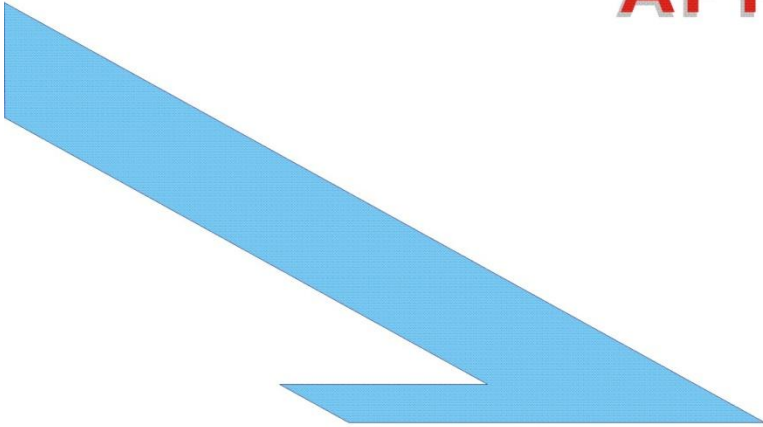




**I T H E A**



**International Journal**  
**INFORMATION THEORIES**  
**&**  
**APPLICATIONS**



**2020**   **Volume 27**   **Number 4**



**International Journal**  
**INFORMATION THEORIES & APPLICATIONS**  
**Volume 27 / 2020, Number 4**

**Editorial board**

Editor in chief: **Krassimir Markov** (Bulgaria)

|                              |            |                            |            |
|------------------------------|------------|----------------------------|------------|
| <b>Alberto Arteta</b>        | (Spain)    | <b>Lyudmila Lyadova</b>    | (Russia)   |
| <b>Aleksey Voloshin</b>      | (Ukraine)  | <b>Martin P. Mintchev</b>  | (Canada)   |
| <b>Alexander Eremeev</b>     | (Russia)   | <b>Natalia Bilous</b>      | (Ukraine)  |
| <b>Alexander Palagin</b>     | (Ukraine)  | <b>Natalia Pankratova</b>  | (Ukraine)  |
| <b>Alfredo Milani</b>        | (Italy)    | <b>Olena Chebanyuk</b>     | (Ukraine)  |
| <b>Avtandil Silagadze</b>    | (Georgia)  | <b>Rumyana Kirkova</b>     | (Bulgaria) |
| <b>Avram Eskenazi</b>        | (Bulgaria) | <b>Stoyan Poryazov</b>     | (Bulgaria) |
| <b>Dimitar Radev</b>         | (Bulgaria) | <b>Tatyana Gavrilova</b>   | (Russia)   |
| <b>Galina Rybina</b>         | (Russia)   | <b>Tea Munjishvili</b>     | (Georgia)  |
| <b>Giorgi Gaganadze</b>      | (Georgia)  | <b>Teimuraz Beridze</b>    | (Georgia)  |
| <b>Hasmik Sahakyan</b>       | (Armenia)  | <b>Valeriya Gribova</b>    | (Russia)   |
| <b>Juan Castellanos</b>      | (Spain)    | <b>Vasil Sgurev</b>        | (Bulgaria) |
| <b>Koen Vanhoof</b>          | (Belgium)  | <b>Vitalii Velychko</b>    | (Ukraine)  |
| <b>Krassimira B. Ivanova</b> | (Bulgaria) | <b>Vitaliy Lozovskiy</b>   | (Ukraine)  |
| <b>Leonid Hulianytskyi</b>   | (Ukraine)  | <b>Vladimir Jotsov</b>     | (Bulgaria) |
| <b>Levon Aslanyan</b>        | (Armenia)  | <b>Vladimir Ryazanov</b>   | (Russia)   |
| <b>Luis F. de Mingo</b>      | (Spain)    | <b>Yevgeniy Bodyanskiy</b> | (Ukraine)  |

**International Journal "INFORMATION THEORIES & APPLICATIONS" (IJ ITA)**  
**is official publisher of the scientific papers of the members of**  
**the ITHEA International Scientific Society**

IJ ITA welcomes scientific papers connected with any information theory or its application.

IJ ITA rules for preparing the manuscripts are compulsory.

The **rules for the papers** for IJ ITA are given on [www.ithea.org](http://www.ithea.org).

Responsibility for papers *published in* IJ ITA belongs to authors.

**International Journal "INFORMATION THEORIES & APPLICATIONS" Vol. 27, Number 4, 2020**

Edited by the **Institute of Information Theories and Applications FOI ITHEA**, Bulgaria, in collaboration with:

University of Telecommunications and Posts, Bulgaria,  
V.M.Glushkov Institute of Cybernetics of NAS, Ukraine,  
Universidad Politécnica de Madrid, Spain,  
Hasselt University, Belgium,  
University of Perugia, Italy,  
Institute for Informatics and Automation Problems, NAS of the Republic of Armenia  
St. Petersburg Institute of Informatics, RAS, Russia,

**Printed in Bulgaria**  
**Publisher ITHEA®**

Sofia, 1000, P.O.B. 775, Bulgaria. [www.ithea.org](http://www.ithea.org), e-mail: [office@ithea.org](mailto:office@ithea.org)  
Technical editor: **Ina Markova**

**Copyright © 2020 All rights reserved for the publisher and all authors.**

® 1993-2020 "Information Theories and Applications" is a trademark of ITHEA®

® ITHEA is a registered trade mark of FOI-Commerce Co.

**ISSN 1310-0513 (printed)**

**ISSN 1313-0463 (online)**

## STEGANALYSIS OF ADAPTIVE EMBEDDING METHODS BY MESSAGE RE-EMBEDDING INTO STEGO IMAGES

Dmytro Progonov, Vladymir Lucenko

**Abstract:** *Counteraction to sensitive information leakage via hidden (steganographic) channels is topical task today. Special interest is taken to the case of hidden message (stego image) revealing under limited a priori information about used embedding methods. The paper is devoted to the performance analysis of image calibration methods, namely by message re-embedding. The case of adaptive message hiding into cover image according to HUGO, S-UNIWARD, MG and MiPOD embedding methods is considered. It is revealed that message re-embedding allows significantly (up to 30%) reduce detection errors for stego images formed according to HUGO and S-UNIWARD methods. For improving stegdetector's performance for MG and MiPOD methods, it is proposed to use linearly transformed features. These features allow reduce classification error even for cover image low payload (less than 10%) in comparison with features for non-calibrated images. Important peculiarity of proposed features is low sensitivity to the number of cover-stego pairs into stegdetector's training set. This makes it possible to apply these features in real cases when steganalytics have limited access to embedding method.*

**Keywords:** *information security, digital images steganalysis, adaptive embedding methods, cover calibration, security level.*

**ACM Classification Keywords:** *Security and privacy – Systems security – Information flow controls*

---

## Introduction

---

The leakage of sensitive information of private corporations and governmental agencies is topical problem today. In most cases, unauthorized transmission of such information is done via hidden (steganographic) channels by message hiding within innocuous files. Reliable detection of embedded messages (stego files) requires comprehensive analysis of digital media, such as digital images – revealing of negligible anomaly changes of cover files caused by message hiding. Special interest is taken to the case of stego image detection under limited a priori information about used embedding method (zero-day problem).

The majority of modern stegdetectors is based on analysis the differences between current (analyzed) images and used statistical model [Fridrich, 2009; Konahovych et al, 2018]. Comprehensive analysis of these differences allows reveals features of statistical models that are sensitive to negligible changes of cover image caused by stego formation. Small changes of revealed features require taking of special classification methods for providing high detection accuracy. This task becomes non-trivial if there is no information of used embedding method (blind steganalysis).

One of possible solution for digital image blind steganalysis is cover image calibration methods. These methods are aimed to increase stego-to-cover ratio either by suppression of cover image context (for example, high-pass filtering), or increasing contribution of distortions caused by message hiding. There is proposed wide range of calibration methods that based on JPEG re-compression, image filtering to name but a few. These methods allows considerably improve detection accuracy for well-known embedding methods while preserving relatively low accuracy for state-of-the-art adaptive methods. Therefore, it is topical task to develop calibration methods that provides high stego-to-cover ratio even for advanced embedding methods.

---

## Related works

---

During last decade it was proposed wide range of digital image steganalysis methods. These methods can be divided into two parts – signature-based and model-based methods [Fridrich, 2009]. The former ones are based on usage of

known signature of message embedding – distinctive alterations of cover due to stego image formation. Nevertheless, practical usage of these methods is limited due to impossibility of obtaining the signature for unknown embedding methods (zero-day problem).

The model-based methods are aimed on analysis the differences between current image and model of cover image. The image model can be based on statistical, spectral and structural features of cover images [Konahovych et al, 2018]. One of the most widespread models is SPAM [Pevny et al, 2010], DCTR [Holub et al, 2014a] to name but a few.

Design of cover image model is non-trivial task that needs high level of expertise in domain of signal processing and statistical modeling. The state-of-the-art models of cover images, such as SRM [Fridrich et al, 2012], incorporate 34,671 features that allows achieving high detection accuracy for wide range of embedding methods. On the other hand, usage of such enormous number of features requires imposes high requirements on volume of used dataset for stegdetector to be tuned. It can be overcome by taking power of artificial neural networks, such as convolutional neural network (CNN). These networks allow learning of distinctive features directly from stego images, for instance, SR-Net [Boroumand et al, 2018]. At the same time, tuning of such networks is compute-intensive procedure that may be inappropriate for real cases.

Performance and computation complexity of model-based and CNN-based steganalysis methods are highly depends on pre-processing (calibration) of analyzed image. The calibration is used for increasing stego-to-cover ratio by amplification of negligible alterations of cover image caused by message hiding. Thorough choosing of calibration method allows significantly improving of stegdetector performance even in case of usage the relatively simple cover image model [Kodovsky et al, 2009].

The state-of-the-art Cartesian calibration method was proposed by Fridrich and based on usage of features of initial and pre-processed (filtered) images [Kodovsky et al, 2009]. It is proposed to pre-process image by applying of high-pass filters for suppression of cover image context. Nevertheless, choosing of optimal calibration transformation of analyzed image for maximization of

detection accuracy is open question today. One of the possible solutions of this problem is message re-embedding into analyzed image. It allows amplifying of cover image distortions caused by message hiding.

---

### Task and challenges

---

The paper is devoted to analysis of statistical stegdetectors performance in case of stego images calibration via message re-embedding. The case of adaptive message hiding into cover image according to state-of-the-art methods HUGO, S-UNIWARD and MiPOD is considered.

---

### Notations

---

High-dimensional arrays, matrices, and vectors will be typeset in boldface and their individual elements with the corresponding lower-case letters in italics.

The symbols

$$\mathbf{U} = (u_{ij}) \in \mathfrak{I}^{N \times M}, \mathbf{X} = (x_{ij}) \in \mathfrak{I}^{N \times M} \text{ and } \mathbf{Y} = (y_{ij}) \in \mathfrak{I}^{N \times M}, \mathfrak{I} = \{0, 1, \dots, 255\}, \quad (1)$$

will always represent pixel values of 8-bit grayscale initial (unprocessed), cover and stego images with size  $N \times M$  pixels. The image's feature vector is denoted as  $\mathbf{F}$ , while the embedding binary message is represented as  $\mathbf{M}$ .

The Iverson bracket  $[a]_I$  equals to one if the Boolean expression  $a$  is true, and zero otherwise. The notation  $\|\cdot\|$  corresponds to Euclidean norm for scalar values, and Frobenius norm for matrices.

---

### Digital images steganalysis via image calibration

---

Improving performance of stegdetectors can be achieved by increasing of stego-to-cover ratio, namely by calibration of analyzed images [Fridrich, 2009]. It was proposed wide range of digital image calibration methods, such as by message re-embedding [Miche et al, 2010], JPEG image re-compression [Kodovsky et al, 2009], image filtering [Kodovsky et al, 2009] to name but a few. Fridrich proposed next classification of calibration methods [Kodovsky et al, 2009]:

1. **Parallel reference** – calibration can be seen as a (near) constant feature space shift. Therefore, applying calibration causes a complete failure of steganalysis because the classes of cover and stego images become indistinguishable.
2. **Eraser** – is achieved by applying transformation that is robust with regards to embedding changes, for instance it erases embedding changes.
3. **Cover image estimate** – calibration transformation maps each stego image to an image whose features approximates the cover image feature. This idea stood behind the original idea of calibration – to come up with a good cover image estimate;
4. **Stego image estimate** – is complementary case to cover image estimation. Here, the calibration provides an estimate of the stego feature instead of the cover ones. A practical form of this approach may be realized by repetitive embedding (re-embedding), when the feature values changes significantly when applied to the cover image while it has a much smaller effect on initial stego image.
5. **Divergent reference** – the action of the reference mapping can be interpreted as a shift of cover's and stego's images features to a different direction.

For improving the stegdetector's performance for wide range of embedding methods, Fridrich proposed the Cartesian calibrated features obtained by dot product between features of initial image  $F(I)$  and its estimated reference  $F_{cal}(I)$  [Kodovsky et al, 2009]. Proposed features allows significantly reduce classification errors for both spatial [Fridrich et al, 2012] and JPEG [Pevny et al, 2007] domains based steganalysis. On the other hand, usage of Cartesian calibrated features leads to doubling of feature space dimensionality that requires corresponding augmentation of dataset. It may be impractical in real cases when steganalytics have limited opportunity to generate stego images.

The alternative approach is linear transformation of features for initial and reference images [Kodovsky et al, 2009]. Fridrich hypothesized that this approach may be ineffective while it may remove potentially useful information that might help us distinguish between cover and stego features. Nevertheless, performance of steganalysis in case of usage the linear transformed features

has not been investigated yet. In this work we analyzed the stegdetector’s performance by applying of diff-features, obtained by taking difference between features for initial and calibrated images.

---

### Adaptive embedding methods

---

The majority of state-of-the-art embedding methods are based on minimizing the empirical distortion estimation function  $D(\mathbf{X}, \mathbf{Y})$  during forming a stego image  $\mathbf{Y}$  from a cover image (CI)  $\mathbf{X}$  [Filler et al, 2011]:

$$D(\mathbf{X}, \mathbf{Y}) = \sum_i \rho_i(\mathbf{X}, \mathbf{Y}) \rightarrow \min, |\mathbf{M}| = \text{const}, \quad (2)$$

where  $\rho_i(\cdot)$  – function for estimation the alteration of cover image’s statistical characteristics by embedding of  $i^{\text{th}}$  stegobit;  $|\mathbf{M}|$  – size of embedded message  $\mathbf{M}$ , bits. Minimization of function (2) allows adapting the embedding process to a cover image, thus corresponding steganographic methods called adaptive.

In most cases, the choice of function  $D(\mathbf{X}, \mathbf{Y})$  is done under the assumption of independency of distortions caused by embedding of individual stegobits (distortions additivity) that simplifies the choice of function (2). Nevertheless, this approach does not taking into account interactions between distortions that may lead to non-linear changes of CI parameters.

In this paper we considered the case of usage the state-of-the-art adaptive embedding methods, namely HUGO [Filler et al, 2010], S-UNIWARD [Holub et al, 2014b], MG [Sedighi et al, 2015] and MiPOD [Sedighi et al, 2016]. These methods are aimed on message in spatial domain of CI with size  $N \times M$  pixels – by manipulation with brightness of individual pixels. Let us consider these methods in details.

The HUGO embedding method is based on solution of next optimization problem [Filler et al, 2010]:

$$\min_{\pi_i} E_{\pi} [D] = \sum_{y \in Y} \pi(y) \cdot D(y), H(\pi) = |M|, \quad (3)$$

where  $y \in Y$  – stego images  $y$  from the set of possible stego images  $Y$ ;  $\pi$  – probability distribution function of the choice of a certain  $y$  from set  $Y$ ;  $E_\pi[D]$  – averaging operator for  $D(X, Y)$  over distribution  $\pi$ ;  $H(\pi) = -\sum_{y \in Y} \pi(y) \cdot \log(\pi(y))$  – entropy function.

The optimal type of distribution  $\pi$  for solving problem (3) is the Gibbs distribution [Filler et al, 2010]:

$$\pi_\lambda(y) = \exp(-\lambda D(y)) / Z(\lambda), \quad (4)$$

$$Z(\lambda) = \sum_{y \in Y} \exp(-\lambda D(y)),$$

where  $Z(\lambda)$  – normalizing constant. The value of the scalar parameter  $\lambda > 0$  is determined by solving of equation (4) [Filler et al, 2010]. If function  $D(\cdot)$  is additive, the equation (4) can be rewritten as [Filler et al, 2010]:

$$\pi_\lambda(y) = \prod_i \pi_\lambda(y_i) = \frac{\prod_i \exp(-\lambda \rho_i(y_i))}{\sum_{t \in I} \exp(-\lambda \rho_t(y_t))},$$

where  $I$  – the range of pixels brightness for cover image.

In seminal paper [Filler et al, 2011], it is proposed to use a limited function, namely local potential  $V_c(y)$ , for cover image distortion estimation. The values of  $V_c(\cdot)$  depend on adjacent pixels brightness correlation in a given pixel's neighborhood (clique)  $c \in C$ . The correlation maybe estimated with usage of adjacency matrix  $C_{k,l}(X)$ :

$$C_{k,l}(X) = \sum_i \sum_j [x_{i,j} = k] [x_{i,j+1} = l], \quad (5)$$

where  $x_{i,j}$  – cover image's pixel brightness value with coordinates  $(i, j)$ .

In the case of CI row-wise processing during message hiding, the matrix (5) can be estimated in the next way [Filler et al, 2011]:

$$\mathbf{A}_{k,l}^{\rightarrow}(\mathbf{X}) = \frac{1}{N(M-2)} \sum_{i,j} \left[ (\mathbf{D}_{i,j}^{\rightarrow}, \mathbf{D}_{i,j+1}^{\rightarrow})(\mathbf{X}) = (k, l) \right]_I,$$

$$(\mathbf{D}_{i,j}^{\rightarrow}, \mathbf{D}_{i,j+1}^{\rightarrow})(\mathbf{X}) = (k, l) \Leftrightarrow \mathbf{D}_{i,j}^{\rightarrow}(\mathbf{X}) = k \& \mathbf{D}_{i,j+1}^{\rightarrow}(\mathbf{X}) = l, \mathbf{D}_{i,j}^{\rightarrow}(\mathbf{X}) = \mathbf{X}_{i,j+1} - \mathbf{X}_{i,j}.$$

If pixels' brightness is changed by  $(\pm 1)$  during stegobit embedding, the normalized adjacency matrix  $\mathbf{A}_{k,l}^{\rightarrow}(\mathbf{X})$  is converted to  $\mathbf{A}_{k,l}^{\rightarrow}(\mathbf{Y})$  [Filler et al, 2010]:

$$|\mathbf{A}_{k,l}^{\rightarrow}(\mathbf{Y}) - \mathbf{A}_{k,l}^{\rightarrow}(\mathbf{X})| = \sum_{c \in C} \mathbf{H}_c^{\rightarrow(k,l)}(\mathbf{Y}),$$

$$\mathbf{H}_c^{\rightarrow(k,l)}(\mathbf{Y}) = \frac{1}{N(M-2)} \left| \left[ (\mathbf{D}_{i,j}^{\rightarrow}, \mathbf{D}_{i,j+1}^{\rightarrow})(\mathbf{Y}) = (k, l) \right]_I - \left[ (\mathbf{D}_{i,j}^{\rightarrow}, \mathbf{D}_{i,j+1}^{\rightarrow})(\mathbf{X}) = (k, l) \right]_I \right|,$$

for all horizontal cliques of a given pixel  $C^{\rightarrow} = \{c : c = \{(i, j), (i, j+1), (i, j+2)\}\}$ .

Similarly, the adjacency matrices for other types of cliques  $C$  ( $\mathbf{A}_{k,l}^{\leftarrow}(\mathbf{Y})$ ,  $\mathbf{A}_{k,l}^{\uparrow}(\mathbf{Y})$  and  $\mathbf{A}_{k,l}^{\downarrow}(\mathbf{Y})$ ) can be calculated.

Thus, message hiding according to HUGO method is carried out by solving of next optimization problem [Filler et al, 2010]:

$$D(\mathbf{Y}) = \sum_{c \in C} \sum_{k,l} w_{k,l} \mathbf{H}_c^{(k,l)}(\mathbf{Y}),$$

where  $C = C^{\rightarrow} \cup C^{\leftarrow} \cup C^{\uparrow} \cup C^{\downarrow}$  – set of three-elements cliques for four-pixels adjacency directions;  $w_{k,l} > 0$  – weighting factors. The HUGO method is widely used in researches as typical adaptive embedding methods.

The S-UNIWARD embedding method is based on spectral transformation of CI. The transformation is used for estimation the cover image distortions caused by embedding of individual stegobits. Similarly to HUGO method, S-UNIWARD method takes additive empirical distortion estimation function [Holub et al, 2014b]:

$$D(\mathbf{X}, \mathbf{Y}) = \sum_{k,u,v} \frac{|\mathbf{W}_{uv}(\mathbf{X}, k) - \mathbf{W}_{uv}(\mathbf{Y}, k)|}{\sigma + |\mathbf{W}_{uv}(\mathbf{X}, k)|}, \quad (6)$$

where  $\mathbf{W}_{uv}(\mathbf{X}, k)$ ,  $\mathbf{W}_{uv}(\mathbf{Y}, k)$  – coefficients of two-dimensional discrete wavelet transform (2D-DWT) of the cover  $\mathbf{X}$  and stego  $\mathbf{Y}$  images with coordinates  $(u, v)$  in the  $k^{\text{th}}$  sub-band;  $\sigma > 0$  – stabilizing constant.

Variation of 2D-DWT basis functions in (6) allows analyzing specific distortions of CI caused by message hiding. Also, usage of empirical distortion estimation function (6) makes it possible to message hiding in spatial (alteration of CI pixels brightness) and transformation (by changing of a CI decomposition coefficients) domains in the uniform way.

The alternative approach to design an empirical distortion estimation function (2) is based on minimization both cover image distortions and statistical detectability of formed stego images [Ker et al, 2013]. As an example of such embedding method, the MG [Sedighi et al, 2015] and MiPOD [Sedighi et al, 2016] embedding methods can be taken. Feature of these methods is usage of locally-estimated multivariate Gaussian cover image model. It allows achieving the stego image's empirical security that is comparable with advanced steganographic methods.

Formation of stego images according to MiPOD method is carried out in several steps [Sedighi et al, 2016]. Firstly, it is suppressed the cover image context  $\mathbf{X} = (x_1, \dots, x_{M \cdot N})$ , using denoising filter  $F$ :

$$\mathbf{r} = \mathbf{X} - F(\mathbf{X}),$$

where  $\mathbf{X}$  is represented in column-wise order. Then, it is measured pixels residual variance  $\sigma_l^2$  using Maximum Likelihood Estimation:

$$\mathbf{r}_l = \mathbf{G}\mathbf{a}_l + \xi_l, \quad (7)$$

where  $\mathbf{r}_l$  – represents the value of the residuals  $\mathbf{r}$  inside the  $p \times p$  block surrounding the  $l^{\text{th}}$  residual put into a column vector;  $\mathbf{G}_{p^2 \times p}$  – the matrix defines the parametric model of remaining expectation;  $\mathbf{a}_{p \times 1}$  – the vector of linear model's parameters;  $\xi_{p^2 \times 1}$  – the signal whose variance is need to be estimated.

At the second step, the pixels residual variance  $\sigma_l^2$  is estimated according to formula:

$$\sigma_l^2 = \|\mathbf{P}_G^\perp \mathbf{r}_l\|^2 / (p^2 - q), \quad (8)$$

where  $\mathbf{P}_G^\perp = \mathbf{I}_l - \mathbf{G}(\mathbf{G}^T \mathbf{G})^{-1} \mathbf{G}^T$  – the orthogonal projection of residual  $\mathbf{r}_l$  (7) to the  $p^2 - q$  dimensional subspace spanned by the left null space of  $\mathbf{G}$ ;  $\mathbf{I}_{l \times l}$  – the unity matrix.

Thirdly, it is determined the probability of  $l^{\text{th}}$  embedding change  $\beta_l, l \in \{1, 2, \dots, L\}$  that minimize the deflection coefficient  $\zeta^2$  between cover and stego image distributions:

$$\zeta^2 = 2 \sum_{l=1}^{M \cdot N} \beta_l^2 \sigma_l^{-4}, \quad (9)$$

under payload constrain

$$R = \sum_{l=1}^{M \cdot N} H(\beta_l),$$

where  $H(z) = -2z \log z - (1-2z) \log(1-2z)$  – ternary entropy function;  $R$  – cover image payload in nats.

