi-i-effektivnyy-sposob-obhoda-blokirovok-saytov-po-ip-adresu-goodbyedpi-i-reqrypt/>
- [10] Вікіпедія Кластерний аналіз (Електронний ресурс)/Режим доступу: <https://ru.wikipedia.org/wiki>
- [11] Nadoor: что, где и зачем (Електронний ресурс)/Режим доступу: <https://habr.com/ru/post/240405/>
- [12] 4G. Тренды и перспективы (Електронний ресурс)
- [13] LTE: взгляд изнутри (Електронний ресурс)
- [14] Сотовая связь HSS (Home Subscriber Server) (Електронний ресурс)
- [15] Эволюция современных сетей мобильной связи 2G/3G/4G (Електронний ресурс)
- [16] Вікіпедія Великі дані (Електронний ресурс)
- [17] Обзор алгоритмов кластеризации данных (Електронний ресурс)
- [18] Вікіпедія Deep packet inspection (Електронний ресурс)
- [19] Автономный способ обхода DPI и эффективный способ обхода блокировок сайтов по IP-адресу/Режим доступу:



Research of serverless computing and its use in mobile communication networks

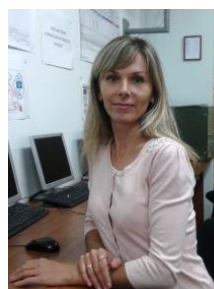
Serverless computing and their use in mobile networks are studied. Today it is impossible to imagine the whole world without calculations. They surround us in all spheres of life - from the usual payment by card in the supermarket to scientific calculations. The main problem of calculations is the growing number of requests and the complexity of providing services. Computing companies that need to constantly increase server bandwidth need to work on network infrastructure and hire professionals to support such systems. All the costs and complexity of developing server systems fall on the developer, which is a very expensive factor for new companies whose activities are related to computing and data processing. To solve the problems, the technology of "serverless computing" is being developed and applied, which is a model of cloud computing. This technology allows you to use resources that do not belong to them physically. All problems with setting up and maintaining the equipment are the responsibility of the service providers. Also, with intelligent load tracking, the system automatically allocates the necessary capacity and resources to perform calculations at the moment. The main problem with server-free computing technology is not the adaptation to the needs of many industries and the problems with certification and standardization, but they are gradually being solved. Research on the use of such technologies in telecommunications systems is a topical issue in the development of science and companies in the country. The structure of the mobile operator, LTE technology, the ability to receive and analyze data from subscribers, methods of tracking subscribers in the network by unique identifiers, as well as the ability to transfer the operator's infrastructure to serverless technologies. The issues of data processing, existing methods and technologies of Big Data data processing, systems of processing / analysis of big data by operators in Ukraine and systems of tracking and analysis of subscriber traffic are given. Complex solutions in the form of serverless (cloud) technologies, classification of such solutions, review of existing services are considered.

Keywords: serverless technologies; cloud technologies; LTE network; big data; mobile network operator.



Марина Завалій,
студентка, Державний університет телекомунікацій, Київ, Україна.

Marina Zavaliy,
student, State University of Telecommunications, Kyiv, Ukraine.



Лариса Дакова,
кандидат технічних наук, посада доцент. Державний університет телекомунікацій

Larisa Dakova,
Candidate technical sciences, position of associate professor. State University of Telecommunications



Антон Завалій,
студент, Державний університет телекомунікацій, Київ, Україна.

Anton Zavaliy,
student, State University of Telecommunications, Kyiv, Ukraine.



Сергій Даков,
кандидат технічних наук, асистент кафедри Київський національний університет імені Тараса Шевченка.

Serhii Dakov,
candidate of technical sciences, assistant department Taras Shevchenko National University of Kyiv.



Іван Пархоменко,
кандидат технічних наук, доцент кафедри кібербезпеки та захисту інформації Київського національного університету імені Тараса Шевченка.

Ivan Parkhomenko,
Candidate of Technical Sciences, Associate Professor of the Department of Cybersecurity and Information Protection, Taras Shevchenko National University of Kyiv.



УДК 004.415.056.5

DOI <https://doi.org/10.17721/ISTS.2020.4.71-80>

С. В. Толюпа, orcid.org/0000-0002-1919-9174, tolupa@i.ua
В. С. Наконечний, orcid.org/0000-0002-0247-5400, nvc2006@i.ua
В. Г. Сайко, orcid.org/0000-0002-3059-6787, vgsaiko@gmail.com
М. Котов, orcid.org/0000-0003-1153-3198, shunki2001@ukr.net
В. Солодовник, orcid.org/0000-0002-8844-4348, valeriamin6@gmail.com
Київський національний університет імені Тараса Шевченка, Київ, Україна

ВИКОРИСТАННЯ СИМЕТРИЧНИХ АЛГОРИТМІВ ШИФРУВАННЯ ДЛЯ ПЕРЕДАЧІ СИГНАЛІВ У ПРИСТРОЯХ БЕЗДРОТОВОГО ВВЕДЕННЯ ДАНИХ

Нині існує безліч обчислювальних систем, які призначені для покращення, полегшення та вдосконалення життя людини. Оскільки з'являється все більше дослідників, які зацікавлені у сфері оброблення інформації, розвиток обчислювальних систем щороку набирає обертів. З розвитком інформаційних систем, не менш швидкими темпами розвиваються і можливі загрози, такі як порушення конфіденційності, цілісності та доступності інформації, що обробляється. З метою запобігання можливим втратам новітні системи захисту інформації постійно оновлюють і вдосконалюють. Оскільки неможливо створити абсолютно захищену систему, завжди є імовірність викрадення даних, тому проблема захисту інформаційних і телекомунікаційних систем набирає все більшої актуальності. Враховуючи те, що жодний захист не може бути досконалим, розроблено спосіб значного зменшення порушення конфіденційності даних. На сьогодні криптосистеми найбільше використовуються для захисту інформаційно-телекомунікаційних систем та інших технологій, зокрема і для захисту надважливої інформації держави, підприємства, особи або інших критично важливих даних, зокрема корпоративних таємниць, розвідувальних даних чи комерційних таємниць.

