

НЕЙРОИНФОРМАТИКА ТА ІНТЕЛЕКТУАЛЬНІ СИСТЕМИ

NEUROINFORMATICS AND INTELLIGENT SYSTEMS

НЕЙРОИНФОРМАТИКА И ИНТЕЛЛЕКТУАЛЬНЫЕ СИСТЕМЫ

UDC 004.932.2:004.93'1

STATISTICAL DATA ANALYSIS TOOLS IN IMAGE CLASSIFICATION METHODS BASED ON THE DESCRIPTION AS A SET OF BINARY DESCRIPTORS OF KEY POINTS

Gadetska S. V. – PhD, Associated Professor of Dep. of Higher Mathematics, Kharkiv National Automobile and Road University, Kharkiv, Ukraine.

Gorokhovatskyi V. O. – Dr. Sc., Professor of the Department of Informatics, Kharkiv National University of Radio Electronics, Kharkiv, Ukraine.

Stiahlyk N. I. – PhD, Head of the Department of Information Technologies and Mathematical Modeling, Educational and Scientific Institute “Karazin Banking Institute” V. N. Karazin Kharkiv National University., Kharkiv, Ukraine.

Vlasenko N. V. – PhD, Associated Professor of Dep. of Informatics and computer Technologies, Simon Kuznets Kharkiv National University of Economics, Kharkiv, Ukraine.

ABSTRACT

Context. Modern computer vision systems require effective classification solutions based on the research of the the processed data nature. Statistical distributions are currently the perfect tool for representing and analyzing visual data in image recognition systems. If the description of a recognized object is represented by a set of vectors, the statistical apparatus becomes fundamental for making a classification decision. The study of data distributions in the feature blocks systems for key point descriptors has shown its effectiveness in terms of achieving the necessary quality of classification and processing speed. There is a need for in-depth study of the descriptor sets statistical properties in terms of the main aspect – the multidimensional data separation for classification. This task becomes especially important for constructing new effective feature spaces, for example, by aggregating a set of descriptors by their constituent components, including individual bits. To do this, it is natural to use the apparatus of statistical criteria designed to compare the parameters of the distribution of the studied samples. Despite the widespread use and applied effectiveness of the feature descriptors apparatus for image classification, the statistical basis of these methods in their implementation in aggregate visual data systems and the choice of effective means to assess their effectiveness for distinguishing real images in application databases remains insufficiently studied.

Objective. Development of an effective images classification method by introducing aggregate statistical features for the description components.

Method. A metric image classifier based on feature aggregation for a set of image descriptors using statistical criteria for assessing the classification decision significance is proposed.

Results. The synthesis of the classification method on the basis of the introduction of aggregated statistical features for a set of image description descriptors is carried out. The efficiency and effectiveness of the developed classifier are confirmed. On examples of application of a method for system of real images features its efficiency is experimentally estimated.

Conclusions. The study makes possible to evaluate the applied effectiveness of the key points descriptors apparatus and build on its basis an aggregate features system for the effective visual objects classification implementation. Our research has shown that the available information in the form of a bit descriptors representation is sufficient for a significant statistical distinction between visual objects descriptions. Analysis of pairs and other blocks for descriptor bits provides a promising opportunity to reduce processing time.

The scientific novelty of the study is the development of a method of image classification based on an integrated statistical features system for structural description, confirmation of the effectiveness of the method and the importance of the created features classification system in the image database.

The practical significance of the work is to confirm the efficiency of the proposed methods on the real image descriptions examples.

KEYWORDS: computer vision, key point, descriptor, data aggregation, statistical distribution, significance of classification decision, processing speed.

ABBREVIATIONS

KP – key point;
ORB (Oriented FAST and Rotated BRIEF) – detector that forms the descriptors of key points;
BRISK (Binary Robust Invariant Scalable Key-points) – detector that forms the descriptors of key points.

NOMENCLATURE

n – dimension of the descriptor KT;
 B^n – space of binary dimension vectors n ;
 E_j – reference descriptor with the number j ;
 E – reference set;
 Z – description of the visual object;
 $\{z_{v,i}\}_{v=1}^S\}_{i=1}^n$ – binary description matrix;
 E_1, E_2, E_3, E_4 – descriptions of reference images in the experiment;
 S – the number of description descriptors;
 m – number of classes;
 θ – aggregate feature vector;
 $\theta^{(1)}$ – aggregate feature vector for the first processing method;
 $\theta^{(2)}$ – aggregate feature vector for the second processing method;
 $P_s^{(i)}(v)$ – probability of occurrence v units in place of the i -th bit in the description of the s vectors;
 p_i – the probability of occurrence of one in place of the i -th bit in the set of descriptor descriptors;
 $D^{(j)}$ – normalized measure of similarity of vectors;
 $d^{(j)}$ – Manhattan distance between vectors $\theta^{(1)}$ for j -th reference and object;
 k – fragment number;
 u_k – the value of the highest level attribute for the fragment with the number k ;
 b – fragment size in bits;
 q – dimension of the vector u_k ;
 K – classifier.