Minimization of (9) can be achieved by using the method of Lagrange multipliers. The change rate  $\beta_l$  and the Lagrange multiplier  $\lambda$  can be determined by numerically solving of next  $(l+1)$  equations:

$$\beta_l \sigma_l^{-4} = \frac{1}{2\lambda} \ln \left( \frac{1-2\beta_l}{\beta_l} \right), l \in [1; M \cdot N],$$

$$R = \sum_{l=1}^{M \cdot N} H(\beta_l).$$

Then, the change rate  $\beta_l$  is converted to the cost  $\rho_l$ :

$$\rho_l = \ln(1/\beta_l - 2). \quad (10)$$

Finally, the desired payload  $R$  is embedded using syndrome-trellis codes (STCs) with pixel costs determined according to (10).

The MG embedding method [Sedighi et al, 2015] is similar to MiPOD algorithm, but it uses the simplified variance estimator:

$$\sigma_l^2 = \|\mathbf{r}_n - \hat{\mathbf{r}}_n\|^2 / (p^2 - q), \quad (11)$$

$$\hat{\mathbf{r}}_n = \mathbf{G}(\mathbf{G}^T \mathbf{G})^{-1} \mathbf{G}^T \mathbf{r}_n.$$

Applying the locally-estimated multivariate Gaussian cover model in MiPOD algorithm gives opportunity to derive a closed-form expression for the performance of the stegdetector and capture the non-stationary character of natural images [Sedighi et al, 2016].

---

## Experiments

---

Performance analysis of image calibration methods was done on standard BOSS dataset. The stegdetector was based on standard SPAM statistical model [Pevny et al, 2010] and ensemble classifier [Kodovsky et al, 2012]. The SPAM model allows estimating correlation of adjacent pixels brightness with usage of 2<sup>nd</sup> and 3<sup>rd</sup> order Markov chains.

We consider the case of analyzed image calibration via message re-embedding. It is performed by applying same embedding method and similar payload as it is for stego images. Therefore, the stegdetector was tuned with usage of next types of image features:

1. **Non-calibrated features** – corresponds to features of initial (non-processed) image  $\mathbf{U}$ :

$$\mathbf{F}_{nc} = \mathbf{F}(\mathbf{U}),$$

2. **Features after message re-embedding** – corresponds to features of calibrated image, obtained after message re-embedding:

$$\mathbf{F}_{re-embed} = \mathbf{F}_{cal}(\mathbf{U}),$$

3. **Cartesian calibrated features** – corresponds to merged features of initial and calibrated images:

$$\mathbf{F}_{CC} = [\mathbf{F}(\mathbf{U}); \mathbf{F}_{cal}(\mathbf{U})],$$

4. **Calibrated features after linear transformation** – corresponds to the differences between features for calibrated and initial images:

$$\mathbf{F}_{DF} = \mathbf{F}_{cal}(\mathbf{U}) - \mathbf{F}(\mathbf{U}).$$

Analysis of stegdetector's detection accuracy was done according to cross-validation procedure. The total error  $P_E$  is used as the performance index [Kodovsky et al, 2012]:

$$P_E = \min_{P_{FA}} \frac{1}{2} (P_{FA} + P_{MD}(P_{FA})),$$

where  $P_{FA}$  and  $P_{MD}$  are probability of false alarm and missed detection, respectively. During testing it was considered two cases:

1. Stegdetectors is tuned with pairs of processed cover and stego images – corresponds to standard practice during image steganalysis;
2. There are no pair of cover and stego images in training subset during stegdetector tuning – corresponds to the real cases when steganalytic has no pairs of cover and corresponding stego images.

The case of adaptive message embedding according to HUGO, S-UNIWARD, MG and MiPOD methods was considered. The CI payload was changed within range 3%, 5%, 10%, 20%, 30%, 40% and 50%.

The dependencies of total error  $P_E$  on CI payload for HUGO embedding method are represented at Fig.1. The case of usage the non-calibrated features  $\mathbf{F}_{nc}$

(solid lines), features after message re-embedding  $\mathbf{F}_{re-embed}$  (dashed lines), calibrated features after linear transformation  $\mathbf{F}_{DF}$  (dotted lines) and Cartesian calibrated features  $\mathbf{F}_{CC}$  (dash-dot lines) is considered.

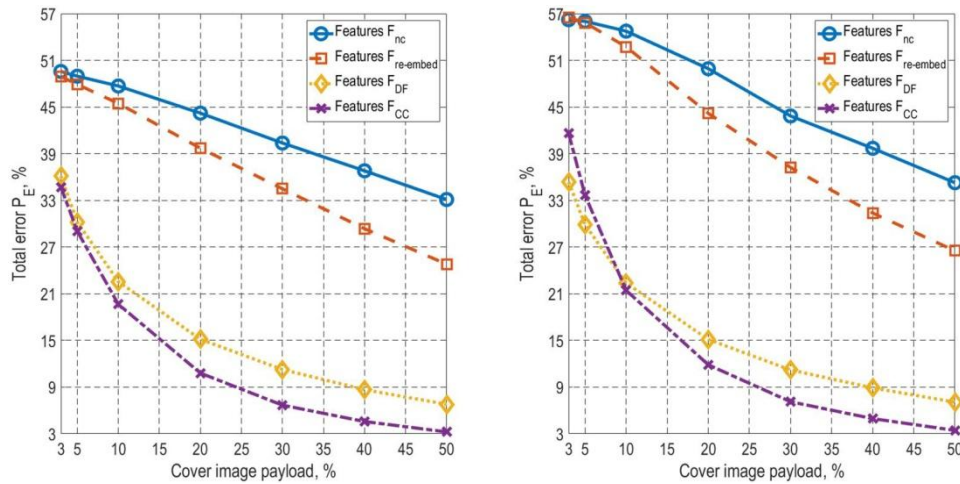


Figure 1. Dependency of total error  $P_E$  on cover image payload for HUGO embedding method by full (left) or none (right) alignment of cover-image pairs during stegdetector testing.

Usage of  $\mathbf{F}_{re-embed}$  features provides relatively small improvement on total error in case of low payload (less than 10%) – the differences achieves up to 2.5% (Fig. 1). For the bigger payloads we obtained up to 7%, especially for high payload (near 50%).

It is revealed significant increasing of total error (approximately 5%-7.5%) in case of none alignment of cover-image pairs during stegdetector testing (Fig. 1). It confirms that usage of cover-stego images pairs during stegdetector testing gives opportunity to achieve the lowest values of total error.

Applying of Cartesian calibrated features  $\mathbf{F}_{CC}$  leads to considerably reducing of total errors  $P_E$  – from 15% for low payload to 30% for high payload (Fig. 1). Usage of  $\mathbf{F}_{DF}$  features allows achieving comparable low values of  $P_E$  only on case of low payload (Fig. 1). On the other hand, it is revealed that  $\mathbf{F}_{DF}$  features

weakly depends on cover-image pairs alignment during stegdetector testing – changing of total errors  $P_E$  values are near 1-1.5% in this case. Therefore,  $F_{DF}$  features allows outperform the Cartesian calibrated features  $F_{CC}$  up to 5-7.5% for the low payload and case on none alignment (Fig. 1).

The dependencies of total error  $P_E$  on cover image payload for S-UNIWARD embedding method are represented at Fig.2.

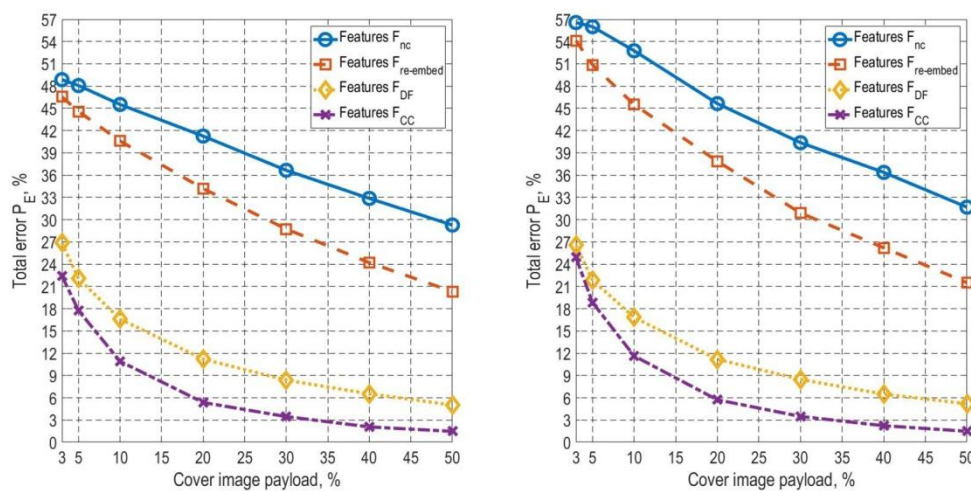


Figure 2. Dependency of total error  $P_E$  on cover image payload for S-UNIWARD embedding method by full (left) or none (right) alignment of cover-image pairs during stegdetector testing.

Similarly to HUGO embedding method (Fig. 1), a message re-embedding according to S-UNIWARD method allows reducing of total error  $P_E$  up to 3% even in case of low cover image payload (less than 10%). The total error's reducing may achieve up to 10% by increasing of cover image payload.

Usage of  $F_{DF}$  features does not allow lower values of  $P_E$  in comparison with Cartesian calibrated features  $F_{CC}$  (Fig. 2) for S-UNIWARD embedding method – the difference between  $P_E$  for both cases achieves up to 5% for full alignment and 3% for none alignment cover-stego images pairs.

Dependencies of total error  $P_E$  on cover image payload for MG embedding method are represented at Fig.3.

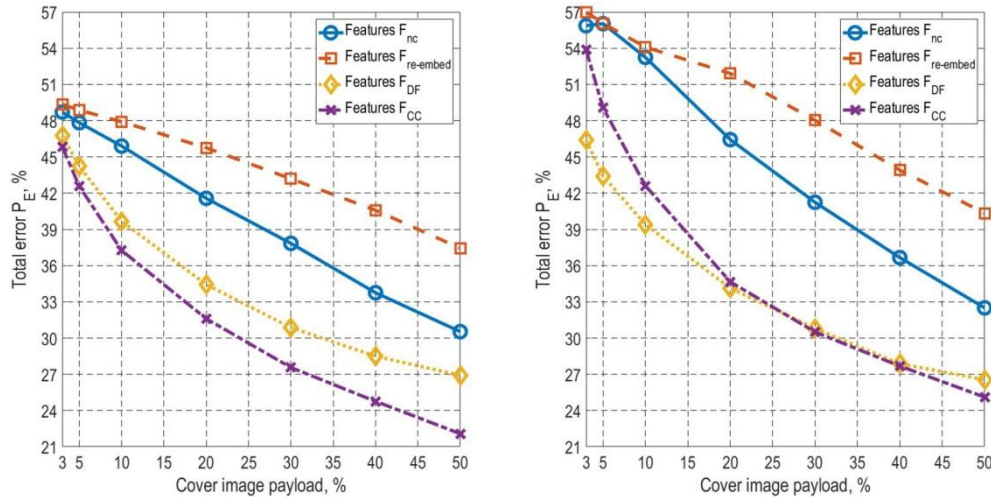


Figure 3. Dependency of total error  $P_E$  on cover image payload for MG embedding method by full (left) or none (right) alignment of cover-image pairs during stegdetector testing.

It should be noted that usage of  $F_{re-embed}$  features leads to considerably reducing of stegdetector performance for MG method (Fig. 3) – the values of total error  $P_E$  increase from 1.5% for low cover image payload to 8% for high payload cases in comparison with case of usage the features of initial (un-processed) images. Therefore, we conclude that message re-embedding can mask of initial stego data for MG method and improve robustness of obtained stego image to steganalysis.

Applying of linearly transformed features  $F_{DF}$  allows reducing  $P_E$  values but only in the case of absence the cover-stego images pairs in stegdetector's training set (Fig.3). The biggest reducing is achieved in the case of low cover image payload (up to 7% reducing of total error values) that is the most difficult for image steganalysis. By increasing cover image payload, usage of  $F_{DF}$  and  $F_{CC}$  features leads to similar values of  $P_E$ .

Dependencies of total error  $P_E$  on cover image payload for MiPOD embedding method are represented at Fig.4.

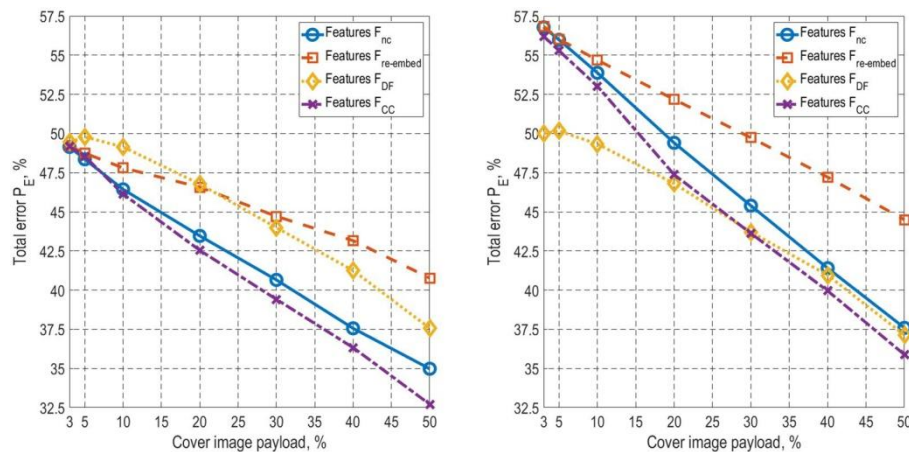


Figure 4. Dependency of total error  $P_E$  on cover image payload for MiPOD embedding method by full (left) or none (right) alignment of cover-image pairs during stegdetector testing.

Calibration of stego images formed according to MiPOD method allows improve stegdetector performance only for  $F_{CC}$  features (Fig. 4, full alignment case). For the none alignment case, we obtained that  $F_{DF}$  features allows significantly (up to 7%) reducing values of  $P_E$  values. Similarly to MG method (Fig. 3), the values of total error remains almost the same for  $F_{DF}$  and  $F_{CC}$  features by increase of cover payload (Fig. 4, none alignment case).

## Discussion

During solving the problem of improving the steganography and cryptography methods, the next question arises –what is the limit of encryption tool’s performance as information protection method. The limit can be achieved by enhancement of known cryptographic protocols, development of interception-proof bit-quantum communication systems, extension the length of encryption keys, improvement of cryptanalysis methods to name but a few. In general,

these approaches are aimed on increasing the cryptographic security level (CSL) for mathematical and technical tools that are used in information security systems. Therefore, it is needed to determine the threshold for CSL when negligible increasing of cryptographic security level requires usage of enormous amount of computational resources. As an example, it can be mentioned the case of achievements the long-term (strategic) cryptographic security level for sensitive information – any additional increasing of CSL requires involvement of tremendous amount of computational resources that may be impractical for real systems. The task of determination the threshold value of CSL can be solved by analysis of performance of known data processing methods.

In most cases, the modeling of physical phenomena and processes requires usage of continuous mathematics. Perturbations (singularities) of processes are represented frequently as superposition of several continuous short-time processes that take places on small scales. These processes are represented as continuous functions over discrete numbers (samples) with fixed bit depth. As a result, researches may face with "mysterious" artefacts (phenomena), such as specific elementary particles (for example, quarks, mesons, neutrinos, muons, leptons), dark matter, wormholes within spacetime to name a few.

It should be noted that usage of discrete mathematic for operations over discrete numbers is not a new approach. This approach appeared recently with development of the complex analysis and integral calculus. The classical discrete mathematic provides concepts of infinitely large and infinitely small actual numbers. Usage of these numbers makes it possible to determine quantities such as the shortest time interval, the maximum size of the Universe, the size of the smallest particle, the longest encryption key, the highest accuracy of cryptanalysis and steganalysis, the smallest error in decrypting closed messages, etc. Therefore, the used numeral system allows precisely represent the structure of the analyzed process or environment.

Let us consider the fine-structure constant  $\alpha$  as indicator of achievement both spatial and thermodynamic collapses. In this case, this constant can be represented as infinitely small actual number for discrete mathematic. Since the electron is the single non-collapsing element within thermodynamic space  $\alpha$ ,

then the minimum measurable size is a quantity  $\alpha = 1/137 \cdot 10^{-57}$  meters. The value  $1/137$  can be calculated from the CODATA's fine-structure constant value  $\alpha = 0.0072973525698 \cdot 10^{-57}$  by finding the harmonic mean (excluding the power factor):

$$\alpha_c = (\alpha^2 - \sigma^2) / \alpha = 0.00729927007299270..., \quad (12)$$

where  $\sigma = 2.2211024289753$  – is coefficient that allows to represent  $\alpha$  as a harmonizing number for numerical series.

The sequence of numbers 0072992700 in (12) has the unique property of symmetry with respect to a pair of zeros and a pair of nines. Moreover, the inverse to such infinite sequence is an integer number that is equal to 137. In what follows, we will designate  $\alpha_M = 1/137$  as the mathematical fine structure constant, while  $\alpha = 137.036$  (according to CODATA) will be denoted as physical fine structure constant. In this case, the value  $\alpha_M = 1/137 \cdot 10^{-57}$  is known as the actual infinitesimal. Therefore, there are no physical phenomena in the observable Universe that can be numerically smaller than  $\alpha_M$ .

Here we presented the concept of an actual infinite quantity in a classic way, i.e. the quantity is an alternative to potential infinite numbers. According to the fractal theory, the value  $\alpha_M$  (represented in meters) is the limiting value of the border between the nested Universes in the SI system. From the point of view of the mathematical description of physical phenomena, the number  $\alpha_M$  is the basis of the numeral system that harmonizes all numeral systems with any basis. Then, the number series in the  $\alpha$  – number system looks like:

$$\alpha_M, 2\alpha_M, 3\alpha_M, \dots, k\alpha_M, k \in \mathbb{N}.$$

From the above, we can conclude that any functions or mathematical constants, which are represented in the decimal numeral system, have no physical meaning if their values are larger than  $1/\alpha = 137 \cdot 10^{57}$  or lesser than  $\alpha$ .

Otherwise, the  $\alpha$  – number system cannot will not describe real physical processes as well as physical and mathematical constants. It has direct effect on achievable limits of security level for cryptography and steganography. For example, it makes no practical sense to use length of encryption keys more than  $1/\alpha$ , to take random number's array with more than  $1/\alpha$  elements, to achieve stego detection error less than  $\alpha$  etc.

---

## Conclusion

---

In this paper, we argue that usage of special type of image calibration method provides additional reducing of classification error. In fact, we recognize that message re-embedding into analyzed stego images allows considerably increasing stegdetector performance even in the most difficult cases – cover images low payload (less than 10%). Our view is supported by obtained results:

1. Usage of features obtained from calibrated images allows considerably improving stegdetector's performance for HUGO and S-UNIWARD embedding methods. It was revealed decreasing of total error up to 30% within all range of image payload (from 3% to 50%).
2. Calibration of stego images formed according to advanced MG and MiPOD embedding methods gives opportunity to reduce classification error up to 8% even for low payload (less than 10%) of cover image. These results were achieved by usage of standard SPAM model, so classification error reducing may be more impressive for rich statistical models, such as SRM and PSRM.
3. The results clearly show the benefit of linearly transformed features of calibrated images. These features allow additionally reduce classification error even for cover image low payload (less than 10%). It is noteworthy that error reducing is achieved without doubling dimensionality of features space as it is performed for state-of-the-art Cartesian calibrated features. Also, performance of stegdetector tuned with linearly transformed features remains almost the same even in case of absence of cover-stego images pairs during training. It makes these features useful for increasing the performance of stegdetectors in real cases.
4. It is proposed hypothesis of finiteness the steganalysis and cryptanalysis performance in real cases. The hypothesis is based on limitations of used numeric systems that represent infinitesimal number

in discrete form. It may introduce additional distortions during stegdetector training.

It should be noted that mentioned results were obtained for the case of message re-embedding according to the same steganographic method is considered. In the future, we would like to investigate stegdetector's performance for message re-embedding by varying of embedding algorithm and image payload.

---

## **Bibliography**

---

- [Boroumand et al, 2018] Boroumand M., Chen M., Fridrich J. Deep Residual Network for Steganalysis of Digital Images. IEEE Trans. Inf. Forensics Security, vol. 14, issue 5, 2018, pp. 1181-1193.
- [Filler et al, 2010] Filler T., Fridrich J. Gibbs Construction in Steganography, IEEE Trans. Inf. Forensics Security, vol. 5, issue 4, 2010, pp. 705-720.
- [Filler et al, 2011] Filler T., Fridrich J. Design of Adaptive Steganographic Schemes for Digital Images. Proc. SPIE, Electronic Imaging, Media Watermarking, Security, and Forensics XIII, 2011, DOI: 10.1117/12.872192.
- [Fridrich, 2009] J. Fridrich Steganography in Digital Media: Principles, Algorithms, and Applications. 1<sup>st</sup> edition. Cambridge University Press, 2009. p. 437. ISBN 978-0-521-19019-0;
- [Fridrich et al, 2012] Fridrich J., Kodovsky J. Rich Models for Steganalysis of Digital Images. IEEE Trans. Inf. Forensics Security, vol. 7, issue 3, 2012, pp. 868-882.
- [Holub et al, 2014a] Holub V., Fridrich J. Low Complexity Features for JPEG Steganalysis Using Undecimated DCT. IEEE Trans. Inf. Forensics Security, vol. 10, issue 2, 2015, pp. 219-228.
- [Holub et al, 2014b] Holub V., Fridrich J., Denemark T. Universal Distortion Function for Steganography in an Arbitrary Domain. EURASIP Journal on Information Security, Vol. 1, 2014, DOI: 10.1186/1687-417X-2014-1.
- [Ker et al, 2013] Ker A. D., Bas P., Böhme R., Cogramne R., Craver S., Filler T., Fridrich J., Pevný T. Moving steganography and steganalysis from the

laboratory into the real world. Proceedings of the first ACM workshop on Information hiding and multimedia security (IH&MMSec '13). New York, 2013;

[Kodovsky et al, 2009] Kodovsky J., Fridrich J. Calibration revisited. Proceedings of the 11<sup>th</sup> ACM workshop on Multimedia and security (MM&Sec'09), 2009, pp. 63-74.

[Kodovsky et al, 2012] Kodovsky J., Fridrich J., Holub V. Ensemble Classifiers for Steganalysis of Digital Media. IEEE Trans. Inf. Forensics Security, vol. 7, issue 2, 2012, pp. 432-444.

[Konahovych et al, 2018] Konachovych G., Progonov D., Puzyrenko O. Digital steganography processing and analysis of multimedia files [In Ukrainian]. 'Tsentr uchbovoi literatury' publishing, 2018, ISBN 978-617-673-741-4.

[Miche et al, 2010] Miche Y., Bas P., Lendasse A. Using multiple re-embeddings for quantitative steganalysis and image reliability estimation. TKK reports in information and computer science. Department of Information and Computer Science, Aalto University. 2010. ISBN 978-952-60-3250-4.

[Pevny et al, 2007] Pevny T., Fridrich J. Merging Markov and DCT features for multiclass JPEG steganalysis. Proceedings SPIE, Electronic Imaging, Security, Steganography, and Watermarking of Multimedia Contents, volume 6505, 2007.

[Pevny et al, 2010] Pevny T., Bas P., Fridrich J. Steganalysis by Subtractive Pixel Adjacency Matrix. IEEE Trans. Inf. Forensics Security, vol. 5, issue 2, 2010, pp. 215-224.

[Sedighi et al, 2015] Sedighi V., Fridrich J., Coganne R. Content-Adaptive Pentary Steganography Using the Multivariate Generalized Gaussian Cover Model. Proc. SPIE, Electronic Imaging, Media Watermarking, Security, and Forensics, 2015, vol. 9409.

[Sedighi et al, 2016] Sedighi V., Coganne R., Fridrich J. Content-Adaptive Steganography by Minimizing Statistical Detectability. IEEE Trans. Inf. Forensics Security. Vol. 11, Iss. 2., 2016. pp. 221-234.

---

## Authors' Information

---



**Dmytro Progonov** – Institute of Physics and Technology, National Technical University of Ukraine "Igor Sikorsky Kyiv Polytechnic Institute"; PhD, Associate Professor; 37, ave. Peremohy, Solomenskiy district, Kyiv, Postcode 03056, Ukraine; e-mail: [progonov@gmail.com](mailto:progonov@gmail.com)

*Major Fields of Scientific Research: digital media steganalysis, digital image forensics, machine learning, advanced signal processing*



**Vladymir Lucenko** – Institute of Physics and Technology, National Technical University of Ukraine "Igor Sikorsky Kyiv Polytechnic Institute"; PhD, Associate Professor; 37, ave. Peremohy, Solomenskiy district, Kyiv, Postcode 03056, Ukraine; e-mail: [lutsenkovn@ukr.net](mailto:lutsenkovn@ukr.net)

*Major Fields of Scientific Research: information security, artificial intelligence, biological and medical cybernetic, neurocomputer technology*

## ДИАГНОСТИКА МЕДИЦИНСКИХ ИЗОБРАЖЕНИЙ ОПУХОЛЕЙ МОЛОЧНОЙ ЖЕЛЕЗЫ С ИСПОЛЬЗОВАНИЕМ ГИБРИДНЫХ СВЕРТОЧНЫХ НЕЧЕТКИХ НЕЙРОННЫХ СЕТЕЙ

Галиб Гамидов, Юрий. П. Зайченко

**Резюме:** Рассмотрена проблема классификации опухолей молочной железы по медицинским изображениям. Для ее решения разработана гибридная нечеткая сеть CNN, в которой сверточная нейронная сеть используется в качестве экстрактора признаков, а нечеткая нейронная сеть NEFClass используется в качестве классификатора. Разработаны и внедрены алгоритмы обучения FNN. Проведены экспериментальные исследования предложенной гибридной сети на стандартном наборе данных BreakHis и проведено сравнение с известными результатами. Рассмотрена проблема уменьшения размерности данных и исследовано применение метода ИКМ.