Презентовано способи використання симетричних алгоритмів для передавання сигналів у пристроях дистанційного введення даних. Подано інформацію про існуючі алгоритми шифрування й використання хеш-функцій. Проведено розмежування між одноключовим і двоключовим методами шифрування інформації. Детально розглянуто алгоритм AES його функції, робота раундів, схеми шифрування даних. Опис кожного алгоритму супроводжується прикладом, в якому пояснюються особливості його використання. Представлено математичну модель і приклад роботи блочного алгоритму шифрування. Висвітлено принципи роботи бездротових пристроїв. Розглянуто вразливості, пов'язані з бездротовими приладами та запропоновано рішення щодо їхнього захисту.

Ключові слова: криптографія, шифрування, симетричне шифрування, блокове шифрування, шифрування сигналів, пристрої бездротового введення даних.

1. ВСТУП

Нині існує безліч обчислювальних систем, які призначені для покращення, полегшення та вдосконалення життя людини. Оскільки з'являється все більше дослідників, які зацікавлені у сфері оброблення інформації, розвиток обчислювальних систем щороку набирає обертів. З розвитком інформаційних систем, не менш швидкими темпами розвиваються і можливі загрози, такі як порушення конфіденційності, цілісності та доступності інформації, що обробляється. З метою

запобігання можливим втратам новітні системи захисту інформації постійно оновлюють і вдосконалюють. Оскільки неможливо створити абсолютно захищену систему, завжди є імовірність викрадення даних, тому проблема захисту інформаційних і телекомунікаційних систем набирає все більшої актуальності. Враховуючи те, що жодний захист не може бути досконалим, розроблено спосіб значного зменшення порушення конфіденційності даних.

© Толюпа С. В., Наконечний В. С., Сайко В. Г., Котов М., Солодовник В., 2020



Можливими способами захисту інформації є стеганографія та криптографія. Основна ціль стеганографії – приховування факту передавання чи зберігання інформації, а криптографії – приховування вмісту переданої інформації. Використовуючи сучасні криптографічні алгоритми шифрування інформації, можливо зберегти її конфіденційність навіть за умови несанкціонованого доступу до неї.

На сьогодні криптосистеми найбільше використовують для захисту інформаційно-телекомунікаційних систем та інших технологій, зокрема і для захисту надважливої інформації держави, підприємства, особи або інших критично важливих даних, зокрема корпоративних таємниць, розвідувальних даних чи комерційних таємниць.

2. ПОСТАНОВКА ПРОБЛЕМИ

Як відомо надійність криптосистем залежить від виконання таких вимог: зберігання ключів у таємниці, генерації псевдовипадкових чисел, алгоритм проведення шифрування тощо [1–5].

У розвиток криптографії значний внесок зробили в своїх працях такі автори: Кузнецов О. О., Горбенко І. Д., Кавальчук Л. В., Dan Boneh, Victor Shoup, Mohamed Barakat, Christian Eder, Timo Hanke, Steven D Galbraith, William Stallings [1–4]. Активно описували симетричне шифрування у своїх наукових працях Joseph Sterling Grah, Nigel Smart, Christof Paar, Jan Pelzl [7, 9, 11].

Питання застосування та функціонування шифру AES свого часу розглядали в своїх роботах Joan Daemen, Vincent Rijmen [12, 13]. Результати досліджень блочних алгоритмів шифрування викладено у працях А. В. Яковлева, А. А. Безбогова, В. В. Родіна, В. Н. Шамкіна, Roberto Avanzi, Bruce Schneier, Lars R, Knudsen Matthew, J.B. Robshaw [16, 18, 19, 20].

Останні події у світі підтвердили особливу важливість протидії порушникам і спробам ведення інформаційної боротьби з метою нанесення втрат. Достатньо детальний аналіз також підтвердив, що поряд з іншими методами і механізмами захисту інформації важливими є такі, що базуються на криптографічних перетвореннях інформації. Дійсно, при правильному виборі та застосуванні, відповідно до вимог криптографічних перетворень, досягається високий рівень безпеки, а також конфіденційність, цілісність, захист від НСД, доступність, неспростовність тощо.

У випадку криптографічного захисту інформації висуваються та повинні бути забезпечені криптографічна стійкість, цілісність, швидкодія криптографічних перетворень і вимоги, що висуваються додатками.

Зрозуміло, що криптографічна стійкість є безумовною вимогою, але поряд із нею швидкодія вже є також безумовною вимогою – потрібно захищати канали зі швидкістю від сотень мегабіт до десятків гігабіт за секунду, виконуючи операції в реальному часі. Указана вимога може бути виконана у разі застосування криптографічних перетворень типу блоковий симетричний шифр.

Тому проблема захисту інформації під час передавання сигналів у пристроях бездротового введення даних (ПБВД) є надзвичайно актуальною. Саме висвітленню та вирішенню цього питання й присвячена ця робота.

3. ВИКЛАД ОСНОВНОГО МАТЕРІАЛУ

Криптографія – наука, що розробляє методи використання складних математичних перетворень для приведення інформації, яка передається через канали поширення, у такій формі, в якій її не зможе отримати ніхто крім авторизованих і дозволених осіб.

Шифрування – ключовий об’єкт дослідження криптографів, це процес, під час якого змінюється форма інформації, що передається через відкриті канали передачі [1, 2].

Алгоритм шифрування – це набір зворотних математичних перетворень, що надають тексту шифрованого вигляду [1].

Існують три базові види шифрування даних [1]:

- шифрування з використанням двох ключів;
- шифрування з використанням одного ключа;
- безключове шифрування даних.

Принципова різниця між безключовими, одноключовими та двоключовими алгоритмами полягає у процесі шифрування вхідної інформації.

Безключові системи не використовують ключі у процесі криптографічного перетворення інформації. До безключових криптосистем належать хеш-функції та генератори псевдовипадкових чисел [5].

Хеш-функція – це функція що здійснює перетворення вхідного масиву даних довільної величини в бітовий рядок фіксованої довжини, та виконується за допомогою певного алгоритму [6, 7].

Як правило хеш-функції використовують [6, 7]:

- для побудови унікальних ідентифікаторів;
- для обчислення контрольних сум від даних для подальшого виявлення в них помилок, що виникають при зберіганні або передаванні даних;
- для пошуку дублікатів в серіях наборів даних;
- для побудови асоціативних масивів;
- для збереження паролів у системах захисту у вигляді хеш-коду. Відновлення такого пароля потребує функцію, яка буде зворотною до тої, що була використана для хешування паролів.