INTRODUCTION

Statistical data science tools usage to build visual objects images classifiers in computer vision systems is aimed at providing the necessary performance based on the study of properties, content, structure of reference data and the introduction of obtained knowledge into the classification process [1–6]. An element of the image space in a vector data environment with real or binary components in the implementation of structural recognition methods is a finite set of key point descriptors (KP) of the image [2]. Recently, BRISK and ORB descriptors with binary components have become popular due to low computational costs [3–6, 13].

Gadetska S. V., Gorokhovatsky V. O., Stiahlyk N. I., Vlasenko N. V., 2021
DOI 10.15588/1607-3274-2021-4-6

Statistical data distributions are perfect tools for representing and analyzing visual data in image recognition systems. If the description of a recognized object is given by a set of vectors, the statistical apparatus becomes fundamental for making a classification decision. Data distributions research in the blocks systems for KP descriptors have shown their effectiveness in terms of providing the required quality of classification and processing speed [2]. There is a need for in-depth study of statistical properties descriptor sets in terms of the main problem – the multidimensional data separation for classification. This task becomes especially important when constructing new effective feature spaces, for example, by aggregating a set of descriptors by vector components [3, 4, 10]. For this purpose, it is natural to use the apparatus of statistical criteria designed to compare the parameters of the distribution of the studied samples.

The aggregator classifier organizes a new data space to describe as a set of descriptors, which evaluates the similarity of the feature vectors of the recognized object and a single reference image, and the classification is done by optimizing the degree of this similarity.

Probabilistic model of generating visual object descriptors vector data is a practical approach to formalize the process of classifier constructing, the essence of which is to build and study statistical distributions of objects or their components with the introduction of aggregation and optimization procedures on multiple classes [1, 6].

Despite the widespread use and practical effectiveness of the apparatus of KP descriptors for the visual objects classification [2–5], there is still remains unexplored statistical basis of these methods and the choice of effective means to assess their effectiveness for real datasets [1, 2].

The object of this research is the introduction of a statistical data analysis apparatus to build the image classifier based on the aggregate data representation as a set of KP descriptors and confirm its effectiveness.

The subject of the research is the synthesis of the classifier on the basis of aggregated features and statistical proof of the separation properties of these features for reference classes and examples of input images.

The aim is to develop a performance-efficient method of image classification by introducing aggregate features for the composition of the description components.

1 PROBLEM STATEMENT

Consider a multidimensional space B^n of any binary vectors of dimension n , where we will construct the object descriptions and reference images. Description Z is defined on the basis of the KP descriptor set of the visual object in the form of a finite set of binary vectors of dimension S

$$Z = \{z_v\}_{v=1}^S, z_v \in B^n.$$

In a more detailed view, we will consider and analyze the description $Z = \{\{z_v\}_{v=1}^s\}_{i=1}^n$ as a matrix of binary values with $s \times n$ size.

We will traditionally consider classification as a reflection

$$K: Z \rightarrow [1, \dots, m],$$

where each class is represented by a reference descriptor in E_j , $j = 1, \dots, m$, which are available for analysis [2].

Let's study the visual objects classification as assigning their description to one of the reference classes, based on the aggregate representation of the description data using the tools and criteria of mathematical statistics. In a general case, the classification problem is formally reduced to establishing the degree of similarity of two vector sets with binary components of equivalent size. We will build a secondary integrated system of features $\theta = \{\theta_k\}_{k=1}^N$ on the basis of descriptions Z , $\{E_j\}_{j=1}^m$ and implement it in the classification solution.

We use a metric approach to determine the degree of similarity of feature values θ for object and reference images. The introduction of aggregate features contributes to a significant acceleration of the classifier decision process, the gain in comparison with the traditional method of voting descriptors reaches hundreds of times [5, 6]. Another task is to investigate the separation properties of the newly created system of features using traditional statistical criteria.

2 REVIEW OF THE LITERATURE

The formal definition of the classification problem with the description of the image as a set of KP descriptors is formulated in [2–4], which also studies the advantages of implementing a structural description model in the methods of statistical classification [1, 5–9]. It is noted that the primary problem is the excessive computational costs caused by large arrays of vector data. Articles [4–6, 9, 16, 17, 26] investigate statistical models for the synthesis of feature space modifications to reduce the amount of computation, in particular, the application of data aggregation methods by forming distributions and defining statistical data centers. Works [1, 2, 8, 14] are devoted directly to the analysis of learning models for the fixed base of descriptions used in computer vision and the definition of the function of belonging to a fixed system of classes.