**Keywords:** medical diagnostics, breast cancer classification, FNN, CNN, hybrid network

**ITHEA Keywords:** 1. Computing methodologies 12. Artificial intelligence, 16.5. Model development

---

### Введение. Анализ состояния проблемы

---

В настоящее время заболевания раком представляют огромную проблему для системы здравоохранения во всем мире [Boyle, 2012]. Среди различных типов рака рак молочной железы занимает второе место по частоте у женщин. Кроме того, смертность от него очень высока в сравнении с другими раковыми заболеваниями.

В настоящее время в практике медицинской диагностики широко используются информационные технологии. Основная цель медицинских информационных систем - расширение сферы практических задач, которые

можно решать с помощью компьютеров, повышение интеллектуального уровня поддержки принимаемых решений в сфере медицинской диагностики на основе обработки и анализа медицинских изображений тканей органов человека с применением различных источников (МРТ, КТ, рентгеновские снимки и т.п.)

Преимуществами систем медицинской диагностики являются скорость и стабильность работы, которые делают их удобным средством экспресс-анализа. Несмотря на юный возраст медицинской информатики, который не превышает 30 лет, информационные технологии быстро проникают в различные сферы медицины и охраны здоровья (семейная медицина, страховая медицина создание единого информационного пространства в медицине и т.д.). Несмотря на очевидный прогресс, достигнутый системами медицинской диагностики, окончательный диагноз рака молочной железы, включая классификацию опухолей выполняется паталого-анатомами которые используют визуальный анализ гистопатологических образцов под микроскопом.

В медицинской диагностике большую часть проблемы представляет извлечение признаков для последующей обработки изображений и их классификации. С развитием и широким распространением СППР в медицине повышаются требования к алгоритмам обучения таких систем.

Последние достижения в технологиях обработки изображений и машинном обучении позволяют создавать системы автоматического обнаружения и к опухолям, которые классификации могут помочь паталого-анатому делать правильный диагноз и ускорить его работу. Классификация изображений на различных паттернах, которые соответствуют раковым и нераковым состояниям тканей часто является первостепенной целью автоматической диагностики рака молочной железы.

На данный момент имеется несколько моделей и методов для обнаружения рака молочной железы, которые используют различные алгоритмы обучения. Используя методы и технологии ИИ, такие, как нейронные сети и метод опорных векторов в работах [Lakhani, 2012],

[Zhang, 2013] была достигнута точность классификации опухолей от 74 до 94% на наборе данных из 92 изображений.

Большинство последних работ относящихся к области классификации рака молочной железы ориентированы на цифровые изображения. (WSI) [Zhang, 2013], [Zhang, 2014], [Doyle, 2008], [Singh]. Однако широкое применение систем автоматической классификации рака молочной железы наталкивается на такие проблемы, как высокая стоимость реализации, недостаточная производительность для большого объема данных и противодействия со стороны патолого-анатомов (человеческий фактор). До настоящего времени большинство работ, основанных на гистологическом анализе, выполнялись на небольших дата-сетах.

Существенным сдвигом в этом направлении является создание датасета, состоящего из 7909 изображений молочной железы, полученных от 82 пациентов [Spanhol, 2016]. На этом датасете авторы оценивали различные текстурные дескрипторы и классификаторы и провели эксперименты и достигли точности от 82% до 85%.

На протяжении многих лет основным средством получения признаков изображений для последующей классификации являлись текстурные дескрипторы. Однако недостаток их использования состоит в отсутствии универсальности, в следствие чего, что для каждого типа изображений необходимо разрабатывать свои дескрипторы. Другой подход базируется на системах машинного обучения презентациям. Этот подход не нов, но стал реализуем только с появлением графических процессоров GPU (Graphic Processing Units), которые способны обеспечить сверхвысокую производительность с относительно низкой стоимостью, благодаря параллельной архитектуре [Bengio, 2013].

Альтернативой к этим подходам является использование сверточных сетей (CNN) для обработки медицинских изображений и диагностики. В ряде работ было показано, что сверточные сети способны превзойти обычные текстурные дескрипторы [LeCun, 2015], [Krizhevsky, 2012]. Кроме того, традиционный подход к выделению признаков на основе дескрипторов требует больших усилий и высокого уровня знаний

экспертов и обычно является специфичным для каждой задачи, что препятствует его непосредственное применение к другим задачам.

Поэтому в настоящей статье предлагается новая гибридная сверточная нечеткая сеть CNN- FNN, в которой CNN используется для выделения признаков на медицинских, а нечеткая нейронная сеть FNN используется для классификации обнаруженных опухолей на изображениях на два класса: доброкачественные и злокачественные.

Основная цель настоящей работы состоит в разработке и исследовании гибридной сети CNN-FNN и исследовании ее эффективности в задаче распознавания медицинских изображений молочной железы и классификации опухолей, а также сравнении ее эффективности с известными работами на стандартном датасете медицинских изображений [Spanhol, 2016].

---

#### **Описание базы медицинских данных**

---

Датасет медицинских данных BreakHis [Spanhol, 2016] содержит данные биопсии доброкачественных и злокачественных опухолей тканей молочной железы. Изображения были получены в результате клинических исследований с января по декабрь 2014. Датасет медицинских данных BreakHis состоит из 7909 клинически представительных изображений под микроскопом опухолей молочной железы от 82 пациентов с различным масштабом увеличения 40X, 100X, 200X, and 400X. Все пациенты, которые в течение этого периода исследовались в медицинской лаборатории, были приглашены принять участие в исследовании. Все данные являлись анонимными. Образцы генерировались путем хирургических срезов опухоли молочной железы, специально раскрашивались и отмечались патолого-анатомом.

Основная цель -сохранить оригинальную структуру опухоли и ее молекулярный состав, что позволит наблюдать ее под оптическим микроскопом. Для исследований все образцы разрезались на слайды размером 3 мкм.

Окончательный диагноз каждого случая делался опытным патолого-анатомом, который затем подтверждался дополнительным исследованием, таким как иммунная гистология.

В исследованиях использовался микроскоп системы Olympus BX-50 с увеличением 3.3, связанный с цифровой камерой Samsung SCC-131AN, которая использовалась для получения цифровых изображений молочной железы. Изображения получались в 3 каналах (8-битные каналы RGB) с коэффициентами увеличения 40X, 100X, 200X, and 400X.

На рис. 1-4 представлены 4 изображения с соответствующими коэффициентами увеличения: (a) 40 ×, (b) 100 ×, (c) 200 × и (d) 400 × - полученными из одного слайда ткани молочной железы, содержащего злокачественную опухоль (breast cancer). Выделенный прямоугольник-регион интереса (ROI), который выбирается патолого-анатомом и будет описан в следующем разделе.

В настоящий момент датасет BreakHis содержит 7909 изображений, разбитых на доброкачественные и злокачественные опухоли.

В табл.1 представлено распределение изображений датасета BreakHis [Spanhol, 2016].

**Таблица 1. Распределение изображений по классам и коэффициентам усиления**

| Коэффициент усиления | Доброкачественные | Злокачественные | Всего образцов |
|----------------------|-------------------|-----------------|----------------|
| 40x                  | 625               | 1370            | 1995           |
| 100x                 | 644               | 1437            | 2081           |
| 200x                 | 623               | 1390            | 2013           |
| 400x                 | 588               | 1232            | 1820           |
| Total                | 2480              | 5429            | 7909           |
| Число пациентов:     | 24                | 58              | 82             |

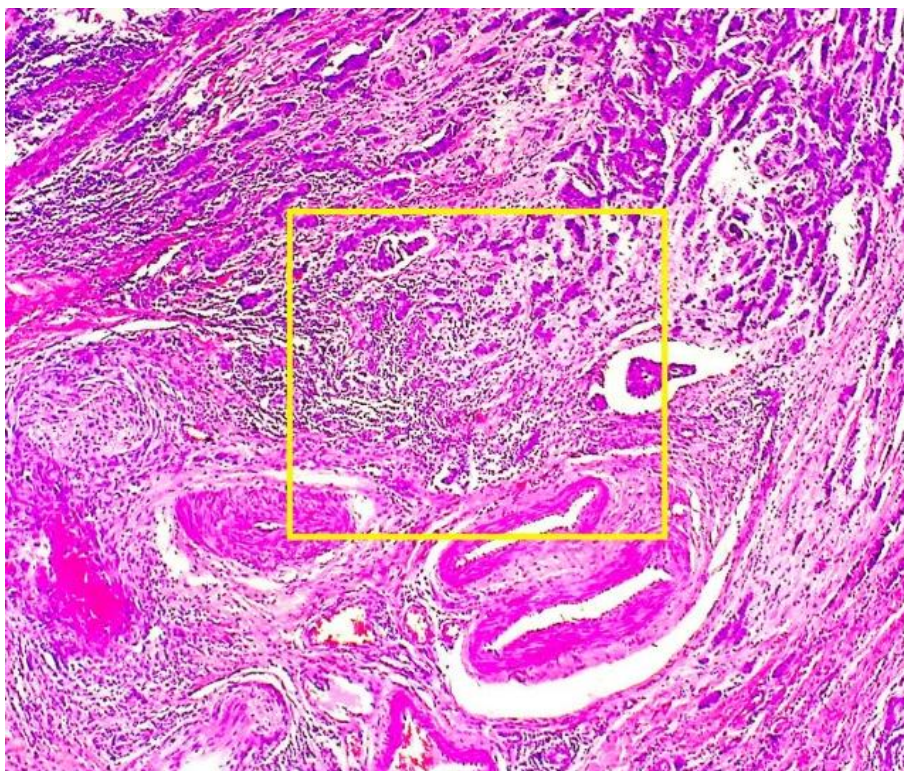


Рис. 1. Слайд злокачественной опухоли с коэффициентом 40X

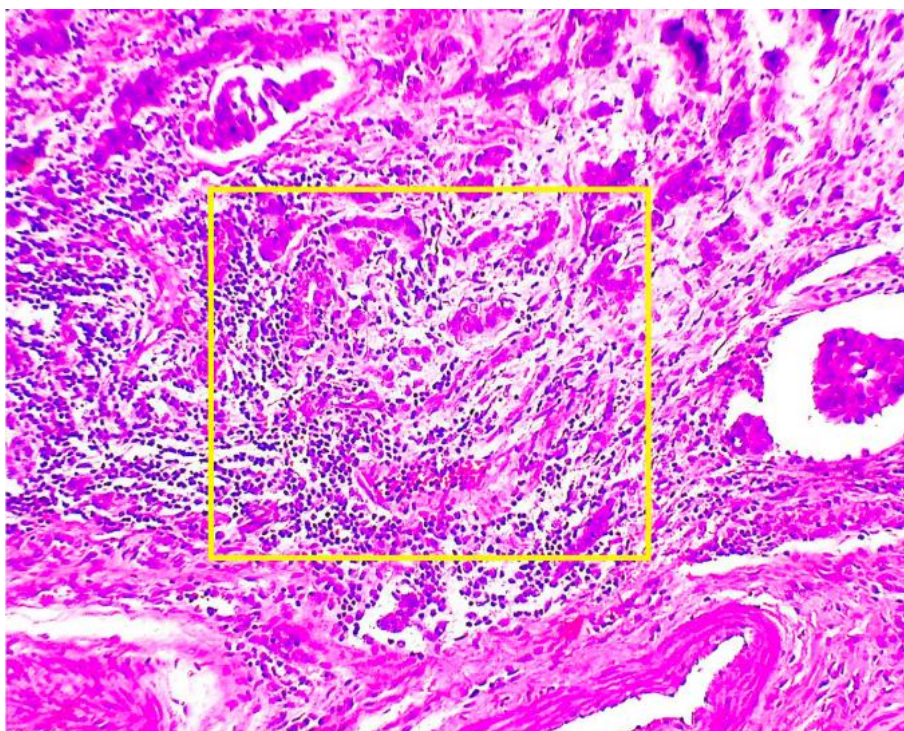


Рис. 2. Слайд злокачественной опухоли с коэффициентом 100X

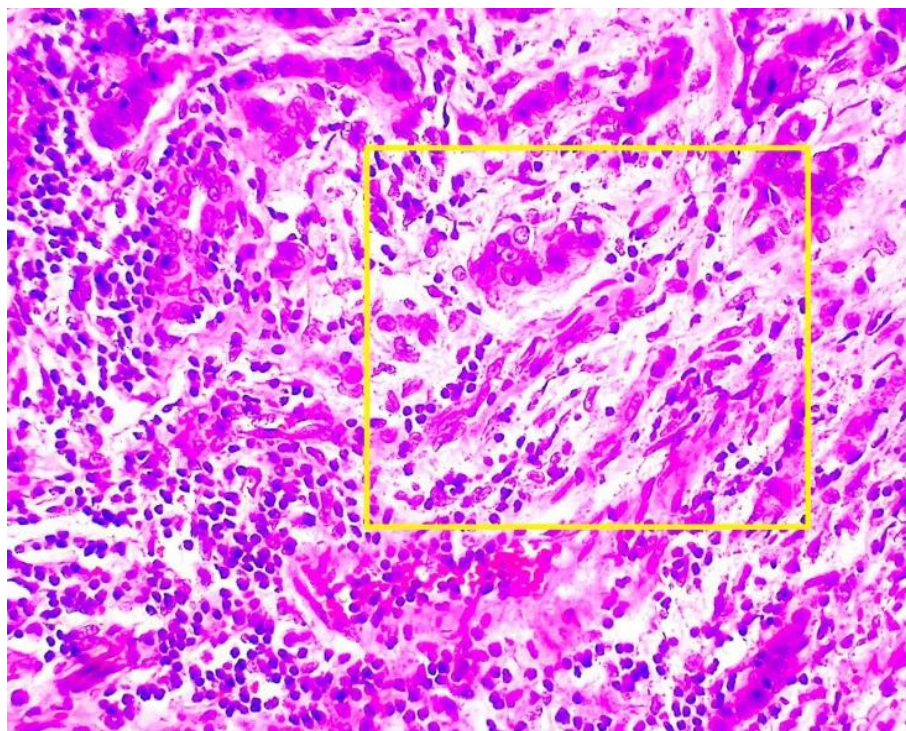


Рис. 3. Слайд злокачественной опухоли с коэффициентом 200X

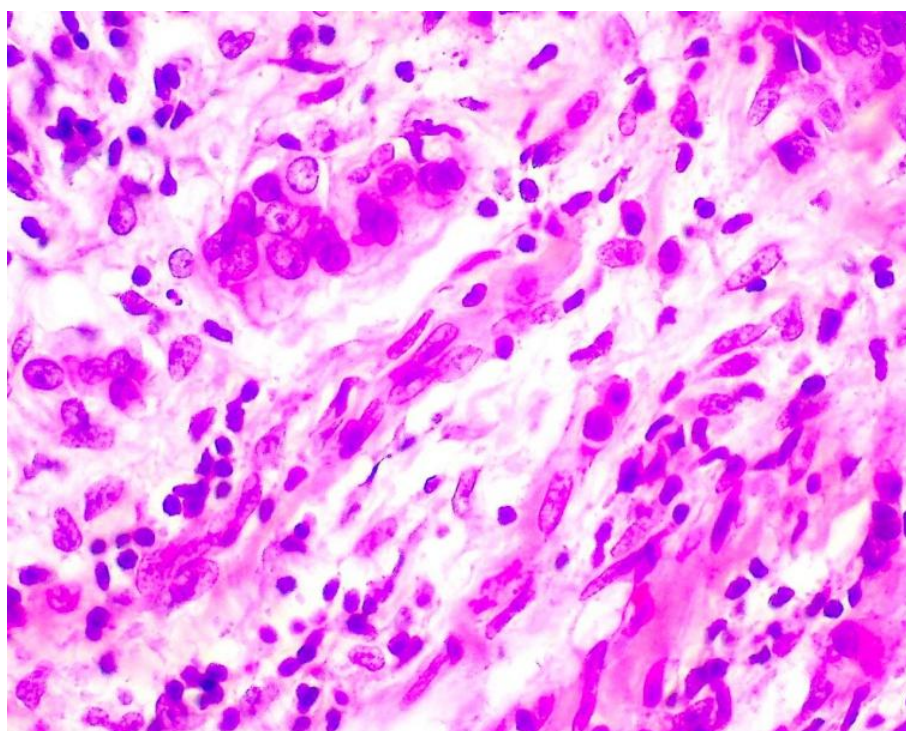


Рис. 4. Слайд злокачественной опухоли с коэффициентом 400X

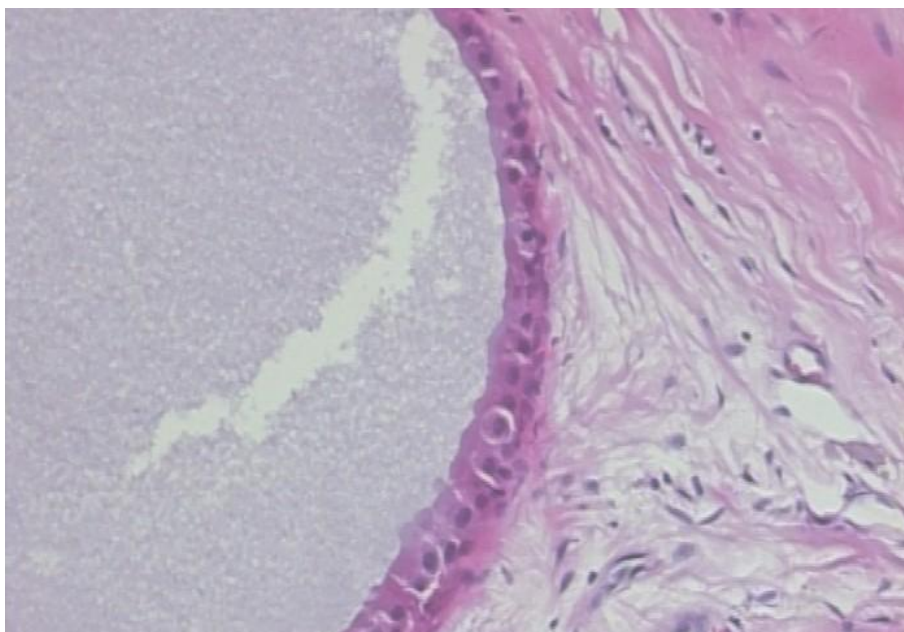


Рис. 5. Слайд доброкачественной опухоли с коэффициентом 100X

---

### Сверточные нейронные сети (краткое описание)

---

Сверточные сети (CNN) – современный метод, широко используемый для обработки изображений. Они обладают способностью извлекать глобальные признаки изображения в иерархическом виде. CNN включает следующие слои [Krizhevsky, 2012].

- **Сверточный слой.** Он рассматривается как основной рабочий компонент сверточной сети и играет жизненно важную роль в этой сети. Ядро свертки (фильтр), обычно представляющий матрицу  $n \times n$ , последовательно проходит по всем пикселям изображения и выделяет информацию из него путем операции свертки.
- **Шаги (Stride) и дополнение (Padding).** Матрица фильтра (ядро) движется по изображению с шагом **Stride**, определяемым его размером, по умолчанию его размер принимается равным 1. Если размер изображения  $5 \times 5$  сканируется фильтром размером  $3 \times 3$  и шагом 1, то мы получим после свертки выходную матрицу  $3 \times 3$ . Однако если мы используем шаг 2, то выходное изображение будет размером  $2 \times 2$ .

Таким образом размер выходного изображения определяется размерами фильтра(ядра) и шага. Чтобы обойти это явление, можно дополнить исходное изображение столбцами и строками, содержащими 0 (Padding). В таком случае размер выходного изображения после свертки будет равным исходному. Такое добавление строк и столбцов, содержащих только 0, называется *zero padding* [Krizhevsky, 2012].

- **Нелинейное преобразование (Nonlinear Performance).** Каждый слой нейронной сети(NN) дает линейный выход и добавление двух таких слоев также осуществляет линейное преобразование. Поэтому увеличение числа слоев нейронной сети не изменяет характера сети. Для преодоления этого недостатка используется нелинейное преобразование в виде следующих функций:
- Rectified Linear Unit (ReLU), Leaky ReLU, TanH, Sigmoid, и т.д.
- Операция пулинга (**Pooling Operation**). Сверточная нейронная сеть производит большой объем информации. Поэтому чтобы уменьшить размерность пространства признаков, используется операция пулинг. Известны несколько стандартных операций пулинга такие, как MaxPooling, AveragePooling.

В нашей работе используется операция MaxPooling, которая выбирает максимальное значение из некоторой подматрицы после свертки.

- Операция **Drop-Out**.  
В результате обучения сети возможно явление переобучения на тестовой выборке, известное как over-fitting. Это явление можно устранить, используя процедуру дропаут (Drop-Out), которая состоит в выключении некоторых нейронов из сети при обучении.
- Решающий слой **Decision Layer**. Для классификации изображений на конце сверточной сети используется решающий слой, обычно в виде многослойного перцептрона (usually MLP). С этой целью используется слой Softmax layer или слой метода опорных векторов (SVM layer). Этот слой реализует нормализованную экспоненциальную функцию и подсчитывает функцию потерь (loss function) для данных классификации.

---

### **CNN модель для классификации изображений**

---

На рис. 6 приведена архитектура сверточной сети VGG-16, которая используется в нашей работе в качестве детектора информативных признаков изображения. Она обучалась различными алгоритмами: стохастическим градиентным методом, методом дифференциальной эволюции (SCD) [Zaychenko, 2009], [Zgurovsky, 2016] и алгоритмом basin hopping [Olson, 2012].

В качестве классификатора полученных признаков в данной работе предложено использовать нечеткую нейронную сеть FNN NEFClass в отличие от известных работ, в которых применялся многослойный перцептрон и слой машины опорных векторов (SVM).

FNN NEFClass была первоначально предложена Д. Науком и В. Крузе (W. Kruse) в [Nauck, 1997]. Она была модифицирована и развита в работах [Zaychenko, 2004], [Zaychenko, 2009] (так называемая FNN NEFClass M). Для нее были разработаны алгоритмы обучения: градиентного спуска, сопряженных градиентов и генетический алгоритм [Zaychenko, 2009] и использованы для распознавания оптических изображений, полученных с использованием мультиспектральной системы.

ННС NEFCLASS M была ранее успешно применена для анализа и распознавания медицинских изображений шейки матки и диагностики опухолей [Zaychenko, 2015]. Основное преимущество ННС NEFClass — это возможность работы с неполными и нечеткими входными данными и осуществлять нечеткую классификацию входных изображений, используя функции принадлежности; высокая скорость и точность классификации. [Zaychenko, 2009], [Zgurovsky, 2016] в сравнении с традиционными методами классификации.