Одноключові системи у процесі криптографічного перетворення інформації використовують криптографічний ключ для шифрування та розшифрування даних. Ключ шифрування може складатись із чисел, слів або символів. Оскільки кожен, хто має доступ до ключа, може розшифрувати інформацію, то такий ключ має залишатися в таємниці, бути відомим лише для відправника й отримувача, щоб забезпечити стійкість шифрування [5].

Схему використання одноключової системи показано на рис. 1.

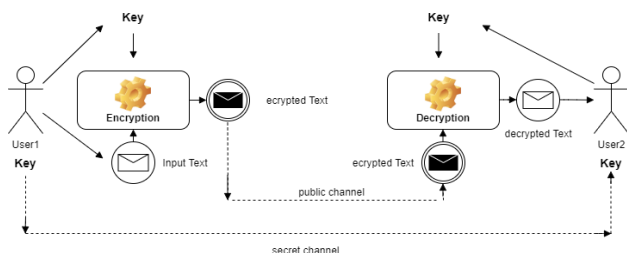


Рис. 1. Одноключове шифрування інформації

Двоключове шифрування інформації – спосіб шифрування даних, при якому використовуються два ключі шифрування: публічний і приватний. Головне досягнення асиметричного шифрування полягає у тому, що необхідність передавання ключа через спеціальний засекречений канал відпадає [5].

Схему використання асиметричної системи показано на рис. 2.

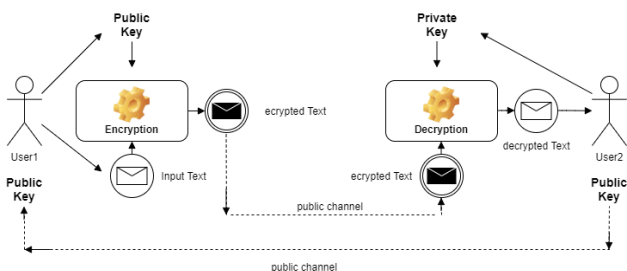


Рис. 2. Асиметричне шифрування інформації

Загалом суть асиметричних алгоритмів полягає в такому [5]: отримувач генерує два ключі – публічний і приватний. Після цього отримувач передає публічний ключ тому, хто надсилає повідомлення, а приватний залишає собі. Отримавши публічний ключ, перший користувач зашифровує інформацію, використавши цей ключ, та надсилає зашифрований текст.

Розшифрувати це повідомлення зможе лише той, хто має приватний ключ. За допомогою пуб-

лічного ключа розшифрувати дане повідомлення – неможливо. Саме тому не потрібно використовувати засекречені канали зв'язку для передавання ключа.

Як зазначалося вище, алгоритми симетричного шифрування використовують лише один приватний ключ для шифрування вхідних даних. Якщо алгоритм шифрування достатньо криптостійкий, то єдиний спосіб зломиснику розшифрувати інформацію – отримати секретний ключ.

Криптостійкість алгоритму шифрування визначається складністю розшифрування тексту методом грубої сили, тобто перебиранням усіх можливих комбінацій ключів. Алгоритм вважається нестійким, якщо вдалося перебрати половину усіх можливих комбінацій ключа. Ключ довжиною 128 біт, сучасні комп'ютери зможуть зламати лише через мільярди років, у свою чергу ключі довжиною 256 біт вважаються безпечними, теоретично спроможні встояти проти грубої атаки квантових комп'ютерів [8–11].

Як правило, для систем криптографічного захисту інформації, що практично використовуються, довжина повідомлення, що захищається за допомогою симетричного блокового шифру, значно перевершує довжину ключа шифрування. У цьому випадку не виконується критерій безумовної стійкості шифру, що використовується, і в таких умовах доцільне введення поліноміального критерію, що припускає наявність обмежень для обчислювальних ресурсів зломисника та часу, протягом якого шифр залишається стійким. Поліноміальний критерій приводить до практичного критерію стійкості – неможливості реалізації атаки на шифр в умовах сучасної обчислювальної бази протягом тривалого строку.

Додатково, з огляду на можливість удосконалення криптоаналітичних методів, вводиться критерій "запасу стійкості" до аналітичних атак – складність атаки на весь алгоритм має бути значно вищою складності силових атак.

Зазвичай цей критерій розглядає версію симетричного блокового алгоритму шифрування зі зменшеною кількістю циклів, що є уразливою проти криптографічного аналізу. Різниця в кількості циклів визначає запас стійкості алгоритму до конкретної криптоаналітичної атаки.

Для оцінювання криптографічної стійкості загальної конструкції шифру доцільно ввести ще один критерій, що розглядає можливість виключення яких-небудь операцій або заміни їх менш складними операціями (наприклад, на деяких наборах вхідних даних операція додавання за модулем 2^{32} близька або еквівалентна операції



додавання за модулем 2). У цьому випадку повноцикловий варіант спрощеного шифру повинен залишатися стійким до аналітичних атак.

Необхідно також враховувати, що більшість сучасних аналітичних атак, насамперед, таких як диференціальний і лінійний криптоаналіз, є статистичними. У процесі криптоаналізу для одержання ключа виконується велика кількість шифрувань і на підставі шифротекстів формуються варіанти підключів. Під час оброблення досить великої вибірки шифротекстів, сформованих на одному ключі, правильне значення ключових бітів зустрічається частіше інших варіантів.

Очевидно, що ймовірність знаходження правильної пари (що пропонує коректне значення ключа) залежить від статистичних властивостей шифру, і для збільшення складності криптоаналізу властивості криптограми повинні бути близькі до властивостей випадкової послідовності. Тому необхідною (але не достатньою) умовою стійкості шифру до аналітичних атак є забезпечення гарних статистичних властивостей вихідної послідовності (шифротекстів).

Для захисту шифру від алгебраїчних атак необхідно, щоб не існувало способу практичної побудови системи рівнянь, що зв'язують відкритий текст, криптограму та ключ шифрування, або не існувало способу розв'язання таких систем у поліноміальний час.