Articles [9, 13, 19, 21, 25] discuss the principle of construction feature detectors for the binary descriptors of KP.

Studies [1, 2, 8, 15, 23] contain results on the applied implementation of statistical approaches to the visual images classification using an ensemble processing. In [1, 6–8, 17–20] methods of evaluating the effectiveness of intelligent systems using statistical and metric measures of similarity are described. The advantages of statistical solutions such as high processing speed, sufficient distortion resistance and ensuring the required level of classification efficiency are discussed.

Works [11, 12] are used as sources of traditional and modern methods of statistical evaluation, the book [15] contains a description of applied features of software modeling, and sources [2–6, 10, 23] include the results of authors' research in implementing statistical approaches to develop structural methods image classification. In particular, [2] proposed technologies of component analysis and spatial processing for the classification of visual objects using statistical characteristics of the structural description of the image.

3 MATERIALS AND METHODS

We introduce a mapping $Z \rightarrow \theta, Z \subset B^n$, from a fixed set Z of binary vectors – KP descriptors for a given object into an integer vector $\theta = \{\theta_k\}_{k=1}^N$, the components of which will be calculated by some rule, according to which $N = n$ or $N = s$. This will make it possible to identify and distinguish visual objects on the basis of smaller data, as set of vectors is transformed into a single vector [4].

We will classify on the basis of estimating the differences in the values of vectors θ for different descriptions, the calculation of which is proposed in two different ways, which aims to take into account the structural features of the studied data and, as a result, ensure the efficiency of the recognition process.

According to the first method of determining vectors θ , we find the sum of binary values (number of units) consecutively for each bit with the number separately, based on the complete set of object descriptions Z . For a fixed description we obtain vectors of the form:

$$\theta^{(1)} = \{\theta_i^{(1)}\}_{i=1}^n, \theta_i^{(1)} = \sum_{v=1}^s z_{v,i}, 0 \leq \theta_i^{(1)} \leq s. \quad (1)$$

The vector (1) is an aggregate parameter formed on a set of descriptor descriptors by bitwise analysis of data in the form of adding the values of the corresponding bits (matrix columns $Z = \{\{z_{v,u}\}_{v=1}^s\}_{u=1}^n$).

If we consider the distribution of values by the i -th bit from the description of the object close to binomial, which is determined by the Bernoulli formula:

$$P_s^{(i)}(v) = C_s^v p_i^v (1 - p_i)^{s-v}$$

for which we calculate $p_i = \frac{\theta_i^{(1)}}{s}$, the value

$\theta_i^{(1)}, i = 1, \dots, n$ as it known in mathematical statistics [11], can be interpreted as the average value of the appearance of the corresponding number of units in place of the i -th bit: $\theta_i^{(1)} = p_i s$.

In general case, the tuple of values $\theta^{(1)} = (\theta_1^{(1)}, \dots, \theta_n^{(1)})$ obtained on the basis of the description Z can be considered as an aggregated parametric representation of the description, where the parameters are the probabilities p_i of occurrence of single bits for the i -th component in the set Z . We introduce the representation $\theta^{(1)}$ into the classification process, as it significantly reduces computational costs by transforming data from a set of vectors into a single parameter-vector [2–4].

Consider $\theta^{(1j)}, j=1, \dots, m$ – a vector aggregated by columns of the matrix for the binary description of the reference E_j with the number $j=1, \dots, m$, according to (1) and $\theta^{(1O)}$ – an aggregate vector for the description of the studied object O .

To compare the aggregate descriptions of objects of type (1) and, accordingly, to solve the classification problem, we introduce the classifier [2]

$$K : k = \arg \max_{j \in [1; m]} D^{(j)}, \quad (2)$$

$$D^{(1)} = 1 - \frac{d^{(j)}}{s \cdot n}, \quad d^{(j)} = \sum_{i=1}^n \left| \theta_i^{(1j)} - \theta_i^{(1O)} \right|, j=1, \dots, m. \quad (3)$$

Classifier (2) implements the principle of analysis “object – reference image” based on the aggregate vector representation θ [2]. We emphasize that expression (2) can be considered as a decisive rule, which is based on the likelihood function, presented, in contrast to its classical probabilistic representation [1], in terms of metrics, in particular, Manhattan.