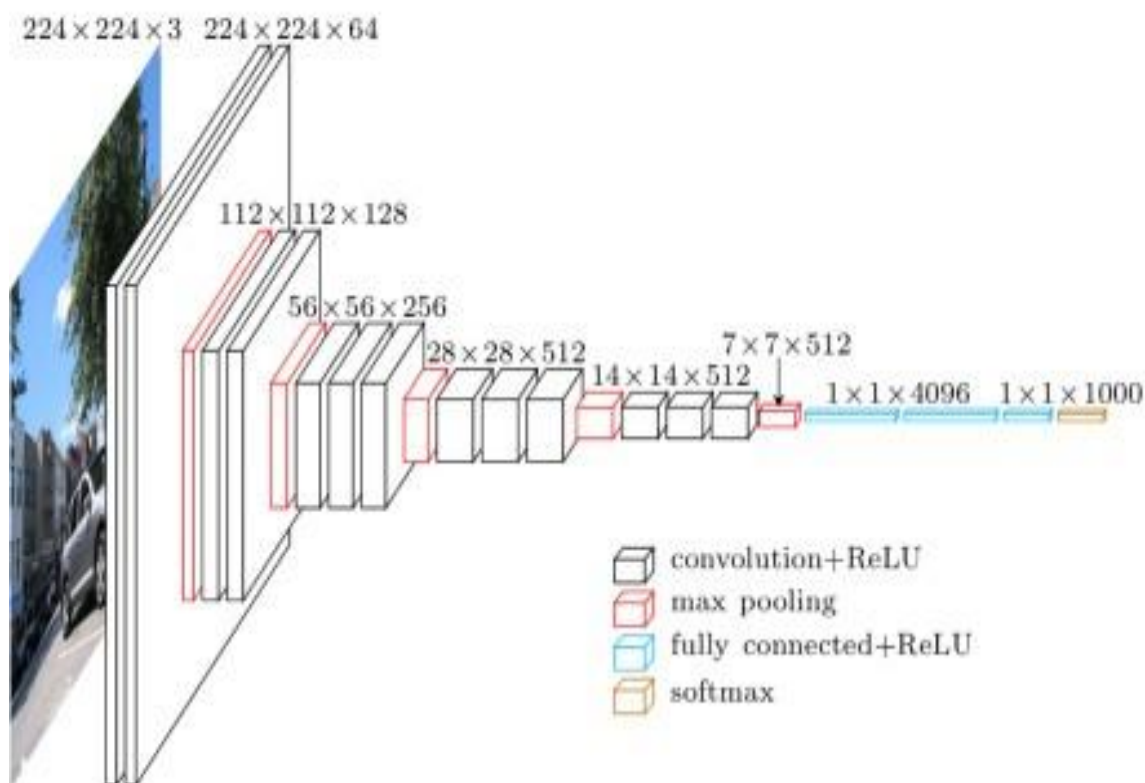


Рис.6. Сверточная нейронная сеть VGG-16

### Экспериментальные исследования и анализ полученных результатов

Как уже указывалось, в данном исследовании было использовано предварительное обучение сверточной сети CNN VGG 16. Существует два основных сценария обучения НС.

**Извлечение признаков.** В этом случае полносвязные слои отключаются, а оставшаяся часть сети используется как экстрактор признаков в новых данных (датасетах).

**Тонкая настройка.** В таком случае новый датасет используется для тонкой настройки предварительно обученной НС.

В настоящем исследовании сеть CNN VGG- 16 была использована для извлечения признаков в медицинских изображениях опухолей молочной железы. После этого найденные признаки подавались на вход ННС

NEFClass. Для обучения использовались 3 алгоритма: basin hopping, стохастический градиентный спуск (СГС) и дифференциальная эволюция (differential evolution) [Zgurovsky, 2016].

Были проведены серии экспериментов и сравнение с результатами предшествующих работ.

В следующих таблицах 2, 3 приводятся результаты классификации с различными параметрами. Все образцы были разбиты на обучающую и тестовую подвыборки в соотношении 80% / 20%.

В первом эксперименте варьировалось число лингвистических переменных и число правил целью определения оптимальных параметров (табл.2).

Таблица 2. Результаты классификации FNN NEFClass

| Число нечетких<br>множеств<br>(лингвистических<br>переменных)/<br>число правил | 40X   | 100X  | 200X  | 400X  |
|--------------------------------------------------------------------------------|-------|-------|-------|-------|
| A                                                                              | 73%   | 74%   | 74.2% | 73.5% |
| 4,2                                                                            | 75.3% | 74.8% | 75.7% | 75.4% |
| 6,2                                                                            | 78.2% | 79%   | 78.4% | 78%   |
| 8,2                                                                            | 76%   | 75.4% | 76.5% | 75.8% |
| 2,4                                                                            | 75%   | 74%   | 73.8% | 73%   |
| 4,4                                                                            | 78.3% | 76.3% | 75.7% | 75.4% |
| 6,4                                                                            | 82%   | 83%   | 82.4% | 83.2% |
| 8,4                                                                            | 82.2% | 81.5% | 81.5% | 83.8% |

| Число нечетких<br>множеств<br>(лингвистических<br>переменных)/<br>число правил | 40X        | 100X       | 200X         | 400X       |
|--------------------------------------------------------------------------------|------------|------------|--------------|------------|
| 2,6                                                                            | 75.4%      | 73.8%      | 74.4%        | 73.2%      |
| <b>4,6</b>                                                                     | <b>90%</b> | <b>91%</b> | <b>90.5%</b> | <b>90%</b> |
| 6,6                                                                            | 89%        | 89.7%      | 90.2%        | 89.5%      |
| 8,6                                                                            | 90.3%      | 90.5%      | 92%          | 91.2%      |
| 4,8                                                                            | 89.3%      | 89.8%      | 89.7%        | 89.3%      |
| 6,8                                                                            | 89.2%      | 88%        | 89.4%        | 88.4%      |
| 8,8                                                                            | 88%        | 87.2%      | 87.2%        | 87%        |

Из анализа этой таблицы следует, что начиная с 6 нечетких множеств на переменную и 6 правил вывода, точность перестает повышаться, а сложность обучения растет. Как следует из этой таблицы, оптимальные значения параметров – 4 нечетких множества (на переменную) и 6 правил.

Для сравнения эффективности предложенной гибридной сверточной сети представим результаты предшествующей работы, полученные с помощью других классификаторов: линейной и полиномиальной машин опорных векторов (SVM) и метода Random forest [Singh, ] ( см. табл.3)

Таблица 3. Сравнение точности различных классификаторов

| классификатор/<br>коэффициент усиления | 40X    | 100X | 200X   | 400X |
|----------------------------------------|--------|------|--------|------|
| Линейная SVM                           | 89%    | 89%  | 88%    | 88%  |
| Полиномиальная SVM                     | 88%    | 90%  | 89%    | 85%  |
| Random forest                          | 89.18% | 88%  | 87.74% | 80%  |
| NEFClass                               | 90%    | 91%  | 90.5%  | 90%  |

Как легко можно видеть из табл.3, ННС NEFClass показывает лучшие результаты, чем альтернативные классификаторы.

В данной работе были применены и исследованы три алгоритма обучения FNN NEFClass: *baisin hopping* [Olson, 2012], стохастический градиентный спуск (СГС) и дифференциальная эволюция [Zaychenko, 2009]. Как показали эксперименты, алгоритмы *baisin hopping* и СГС (*stochastic gradient descent*) дают примерно одинаково хорошие результаты, а алгоритм дифференциальной эволюции оказался значительно хуже.

Следует заметить, что в данной проблеме число признаков, извлекаемых сетью CNN VGG-16, оказалось очень большим – 4096 (проблема Big Data). Поэтому была поставлена задача сокращения числа признаков. Для этих целей был применен метод главных компонент (ГК) [Jindal, 2013]. В таблице 4 представлены результаты такого сокращения признаков.

Таблица 4. Зависимость величины общей вариации от числа главных компонент и примерное время обучения

| Число главных компонент | Нормированная общая вариация (дисперсия) | Примерное время обучения (час) |
|-------------------------|------------------------------------------|--------------------------------|
| 100                     | 0.840587442159                           | ~2 час                         |
| 200                     | 0.897366730496                           | ~3 час                         |
| 250                     | 0.912324353994                           | ~4 час                         |
| 500                     | 0.954868534936                           | ~9 час                         |

Из табл.4 следует, что наиболее приемлемым является сокращение числа главных компонент до 250.

Вследствие лимита времени последующие эксперименты проводились с коэффициентом усиления (magnificence factor) 100X (2081 образцов).

В Таблице 5. приводится точность классификации для различных варьируемых параметров ННС NEFClass.

Таблица 5. Точность классификации для 250 признаков

| Число нечетких множеств / число правил | 100X   |
|----------------------------------------|--------|
| 4/4                                    | 80.64% |
| 4/6                                    | 87.24% |
| 4/8                                    | 88.18% |

В таблице 6 приводится зависимость точности классификации от числа признаков. Легко можно видеть, что точность при сокращении числа признаков до 250 (примерно 16 раз) точность уменьшилась на несколько процентов, но при этом существенно сократилось время обучения.

Таблица 6. Точность классификации с разным числом признаков

| Число нечетких множеств, число правил /число признаков | 100    | 250    | 4096  |
|--------------------------------------------------------|--------|--------|-------|
| 4,4                                                    | 75.23% | 80.64% | 86.3% |
| 4,6                                                    | 83.34% | 87.24% | 91%   |
| 4,8                                                    | 84.21% | 88.18% | 89.8% |

Из данной таблицы можно видеть, что точность с уменьшением числа признаков падает всего на 3-5%, если сравнивать 100 и 250 признаков. Для сравнения была проведена классификация с полным набором 4096 признаков и было установлено, что с уменьшением числа признаков в 16 раз точность упала всего, в среднем на 3-5%.

Это свидетельствует об эффективности применения метода главных компонент для сокращения числа признаков в задачах классификации медицинских изображений.

---

## Выводы

- 1 Рассмотрена проблема анализа медицинских изображений молочной железы и классификации обнаруженных опухолей на два класса: доброкачественные и злокачественные.
- 2 Для распознавания опухолей была разработана гибридная сверточная нечеткая сеть CNN- FNN, в которой CNN VGG 16 была

использована для выделения признаков на изображениях, а ННС NEFClass– для классификации обнаруженных опухолей на основе этих признаков.

- 3 Для обучения ННСNEFClass были предложены и реализованы алгоритмы обучения baissin hopping, стохастический градиентный спуск (СГС) и дифференциальная эволюция.
- 4 Проведены экспериментальные исследования предложенной гибридной CNN-FNN сети в задаче классификации реальных изображений на специальном датасете BreakHis.
- 5 Сравнение точности классификации предложенной гибридной CNN-FNN сети с известными результатами для сверточной сети с алгоритмами классификации SVM и Random forest показало целесообразность использования гибридной сети.
- 6 Рассмотрена и проанализирована проблема снижения размерности признаков в задачах классификации медицинских изображений с использованием метода главных компонент и оценена его эффективность.

---

## Библиография

---

- [Boyle, 2012] P. Boyle and B. Levin, Eds., World Cancer Report 2012. Lyon: IARC, 2012. [Online]. Available:[http://www.iarc.fr/en/publications/pdfs-online/wcr/2008/wcr\\_2012.pdf](http://www.iarc.fr/en/publications/pdfs-online/wcr/2008/wcr_2012.pdf)
- [Lakhani, 2012] S. R. Lakhani, E. I.O., S. Schnitt, P. Tan, and M. van de Vijver, WHO classification of tumours of the breast, 4th ed. Lyon: WHO Press,2012.
- [Zhang, 2013] Y. Zhang, B. Zhang, F. Coenen, and W. Lu, “Breast cancer diagnosis from biopsy images with highly reliable random subspace classifier ensembles,” Machine Vision and Applications, vol. 24, no. 7, pp. 1405–1420,2013.
- [Zhang, 2014] Y. Zhang, B. Zhang, F. Coenen, J. Xiau, and W. Lu, “One-class kernel subspace ensemble for medical image classification,” EURASIP Journal on Advances in Signal Processing, vol. 2014, no. 17, pp. 1–13,2014.

- [Doyle, 2008] S. Doyle, S. Agner, A. Madabhushi, M. Feldman, and J.Tomaszewski, “Automated grading of breast cancer histopathology using spectral clustering with textural and architectural image features,” in Proceedings of the 5th IEEE International Symposium on Biomedical Imaging (ISBI): From Nano to Macro, vol. 61. IEEE, May 2008, pp.496–499.
- [Singh, ] Aditi Singh, Hadi Mansourifar, Hasnain Bilgrami, Nikhil Makkar, Tanay Shah. Classifying Biological Images Using Pre-trained CNNs. <https://docs.google.com/document/d/1H7xVK7nwXcv11CYh7hl5F6pM0m218FQloAXQODP-Hsg/edit?usp=sharing>
- [Spanhol, 2016] F. Spanhol, L. S. Oliveira, C. Petitjean, and L. Heutte, “A dataset for breast cancer histopathological image classification,” IEEE Transactions of Biomedical Engineering, 2016.
- [Bengio, 2013] Y. Bengio, A. Courville, and P. Vincent, “Representation learning:A review and new perspectives,” IEEE Transactions on Pattern Analysis and Machine Intelligence, vol. 35, pp. 1798–1828,2013.
- [LeCun, 2015 ] Y. LeCun, Y. Bengio, and G. Hinton, “Deep learning,” Nature, vol. 521, pp. 436–444, 2015.
- [Krizhevsky, 2012] A. Krizhevsky, I. Sutskever, and G. E. Hinton, “Imagenet classification with deep convolutional neural networks,” in Advances in NeuralInformation Processing Systems 25, 2012, pp.1097–1105.
- [Olson, 2012] Olson,B.,Hashmi,I.,Molloy,K.,andShehu1,A.,Basin Hopping as a General and Versatile Optimization Framework for the Characterization of Biological Macromolecules, Advances in Artificial Intelligence, Volume 2012 (2012), Article ID674832.
- [Nauck, 1997] Detlef Nauck and Rudolf Kruse. New learning strategies for NEFCLASS. In Proc. Seventh International Fuzzy Systems Association World Congress IFSA’97, Vol. IV, pp. 50-55, Academia Prague, 1997.
- [Zaychenko, 2004] Zaychenko Yu. P. , Sevae Fatma, Matsak A.V. Fuzzy neural networks for economic data classification// Vestnik of National Technical University of Ukraine “KPI”, section “Informatic, control and computer engineering. Vol. 42.-2004.-pp.121-133. (rus)

- [Zaychenko, 2009] Zaychenko Yu.P. , Petrosyuk I.M., Jaroshenko M.S. The investigations of fuzzy neural networks in the problems of electro-optical images recognition// System research and information technologies. -2009.- №4.- pp. 61-76. (rus).
- [Zgurovsky, 2016] M. Zgurovsky, Yu. Zaychenko. The Fundamentals of Computational Intelligence: System Approach. Springer International Publishing AG, Switzerland.-2016.-308p.
- [Zaychenko, 2015] Yuriy Zaychenko, Vira Huskova. Recognition of objects on Optical Images in Medical Diagnostics Using Fuzzy Neural Network NEFClass. Intern. Journal Information Models and Analysis, 2015. Vol .4, Number 1, pp. 13-22.
- [Jindal, 2013] N. Jindal. Enhanced Face Recognition Algorithm using PCA with Artificial Neural Networks. / N Jindal, V Kumar // International Journal of Advanced Research in Computer Science and Software Engineering.– 2013. – Vol. 3, pp. 864-872.

---

#### Информация об авторах

---



**Yuri Zaychenko** – Professor, doctor of technical sciences, Institute for applied system analysis, NTUU “KPI”, 03056, Ukraine, Kyiv, Peremogi pr. 37, Corpus 35; e-mail: [baskervil@voliacable.com](mailto:baskervil@voliacable.com), [zaychenkoyuri@ukr.net](mailto:zaychenkoyuri@ukr.net)

*Major Fields of Scientific Research: Information systems, Fuzzy logic, Decision making theory*



**Galib Hamidov**- PhD, doctor philosophy of technical sciences, “Bakielektrikshebeke” JSC, Head of the Information technologies department, AZ 1065, Azerbaijan, Baku, Kazimzade st. 20, e- mail: [galib.hamidov@gmail.com](mailto:galib.hamidov@gmail.com)

*Major Fields of Scientific Research: Information technologies, Data Mining*

## Medical Images of Breast Tumor Diagnostics with Application of Hybrid Fuzzy Convolutional Networks

Galib Hamidov, Yuriy Zaychenko

**Abstract:** *The problem of classification of breast tumors on medical images is considered. For its solution hybrid fuzzy CNN network is developed in which convolutional neural network is used as feature extractor while fuzzy neural network NEFClass is used as classifier. Training algorithms of FNN were developed and implemented. The experimental investigations of the suggested hybrid network on the standard data set BreakHis were carried out and comparison with known results was performed. The problem of data dimensionality reduction is considered and application of PCM method is investigated.*

**Keywords:** *medical diagnostics, breast cancer classification, FNN, CNN, hybrid network*

**ITHEA Keywords:** *1. Computing methodologies 12. Artificial intelligence, 16.5. Model development*

## ИССЛЕДОВАНИЕ НЕЧЕТКИХ НЕЙРОННЫХ СЕТЕЙ С РАЗЛИЧНЫМИ АЛГОРИТМАМИ ВЫВОДА В ЗАДАЧАХ ПРОГНОЗИРОВАНИЯ НА РЫНКАХ ЦЕННЫХ БУМАГ

Тофик Кязимов, Юрий Зайченко

**Аннотация.** В статье рассматривается проблема прогнозирования на финансовых рынках с использованием нечетких нейронных сетей (FNN). Для этого предлагаются нейронные сети с различными алгоритмами Мамдани, Цукамото и Сугено. Рассмотрены структура этих сетей и алгоритмы вывода. Проведены экспериментальные исследования точности прогнозов FNN, представлено и проанализировано сравнение их эффективности при прогнозировании цен акций и индексов рынков на фондовых биржах. Были разработаны и исследованы алгоритмы обучения FNN. Чувствительность точности прогноза в зависимости от алгоритма вывода, количества входов, правил и количества итераций была исследована и сравнена для различных FNN.

**Key words:** financial markets, forecasting, fuzzy neural networks

**ITHEA Keywords:** 1. Computing methodologies 12. Artificial intelligence, 16.5. Model development

---

### Введение

В последние годы нечеткие нейронные сети получили широкие применения в различных задачах вычислительного интеллекта. Их использование определяется их следующими важными свойствами [Зайченко, 2008], [Згуровский, 2013]:

1. Возможность работы в нечеткой и неполной информации, а также качественной информации;

2. Использование экспертной информации в виде нечетких правил вывода;
3. Они являются универсальными аппроксиматорами, т.е. способны аппроксимировать с произвольной точностью непрерывные ограниченные функции от  $n$  переменных.

Вместе с тем они обладают рядом недостатков, а именно:

- база нечетких правил может оказаться неполной или противоречивой;
- функции принадлежности и их параметры, которые задаются экспертом могут оказаться неадекватными реальным моделируемым процессам.

Для устранения указанных недостатков их необходимо обучать. Для этого используется целый арсенал алгоритмов обучения многие из которых были разработаны для обычных нейронных сетей [Згуровский, 2013], [Bodyanskiy, 2009], [Зайченко, 2017].

В настоящее время разработаны различные системы нечеткого логического вывода и ННС, отличающиеся алгоритмами нечеткого вывода и обладающие разными прогнозирующими возможностями. Целью настоящей работы является исследование и сравнительный анализ эффективности различных алгоритмов нечеткого вывода в задачах прогнозирования на рынках ценных бумаг.

---

### Основные этапы нечеткого логического вывода

---

Используемый в разных экспертных и управляющих системах механизм нечетких выводов в своей основе имеет базу знаний, формируемую специалистами предметной области в виде совокупности нечетких предикатных правил вида [Зайченко, 2008]

$\Pi_1$  : если  $x$  есть  $A_1$ , то  $y$  есть  $B_1$ ;

$\Pi_2$  : если  $x$  есть  $A_2$ , то  $y$  есть  $B_2$ ;

...

$$\Pi_n : \text{если } x \text{ есть } A_n, \text{ то } y \text{ есть } B_n,$$

где  $x, x \in X$  — входная переменная (имя для известных значений данных);  $y, y \in Y$  — переменная вывода (имя для значения данных, которое будет вычислено);  $A_i$  и  $B_i$  — функции принадлежности, заданные соответственно на множествах  $X$  и  $Y$ .

Приведем более детальное пояснение. Знания эксперта  $A \rightarrow B$  отражают нечеткое причинное отношение предпосылки (антецедент) и следствия (консеквент), поэтому его можно назвать нечетким отношением и обозначить через  $R$ :

$$R : A \rightarrow B,$$

где « $\rightarrow$ » называют нечеткой импликацией.

Отношение  $R$  можно рассматривать как нечеткое подмножество прямого произведения  $X \times Y$  полного множества предпосылок  $X$  и выводов  $Y$ . Таким образом, процесс получения (нечеткого) результата вывода  $B'$  с использованием данного наблюдения  $A'$  и знания  $A \rightarrow B$  можно представить в виде композиционного правила — нечеткий «modus ponens»:

$$B' = A' \bullet R = A' \bullet (A \rightarrow B),$$

где « $\bullet$ » — операция свертки.

Как операцию композиции, так и операцию импликации в алгебре нечетких множеств можно реализовывать по-разному (при этом будет различаться и получаемый результат), но в любом случае общий логический вывод осуществляется за следующие четыре этапа [Зайченко, 2008],

### **Введение нечеткости (фаззификация —fuzzification).**

Функции принадлежности, определенные на входных переменных, применяются к их фактическим значениям для определения степени истинности каждой предпосылки каждого правила.

### **2. Логический вывод.** Он состоит из двух подэтапов.

2.1. На первом из них определяется степень истинности всей совокупности предпосылок правила. Для этого используется операция пересечения.

2.2. На втором подэтапе определяется выход каждого правила используя нечеткий логический вывод и определяется степень его выполнения (истинности). Для этого используется операция нечеткого пересечения условий и следствия правила.

В качестве правил логического вывода обычно используются только операции  $\min$  (МИНИМУМ) или  $\text{prod}$  (УМНОЖЕНИЕ). В логическом выводе МИНИМУМА функция принадлежности вывода «отсекается» по высоте, соответствующей вычисленной степени истинности предпосылки правила (нечеткая логика «И»). В логическом выводе УМНОЖЕНИЯ функция принадлежности вывода масштабируется при помощи вычисленной степени истинности предпосылки правила.

**3. Композиция.** Все нечеткие подмножества, назначенные к каждой переменной вывода (во всех правилах), объединяются вместе, чтобы сформировать одно нечеткое подмножество для всех переменных вывода. При подобном объединении обычно используются операции  $\max$  (МАКСИМУМ) или  $\text{sum}$  (СУММА). При композиции МАКСИМУМА комбинированный вывод нечеткого подмножества конструируется как поточечный максимум по всем нечетким подмножествам (нечеткая логика «ИЛИ»). При композиции СУММЫ комбинированный вывод нечеткого подмножества формируется как поточечная сумма по всем нечетким подмножествам, назначенным переменной вывода правилами логического вывода.

#### **4. Приведение к четкости (дефаззификация—*defuzzification*).**

Используется, если нужно преобразовать нечеткий набор выводов в четкое число. Существует значительное количество методов приведения к четкости, некоторые из которых рассмотрены ниже.

**Пример.** Пусть некоторая система описывается следующими нечеткими правилами:

$\Pi_1$  : если  $x$  есть  $A$ , то  $z$  есть  $D$  ;

$\Pi_2$  : если  $y$  есть  $B$ , то  $z$  есть  $E$  ;

$\Pi_3$  : если  $w$  есть  $C$ , то  $z$  есть  $F$  ,

где  $x, y$  и  $w$  — имена входных переменных;  $z$  — имя переменной вывода, а  $A, B, C, D, E, F$  — заданные функции принадлежности (треугольной формы).

Процедура получения логического вывода показана на рис. 1. Предполагается, что заданы конкретные (четкие) значения входных переменных:  $x_0, y_0, z_0$ .

На первом этапе на основании данных значений и исходя из функций принадлежности  $A, B, C$  находят степени истинности  $\alpha(x_0), \alpha(y_0)$  и  $\alpha(w_0)$  для предпосылок каждого из трех приведенных правил. На втором этапе происходит «отсечение» функций принадлежности выводов правил ( $D, E, F$ ) на уровнях  $\alpha(x_0), \alpha(y_0)$  и  $\alpha(w_0)$ . На третьем этапе рассматриваются функции принадлежности, усеченные на предыдущем этапе, и проводится их объединение с использованием операции  $\max$ , в результате чего получается комбинированное нечеткое подмножество, описываемое функцией принадлежности  $\mu_\Sigma(z)$  и соответствующее логическому выводу для выходной переменной  $z$ . Наконец, на четвертом этапе при необходимости находят четкое значение выходной переменной, например, с применением центроидного метода — четкое значение выходной переменной определяется как центр тяжести для кривой  $\mu_\Sigma(w)$ :

$$z_0 = \frac{\int_{\Omega} z \mu_\Sigma(z) dz}{\int_{\Omega} \mu_\Sigma(z) dz}.$$

---

### Алгоритмы нечеткого вывода

---

Рассмотрим наиболее применяемые алгоритмы нечеткого вывода, для простоты считая, что база знаний содержит два нечетких правила вида

$\Pi_1$  : если  $x$  есть  $A_1$  и  $y$  есть  $B_1$ , то  $z$  есть  $C_1$ ;

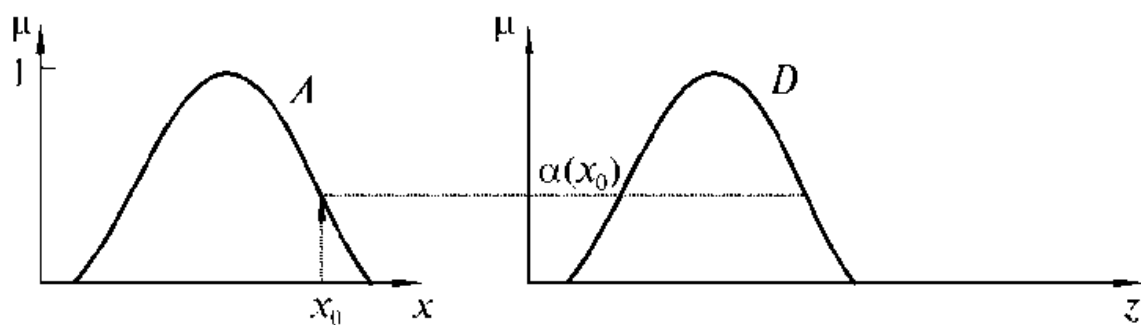
$\Pi_2$  : если  $x$  есть  $A_2$  и  $y$  есть  $B_2$ , то  $z$  есть  $C_2$ ,

где  $x$  и  $y$  — имена входных переменных;  $z$  — имя переменной вывода;  $A_1, B_1, C_1, A_2, B_2, C_2$  — некоторые заданные функции принадлежности.