У побудові засобів криптографічного захисту необхідно враховувати можливість організації атак на реалізацію (зміна температурного режиму електронного пристрою, вхідної напруги, поява іонізуючого випромінювання, вимірювання струмів, що споживаються, ПЕМІН, час виконання тощо).

Такі атаки можуть бути ефективні проти всіх криптографічних алгоритмів, і захист від таких атак вимагає інженерних рішень уже на етапі проектування засобів криптографічного захисту інформації.

4. РЕЗУЛЬТАТИ

Нині є два види симетричного шифрування [11]:

- блочний;
- поточний.

Свої назви, ці види дістали від способу оброблення вхідного тексту.

Блочні алгоритми шифрують текст блоками по 128 біт, шляхом застосування алгоритмів шифрування, до яких у свою чергу додається ключ довжиною 128, 192 або 256 біт.

Схему блочного шифру показано на рис. 3.

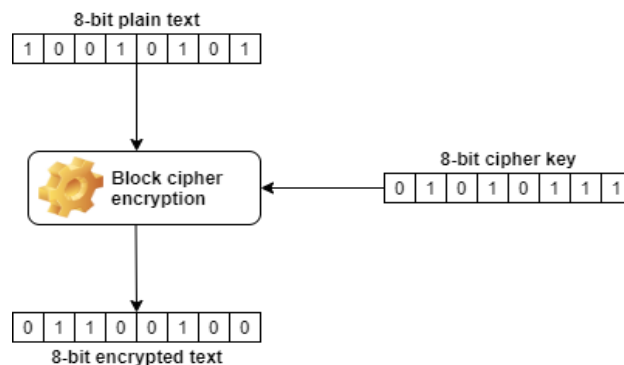


Рис. 3. Схема блочного шифру

Поточні шифри, у свою чергу, працюють із кожним бітом вхідного тексту та видають по біту вихідного, шифрованого тексту.

Схему поточного шифру показано на рис. 4.

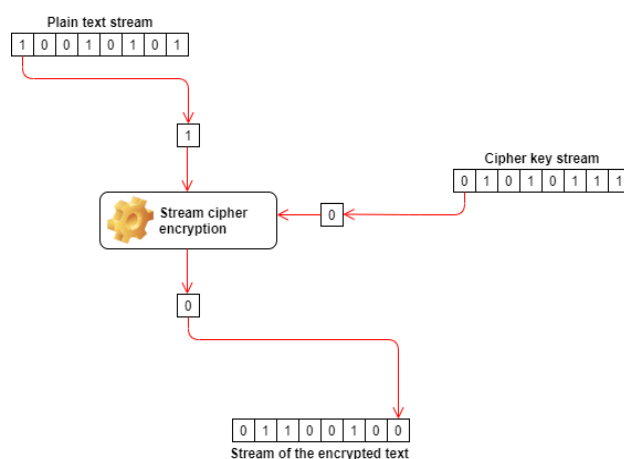


Рис. 4. Схема поточного шифру

Схема роботи AES показана на рис. 5.

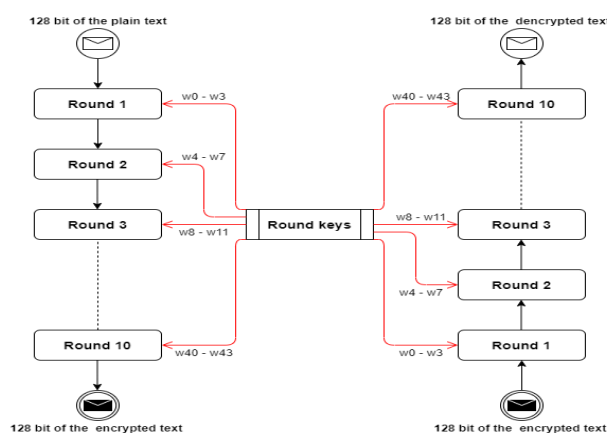


Рис. 5. Схема шифрування алгоритму AES

Відповідно до схеми на рис. 5 шифрування та дешифрування тексту відбувається протягом



десяти раундів, на кожному з раундів додається ключ раунду.

Нехай маємо 128 біт вхідного тексту A та 128 біт раундового ключа K , перед першим раундом, спочатку виконується операція $AddRoundKey$, це операція XOR вхідного тексту з ключем, тобто:

$$B = A \text{ XOR } K,$$

де B – 128 біт тексту після операції $AddRoundKey$ [12].

Кожний раунд шифрування та дешифрування AES складається з чотирьох перетворень вхідної матриці [12]:

1. Для раунду шифрування виконують такі перетворення:

1) $SubBytes$ – побайтова підстановка в S-BOX, з фіксованою таблицею заміन.

2) $ShiftRows$ – побайтове зрушення рядків матриці $State$.

3) $MixColumns$ – перемішування байтів у колонках.

4) $AddRoundKey$ – додавання ключа раунду.

2. Для раунду дешифрування виконуються такі перетворення:

1) $InvShiftRows$ – зворотна операція до $ShiftRows$.

2) $InvSubBytes$ – зворотна операція до $SubBytes$.

3) $AddRoundKey$ – додавання ключа раунду.

4) $InvMixColumns$ – зворотна операція до $MixColumns$.

Останній раунд шифрування відрізняється від інших тим, що не активізує перетворення $MixColumns$ [13].

На рис. 6 представлено схему перетворень у кожному раунді шифрування та дешифрування.

Відповідно до принципу Керкгофса, оцінювати криптографічну стійкість алгоритму потрібно з урахуванням того, що зломисник знає всі процедури перетворення вхідного тексту, які виконуються даним алгоритмом. Тобто в секреті залишається лише такий параметр – ключ шифрування [14, 15].

Зі сторони зломисника, довжина ключа визначає доцільність виконання повного перебору всіх можливих його комбінацій, оскільки інформацію, що зашифрована довшим ключем, набагато складніше зламати атакою грубої сили.

Для прикладу візьмемо ключ довжиною 4 біти, використовуючи основні закони комбінаторики, можливо обчислити всі можливі варіанти перебору: $S = 2 * 2 * 2 * 2 = 16$. Таким чином, кількість можливих варіантів перебору становить $S = 16$.