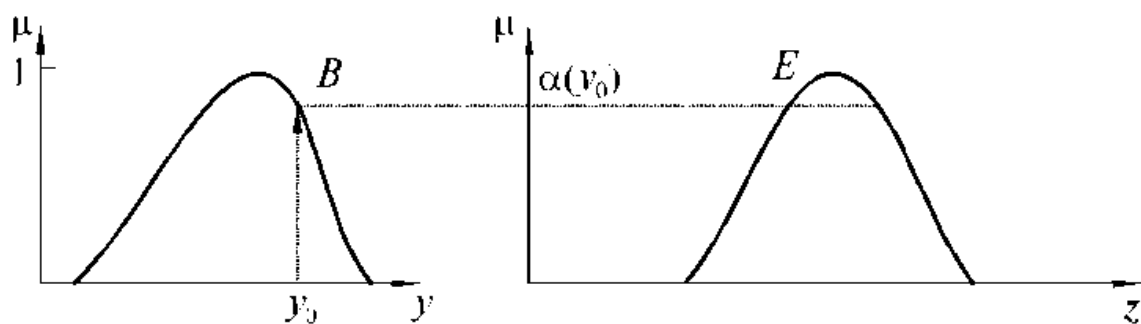
### Алгоритм нечеткого вывода Мамдани

Данный алгоритм соответствует рассмотренному примеру на рис. 1.

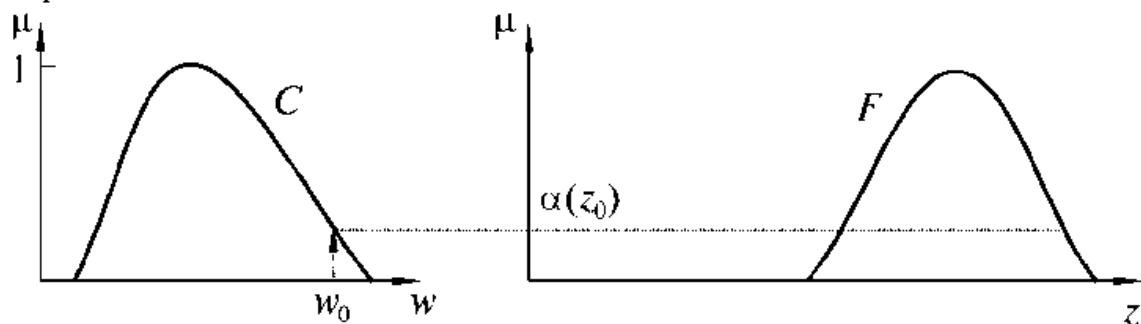
Правило 1:



Правило 2:



Правило 3:



Композиция  
и приведение к четкости:

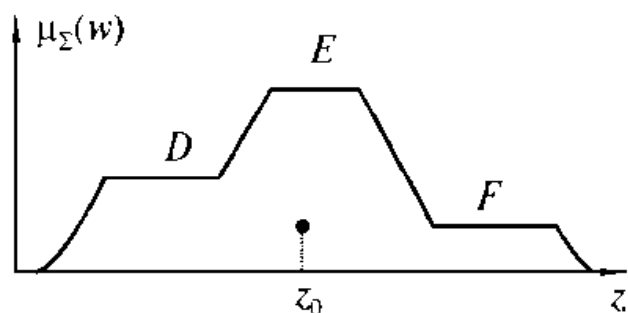


Рис.1. Алгоритм нечеткого вывода Мамдани

При этом он математически может быть описан следующим образом [Зайченко, 2008], [Згуровский, 2013].

1. Введение нечеткости. Находят степени истинности для предпосылок каждого правила:

$$A_1(x_0), A_2(x_0), B_1(y_0), B_2(y_0).$$

2. Логический вывод. Находят уровни «отсечения» для предпосылок каждого из правил (с использованием операции МИНИМУМ):

$$\alpha_1 = A_1(x_0) \wedge B_1(y_0);$$

$$\alpha_2 = A_2(x_0) \wedge B_2(y_0),$$

где через « $\wedge$ » обозначена операция логического минимума (min). Затем находят «усеченные» функции принадлежности выходов правил:

$$C'_1 = (\alpha_1 \wedge C_1(z)); C'_2 = (\alpha_2 \wedge C_2(z)).$$

3. Композиция. Проводится объединение найденных усеченных функций с использованием операции МАКСИМУМ (max, обозначенная далее как « $\vee$ »), что приводит к получению итогового нечеткого подмножества для переменной выхода с функцией принадлежности:

$$\mu_\Sigma(z) = C(z) = C'_1(z) \vee C'_2(z) = (\alpha_1 \wedge C_1(z)) \vee (\alpha_2 \wedge C_2(z)). \quad (1)$$

4. Приведение к четкости. Проводится для нахождения  $z_0$ , например, центроидным методом.

$$Z_0 = \frac{\int_{\Omega} z \cdot \mu_\Sigma(w) dw}{\int_{\Omega} \mu_\Sigma(w) dw} \quad (2)$$

## 2.2 Алгоритм нечеткого вывода Цукамото

Исходные посылки — как у предыдущего алгоритма, но здесь предполагается, что функции  $C_1(z), C_2(z)$  являются монотонными (рис. 2) [Зайченко, 2008], [Згуровский, 2013]:

1. Введение нечеткости (как в алгоритме Мамдани).

2. Нечеткий вывод. Сначала находят уровни «отсечения»  $\alpha_1$  и  $\alpha_2$  (как в алгоритме Мамдани), а затем решают уравнения:

$$\alpha_1 = C_1(z_1) \text{ и } \alpha_2 = C_2(z_2) .$$

При этом определяют четкие значения ( $z_1$  и  $z_2$ ) для каждого исходного правила.

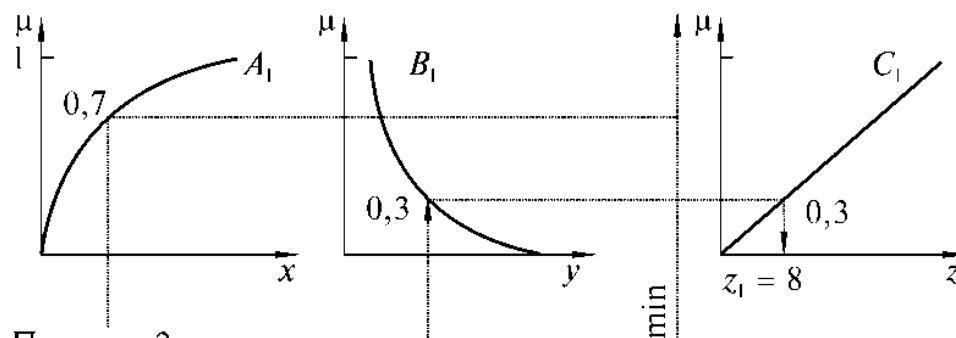
3. Определяют четкое значение переменной вывода (как взвешенное среднее  $z_1$  и  $z_2$ ):

$$z_0 = \frac{\alpha_1 z_1 + \alpha_2 z_2}{\alpha_1 + \alpha_2} .$$

В общем случае дискретный вариант центроидного метода реализуется так:

$$z_0 = \frac{\sum_{i=1}^n \alpha_i z_i}{\sum_{i=1}^n \alpha_i} . \quad (3)$$

Правило 1:



Правило 2:

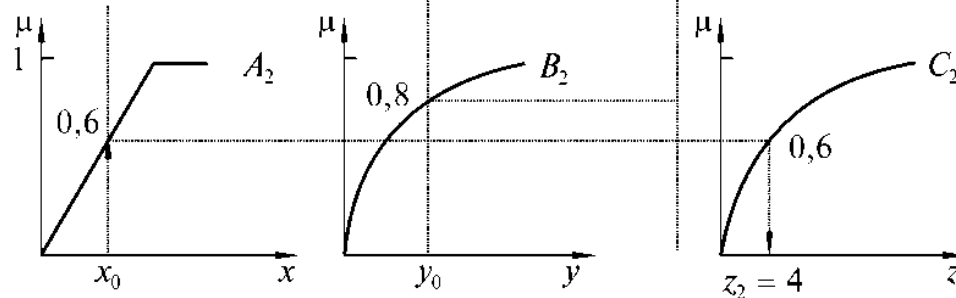


Рис. 2. Алгоритм Цукамото

### 2.3 Алгоритм нечеткого вывода Сугено

Сугено и Такаги использовали набор правил в следующей форме (как и ранее, приведем пример двух правил):

$\Pi_1$ : если  $x$  есть  $A_1$  и  $y$  есть  $B_1$ , то  $z_1 = a_1x + b_1y$ ;

$\Pi_2$ : если  $x$  есть  $A_2$  и  $y$  есть  $B_2$ , то  $z_2 = a_2x + b_2y$ .

Описание алгоритма (рис. 3.):

1. Введение нечеткости (как в алгоритме Мамдани).

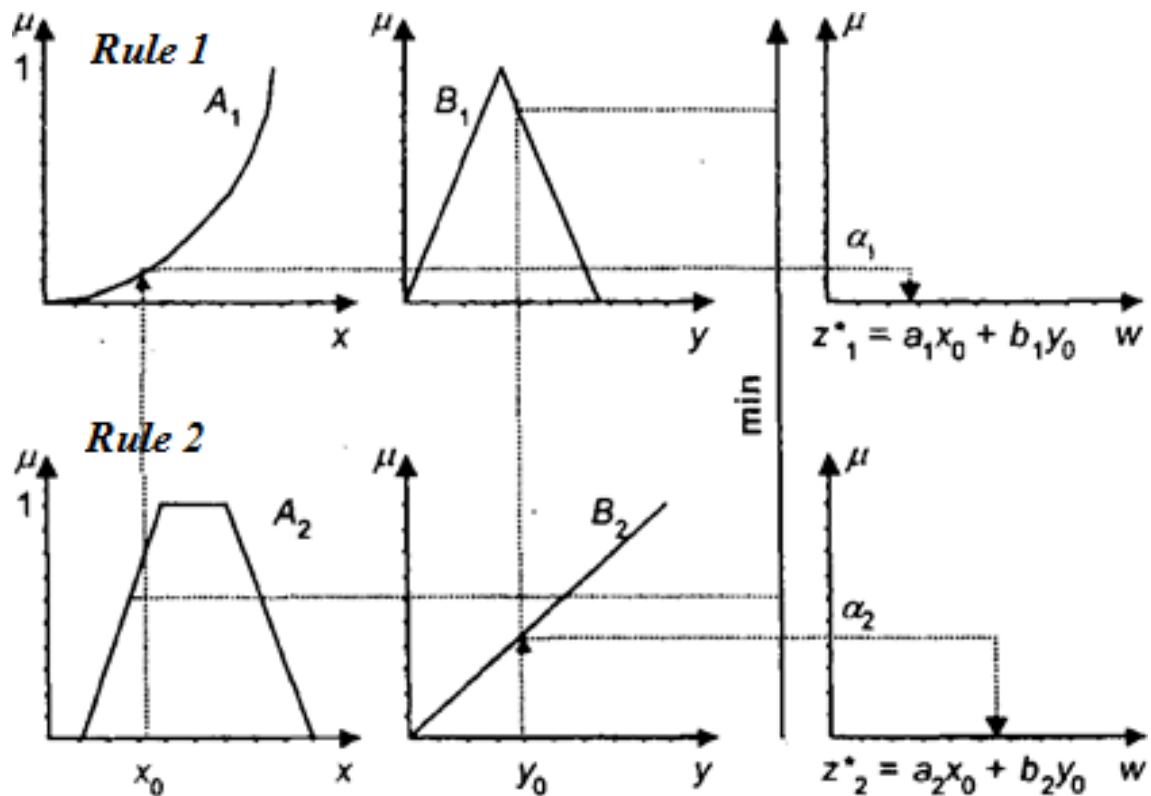


Рис. 3.

2. Нечеткий вывод. Находят  $\alpha_1 = A_1(x_0) * B_1(y_0)$  ,  $\alpha_2 = A_2(x_0) * B_2(y_0)$  и индивидуальные выходы правил:

$$\dot{z}_1 = a_1 x_0 + b_1 y_0 ;$$

$$\dot{z}_2 = a_2 x_0 + b_2 y_0 .$$

3. Определяют четкое значение переменной вывода центроидным методом:

$$z_0 = \frac{\alpha_1 \dot{z}_1 + \alpha_2 \dot{z}_2}{\alpha_1 + \alpha_2} . \quad (4)$$

#### 2.4. Алгоритм нечеткого вывода Ларсена

В алгоритме Ларсена (Larsen) нечеткая импликация моделируется с использованием оператора умножения.

Описание алгоритма (рис.4):

1. Введение нечеткости (как в алгоритме Мамдани).
2. Нечеткий вывод. Сначала, как в алгоритме Сугено, находят значения:

$$\alpha_1 = A_1(x_0) \wedge B_1(y_0) = A_1(x_0) * B_1(y_0)$$

$$\alpha_2 = A_2(x_0) \wedge B_2(y_0) = A_2(x_0) * B_2(y_0)$$

а затем определяют частные выходные нечеткие подмножества:

$$\alpha_1(C_1(z)) , \alpha_2(C_2(z)) .$$

3. Композиция. Проводится композиция выходов правил, как в алгоритме Мамдани с использованием операции МАКСИМУМ (max, обозначенная далее как « $\vee$ »), что приводит к получению итогового нечеткого подмножества для переменной выхода с функцией принадлежности:

$$\mu_z(z) = C(z) = C_1'(z) \vee C_2'(z) = (\alpha_1 \wedge C_1(z)) \vee (\alpha_2 \wedge C_2(z)) . \quad (5)$$

4. Приведение к четкости. Проводится для нахождения  $z_0$ , например, центроидным методом, как в алгоритме Мамдани, формула (2).

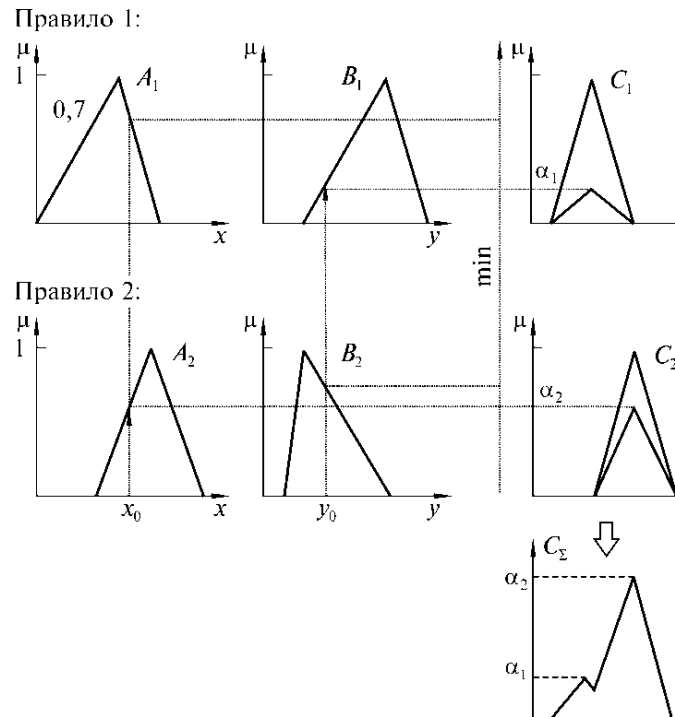


Рис. 4.

### Градиентный метод обучения нечетких нейронных сетей

Рассмотрим ННС с алгоритмом вывода Цукамото. Структура этой сети приведена на рис.5.

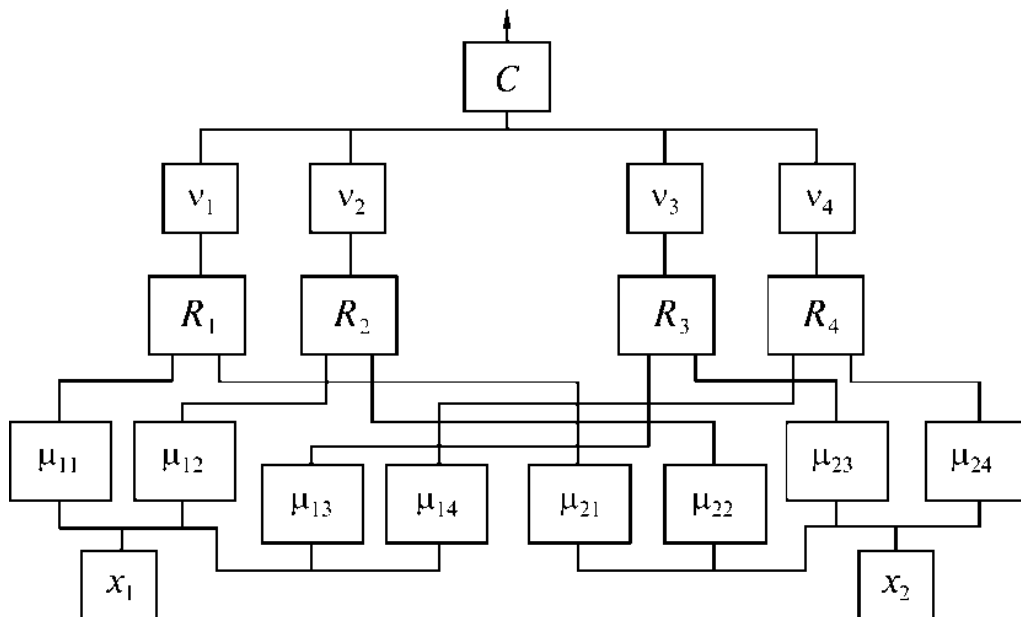


Рис. 5. Структура ННС с выводом Цукамото

Модули  $x_1$  и  $x_2$  здесь представляют входные переменные, и они посылают свои значения в свои  $\mu$ -модули, которые содержат соответствующие ФП.  $\mu$ -модули связаны с  $R$ -модулями, которые представляют собой нечеткие правила «если, то».

Каждый  $\mu$ -модуль передает всем связанным с ним  $R$ -модулям значения ФП  $\mu_{ij}(x_i)$  ее входной величины  $x_i$ .  $R$ -модуль использует операцию пересечения и находит  $\min_i \{\mu_{ij}(x_i)\}$  и передает это значение дальше в  $v$ -модуль, который содержит ФП, описывающую выходные значения.  $v$ -модуль, используя монотонные функции принадлежности, вычисляет  $r_i$  и  $v^{-1}(r_i)$  и передает их в  $C$ -модуль, который вычисляет итоговую выходную переменную — управляющее воздействие  $C$  используя алгоритм центра масс (COA),

где  $n$  — число правил вывода;  $r_i$  — степень, с которой правило  $R_i$  выполняется. Как нетрудно заметить, система на рис. 4 напоминает последовательную многослойную НС, где  $x$ -,  $R$ - и  $C$ -модули выполняют функции нейронов, а  $\mu$ - и  $v$ -модули играют роль адаптируемых весов связей сети.

Пусть ФП  $i$ -го  $\mu$ -модуля, который связан с правилом  $R_k$  описывается таким образом

$$\mu_{ik}(x_i) = \exp \left\{ -\frac{1}{2} * \frac{(x_i - a_{ik})^2}{\sigma_{ik}^2} \right\} \quad (6)$$

где  $a_{ik}$ ,  $\sigma_{ik}$  - параметры ФП, которые определяются в процессе обучения.

$v_k$  - соответствующий вес выходной связи ( $R_k, C$ ) с следующей ФП

$$\mu_k(y_i) = \exp \left\{ -\frac{1}{2} * \frac{(y_i - a_k)^2}{\sigma_k^2} \right\},$$

Пусть пересечение условий правил определяется в форме произведения:

$$\alpha_k = \prod_{i=1}^n \mu_{ik}(x_i) = \exp \left\{ - \sum_{i=1}^n \frac{1}{2} \cdot \frac{(x_i - a_{ik})^2}{\sigma_{ik}^2} \right\}. \quad (7)$$

Допустим, что мы используем центроидный метод фаззификации. Тогда общий выход сети определяется так:

$$z_0 = \frac{\sum_k z_k \alpha_k}{\sum_k \alpha_k}$$

Допустим, что в выходах правил используется монотонная ФП  $C(z)$ . Тогда  $z_k$  определяется как решение следующего уравнения

$$C_k(z_k) = \alpha_k, \quad (8)$$

где

$$C_k(z_k) = \exp \left\{ - \frac{1}{2} * \frac{(z_k - a_k)^2}{\sigma_k^2} \right\}, \quad (9)$$

Решая уравнение:

$$\exp \left\{ - \frac{1}{2} * \frac{(z_k - a_k)^2}{\sigma_k^2} \right\} = \alpha_k,$$

Находим два корня:

$$z_k = a_k \pm \sqrt{2 \ln \frac{1}{\alpha} * \sigma_k}.$$

Первый корень:

$$z_{1k} = a_k - \sqrt{2 \ln \frac{1}{\alpha} * \sigma_k}.$$

Он находится на части кривой где график  $z_r(a_k)$  монотонно возрастает, а второй -  $z_{2k}$ , где график монотонно спадает.

Пусть критерий обучения имеет вид  $E(z) = \frac{1}{2} (z_0 - z^*)^2 \rightarrow \min$ ,

где  $z^*$  - желаемый выход;  $z_0$  – выход сети,

Тогда мы находим следующие производные критерия по параметрам:

$$\frac{\partial E}{\partial a_k} = \frac{\partial E}{\partial z_0} \frac{\partial z_0}{\partial z_k} \frac{\partial z_k}{\partial a_k} = + (z_0 - z^*) \frac{\alpha_k}{\sum_{k=1}^K \alpha_k}, \quad (10)$$

На монотонно возрастающей части кривой:

$$\frac{\partial E}{\partial \sigma_k} = \frac{\partial E_0}{\partial z_0} \frac{\partial z_0}{\partial z_k} \frac{\partial z_k}{\partial \sigma_k} = - (z_0 - z^*) \frac{\alpha_k}{\sum_{k=1}^K \alpha_k} \cdot \sqrt{2 \ln \frac{1}{\alpha}}, \quad (11)$$

На монотонно спадающей части:

$$\frac{\partial E}{\partial \sigma_k} = + (z_0 - z^*) \frac{\alpha_k}{\sum_{k=1}^K \alpha_k} \sqrt{2 \ln \frac{1}{\alpha}}, \quad (12)$$

Для входных  $\mu$  - модулей находим

$$\begin{aligned} \frac{\partial E}{\partial \alpha_{ik}} &= \frac{\partial E_0}{\partial z_0} \frac{\partial z_0}{\partial \alpha_k} \frac{\partial \alpha_k}{\partial \alpha_{ik}} = + (z_0 - z^*) \frac{z_k \sum_k \alpha_k - \sum_k z_k \cdot \alpha_k}{\left( \sum_k \alpha_k \right)^2} \cdot \prod_{k=1}^K \exp \left\{ - \frac{(x_i - a_{ik})^2}{2 \cdot \sigma_{ik}^2} \right\} \cdot \frac{(x_i - a_{ik})}{\sigma_{ik}^2} = \\ &= (z_0 - z^*) \cdot \alpha_k \cdot \frac{z_k \sum_k \alpha_k - \sum_k z_k \cdot \alpha_k}{\left( \sum_k \alpha_k \right)^2} \cdot \frac{(x_i - a_{ik})}{\sigma_{ik}^2} \end{aligned}$$

$$\begin{aligned} \frac{\partial E}{\partial \sigma_{ik}} &= \frac{\partial E}{\partial z_0} \frac{\partial z_0}{\partial \alpha_k} \frac{\partial \alpha_k}{\partial \sigma_{ik}} = \\ &= \left( (z_0 - z^*) \alpha_k \frac{z_k \cdot \sum_{k=1}^K \alpha_k - \sum_k z_k \cdot \alpha_k}{\left( \sum_k \alpha_k \right)^2} \cdot \frac{(x_i - a_{ik})^2}{\sigma_{ik}^3} \right) \end{aligned} \quad (13)$$

Тогда градиентный алгоритм обучения ННС Цукамото имеет следующий вид:

Для связей выходного слоя

$$a_k(n+1) = a_k(n) - \gamma_n \frac{\partial E}{\partial \alpha_k} = a_k(n) - \gamma_n (z_0 - z^*) \frac{\alpha_k}{\sum_k \alpha_k}, \quad (14)$$

$$\begin{aligned}\sigma_i(n+1) &= \sigma_k(n) - \gamma'_n \frac{\partial E}{\partial \sigma_k} = \\ &= \sigma_k(n) \pm \gamma_n (z_0 - z_\Sigma^*) \frac{\alpha_k}{\sum_k \alpha_k} * \sqrt{2 \ln \frac{1}{\alpha_k}}\end{aligned}\quad (15)$$

для входных  $\mu$  - модулей

$$\begin{aligned}a_{ik}(n+1) &= a_{ik}(n) - \gamma_{n_2} \frac{\partial E}{\partial a_{ik}} = a_{ik}(n) - \gamma_n (z_0 - z_\Sigma^*) \cdot \\ &\cdot \frac{\alpha_k \left( z_k \cdot \sum_k \alpha_k - \sum_k z_k \cdot \alpha_k \right)}{\left( \sum_k \alpha_k \right)^2} \cdot \frac{(x_i - a_{ik})}{\sigma_{ik}^2}\end{aligned}\quad (16)$$

$$\sigma_{ik}(n+1) = \sigma_{ik}(n) - \gamma_{n_3} (z_0 - z_\Sigma^*) \frac{\alpha_k \left( z_k \cdot \sum_k \alpha_k - \sum_k z_k \cdot \alpha_k \right)}{\left( \sum_k \alpha_k \right)^2} \cdot \frac{(x_i - a_{ik})^2}{\sigma_{ik}^3}.\quad (17)$$

где  $\gamma_n, \gamma_{n_1}, \gamma_{n_2}, \gamma_{n_3}$  - размер шага спуска.

---

### Экспериментальные исследования и анализ результатов

---

С целью проведения исследований был разработан программный продукт на языке программирования Python в среде JetBrains PyCharm. Он позволяет моделировать работу нечетких нейросетей Мамдани, Цукамото и а Сугено, осуществлять их обучение и выбирать параметры ФП.

В качестве исходных данных для прогнозирования был выбран промышленный индекс Доу Джонса (Dow Jones Industrial Average (DJII)). Данные были взяты из веб-ресурсу [https:// www.investing.com/indices/us-30](https://www.investing.com/indices/us-30). Всего были выбраны 80 значений индекса в период с 2019-01-11 по 2019-05-07.

Прогноз осуществлялся по переменной Price. На этапе подготовки данных значения индекса (Price) нормируются в интервале от 0 до 1. Критериями оценки качества прогноза выбраны MSE (Mean Squared Error) и MAE (Mean Absolute Error).

---

### Анализ и сравнение полученных результатов

---

Сравнение различных алгоритмов проходило при варьировании следующих параметров:

- соотношение обучающей выборки к проверочной ;
- число входов 6, число итераций: 100.

Результаты прогноза для разных методов представлены на рис.6 и в табл.1. А в табл.2 приведены значения метрик MSE та MAE соответствующих алгоритмов

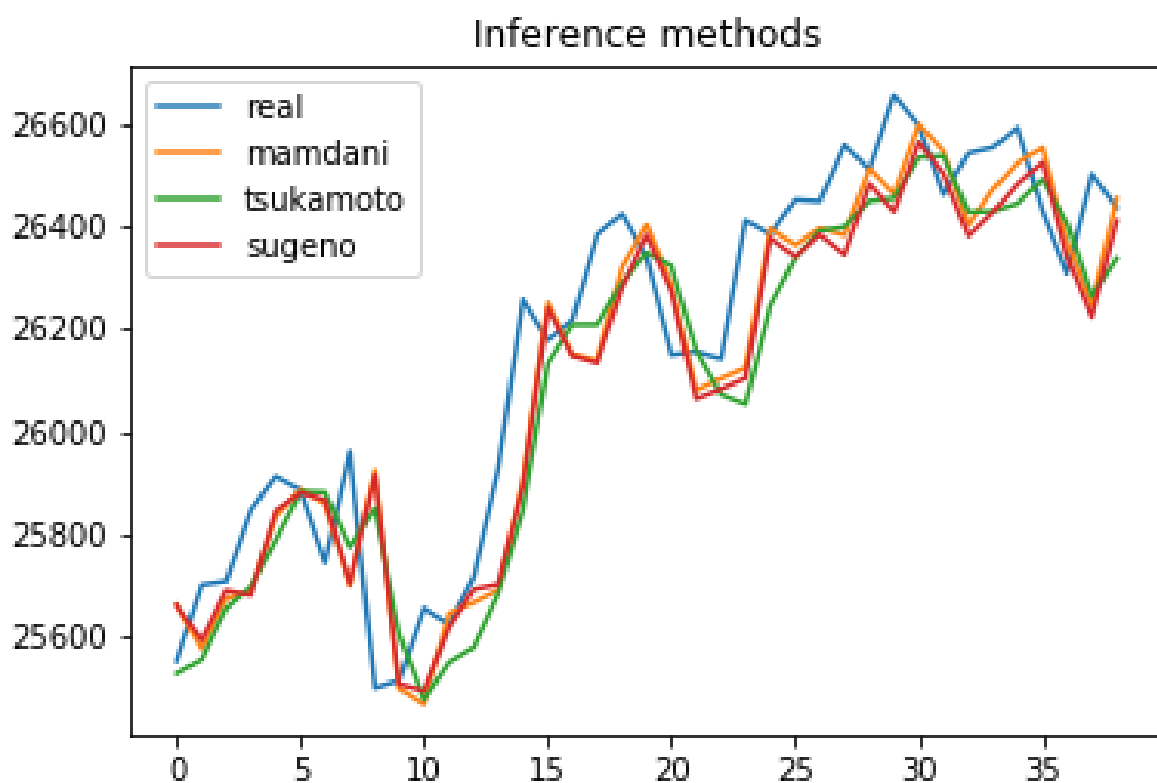


Рис. 6. Сравнение алгоритмов Мамдани, Цукамото та Сугено

**Табл. 1. Реальные та прогнозируемые результаты индекса**

|                   |          |          |          |          |          |          |          |          |
|-------------------|----------|----------|----------|----------|----------|----------|----------|----------|
| Реальные значения | 25554.66 | 25702.89 | 25709.94 | 25848.87 | 25914.1  | 25887.38 | 25745.67 | 25962.51 |
| Мамдани           | 25663.93 | 25578.39 | 25676.91 | 25690.09 | 25835.78 | 25888.9  | 25861.08 | 25701.35 |
| Цукамото          | 25530.06 | 25557.51 | 25657.21 | 25699.97 | 25789.46 | 25884.25 | 25883.06 | 25776.07 |
| Сугено            | 25663.13 | 25594.98 | 25691.98 | 25682.7  | 25847.02 | 25882.16 | 25868.09 | 25704.95 |

**Табл. 2. Значения MSE та MAE в зависимости от алгоритма вывода**

|     | Мамдани  | Цукамото | Сугено   |
|-----|----------|----------|----------|
| MSE | 0.000944 | 0.000999 | 0.001003 |
| MAE | 0.023519 | 0.025614 | 0.024212 |

Сравним работу Мамдани при соотношении обучающей к проверочной выборке 80%/20% и различном числе итераций. Число входов равно 6.

В табл.3 показаны результаты прогнозирования, а в табл.4 соответствующие значения критериев MSE и MAE.

**Табл. 3. Сравнение реальных и прогнозных значений**

|                   |         |         |         |         |         |         |         |         |
|-------------------|---------|---------|---------|---------|---------|---------|---------|---------|
| Реальные значения | 25554.6 | 25702.8 | 25709.9 | 25848.8 | 25914.1 | 25887.3 | 25745.6 | 25962.5 |
| 10 итераций       | 25463.4 | 25504.3 | 25531.8 | 25646.8 | 25676.3 | 25807.5 | 25797.4 | 25718.1 |
| 50 итераций       | 25662.6 | 25565.4 | 25684.3 | 25685.7 | 25837.4 | 25886.1 | 25859.1 | 25702.0 |
| 100 итераций      | 25663.9 | 25578.3 | 25676.6 | 25690.0 | 25835.7 | 25888.9 | 25861.0 | 25701.3 |

**Табл. 4. Значения MSE и MAE при соотношении 80/20**

|     | 10 итераций | 50 итераций | 100 итераций |
|-----|-------------|-------------|--------------|
| MSE | 0.001653    | 0.000947    | 0.000953     |
| MAE | 0.035366    | 0.023586    | 0.023841     |

Следующий эксперимент имел аналогичные входные данные, а соотношения обучающей к проверочной выборке 50/50. В табл.5 приведены результаты прогнозирования, а в табл.6 значения MSE и MAE.

**Табл. 5. Результаты прогнозирования**

|                   |          |          |          |          |          |          |          |          |
|-------------------|----------|----------|----------|----------|----------|----------|----------|----------|
| Реальные значения | 24100.51 | 23592.98 | 23675.64 | 23323.66 | 22859.6  | 22445.37 | 21792.2  | 22878.45 |
| 10 итераций       | 25456.96 | 24456.96 | 24066.91 | 23686.48 | 23391.21 | 23227.27 | 22795.13 | 22185.08 |
| 50 итераций       | 24692.04 | 24160.58 | 23578.21 | 23768.85 | 23577.18 | 23017.44 | 22575.91 | 22028.06 |
| 100 итераций      | 24687.94 | 24139.37 | 23584.02 | 23790.99 | 23548.44 | 23011.34 | 22578.6  | 22012.83 |

**Табл. 6. Значения MAE та MSE при соотношении 50/50**

|     | 10 итераций | 50 итераций | 100 итераций |
|-----|-------------|-------------|--------------|
| MSE | 0.003073    | 0.002744    | 0.002720     |
| MAE | 0.039513    | 0.035193    | 0.035167     |

В следующем эксперименте соотношение обучающей к проверочной выборке было: 30/70. Все остальные параметры оставались прежними.

В табл.7 показаны результаты прогнозирования ННС Мамдани, а в табл.8 соответствующие значения MSE та MAE.

**Табл. 7. Результати прогнозування**

|                   |          |          |          |          |          |          |          |          |
|-------------------|----------|----------|----------|----------|----------|----------|----------|----------|
| Реальные значения | 25379.45 | 25444.34 | 25317.41 | 25191.43 | 24583.42 | 24984.55 | 24688.31 | 24442.92 |
| 10 итераций       | 25577.77 | 25520.29 | 25399.68 | 25466.77 | 25358.33 | 25012.06 | 24939.12 | 24909.9  |
| 50 итераций       | 25712.71 | 25303.38 | 25437.78 | 25400.67 | 25252.11 | 24592.77 | 25139.71 | 24846.86 |
| 100 итераций      | 25710.87 | 25356.81 | 25352.43 | 25408.75 | 25294.8  | 24616.37 | 25091.68 | 24875.04 |

**Табл. 8. Значения MSE и MAE при соотношении выборок 30/70**

|     | 10 итераций | 50 итераций | 100 итераций |
|-----|-------------|-------------|--------------|
| MSE | 0.004609    | 0.003246    | 0.003428     |
| MAE | 0.052707    | 0.040476    | 0.040608     |

В последней серии экспериментов проведено варьировании числа входов при различных соотношениях обучающей и проверочной выборок.. В таблице 9 приведены значения критериев MSE та MAE, а их зависимость показана на Рис.7 и Рис.8.

**Табл. 9. Значения MSE та MAE при разных соотношениях выборок и разном числе входов**

| %     | 80/20    |          |          | 50/50    |          |          | 30/70    |          |          |
|-------|----------|----------|----------|----------|----------|----------|----------|----------|----------|
| Входи | 2        | 5        | 7        | 2        | 5        | 7        | 2        | 5        | 7        |
| MSE   | 0.000967 | 0.000943 | 0.000934 | 0.002499 | 0.003059 | 0.002958 | 0.003139 | 0.003779 | 0.003561 |
| MAE   | 0.023229 | 0.023613 | 0.023613 | 0.033441 | 0.038011 | 0.037819 | 0.039286 | 0.043735 | 0.042632 |

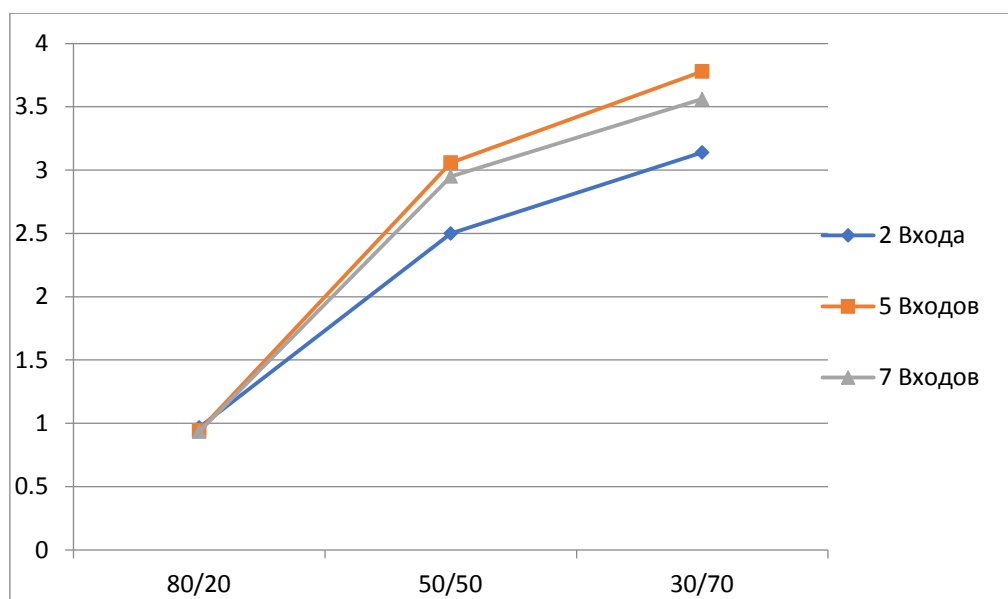


Рис.7. Значения MSE при разных входах и соотношениях выборок (x1000)

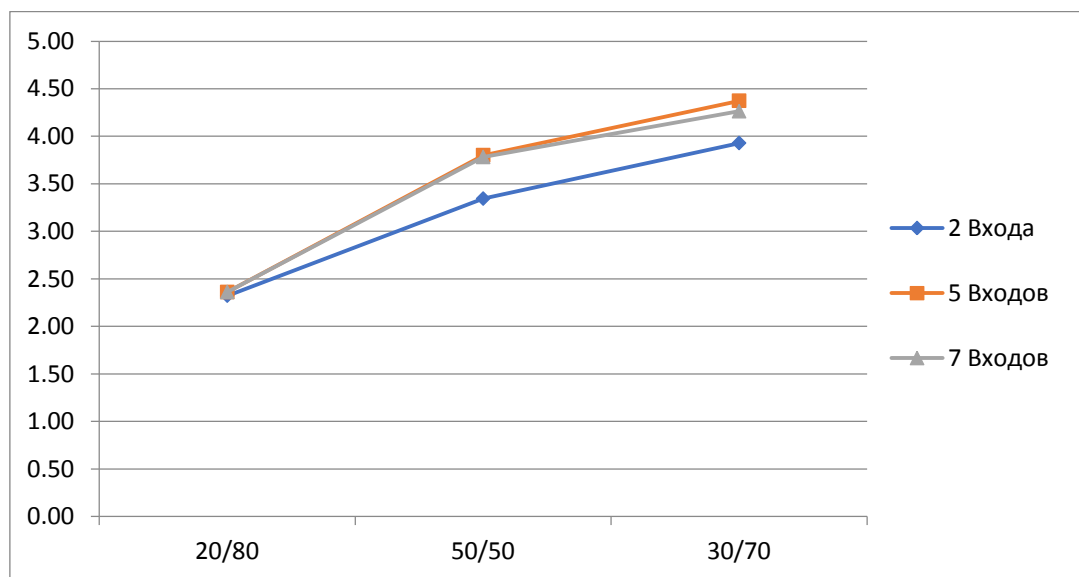


Рис.8. Значения MAE при разном числе входов и соотношениях выборок

Анализируя эти результаты на рисунках можно сделать вывод, что при уменьшении доли обучающей выборки значения критериев растут (точность падает).

Кроме того, наилучшие результаты показали нейросети с наименьшим числом входов. Этот факт можно объяснить тем, что при увеличении числа входов сложность сети возрастает и для ее обучения требуется большее число итераций.

Анализируя результаты прогнозирования различных нейронных сетей можно сделать вывод, что наиболее высокую точность прогнозирования имеет ННС с выводом Мамдани, далее с примерно одинаковыми результатами следуют ННС Цукамото и Сугено.

---

## **Выводы**

---

В работе рассмотрены ННС с выводами Мамдани, Цукамото и Сугено в задаче прогнозирования на рынке ценных бумаг. Рассмотрены алгоритмы их обучения.

Проведены экспериментальные исследования нечетких нейросетей в задаче прогнозирования индекса Доу-Джонса на рынке ценных бумаг NYSE. В процессе экспериментов варьировалось соотношение обучающей к проверочной выборке, число входов. В качестве критериев точности прогнозов использовались MAPE и MAE. Исследования показали достаточно высокую точность прогнозов для всех типов сетей.

При этом наиболее высокую точность прогноза показала ННС с выводом Мамдани.

В целом по результатам исследований можно заключить, что нечеткие нейронные сети являются хорошим инструментом прогнозирования курсов акций и биржевых индексов на рынках ценных бумаг.

---

## Список литературы

---

- [Зайченко, 2008] Зайченко Ю.П. Нечеткие модели и методы в интеллектуальных системах. — К.: Изд. дом «Слово», 2008. — 354 с.
- [Згуровский, 2013] Згуровский М.З., Зайченко Ю.П. Основы вычислительного интеллекта. — К.: Наукова думка, 2013. — 406 с.
- [Bodyanskiy, 2009] Ye. Bodyanskiy, Yu. Zaychenko, E. Pavlikovskaya, M. Samarina and Ye. Viktorov, The neo-fuzzy neural network structure optimization using the GMDH for the solving forecasting and classification problems, Proc. Int. Workshop on Inductive Modeling, Krynica, Poland, 2009, pp. 77-89.
- [Bodyanskiy, 2013] Ye. Bodyanskiy, O. Vynokurova, A. Dolotov and O. Kharchenko, Wavelet-neuro-fuzzy network structure optimization using GMDH for the solving forecasting tasks, Proc. 4th Int. Conf. on Inductive Modelling ICIM 2013, Kyiv, 2013, pp. 61-67.
- [Зайченко, 2017] Ю.П. Зайченко, Галиб Гамидов. Каскадные нейронечеткие сети в задачах прогнозирования на рынках ценных бумаг. Системні дослідження та інформаційні технології, 2017. №2,- с.92-102.

---

## Информация об авторах

---



**Yuri Zaychenko** – Professor, doctor of technical sciences, Institute for applied system analysis, NTUU “KPI”, 03056, Ukraine, Kyiv, Peremogi pr. 37, Corpus 35; e-mail: [baskervil@voliacable.com](mailto:baskervil@voliacable.com), [zaychenkoyuri@ukr.net](mailto:zaychenkoyuri@ukr.net)

*Fields of Scientific Research: Information systems, Fuzzy logic, Decision making theory*



**Tofig Kazimov**– Professor, Doctor of Physics and Mathematics of sciences, Institute of Information Technology of ANAS, AZ1141, Azerbaijan Republic, Baku, B.Vahabzade str., 9A, e-mail: [tofig@mail.ru](mailto:tofig@mail.ru)

*Major Fields of Scientific Research: Information systems, Fuzzy logic, Data Mining, software engineering*

## **Fuzzy Neural Networks Investigations in Forecasting Problems at Stock Exchanges**

**Tofiq Kazimov, Yuriy Zaychenko**

**Abstract.** *The problem of forecasting at financial markets with application of fuzzy neural networks (FNN) is considered in this paper. For this aim neural networks with different algorithms Mamdani, Tsukamoto and Sugeno are suggested. The structure of these networks and inference algorithms are considered. The experimental investigations of forecasting accuracy of FNN were carried out and comparison of their efficiency in predicting of share prices and markets indices at stock exchanges is presented and analyzed. The training algorithms for FNN were developed and explored. The sensitivity of forecasting accuracy in dependence on inference algorithm, number of inputs, rules and iterations number was investigated and compared for different FNN.*

**Key words:** *financial markets, forecasting, fuzzy neural networks*

**ITHEA Keywords:** *I. Computing methodologies I2. Artificial intelligence, I6.5. Model development*

## **УНИВЕРСАЛЬНЫЕ РЕШАЮЩИЕ ПРАВИЛА: ПОДДЕРЖКА И ПРИНЯТИЕ РЕШЕНИЙ В ЗАДАЧАХ РАСПОЗНАВАНИЯ, ДИАГНОСТИКИ, ВЫБОРА, КЛАССИФИКАЦИИ**

**Крислов А.Д., Крислов В.А., Чумаченко В.В.**

---

### **ПОСТАНОВКА ЗАДАЧИ И ИСХОДНЫЕ ПОЛОЖЕНИЯ**

---

Где ни просвищет грозный меч,  
Где конь горячий ни промчится, ...  
(А. Пушкин)

Принятие и обоснование решений является неотъемлемой частью целого ряда задач, относимых к области искусственного интеллекта. В задачах классификации и распознавания образов, в задачах выбора и оценки оказывается необходимым принимать решения, обосновывать выбор и т. д. (см., напр., [1, 2] и др.) В подавляющем числе случаев решающие правила строятся в предположении о независимости признаков, составляющих описание распознаваемых или анализируемых объектов. Желание учесть зависимости между элементами описания требует весьма больших затрат времени и памяти и в итоге приводит, как правило, к учету лишь парных зависимостей. Мало того, достаточно редкими являются случаи, когда разработчики или пользователи обладают информацией о реальных величинах условных вероятностей появления признаков (друг от друга) в различных процессах, в распознаваемых классах etc.

В то же время довольно трудно представить задачу, в которой описание распознаваемых объектов складывалось бы из признаков, независимых в совокупности. В медицинской или технической диагностике, при машинном чтении текста или в природо-охранных задачах - всюду признаки,

составляющие описание анализируемых объектов, связаны друг с другом довольно большими условными вероятностями появления-непоявления (см. напр., [3, 4] и др.). Применение в этих случаях решающих правил, работающих с признаками как с независимыми, приводит к ошибкам, тем большим, естественно, чем больше признаки зависят друг от друга.

Таким образом, существует задача: построить такие решающие правила, кото-рые в возможно большей мере учитывали бы связи и зависимости между признаками - без привлечения таких громоздких методов, как полный перебор или сходные с ним.

Построение требуемой решающей функции (или набора функций) удобно провести на примере задачи распознавания/классификации [5].

Рассмотрим множество подлежащих распознаванию ситуаций, состоящее из  $S$  подмножеств или классов:

$$M = \{M_j\}; j = 1, \dots, S \quad (1)$$

В нашем случае понятие "класс" объединяет некоторую совокупность объектов (ситуаций), для которых определено понятие близости или сходства.

Каждый класс описывается совокупностью параметров или признаков общим числом  $n$ . Это означает, что в общем случае описание данного класса (и данной ситуации) включает в себя  $n$  компонент:

$$Q = \{Q_i\}; i = 1, \dots, n \quad (2)$$

Каждый признак  $Q_i$  в том или ином объекте характеризуется определенной степенью выраженности  $v_i$ ; для простоты будем считать признаки двоичными, на общность наших построений это не повлияет.

Обозначим каждую анализируемую ситуацию символом  $f_k$ , где  $k$  - текущий номер данной рассматриваемой ситуации или объекта. Данная ситуация в терминах нашего описания (2) выражена следующим образом:

$$f_k = \{v_1, \dots, v_i, \dots, v_n\}. \quad (3)$$

Для принятого упрощения объект  $f_k$  представлен в пространстве признаков двоичным  $n$ -мерным вектором.

Задача состоит в том, чтобы построить распознающую систему  $W$  (в виде автономного устройства или программы), могущую принимать решения об отнесении очередной анализируемой ситуации к определенному подмножеству  $M_j$  из выражения (1).

Это означает, что нужно найти, сформировать такое решающее правило, такое количественное обоснование, которое бы наилучшим образом позволяло отнести каждую из рассматриваемых ситуаций к тому или иному классу. Слова «наилучшим образом отнести» означают, в частности, что оптимальность решающих правил является понятием относительным, и нужно признать определенную зависимость выбираемых (конструируемых) решающих правил от целого ряда внешних условий. Среди таких условий могут быть требование минимального риска или максимальной близости к эталону и т. д. Наиболее распространенным является требование сохранения минимальной ошибки, то есть максимальной достоверности даваемых системой ответов. В настоящей работе уделяется внимание именно таким решающим правилам.