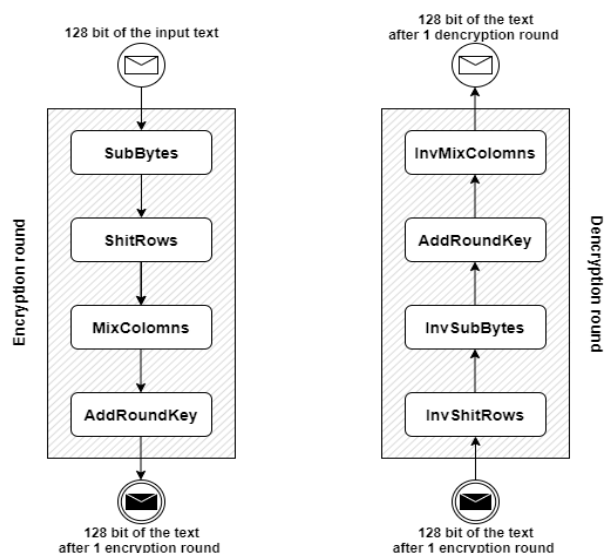


Рис. 6. Схема раундів алгоритму AES

Нехай маємо 10 000 машин, кожна перевіряє 1 000 000 комбінацій ключа шифрування за секунду, алгоритм вважається зламаном, якщо вдалось перебрати половину варіантів, на цій основі виконано табл. 1.

Таблиця 1

Таблиця переборів ключів

Key size (bit)	Combinations	Time to check all combinations
1	2	$2 * 10^{(-10)}$ sec
2	4	$4 * 10^{(-10)}$ sec
4	16	$1.6 * 10^{(-9)}$ sec
8	256	$2.56 * 10^{(-8)}$ sec
16	65 536	$6.5536 * 10^{(-6)}$ sec
32	4 294 967 296	0.4294967296 sec
56 (DES)	$7.20575940 * 10^{16}$	83.3999930995 days
64	$1.84467441 * 10^{19}$	58.4942417355 years
128 (AES)	$3.40282367 * 10^{38}$	$1.07902831 * 10^{21}$ years
192 (AES)	$6.27710174 * 10^{57}$	$1.99045590 * 10^{40}$ years
256 (AES)	$1.15792089 * 10^{77}$	$3.67174306 * 10^{59}$ years

Оцінивши результати обчислень з табл. 1, можна зробити висновок, що ключ довжиною 128 біт надійно шифрує інформацію. Ключ довжиною 256 біт – теоретично стійкий до лобової атаки квантових комп'ютерів.

Блочним типом шифрування є система шифрування, яка залежно від алгоритму, обраного ще на початку своєї роботи, використовує його у кожному такті своєї роботи.

Відомо, що блочний алгоритм шифрування, по суті є шифром заміни, проте основна його особливість полягає в тому, що цей шифр має величезний алфавіт. Кожний алфавітний знак



представляє собою певну послідовність двійкових чисел 0 і 1, тобто двійкових даних, визначеного розміру [16–21].

Щоб краще уявити, як працює ця система, візьмемо блок розміру N (послідовність 0 і 1): $p = (p_0, p_1, p_2, \dots, p_{N-1}) \in Z_{2,N}$, де p в $Z_{2,N}$ – це вектор, представлений у вигляді фіксованої довжини 0 і 1. Тоді двійкове подання числа ми отримаємо, підставивши дані у формулу

$$|p| = \sum_{i=0}^{N-1} 2^{N-i-1} p_i. \quad (1)$$

Тоді очевидним стає те, що стійкість шифру напряму залежить від розмірів таблиць заміни, які неможливо фізично представити через величину їх обсягів.

Створимо таку таблицю. Для прикладу візьмемо значення $N = 5$, тоді табл. 1 виглядатиме таким чином:

Таблиця 2

Таблиця заміни

(0,0,0,0,0)→0	(0,0,0,0,1)→1	(0,0,0,1,0)→2	(0,0,0,1,1)→3
(0,0,1,0,0)→4	(0,0,1,0,1)→5	(0,0,1,1,0)→6	(0,0,1,1,1)→7
(0,1,0,0,0)→8	(0,1,0,0,1)→9	(0,1,0,1,0)→10	(0,1,0,1,1)→11
(0,1,1,0,0)→12	(0,1,1,0,1)→13	(0,1,1,1,0)→14	(0,1,1,1,1)→15
(1,0,0,0,0)→16	(1,0,0,0,1)→17	(1,0,0,1,0)→18	(1,0,0,1,1)→19
(1,0,1,0,0)→20	(1,0,1,0,1)→21	(1,0,1,1,0)→22	(1,0,1,1,1)→23
(1,1,0,0,0)→24	(1,1,0,0,1)→25	(1,1,0,1,0)→26	(1,1,0,1,1)→27
(1,1,1,0,0)→28	(1,1,1,0,1)→29	(1,1,1,1,0)→30	(1,1,1,1,1)→31

Позначимо, що блоковий шифр виглядає таким чином:

$$\pi \in \text{SYM}(Z_{2,N}); \pi: p \rightarrow q = \pi(p),$$

де $p = (p_0, p_1, p_2, \dots, p_{N-1})$, $q = (q_0, q_1, q_2, \dots, q_{N-1})$.

Як уже зазначалося, блоковий шифр є окремим випадком підстановки, тільки він має великий алфавіт, проте цей вид шифрів розглядається окремо, адже зараз вони широко використовуються і їх набагато легше і зручніше описати за допомогою алгоритмів.

Усі криптосистеми реалізовані за допомогою блокових шифрів, адже методика їхнього використання полягає в тому, що створюється ланцюг з уже зашифрованих даним алгоритмом байтів. Це дозволяє шифрувати інформацію, пакетний об'єм якої є необмеженим [16].

Вихідний текст шифрується за допомогою підстановки π , яку обирають із повної симетричної групи.

Тоді зловмисник, який займається проведенням атаки і робить спроби для знаходження відповідності між зашифрованим і вихідним текстом $p_i \leftrightarrow q_i$, $0 \leq i \leq m$, не може, навіть маючи ці відомості про шифр, визначити початковий текст, який має таку відповідність: $q \notin \{q_i\}$ [20].