---

## СОСТАВЛЕНИЕ ОПИСАНИЯ ОБЪЕКТА И ПРИНЯТИЕ РЕШЕНИЙ

---

Сперва производится обучение системы  $W$  различению предполагаемых классов, например, путем показа реальных известных объектов или ситуаций и запоминания их характеристик, т. е. формирования эталонов.

В результате, описание внешнего мира в достаточно простом случае может быть представлено в нашей распознающей системе набором или одной из матриц вида

$$C_n = \|p_{ji}\|, \quad (4)$$

где строки пронумерованы по  $j$  (классы), а столбцы – по  $i$  (признаки).

Здесь  $p_{ji}$  представляет собой число, могущее отражать различные величины: информационный, частотный, диагностический (экспертный)

или другой вес  $i$ -того признака для  $j$ -того класса, его нечеткую оценку и т. д. Каждая строка этой матрицы представляет собой эталонное описание класса  $M_j$ , которое и нужно хранить в памяти устройства и с которым будет сравниваться неизвестный распознаваемый объект.

Этап принятия решения об отнесении неизвестной ситуации к тому или иному классу, состоит в следующем.

Входной вектор  $f_k$ , (состоящий из единиц и нулей, (выражение (3)), оценивается по степени сходства с каждой из строк матрицы  $C_n$ , выражение (4).

Эта оценка (обозначим ее  $X_{kj}$  – оценка  $k$ -того входного объекта на  $j$ -том эталоне) производится по тому или иному критерию и имеет смысл меры близости к данному эталону, степени сходства с ним, вероятности порождения данного вектора данным классом и т. д., - в зависимости от выбранного характера описания и других обстоятельств. Назовем это число интегральной характеристикой данного входного объекта относительно данного эталона. В упоминавшемся ранее пространстве признаков эта оценка есть расстояние между концом данного входного вектора и точкой (или областью), соответствующей описанию данного класса (см., например, [5, 6, 7] и др.).

После вычисления  $X_{kj}$  для всех  $j$  оказывается возможным реализовать следующее решающее правило:

$$f_k \in M_j / R_f = \text{extr}_j X_{kj}, \quad (5)$$

то есть «данный объект принадлежит подмножеству  $M_j$ , если оценка принимает экстремальное значение ( $R_f$ ) на эталоне с номером  $j$ ». Если мы оцениваем в указанном пространстве расстояние до «идеального» представителя данного класса, экстремум должен иметь смысл минимума; если ищем вероятность порождения данного вектора данным классом – смысл максимума и т. д.

Таким образом, решающее правило (5) требует непосредственного численного определения степени сходства данного неизвестного объекта

или ситуации с каждым из эталонов, запасенных в памяти системы, и указывает в качестве ответа тот из классов, с которым это количественно измеренное сходство оказывается наибольшим.

### РАЗНОВИДНОСТИ РЕШАЮЩИХ ПРАВИЛ И УЧЕТ ЗАВИСИМОСТЕЙ МЕЖДУ ПРИЗНАКАМИ

Принятие решения по правилу (5) для случая, когда параметры рассматриваемых ситуаций имеют вероятностную природу, целесообразно производить следующим образом [3, 5]:

$$f_k \in M_j / R_f = \max_j P(M_j / f_k) = \max_j \prod_i^n p_{ji}^{(v_i)}, \quad (6)$$

где:  $P(M_j / f_k)$  – вероятность принадлежности данного вектора  $f_k$  к классу  $M_j$ ;

$p_{ji}^{(v_i)}$  – равняется  $p_{ji}$ , если  $v_i=1$  и равняется  $1-p_{ji}$ , если  $v_i=0$  здесь подразумевается, что символ  $p_{ji}$  из выражения (4) обозначает вероятность появления данного признака в представительной совокупности объектов данного класса.

Отметим, что в решающем правиле (6) величина  $P(M_j / f_k)$ , подсчитываемая как произведение для данного  $f_k$  всех  $p_{ji}^{(v_i)}$  при фиксированном  $j$ , собственно, и представляет собой оценку  $X_{kj}$  из выражения (5). В нашем случае эта величина есть скалярное произведение вектора  $f_k$  (то есть конкретной входной реализации) на  $j$ -тую строку матрицы  $C_n$ . Отметим также, что при вероятностной постановке задачи нет более весомых критериев для принятия правильного решения, чем решающее правило (6).

Итак, в выражении (6) величина  $R_f$  указывает класс, который мог породить данный вектор  $f_k$  с наибольшей вероятностью, - если признаки были независимы.

Для учета зависимостей между признаками введем вторую стадию обучения [5, 6]. Будем снова предъявлять системе  $W$  множество

известных объектов класса  $M_j$  и для каждого объекта вычислять  $X_{kj}$ , то есть его оценку на своем эталоне. Далее будем фиксировать частоту (или вероятность) появления различных величин  $X_{kj}$ .

Результатом второй стадии обучения для данного  $M_j$  будет кривая распределения вероятностей (КРВ)  $P(X_{kj}/M_j)$  различных оценок  $X_{kj}$  всех предъявленных  $f_k \in M_j$  на эталоне своего класса (здесь  $\epsilon$  – символ принадлежности).. Поступив так же с другими классами, получим набор КРВ для всех подлежащих анализу классов решений.

Теперь для принятия решения может быть предложено следующее правило:

$$f_k \in M_j / R_\Sigma = \max_j P(X_{kj} / M_j). \quad (7)$$

Здесь решение принимается не по максимуму оценки данного  $f_k$  на данном эталоне, а по максимуму вероятности данной оценки для соответствующего класса решений [5].

Анализируя процедуру обучения, можно сказать, что механизм формирования КРВ выявляет и закрепляет количественные связи между признаками, имеющие место внутри данного класса ситуаций: если какие-то группы признаков являются связанными, более частыми (или наоборот), то это отразится на форме соответствующей КРВ и, что, собственно, и важно, - будет учтено решающим правилом (7). Разработчик может не знать существующих условных вероятностей между признаками, - вводимая стадия обучения (и – соответствующее правило принятия решения) снимают эту проблему.

По характеру немонотонности этих кривых можно судить о величине зависимостей между признаками (и, кстати, использовать КРВ как инструмент для работы с ними – отбора, ранжирования и т. д.).

Однако этого мало. Мы можем продолжить вторую стадию обучения: будем взвешивать объекты класса  $M_j$  на эталонах/строках других классов (обозначим для этого случая текущие номера других классов через  $r$ :  $r = 1, \dots, s$ ) и сформируем соответствующие КРВ. Теперь описание каждого

класса построено как бы по принципу дополнителности, - в памяти решающей системы имеются слепки событий, пропущенные не только через «свои» входы (в нашем случае – это взвешивание на своих эталонах), но и через входы/эталонны других классов, как бы через другие рецепторные каналы. При этом ситуация является симметричной.

Теперь эталоном класса является совокупность

$$B_j = \{P(X_{kr} / M_j)\} \quad (8)$$

кривых (функций) распределения вероятностей оценок  $X_{kr}$  векторов этого класса (т. к.  $j$  фиксировано) на каждой строке матрицы  $C_n$ .

Многомерный вариант решающей функции теперь принимает вид:

$$f_k \in M_j / R = \max_j \prod_r [P(X_{kr} / M_j)] \quad (9)$$

Согласно этому решающему правилу процедура распознавания неизвестного вектора  $f_k = \{v_1, \dots, v_i, \dots, v_n\}$  состоит в следующем:

- а) находятся интегральные оценки  $X_{kj}$  для данного вектора  $f_k$  для всех  $j$ ;
- б) для каждого  $j$  определяем по кривым  $B_j$  (8) величины  $P(X_{kr})$  для всех  $r$ ;
- в) для каждого  $j$  определяются произведения найденных в п. (б) величин;
- г) максимальное значение произведения указывает номер класса, к которому следует отнести неизвестный вектор.

В силу ряда обстоятельств имеет место тот факт, что разные векторы  $f_k$  получают одну и ту же оценку  $X_{kj}$ . Происходит это вследствие представления вероятностного диапазона конечным числом градаций, из-за наличия  $p_{ji} = 0,5$  и т. п. Вследствие этого даже при достаточно представительной обучающей выборке система не сможет в процессе экзамена узнавать каждый вектор «в лицо», а только по вероятности оценки для группы векторов. Возможные при этом ошибки в распознавании при применении правила (7) являются, таким образом, платой за отказ от полного перебора.

Однако, если эти разные векторы, одинаково оцениваемые на одной строке матрицы  $C_n$ , взвешивать на разных строках, как это предусматривается выражением (9), то они будут получать различные оценки  $X_{kr}$ . Поэтому решающее правило, учитывающее вероятности оценок  $X_{kr}$  (9), обеспечивает меньшее количество ошибок при распознавании, чем правило (7). Это происходит потому, что многомерные распределения (8), получаемые во второй стадии обучения, более полно реализуют информацию о классах, содержащуюся в матрице  $C_n$ , чем это происходит в правиле (7), - более полно используется информация, которая имеется в обучающей последовательности.

Рассмотрим теперь величину

$$Y_{kj} = |X_{kj}^0 - X_{kj}|,$$

которая характеризует отклонение веса, подсчитанного для неизвестного вектора  $f_k$ , от наиболее вероятного в этом классе значения  $X_{kj}^0$ :

$$P(X_{kj}^0) = \max_k P_{ik}(X_{kj} / M_j); \quad k \in j.$$

Тогда, распространив эту оценку на многомерный случай:

$$Y_{kj}^0 = \{ |(X_{kj}^0 / M_r) - (X_{kj} / M_r)| \}; \quad r, j = 1, \dots, s,$$

можно предложить такую модификацию правила (9):

$$f_k \in M_j / R^0 = \min_j \prod_r [(Y_{kj}^0 / M_r) - (Y_{kj} / M_r)] \quad (10)$$

Упрощенным вариантом правила  $R^0$  (10) является выражение:

$$R_r = \min_r Y_{kj}, \quad (11)$$

которое при определенном характере КРВ (например, наличие одного экстремума) может явиться вполне приемлемой заменой правила (10), давая при этом значительную экономию памяти, времени и т. д.

В дальнейшем при конструировании решающего устройства или программного обеспечения задач принятия решений необходимо оценить,

оправданы ли для достижения выигрыша в уменьшении ошибок необходимые усложнения (запоминание кривых, увеличение объема вычислений и т. д.). При использовании данных по выбору числа градаций вероятности можно получить достаточно простую систему, тем более, если реализуется правило (11). В дальнейшем очень большой интерес представляет сравнение предлагаемых модификаций решающих правил в различных задачах на больших массивах исходных данных.

---

### **ЗАМЕЧАНИЯ О ПРАКТИЧЕСКОЙ РЕАЛИЗАЦИИ И ПРИМЕРЫ ПРИМЕНЕНИЯ**

---

При реализации любого из рассмотренных решающих правил представляется целесообразным применить некоторые технические рабочие приемы, повышающие качество решающей системы, облегчающие общение с пользователем и т. д.

Первый из таких приемов связан с желанием предоставить системе возможность отказываться от решения в сомнительных ситуациях, если, например, при чтении текста устройство вместо буквы видит кляксу.

В этом случае можно рекомендовать установление порогов снизу на величину  $X_{kj}$ , требование ощутимого расстояния до ближайшего конкурента и т. д. Введение та-кого порога дает возможность системе еще и оценивать свои ответы: при каких-то значениях  $X_{kj}$  давать определенный ответ, при других – сопровождать его сигналом сомнения, при третьих – отказываться от ответа. Сам такой отказ означает введение «нулево-го» класса, то есть  $j = 0, 1, \dots, s$ . Подобный прием в различных вариантах реализуется сейчас в разных распознающих программах.

Далее, в ряде случаев (при диагностике, при социально-экономическом анализе и др.) полезно, кроме основного ответа, указывать ближайшее конкурирующее решение. Реализация этого приема существенно расширяет возможности системы, особенно, если к тому же выдавать количественные оценки предлагаемых решений.

Отдельно должен быть рассмотрен вопрос о том, какое число  $N$  градаций вероятности появления признака должна различать система. Практика

показывает, что для многих практических применений вполне достаточно трех-пяти градаций с неравно-мерным разбиением всего диапазона вероятностей (среднюю градацию выбирать больше, крайние – поменьше; сами эти величины легко могут быть рассчитаны и т. д.).

В качестве одного из вариантов этого последнего предложения может быть рас-смотрена для решающих правил (7) и (9) ступенчатая аппроксимация полученных КРВ, например, П-образной формой.

Описанные в настоящей работе решающие правила и их различные модифика-ции были применены при решении ряда задач из различных предметных областей.

В задачах медицинской дифференциальной диагностики с применением ре-шающих правил (6) и (7) осуществлялся диагноз туберкулеза легких [8], активного и неактивного ревматизма [9], других нозологических форм. В первом случае достовер-ность распознавания составила (для  $S=3$ ,  $n=31$ , решающее правило (6)) около 90% при 6,3% отказов от распознавания. Во втором (для  $S=5$ ,  $n=26$ , весовые коэффициенты для симптомов по Джонсу – 2:1, экзаменационная выборка – 170 историй болезни) было получено для модифицированного варианта решающего правила (7) 93,3% правильных ответов при 8,3% отказов от постановки диагноза; врачебный диагноз на том же мате-риале дал 82,4% и 11% соответственно.

Значительное количество опытов было связано с распознаванием машинопис-ных и рукописных букв и цифр. Для этих знаков составлялось описание в терминах геометрических признаков (например, вертикальные, горизонтальные и наклонные линии), а затем проводилось обучение и классификация с помощью различных РП. Опишем кратко результаты одного из этих экспериментов. На различных рукописных цифрах, вписанных в габаритную рамку десятью различными людьми, не знавшими о цели эксперимента, без особых ограничений на написание, были опробованы РП (6), (7) и (9) в различных модификациях. Число знаков, взятых для эксперимента, - 1000; число признаков  $n = 36$ ;  $v = 2$ ;  $S = 10$ . При формировании КРВ использовался дискретный вариант представления; производилось выделение признаков, составление эталонов и

последующее распознавание. Результаты распознавания для трех названных РП соответственно составили 84,5%, 85,9% и 93,1% [3, 5].

Кроме того, отдельно видоизмененное решающее правило (5) было опробовано при распознавании слабостилизованных рукописных цифр с применением идеологии и аппарата размытых множеств. Эта идеология была применена на двух уровнях принятия решений: для выделения признаков и для узнавания знаков. Результат распознавания составил 84% правильных ответов [10]. Помимо того, что для слабостилизованных рукописных цифр этот результат хорош и сам по себе, - он еще иллюстрирует возможности применения строгих методов в плохо формализованных областях.

Наряду с описанными были проведены эксперименты по распознаванию с помощью РП (6) и (7) звуковых команд (данные ИППИ АН СССР) и по разработке рекомендаций для капитанов сухогрузных судов по нормативам скорости их движения в рейсе (данные Черноморского Морского Пароходства). И здесь качество ответов составляло 90 - 95% [5].

Среди других областей применения, кроме перечисленных, следует назвать распознавание нефтеносных и водоносных слоев по данным геофизической разведки (145 групп замеров в пластах Бугульминского месторождения; применялись две модификации решающего правила (7)), а также некоторые задачи экологического и социально-экономического анализа, в которых решающие правила, сходные с (6) и (7), использовались в целях квалиметрии [3, 6, 7, 11]. Определенная группа вопросов методологии, в частности, - по построению решающих правил, рассмотрена в работах [11 - 15].

Во всех перечисленных случаях, хоть и в разной степени, введение второй стадии обучения и, соответственно, применение решающих правил, учитывающих зависимости между признаками, - существенно уменьшало число ошибок и повышало качество принимаемых решений.

Полученные результаты показывают, что предложенные алгоритмы количественного обоснования решений успешно могут работать в

различных сложных ситуациях и найдут применение в задачах управления и диагностики, в экономике, природо-охранной деятельности и т. д.

---

## ЛИТЕРАТУРА

---

1. R. Schlaifer. Analysis of Decisions under Uncertainty. McGraw Hill, 1979.
2. R.B. Banerji. A Language for Pattern Recognition. Patt. Rec., 1978, 1.
3. А.Д. Крислов. Модифицированные решающие функции для распознавания сложных ситуаций. Тр. УИ Всесоюзн. Симп. по теоретич. киберн., Т.3, Тбил., 1974.
4. V. Gladun, N. Vashchenko. Analytical Processes in Pyramidal Networks. Proc. of 5-th Intern. Conf. ITA 2000, Foi-Comm., Sofia, 2000.
5. А.Д. Крислов. Алгоритмы количественного обоснования решений для систем, работающих в условиях неопределенности. Знание, Киев, 1984.
6. A. Krissilov, V. Krissilov, A. Shutko. Decision Making Procedures that Operate with Dependent Features and Their Environmental Applications. Proc. of 5-th Intern. Conf. ITA 2000, Foi-Comm., Sofia, 2000.
7. V.A. Krissilov, A.D. Krissilov, R.A.Tarasenko. Transformation of Object Feature Space Under the Goal of Evaluation. Proc. of Conf. IPMU'98, Paris, 1998, pp.1901-.1908.
8. А.И. Штейнберг, А.Д. Крислов. Решение задачи дифференциальной диагностики туберкулеза легких с помощью гибридных решающих функций. «Кибернетика и вычислительная техника», вып. 59, Киев, 1983.
9. A.A. Popov, A.D. Krissilov et al. On Automation of Medical Diagnostic Procedures. Proc. of IFIP-71 Congr., TA-7, Nederl., 1971, pp.841-847.

10. А.Д. Крисиллов, И.В. Хихловская, Н.Н. Гогунская. Экспериментальная проверка некоторых алгоритмов описания и распознавания применительно к стилизованным рукописным знакам. Тр. 5-ой Всесоюзн. конф. «Автоматизация ввода письменных знаков в ЦВМ», Т.2, Каунас, 1984.
11. A.D. Krissilov. Towards a New Econo-Ecological Order for the Black Sea Region.// Int. Leadership Sem., Internat. Ocean Institute – Black Sea Operational Centre, Mamaia, Romania, Sept., 1999, pp.87-96.
12. А. Крисиллов, Е. Соловьева, А. Уемов. Краткий методологический меморандум – ч. I. Information Science & Computing. International Book Series: Knowledge – Dialogue – Solution, N. 15. – ITHEA, Sofia, 2009.
13. А. Крисиллов, А. Уемов. О шагах системного синтеза (продолжение «Методологического меморандума» – часть II). Information Models of Knowledge, v. 19. Proc. of Intern. Confer. "Knowledge – Dialogue – Solution" (KDS – 2010), Kiev, Sept. of 2010, ITHEA, Kiev – Sofia, 2010.
14. А. Крисиллов. «Шаги системного синтеза и построение Интеллектуального Агента в задачах представления знаний», Труды Междун. Конфер. «КДС - 2010». Варна, июнь 2010.
15. А. Крисиллов. «Содержание некоторых этапов системного анализа и синтеза (продолжение «методолог. меморандума» – ч. III)» Proc. of Intern. Confer. "III – 2011"), ITHEA, Sofia, 2011.

---

### Информация об авторах

---

**Крисилов А.Д.** – Институт проблем рынка и экономико-экологических исследований НАН Украины, 65044, Одесса-44, Французский б-р, 29,

e-mail: [adkrissilov@list.ru](mailto:adkrissilov@list.ru)

**Крисилов В.А.** – Одесский Государственный Политехнический Университет, 65044, Одесса-44, пр. Т. Шевченко, 1;

e-mail: [krissilovva@gmail.com](mailto:krissilovva@gmail.com)

**Чумаченко В.В.** - ООО “GERC” Одесса-29, ул. Канатная, 48,

e-mail: [viachesv@chumachenko.net](mailto:viachesv@chumachenko.net)

## UNIVERSAL DECISION RULES: SUPPORT AND DECISION-MAKING IN RECOGNITION, DIAGNOSIS, SELECTION, CLASSIFICATION

**Krisilov A.D., Krisilov V.A., Chumachenko V.V.**

**Abstract:** *The decision-making functions are described, that allow take into account relations and de-pendence among the features in such tasks as classification, choice, pattern reco-gnition and evaluation. The algorytms proposed in this paper provide with high level of decision support and situations recognition under uncertain conditions. Some examples are given.*

## ГРАФО-АНАЛИТИЧЕСКИЙ МЕТОД ОЦЕНКИ СОСТОЯНИЯ ДИАГНОСТИРУЕМОГО ОБЪЕКТА

Игорь Арефиев, Теа Мунджишвили

**Abstract:** Авторы предлагают метод диагностики технического объекта, когда набор данных для оценки его состояния превышает известный экспертам множественный порог его характеристик. о состоянии исследуемого объекта даже при теоретической бесконечности исходных параметров. Метод основан на графо-аналитических представлениях и логико-вероятностной моделирование.

**Keywords:** диагностика, контроль, объект, граф, параметр, показатель, фактор.

---

### Введение

Научно-технический прогресс в конце 20 века привёл к внедрению в практику судостроения и судоремонта новых методов организации производства с учётом прогресса в нанотехнологиях. Этот процесс стал заметен и в формировании судостроительных комплектов: энергетических установок, комплектующих материалов и изделий, систем навигации, деталей корпусов и отделки помещений, судовых средств погрузки и разгрузки. Указанный прогресс требует повышения уровня эксплуатационно-технических показателей качества всех судовых элементов. В связи с этим, расчётный срок службы судов и их элементов так же увеличивается. Согласно условиям эксплуатации, безопасности плавания и сроков гарантии, устанавливаются сроки технического осмотра. Ясно, что элементы судов и судовой техники изнашиваются неравномерно. С другой стороны, технический прогресс постоянно требует

замены оборудования и реновации любой конструкции даже между установленными периодами технических осмотров. Износа оборудования часто не совпадают с установленными сроками осмотров. В этих условиях наблюдается постоянный рост числа контролируемых параметров о состоянии объекта (судна) и его элементов. Следствием такого положения стало резкое увеличение числа контролируемых показателей качества любого элемента судовой техники. Здесь диагностика объекта сталкивается с „проклятием размерности”: любое сочетание исходных показателей вызывает лавинообразное (многосвязное) сочетание малозначных показателей с базовыми (основными). Тогда вывод о значимости конкретного показателя оказывается возможным только на экспертном уровне [6].

Здесь важно заметить, что экспертные заключения об основных и малозначимых показателях носят условный характер. На пример, при повышенных (пониженных) внешних температурах, критерий вязкости масла двигателя внутреннего сгорания важен, а в условиях обычной температур им можно пренебречь. Другой пример: колебания напряжения в сети освещения малозначны, а в системах навигационной автоматики весьма важны и требует специальной аппаратуры стабилизации. Анализ существующих методов контроля объектов по конечному множеству показателей показал их низкую эффективность [4].

Выявление новых показателей и характеристик определяет необходимость в дополнительном инструментарии повышения качества комплексного измерения их параметров. Традиционные средства алгоритмического моделирования подобных процессов (циклические и вложенные циклы) сталкиваются с размерность бесконечности и решить указанной задачи не могут. К этому добавим, что практически все оценки параметров носят вероятностный (статистический) характер. Они требуют постоянной коррекции

Так возникает проблема: разработать методологию оценки состояния диагностируемого объекта для набора его конечных факторов в условиях

выявленного множества его основных (базовых) и текущих (малозначных) показателей.

---

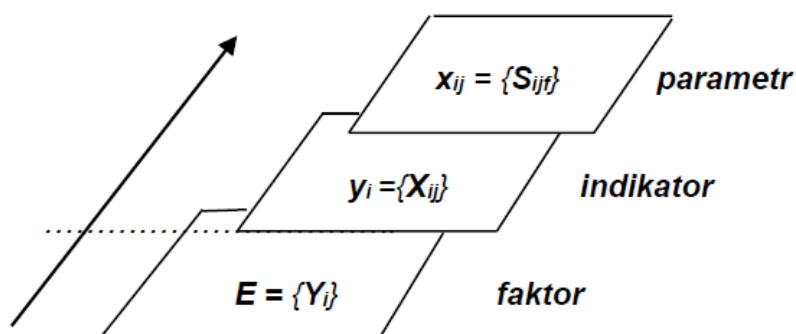
## ПОСТАНОВКА ЗАДАЧИ

---

Анализ любого изделия начинается с выявления множества параметров, по которым необходимо оценить его состояние. Дать такую оценку на бесконечном множестве параметров для всех элементов судна нереально. Этот процесс зависит от степени изученности параметров, технических, теоретических и статистических вычислений, измерений и инструментальных средств определения реальной оценки каждого из них. Степень достоверности измерений и вычислений параметра определено допусками (колебаниями) которые исследователь может принять за возможную ошибку ( $\Delta x_{ijf}$ ). Обычно выбирают группу наиболее значимых параметров с их допусками: 5-7% от измеряемой величины. Их принимают за критерии оценки состояния объекта (показатель). Остальные же трактуются, как ограничения формируемой модели объекта. Такой подход приводит к тому, что нарушение любого ограничения рассматривается как факт неработоспособности объекта (брак, ненадёжность, низкое качество и т.п.). По правилам Теории Систем, такой объект (элемент, изделие) считается неэффективным [7, 8]. Искать на этом пути множество Парето не имеет смысла, т.к. оно само является подмножеством эффективных решений по  $x_{ij}$  (Рис.1).