У цьому випадку правильно буде стверджувати, що блочний шифр є окремим випадком моноалфавітної підстановки, причому його алфавіт можна записати так:

$$Z_2^N = Z_{2,N}. \quad (2)$$

$\Pi[K]$ – ця підмножина є підмножиною симетричної групи $\text{SYM}(Z_{2,N})$ і ключовою системою блочних шифрів. Її індексація відбувається за допомогою параметра $k \in K$, де k є ключем, а K – простір ключів.

Математичним записом $\Pi[K]$ є вираз

$$\Pi[K] = \{\pi\{k\} : k \in K\}. \quad (3)$$

Використовуючи принцип побудови таблиці при $N = 5$, розглянемо інший випадок. Нехай $N = 64$ і в такому разі кожен елемент $\text{SYM}(Z_{2,N})$ потенційно може розглядатися як підстановка, так, щоб $K = \text{SYM}(Z_{2,N})$.

Таким чином звідси випливає, що

– 2^{64} кількість блоків, розрядність яких – 64, проте атакуючий не має жодної можливості для підтримання каталогу, розмір якого становить $2^{64} \approx 1,8 \cdot 10^{19}$ рядків;

– спроба отримати ключ за кількості ключів рівній $(2^{64})!$ неможлива.

Неподільність систем шифрування з блочними шифрами пояснюється тим, що в них наявний алфавіт $Z_{2,64}$, простір ключів $K = \text{SYM}(Z_{2,64})$ і найголовніше те, що розмір каталогу частот появи символів для 64-розрядних блоків за умови того, що кількість ключів дорівнює 2^{64} або ж спроба отримати ключ зловмисником, виходить за межі його можливостей. У цей самий час атакуючий стикається з проблемою недостатньої кількості ресурсів, а це обмежує його можливості, тим самим забезпечуючи правильну роботу системи. Ця ситуація дуже схожа на ту, коли зловмисник намагається зробити криптоаналіз текстових даних.

Прикро усвідомлювати той факт, що і порушник і розробник мають однакову проблему. Причини цієї проблеми пояснюють такими факторами їхнього спричинення [16]:

- сам розробник не має можливості створити систему, де він міг би реалізувати $2^{64}!$ підстановок $\text{SYM}(Z_{2,64})$;
- атакуючий у свою чергу також не може перебрати всі ключі із цієї групи.

Формулювання вимог до блочного шифру можна записати так [16]:

- N – має мінімально дорівнювати 64, а ще краще бути більше за 64. Це робиться для ускладнення створення каталогу;
- простір ключів повинен бути якомога більшим, щоб виключити можливість його підбору;
- $\pi\{k, p\}: p \rightarrow y = \pi\{k, p\}$ – співвідношення вихідного і шифрованого тексту мають бути складними, щоб не можна було застосувати статистичні або аналітичні методи аналізу, на основі відповідності вихідного тексту або ключа.

Отже, бачимо, що використовуваний криптоалгоритм, за допомогою якого шифруються дані, повинен бути ідеально стійким, а це можна реалізувати тільки, якщо читання вихідних даних можливе за умови перебору всіх варіантів ключів, а доти повідомлення не буде таким, яке можна прочитати.

Якщо звернутися до теорії імовірності, тоді шуканий ключ, буде знайдений з імовірністю 0.5 після того, як буде здійснено перебір половини всіх ключів. У цьому випадку на злам такого криптоалгоритму, довжина ключа якого N , потрібно буде в середньому 2^{N-1} перевірок. Із цього випливає, що стійкість блочного шифру залежить виключно від довжини ключа, і це дає можливість експоненційного зростання стійкості шифру одночасно зі зростанням довжини ключа. І навіть, якщо можна було б припустити, що перебір ключів буде здійснюватися на спеціальній багатопроекторній системі, то на злам ключа довжиною 128 біт, знадобиться занадто багато часу. Із цього факту випливає, що розшифрування методом лобової атаки стає не раціональним.

Звісно, усе це стосується виключно ідеальних за стійкістю шифрів, які неможливо реалізувати через різноманітні фізичні чинники [18].

Використання шифрування сигналів для пристроїв введення даних. Нині існує безліч пристроїв дистанційного введення даних. Під приладом дистанційного введення даних мається на увазі – пристрій бездротової передачі введених даних (ПБПВД) до обчислювальної машини. Прикладом таких приладів є: бездротові клавіатури, або миші, бездротова гарнітура, пристрої сенсорного введення даних тощо.

Пристрої ПБПВД передають інформацію за допомогою радіохвиль частотою від 27 MHz до 2.4 GHz. ПБПВД виконують свою функцію за допомогою двох приладів: передавача та приймача [22, 23].

Основною проблемою пристроїв, які передають інформацію на частоті 2.4 GHz є те, що немає єдиного стандарту щодо захисту інформації, яка передається.

У разі підключення таких пристроїв як бездротова комп'ютерна миша, автентифікація пристрою не використовується, що може призвести до перехоплення управління курсором.

У випадку бездротової клавіатури у деяких моделях використовується шифрування сигналів. Однак у більшості випадків захисту сигналів бездротових пристроїв не надається значної уваги. Відтак, наприклад, шифрування сигналів комп'ютерної миші, зазвичай не використовується, шифрування сигналів клавіатури також не проводиться. Завдяки чому зловмисники мають змогу перехопити сигнали, що може надати їм таку інформацію: особисті дані; дані банківських карт; логіни та паролі облікових записів; зміст листів особистого характеру; дані про стан здоров'я (під час онлайн консультацій лікарів, або ведення медичних карток онлайн); дані про фінансовий стан; дані про політичні вподобання; дані про короткострокові та довгострокові плани; інформацію про наукові розробки.

Навіть у разі шифрування сигналів клавіатури, отримання даних усе ще можливо. Прикладом методу отримання даних є MouseJack. Перед цією атакою не зможуть встояти товари навіть таких корпорацій-гігантів як Dell, Logitech, Microsoft, HP, Amazon, Gigabyte та Lenovo [24]. Ця атака може здійснюватися на відстанях до 100 м, а все, що потрібно мати зловмиснику, це USB-dongle [24–27].

Суть цієї атаки така: і бездротова клавіатура, і бездротова миша використовують два прилади для передавання сигналів (приймач та передавач) у ролі приймача виступає USB-dongle. Деякі моделі клавіатур шифрують сигнали, але комп'ютерні миші найчастіше цього не роблять.