Следовательно, найти оптимум качества объекта можно только по группам критериальных характеристик [1].

Остальные характеристики объекта определяют показатели, которые фиксируют факт неработоспособности, ненадёжности, отклонения контролируемой величины от заданной. Это вызывает отбраковку данного изделия. Здесь необходима декомпозиция структуры объекта (судна) и элементаризация изделий по контролируемым параметрам на каждом из уровней (фактор).



**Рис. 1.** Схема факторного анализа состояния объекта

На начальном уровне диагностики рассматриваются факторы ( $E$ ), определяющие общее состояние объекта (судна). Ввод в модель конкретных факторов или их сочетания  $\{Y_i\}$  задаёт Морской Регистр. При этом, все выбранные факторы считаются равно-важными. Такое решение выгодно отличает логико-рефлексивный метод от других методов, в том числе и логико-вероятностных. Он позволяет оперативно исключать (включать) дополнительные факторы при прогнозировании, избегая вероятностные оценки. Следующий этап декомпозиции относится к функциональным комплексам ( $X$ ): силовая установка, палубное оборудование, навигационные системы, движительные системы и т.п. Наконец, диагностируются параметры конкретных устройств, элементов, блоков ( $S$ ). На каждом из уровней допустима процедура замены, реновации, регуляции, если изделие не соответствует заданным техническим условиям. Указанные технические условия (размеры, характеристики, уровень износа) принимаются за эталон либо одного изделия, либо комплекса или конструкции в целом. Здесь важно отметить, что технология выбора характеристик элементов и обоснование их значений для Баз Данных на каждом уровне аналогична. В связи с этим остановимся на исходном уровне оценки состояния исследуемого объекта – организация контроля параметров отдельных изделий и деталей, составляющих комплексы и конструкции судна.

Эффективным способом решения таких задач диагностики можно считать семантическую сеть, когда для пары вершин  $i$  и  $j$  соответствует дуга величины  $\kappa(i,j) \in [0,1]$ . Эта величина определяет степень принадлежности данной пары нечёткому отношению. В реальности, указанный процесс характеризует неполноту знаний субъекта, который осознаёт причинно-следственную связь  $i - j$ , но не может объяснить промежуточные звенья рассуждений. Интуитивно понятно, что между  $i$  и  $j$  существует множество подобных пар, но для их идентификации исследователь не обладает соответствующими знаниями и данными [2]. К тому же, семантическая сеть упрощённо носит линейно-плоский (двумерный) характер. В таких случаях формальная модель принимает вид

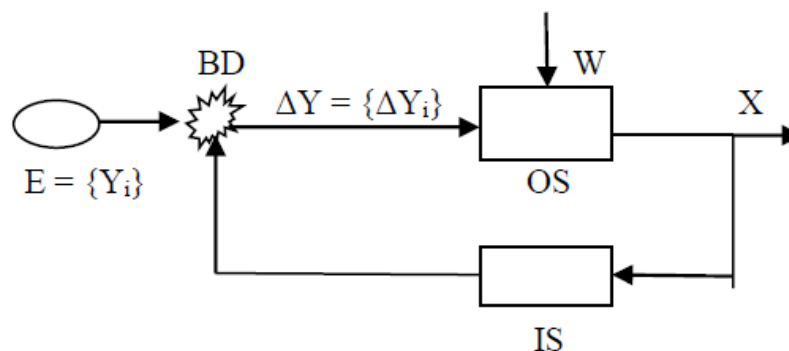
$$F = \langle S, Mr \rangle$$

где  $S$  – выявленное множество свойств объекта с конкретными значениями, а  $Mr$  – нечёткое отношение „причина – следствие” на множестве  $S$ . Семантика заключается в том, что если  $z_i = (q_i, v_i)$ ,  $z_j = (q_j, v_j)$ ,  $(q_i, q_j) \in S$ ,  $i \neq j$ , когда  $v_i, v_j$  – значения свойств  $q_i, q_j$  то для каждой пары  $(z_i, z_j)$  соответствует число  $\kappa(z_i, z_j) \in [0,1]$ , определяющее степень принадлежности пары нечёткому отношению.

Вместе с тем, отказавшись от строгого математического решения при многопараметрической оценке состояния объекта, логично ввести в модель некоторую систему условных кодов. Тогда каждый известный или потенциально допустимый параметр считается равноценным в представлении показателей состояния объекта [5]. Качество состояния объекта будет определено по конечному множеству его показателей в соответствии с обоснованным эталоном. Данное кодирование следует заложить в модель выявленных групп параметров  $x_{ij}$  как некоторую совокупность показателей  $Y_i$ . Следовательно, эталонным можно считать такое состояние объекта (изделия), когда все его факторы  $E$  выявлены и однозначно соответствуют системе заданных показателей  $y_i$ . В основу приведённого вывода положены рассуждения, основанные на модифицированной схеме N.Viniera (Рис.2).

Действительно, если за объект диагностики принимается некоторая система (OS), то не трудно определить эталон её состояния  $E$  в каждый момент времени или на определённом интервале его эксплуатации.

В реальных процессах управления и за такой эталон можно принять набор заранее заданных или рассчитанных факторов  $E = \{Y_i\}$  ( $i = 1, l$ ). Они имеют фиксированные значения:  $l$  – число обоснованных факторов оценки состояния объекта прогнозирования нет смысла давать вероятностную оценку этим значениям [3]. Любое отклонение от эталонного ( $+\Delta y_i, -\Delta y_i$ ) есть ошибка, снижающая качественную характеристику элемента по отношению к эталону. Она должна учитываться как расчётная величина, полученная на основании оценок  $X, S$ . Она принимает абсолютное значение  $y = |y - \Delta y_i|$ . Для этого следует определить группу показателей  $y_i$  ( $y_i = \{X_{ij}\}, j=1, J$ ), формирующих данный объектный фактор и дать ей оценку по каждому из показателей  $y_i$ .



OS – объект управления

E – эталон (модель) объекта управления

BD – блок принятия решения о состоянии объекта

IS – информационно - измерительная система

W – внешние воздействия среды на объект

**Рис. 2.** Модифицированная схема N. Viniera

Проявление любой совокупности  $\Delta Y$ , полученной по данным IS трактуется как нарушение условия E и требует соответствующей реакции BD в оценке состояния OS. Следовательно, необходимо в оценку каждого из показателей  $X_{ij}$  ввести конечное множество групп параметров  $S_{ijf}$ , когда  $x_{ij} = \{S_{ijf}\}$  при  $f = 1, F$ , где  $F$  – число исходных параметров оценки состояния объекта (База Данных). На этом уровне конкретный параметр рассматривается как величина измерения (вычисления). Она оценивается по реальным выходным показателям функционирования объекта на данный момент или за установленный промежуток времени (интервала контроля, периода прогнозирования, времени жизни объекта и т.п). Здесь нужно указать, что показатели о состоянии объекта должны быть обоснованы измерениями параметров, показаниями датчиков или документированной информацией с числовыми значениями.

Примем за меру качества объекта результат сравнения процедур последовательной оценки фиксированных групп параметров, показателей и факторов данного с их эталонными, установленными заказчиком (пользователем).

---

#### ГРАФО-АНАЛИТИЧЕСКИЙ МЕТОД ДИАГНОСТИКИ ОБЪЕКТА

---

С позиции Теории систем, за эталон состояния диагностируемого объекта можно принять комплекс факторов, свободных от ошибок (отклонений) их показателей и параметров. Они прошли контрольные процедуры, т.е.  $\sum \Delta Y_i = 0$ . Утверждение справедливо при  $W=0$  (Рис.2). Здесь сразу возникает смысловая проблема. Элементы, составляющие анализируемую систему (производство, процесс, управление и т.д.), практически всегда имеют разную физическую природу, измерительные характеристики, методы определения и контроля. Это чётко проявляется даже в элементарных системах массового производства микросхем, метиза, крепежа и т.д. Следовательно, схема оценки факторов показатель – параметр по принципу допусков, вероятностных условий, статистических расчётов носит приблизительный характер. При этом, методы многофакторного анализа слишком громоздки и сложны для простых объектов, выпускающих малоценную продукцию. Но уже на следующем

шаге оценки состояния объекта задача усложняется: каждый из параметров может сохранять своё постоянное значение, но только в пределах частной задачи. Переход к другому варианту выпуска продукции или при изменении задачи (переналадка оборудования, смена инструментария), параметр, как правило, меняет своё значение. На пример, чистота поверхности при обработке вала для контрольно - измерительной аппаратуры является основным параметром, а при производстве редукторов носит второстепенный характер.

Понятно, что ошибка измерения присутствует в обоих вариантах приведённого примера и не может быть нулевой ( $\Delta \neq 0$ ). Поэтому она всегда учитывается, как отрицательная ошибка в оценке как данного показателя, так и параметра (фактора). Не трудно сделать вывод, что оценка состояния любого объекта или технологического процесса представляет собой тернарную систему последовательных операций, которую схематично можно представить в виде линейного алгоритма (Рис. 1).

Однако, утверждать, что совокупность показателей  $Y$  есть эталон  $E$  в такой интерпретации не справедливо. Каждый  $Y_i$  характеризует конкретную группу  $Y_i$  ( $i=1, I$ ). Эти группы разномерны по шкале измерений физических, экономических, информационных и др. факторов. Для наиболее эффективного результата принятия решения в BD о состоянии объекта необходимо изменить систему кодирования исходных данных по факторам и привести их к единой логике параметрического представления. В такой интерпретации и рассмотрим двух-шаговую процедуру реализации указанного предложения:

1. Логика системного анализа диктует условие многофакторности объектов и явлений. Тогда исследователь находится в весьма трудном положении. Учесть все возможные факторы оценки объекта нереально из-за невозможности сформировать многофакторную Базу Данных из-за ограниченного объёма знаний, парка контрольно-измерительной техники, отсутствия методов и средств вычислений ( $S \in X, X \in Y, Y \in E$ ). К тому же факторы могут быть различной природы и противоречивы: безопасность,

экономика, экология и т.п. Необходимо ввести относительную меру для каждого из множества известных для данного объекта значения  $S$ ,  $X$ ,  $Y$  по условной шкале измерения ( $0 \leq E \leq 1$ ).

Предположим, что известно (вычислено, установлено) эталонное значение каждого параметра  $S_{ijf}$  в последовательности  $S_{ij1}, S_{ij2}, S_{ij3}, S_{ij4}, \dots, S_{ijF}$  ( $f=1, F$ ). Если значения параметров  $S$  располагать линейно по временной оси абсцисс, а величины показателей  $X$  – на оси ординат, то они сольются в неупорядоченную прямую линию. Их сравнимость на уровне факторов  $E$  окажется случайной. Плоско-парное расположение параметров  $X$  даёт адекватную оценку только для данной пары показателей и не решает задачу оценки состояния объекта через характеристики групп по фактору  $E$  в целом [9].

Будем считать, что в оценке качества объекта каждый выявленный параметр равноважен по отношению к другим при оценке показателя  $j$ . Тогда можно перейти от абсолютной шкалы их представления к относительной (единичной) форме:

$$\begin{array}{ccc} S_{ijf} & \longrightarrow & 1,0 \\ S''_{ijf} & \longrightarrow & Z \end{array}$$

Решив традиционное уравнение перехода к относительной шкале, получим значение каждого параметра на единичной шкале измерения:

$$Z_{ijf} = S''_{ijf} \cdot 1,0 / S_{ijf}$$

2. Представим полученные результаты в табличной форме. Примем её за Базу Данных для оценки состояния объекта по параметрическому анализу (Таблица 1). Данные Таблицы 1 подразумевают конечность числа параметров  $S_{ijf}$  в каждой группе показателей  $X_{ij}$ .

**Table 1. База Данных оценки состояния объекта по параметрам**

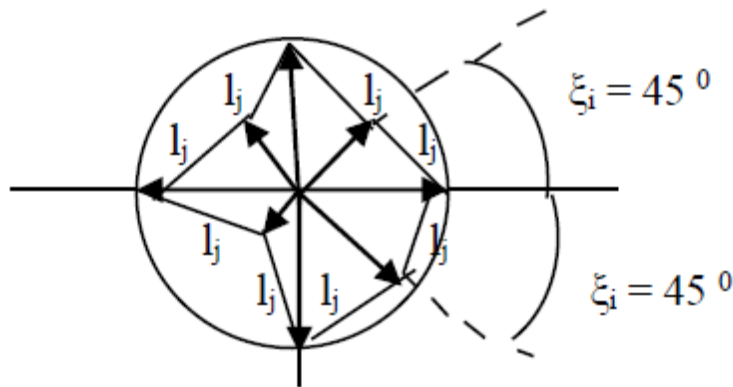
| №   | Параметр ( $s_{ijf}$ )      | Код параметра | Абсолютное значение | Относительное значение |
|-----|-----------------------------|---------------|---------------------|------------------------|
| 1   | 2                           | 3             | 4                   | 5                      |
| 1   | Определение параметра $s_1$ | ij1           | $S_{ij1}$           | $Z_{ij1}$              |
| 2   | Определение параметра $s_2$ | ij2           | $S_{ij2}$           | $Z_{ij2}$              |
| 3   | Определение параметра $s_3$ | ij3           | $S_{ij3}$           | $Z_{ij3}$              |
| ... | .....                       | .....         | .....               | .....                  |
| f   | Определение параметра $s_f$ | ijf           | $S_{ijf}$           | $Z_{ijf}$              |

Следует заметить, что в предлагаемой модели любое отклонение параметра  $X_{ij}$  ( $\pm \Delta X_{ij}$ ) указывает на снижение качества показателя и фактора состояния объекта в целом [7]. Тогда модель состояния объекта может быть представлена набором значений групп параметров и показателей, составляющих систему оценки принятия решения по конечному множеству факторов [3]. На пример, если конкретный показатель оценивается по 8 параметрам, то распределим их равномерно по окружности ( $45^0$ ). Длина каждого вектора соответствует величине параметра по относительной (единичной) шкале измерения (Рис.3). Следовательно:

$$r = 1,0; D = 2,0; L = 2\pi, P = \pi$$

Можно утверждать:

1. Чем больше изучен элемент (изделие, объект), тем больше векторов в единичной окружности.
2. Длина каждого вектора определяется текущим значением параметра.



**Рис. 3.** Пример графо-аналитической модели оценки состояния объекта

3. Чем больше вектор, тем ближе объект по параметру  $S_{ijf}$  к эталонной оценке ( $S_{ijf} = 1$ ).

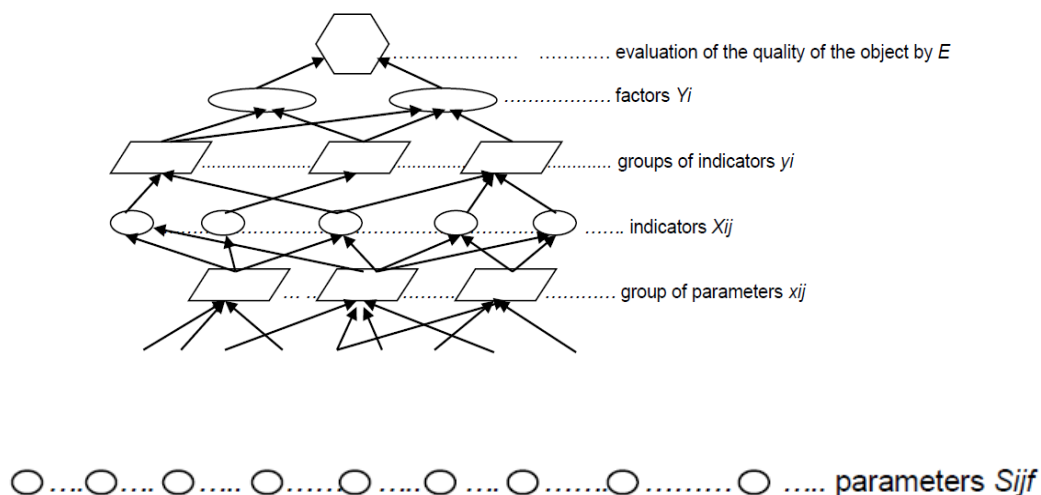
Рассуждая подобным образом, установим, что при  $x_{ij} \rightarrow \infty$  качество состояния элемента объекта теоретически может быть установлено с вероятностью  $P = 1,0$  как по совокупности параметров, так и по производным от него показателям  $y_i$  и фактору  $E$ . Это решение так же эффективно для моделирования состояния объекта при внезапном отказе (полном или частичном) одного из его элементов. Достаточно либо исключить вектор – показатель  $l_j$  данного элемента, либо уменьшить этот вектор на соответствующую величину (Рис. 3).

Для оценки состояния объекта в целом, достаточно соотнести теоретическую длину единичного круга  $L^0$  с длиной ломаной линии, соединяющей вершины векторов реальной оценки параметров ( $L_f$ ):

$$L^0 = \pi D = 3,14 \cdot 2 = 6,28$$

$$L_f = \sum l_{ijf} \quad (f = 1, F)$$

Рассуждая подобным образом, получим гириомат (Рис.4).



**Рис.4.** Гириомат состояния объекта управления

Основываясь на логике гириомата, по единичной шкале диагностируется объект как в целом, так и на каждом из его уровней:

$L^0 = 6,28$  – абсолютный эталон качества исследуемого объекта.

$L^0 - L_{ij} = \Delta L_{ij}$  – параметрическая мера состояния исследуемого объекта.

$L^0 - L_i = \Delta L_i$  – показательная мера состояния исследуемого объекта.

$L^0 - L_E = \Delta L_E$  – факторная мера состояния исследуемого объекта

## ЗАКЛЮЧЕНИЕ

Таким образом, для каждого объектного фактора, учитывая его оригинальность (специфику), можно графо-матричным способом дать вполне обоснованную оценку объекту или изделию по совокупности его

выявленных и вычисленных параметров (показателей). При этом, поступая аналогичным образом, когда удаётся обосновать меры и характер различных факторов состояния объекта, вполне допустимо провести и его многофакторный анализ.

---

## СПИСОК ЛИТЕРАТУРЫ

---

- [1] Арефьев И.Б., Коровяковский Е.К. Принцип распознавания проблемных ситуаций на базе интегральных характеристик // Гибридные и сенергетические интеллектуальные систем.2016. Калининград, Изд-во БФУ им. Канта. С. 206-217.
- [2] Арефьев И.Б. Логико-рефлексивное представление моделей древовидных структур. // Системный анализ в проектировании и управлении. 2014. СПб, Изд-во СПГТУ. С. 56-60.
- [3] Ariefiew I. Forecasting and control object of management in the environment of system PERT. Maritime University, Szczecin., Biblioteka cyfrowa., 2012.. 293 p.
- [4] Арефьев И.Б. Элементы метода логико-рефлексивного моделирования. Из-во „LAMBERT akademik publishing”, Riga. 2019. 68 p.
- [5] Карпов А.В. Психология рефлексивных механизмов деятельности. // М. Изд-во «ИП РАН», 2004.- 31 с.
- [6] Кант Э. Критика чистого разума. М: Мысль. Соч. том 3. 1964,]. [7] Чебышёв П.Л. Полное собрание сочинений.— Изд-во АН СССР, 1948.— Т.III.— С. 404
- [8] Мунджишвили Т., Логико-вероятностная модель оценки финансового состояния предприятия, Труды Тбилисского Университета, прикладная математика и компьютерные науки, #364 (24), стр. 216-229, Тб., 2005
- [9] Jack Heidel; Zhang Fu (1999). “Nonchaotic behaviour in three-dimensional quadratic systems II. The conservative case”. Nonlinearity.12(3): 617 – 633

---

## ИНФОРМАЦИЯ ОБ АВТОРОВ

---

**Igor Ariefjew** – Maritime University. Szczecin, Poland; e-mail: [i.aeriefjew@am.szczecin.pl](mailto:i.aeriefjew@am.szczecin.pl).

**Tea Munjishvili** – Tbilisi state university; e-mail:

### GRAPH-ANALYTICAL METHOD FOR ASSESSING THE STATE OF THE OBJECT BEING DIAGNOSED

**Igor Ariefjew, Tea Munjishvili**

**Abstract:** *The authors propose a method for diagnosing a technical object when the data set for assessing its state exceeds the multiple threshold of its characteristics about the state of the object under study known to experts, even with the theoretical infinity of the initial parameters. The method is based on graph-analytical representations and logical-probabilistic modeling.*

**Keywords:** *diagnostics, control, object, graph, parameter, indicator, factor.*

## TABLE OF CONTENTS OF IJ ITA VOL. 27, NO.1

*Using Visual Analytics and K-Means Clustering for Monetising Logistics Data, a Case Study With Multiple E-Commerce Companies*

Hamzah Qabbaah, George Sammour, Koen Vanhoof ..... 3

*On Logical-Combinatorial Supervised Reinforcement Learning*

Levon Aslanyan, Vladimir Ryazanov, Hasmik Sahakyan ..... 40

*Optimization Problem: Systemic Approach*

Albert Voronin, Yuriy Ziatdinov ..... 52

*Multisensor Prototype for Beverage Quality Control: Principle Scheme and Test Results*

Volodymyr Romanov, Igor Galelyuka, Oleksandr Voronenko, Oleksandra Kovyrova,  
Sergei Dzyadevych, Lyudmyla Shkotova ..... 82

*A Comparative Examination of Convolutional Autoencoder and Densenet Applications for Breast Cancer Classification*

Naderan Maryam, Yuri Zaychenko ..... 93

*Table of contents*..... 100

## TABLE OF CONTENTS OF IJ ITA VOL. 27, NO.2

### *Integration of Chain-Computation Algorithms in Cube-Splitting Technique*

Hasmik Sahakyan, Levon Aslanyan ..... 103

### *Mathematical Logic Problems in Natural Language Processing*

Serhii Kryvyi, Nataliia Darchuk, Hryhorii Hoherchak ..... 121

### *Synthesis of Chat-Bot Responses in the Inflecting Natural Language Based on the Results of Queries to Ontology and Analysis of the Chat Previous Phrase*

Anna Litvin, Vitalii Velychko, Vladislav Kaverinsky ..... 152

Table of contents ..... 200

## TABLE OF CONTENTS OF IJ ITA VOL. 27, NO.3

|                                                                                                                                  |     |
|----------------------------------------------------------------------------------------------------------------------------------|-----|
| <i>Using Apriori Association Rules and Decision Tree Analysis for Detecting Customs Fraud</i>                                    |     |
| Hamzah Qabbaah, George Sammour, Koen Vanhoof .....                                                                               | 203 |
| <i>Substantive Formulation of UAV Groups Routes Optimization Problem</i>                                                         |     |
| Maksym Ogurtsov, Vyacheslav Korolyov .....                                                                                       | 229 |
| <i>Isomorphism of Predicate Formulas as the Base of Logic Ontology Construction</i>                                              |     |
| Tatiana Kosovskaya.....                                                                                                          | 248 |
| <i>Digital Analysis of Genetic Proximity of the Balkan Population</i>                                                            |     |
| Georgi Gachev, Jordan Tabov.....                                                                                                 | 256 |
| <i>Didactical Mapping in Topic-Oriented Didactical Units</i>                                                                     |     |
| Borislav Yordanov Lazarov .....                                                                                                  | 273 |
| <i>Analysis of Approaches for Effective Applications Development in Monolith, Service-Oriented and Microservice Architecture</i> |     |
| Yurii Milovidov .....                                                                                                            | 288 |
| Table of contents.....                                                                                                           | 300 |

## TABLE OF CONTENTS OF IJ ITA VOL. 27, NO.4

### *Steganalysis of Adaptive Embedding Methods by Message Re-Embedding Into Stego Images*

Dmytro Progonov, Vladymir Lucenko ..... 303

### *Диагностика медицинских изображений опухолей молочной железы с использованием гибридных сверточных нечетких нейронных сетей*

Галиб Гамидов, Юрий. П. Зайченко..... 325

### *Исследование нечетких нейронных сетей с различными алгоритмами вывода в задачах прогнозирования на рынках ценных бумаг*

Тофик Кязимов, Юрий Зайченко ..... 345

### *Универсальные решающие правила: поддержка и принятие решений в задачах распознавания, диагностики, выбора, классификации*

Крисиллов А.Д., Крисиллов В.А., Чумаченко В.В. .... 369

### *Графо-аналитический метод оценки состояния диагностируемого объекта*

Игорь Арефиев, Теа Мунджишвили ..... 383

Table of contents of IJ ITA Vol. 27, No.1 ..... 397

Table of contents of IJ ITA Vol. 27, No.2 ..... 398

Table of contents of IJ ITA Vol. 27, No.3 ..... 399

Table of contents of IJ ITA Vol. 27, No.4 ..... 400