У випадку з клавіатурами це працює таким чином [24]: ключ шифрування має лише USB-dongle, тобто тільки він має можливість розшифрувати сигнал, що несе в собі інформацію про введені клавіші, зловмисник, навіть якщо перехопить цей сигнал, то розшифрувати його не зможе.

На рис. 7 показано схему роботи 1–3 фаз указаної атаки.

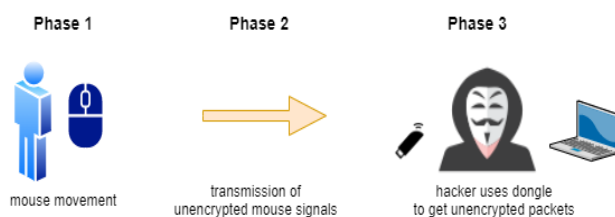


Рис. 7. Атака MouseJack (1–3 фази)



У ході перших трьох фаз відбувається таке: користувач задає зміну (x, y) координат положення миші, у свою чергу передавач, що міститься в миші, передає сигнал до USB-dongle, не шифруючи його, за допомогою радіосигналів. Зловмисник, використовуючи власний, налаштований USB-dongle, перехоплює незашифровані сигнали.

На рис. 8 зображено 4–6 фази атаки MouseJack [24].

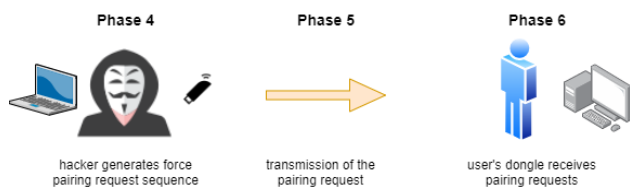


Рис. 8. Атака MouseJack (4–6 фази).

Виконуючи 4–6 фази, зловмисник надсилає послідовність запитів на підключення до USB-dongle користувача, у свою чергу USB-dongle користувача отримує послідовність цих запитів і під'єднується до пристрою користувача [24].

На рис. 9 зображено 7–9 фази атаки MouseJack.

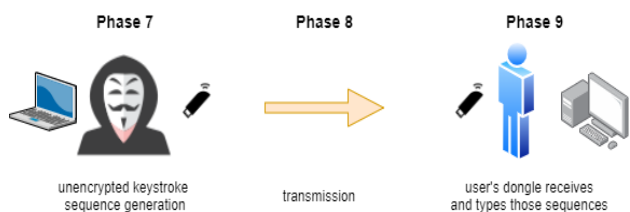


Рис. 9. Атака MouseJack (7–9 фази)

Під час виконання 7–9 фази зловмисник надсилає послідовність символів на комп'ютер користувача. Після виконання всіх зазначених фаз зловмисник має значні можливості маніпулювання над комп'ютером користувача [24].

Використання цієї вразливості можливе на всіх операційних системах, оскільки ця вразливість не відноситься до вразливостей приймача.

Застосовуючи дану вразливість, зловмисник має змогу встановити зловмисне програмне забезпечення, отримати доступ до даних комп'ютера, видаляти дані, порушувати їхню цілісність або доступність.

Отже способом захисту від цієї вразливості є шифрування сигналів ПБПВД, використовуючи надійні алгоритми шифрування, такі як розглянутий вище алгоритм AES. Слід установити однозначну ідентифікацію пристроїв, що намагаються відправляти пакети до USB-dongle.

Недоліком використання шифрування для сигналів ПБПВД є збільшення часу відгуку, тобто часу між введенням даних користувачем та їхнім відображенням, а також зростання ціни приладів.

5. ВИСНОВКИ

Захист даних – першочергова справа для професіоналів у сфері інформаційної та кібербезпеки. Серед різноманіття атак і спроб вилучити конфіденційну інформацію, фахівцям усе важче досягати ефективного захисту критично важливої інформації. Не полегшують ситуацію і нові технології, які створені з одного боку для покращення функцій захисту, а з іншого – нею так само користуються і порушники, намагаючись зазіхнути на конфіденційність, цілісність і доступність даних у системі.

Ця стаття написана для ознайомлення з методами перехоплення та ранжування інформації під час введення її з пристроїв бездротового введення даних.

У роботі дано загальну характеристику алгоритму шифрування, як засобу захисту інформації. Прописано найбільш розповсюджені види шифрувань і відомості про них. Наведено приклади застосування таких методів під час захисту даних.

Ця робота описує симетричні алгоритми шифрування, а саме алгоритм AES, як найбільш широко застосованого алгоритму на теперішній час. Прописано його переваги, недоліки та сферу застосування.

Для порівняння розглянуто блочний алгоритм шифрування та побудовано його математичну модель. У цій частині роботи алгоритм розглядається як такий, що має певні переваги, а саме надійність. Проте він незначно поступається швидкості реалізації порівняно з методом поточного шифрування даних.

Висвітлення проблематики наукової роботи чітко розкриває суть необхідності застосування шифрування під час використання пристроїв дистанційного введення даних.

СПИСОК ВИКОРИСТАНОЇ ЛІТЕРАТУРИ

- [1] Dan Boneh, Victor Shoup, A Graduate Course in Applied Cryptography, version 0.4, Stanford University, September 2017.
- [2] Mohamed Barakat, Christian Eder. An Introduction to Cryptography, September 20, 2018.
- [3] Steven D Galbraith, Mathematics of Public Key Cryptography, version 2.0,
- [4] William Stallings, Cryptography and Network Security Principles and Practices, Fourth Edition, Prentice Hall, November 16, 2005/
- [5] Information on <https://www.pvsm.ru/algoritmy/225093#begin>.



- [6] Information on <https://ru.wikipedia.org/wiki/%D0%A5%D0%B5%D1%88-%D1%84%D1%83%D0%BD%D0%BA%D1%86%D0%B8%D1%8F>.
- [7] Joseph Sterling Grah, Hash Functions In Cryptography, The University of Bergen, June 1, 2008.
- [8] Information on <https://www.binance.vision/security/what-is-symmetric-key-cryptography>.
- [9] Nigel Smart, Cryptography: An Introduction, 3rd edition, McGraw-Hill College, December 30, 2004.
- [10] Information on https://en.wikipedia.org/wiki/Symmetric-key_algorithm.
- [11] Christof Paar, Jan Pelzl, Understanding Cryptography, Springer-Verlag Berlin Heidelberg, 2010.
- [12] Joan Daemen, Vincent Rijmen, AES Proposal: Rijndael, October 1999.
- [13] Joan Daemen, Vincent Rijmen, The Design of Rijndael, Springer-Verlag, November 26, 2001.
- [14] Information on <https://ru.wikipedia.org/wiki/%D0%9F%D1%80%84%D1%81%D0%B0>.
- [15] Information on <https://dic.academic.ru/dic.nsf/ruwiki/12112>.
- [16] А.В. Яковлев, А.А. Безбогов, В.В. Родин, В.Н. Ша-мкин, Криптографическая защита информации. Издательст-во ТГТУ, 2016.
- [17] Information on https://studref.com/403682/informatika/blochnye_shifry.
- [18] Roberto Avanzi, A Salad of Block Ciphers, Munich, August 1, 2017.
- [19] Bruce Schneier, A self-study course in block-cipher cryptanalysis, Cryptologia, January 2000.
- [20] Information on <https://sites.google.com/site/anisimovkhv/learning/kripto/lecture/tema4>.
- [21] Lars R. KnudsenMatthew J.B. Robshaw, The Block Cipher Companion, Springer-Verlag Berlin Heidelberg, 2011.
- [22] Information on <https://techspirited.com/how-does-wireless-keyboard-mouse-work>.
- [23] Information on <https://compfonyk.com/kak-rabotaet-besprovodnaya-klaviatura-dlya-kompyutera/>.
- [24] Information on <https://xakep.ru/2016/02/25/mousejack/>.
- [25] Information on <https://www.securitylab.ru/news/381480.php>.
- [26] Information on <https://www.mousejack.com/faq>.
- [27] Information on <https://ru.wikipedia.org/wiki/%D0%AD%D0%BB%D1%8E%D1%87>.

Стаття надійшла до редколегії

11.11.2020



Use of symmetric encryption algorithms for signal transmission in wireless data input devices

Today, there are many computer systems that are designed to improve, facilitate and improve human life. As more and more researchers become interested in information processing, the development of computer systems is gaining momentum every year. With the development of information systems, possible threats, such as breaches of confidentiality, integrity and availability of processed information, are developing at an equally rapid pace. In order to prevent possible losses, the latest information security systems are constantly updated and improved. Since it is impossible to create a completely secure system, there is always the possibility of data theft, so the problem of protection of information and telecommunications systems is becoming increasingly important. Given that no protection can be perfect, a way has been developed to significantly reduce data breaches. Today, cryptosystems are mostly used to protect information and telecommunications systems and other technologies, including the protection of critical information of the state, enterprise, person or other critical data, including corporate secrets, intelligence or trade secrets. This article presents ways to use symmetric algorithms for signal transmission in remote data input devices. Information on existing algorithms of encryption and use of hash functions is given. A distinction is made between single-key and two-key methods of information encryption. The AES algorithm of its function, work of rounds, schemes of data encryption are considered in detail. The description of each algorithm is accompanied by an example which explains the features of their use. A mathematical model and an example of a block encryption algorithm are presented. The principles of operation of wireless devices are highlighted. Vulnerabilities related to wireless devices are considered and solutions for their protection are proposed.

Keywords: cryptography, encryption, symmetric encryption, block encryption, signal encryption, wireless data entry devices.



Сергій Толопа,

доктор технічних наук, професор, професор кафедри кібербезпеки та захисту інформації Київського національного університету імені Тараса Шевченка.

Serhii Toliupa,

Doctor of Technical Sciences, Professor, Professor of the Department of Cybersecurity and Information Protection, Taras Shevchenko National University of Kyiv.



Володимир Сайко,

доктор технічних наук, професор, професор кафедри прикладних інформаційних систем Київського національного університету імені Тараса Шевченка.

Volodymyr Saiko,

Doctor of Technical Sciences, Professor, Professor of the Department of Applied Information Systems, Taras Shevchenko National University of Kyiv.



Володимир Наконечний,

доктор технічних наук, старший науковий співробітник, професор кафедри кібербезпеки та захисту інформації Київського національного університету імені Тараса Шевченка.

Volodymyr Nakonechniy,

Doctor of Technical Sciences, Senior Research Fellow, Professor of the Department of Cybersecurity and Information Protection, Taras Shevchenko National University of Kyiv.



Максим Котов,

студент кафедри кібербезпеки та захисту інформації Київського національного університету імені Тараса Шевченка.

Maxim Kotov,

student of the Department of Cyber Security and Information Protection of Taras Shevchenko National University of Kyiv.



Валерія Солодовник,

студентка кафедри кібербезпеки та захисту інформації Київського національного університету імені Тараса Шевченка.

Valeria Solodovnik,

student of the Department of Cyber Security and Information Protection of Taras Shevchenko National University of Kyiv.

Наукове видання



БЕЗПЕКА ІНФОРМАЦІЙНИХ СИСТЕМ І ТЕХНОЛОГІЙ

№ 1/2 (3/4) 2020

Автори опублікованих матеріалів несуть повну відповідальність за підбір, точність наведених фактів, цитат, економіко-статистичних даних, власних імен та інших відомостей. Редколегія залишає за собою право скорочувати та редагувати подані матеріали.

Оригінал-макет виготовлено Видавничо-поліграфічним центром "Київський університет"
Виконавець *В. Гаркуша*



Формат 60x84^{1/16}. Ум. друк. арк. 9,3. Наклад 100. Зам. № 221-10348.
Гарнітура Arial. Папір офсетний. Друк офсетний. Вид. № ІТ1.
Підписано до друку 29.12.2020

Видавець і виготовлювач
ВПЦ "Київський університет",
6-р Тараса Шевченка, 14, м. Київ, 01601, Україна
☎ (38044) 239 32 22; (38044) 239 31 72; тел./факс (38044) 239 31 28
e-mail: vpc_div.chief@univ.net.ua; redaktor@univ.net.ua
http: vpc.knu.ua

Свідоцтво суб'єкта видавничої справи ДК № 1103 від 31.10.02