To confirm the significance of the decision, as well as to control the obtained result of the classifier (2) with the involvement of aggregate vectors, we use methods of mathematical statistics, namely, a paired two-sample t-test for averages [11, 12], which provides pairwise (coordinate) comparison of the studied objects – vectors θ for a statistically significant difference in their average values. When using this test, two samples of the same volume are considered, in which the elements have a fixed location (as coordinates).

In the process of testing the null hypothesis regarding the equality of the averages in these samples, Student's statistics is used [11]; a level of significance α is established, equal to the probability of making an error of the I kind, i.e. rejecting the null hypothesis if it was correct; based on the initial data, the p-value is calculated as the maximum possible probability of error of the I kind. Then, if the p-value is less than the established α , then the null hypothesis is rejected, and an alternative hypothesis is accepted regarding the significant difference of the means (at the level of significance α). Otherwise, there is no reason to reject the null hypothesis of no statistically significant differences between the means [11].

Based on the features θ_i , it is possible to calculate higher-level features u_k for data blocks as sets of columns [1]

$$u_k = \sum_{i=k}^{k+b-1} \theta_i. \quad (4)$$

In relation (4) the relation is established $b = n/q$, $k=1, b+1, 2b+1, \dots, n-b+1$.

Features (4) implement cross-correlation processing of the matrix Z with a rectangular mask size $b \times s$ [1, 14]. As a result of calculation (4) we obtain an integer vector u_k of dimension q . The parameter q is a characteristic of the newly created system of fragments, it varies from n to 1 with increasing fragment size from 1 to n .

The values of the vector $u = (u_1, \dots, u_k, \dots, u_q)$ can be used for classification as independent structural features of the statistical type. Given a simple model for calculating functions (4), they are all easy to determine for an arbitrary fragment size (logically or by adding integers).

Based on representation (4), a hierarchical recognition method can also be used, which uses a system of features u_k with different degrees of data integration to compare with reference set [2, 16]. The range of integer values for features u_k can be directly determined by the size of the fragment $u_k \in \{0, \dots, sb\}$ as Model (4) implements the procedure of reducing the information redundancy of the spatial signal due to the allowable resolution reduction of the feature system [2, 14].

Note also that for the sake of universality of the study, it may be appropriate to perform analysis of variance of aggregate vectors constructed from reference E_j , $j=1, \dots, m$, in order to ensure that the reduction of data dimensionality did not affect the difference in the references set. In this case, since the structure of the vectors aggregated by formula (1) involves the consideration of paired samples, in this case it is possible to use only non-parametric analysis of variance, for example, in the form of the Friedman test [11].

The second way to calculate vectors θ is to represent the components of the vector as the sum of unit bits separately for each binary descriptor of the description. In this case, by adding the elements of the rows of the matrix (1) we obtain aggregate description vectors in the form:

$$\theta^{(2)} = \{\theta_v^{(2)}\}_{v=1}^s, \theta_v^{(2)} = \sum_{i=1}^n z_{vi}, 0 \leq \theta_v^{(2)} \leq n. \quad (5)$$

Then $\theta^{(2j)}$, $j=1, \dots, m$ – is a vector aggregated along the rows of the matrix of the description of the reference E_j with the number $j=1, \dots, m$ by expression (5); where

$\theta^{(2O)}$ – aggregated vector according to the description of the studied object O.

Note that the process of formation of aggregate vectors in the form (5) leads to the creation of already independent samples, the study of which does not involve a coordinate comparison. In this case, due to the independent nature of the data, it is proposed to check for a significant difference only by statistical methods. The essence of the introduction in this situation of a two-sample t-test for averages for two independent samples is to compare the averages of two sets of disordered elements, which are the number of units calculated separately for each binary descriptor description, the sequence of which obviously does not matter. The procedure for implementing the test remains the same as for the case of dependent (paired) samples, and differs only in the formula, according to which the relevant statistics are calculated on the basis of the studied data [11].

Note that here it may also be appropriate to perform analysis of variance of aggregate vectors (5), built on the references E_j , $j = 1, \dots, m$; in order to verify the fact of statistically significant differences $\theta^{(2j)}$, $j = 1, \dots, m$ in the set of these vectors. Here, the structure of aggregate vectors involves the consideration of independent samples, which involves the use of methods of both parametric and nonparametric variance analysis (for example, in the form of the Kruskal-Wallis test) [11].

4 EXPERIMENTS

Consider an example with experimental descriptions of three fixed reference images E1, E2, E3, and E4, obtained from E1 by rotation. Examples of images based on the results of software modeling with the formed coordinates of the BRISK KP descriptors are shown in Fig. 1 [13, 15, 21]. For the descriptions of these images in the form of a set of descriptors, our calculations are performed.

In the calculation example, $n = 512$ is the dimension of the descriptor, $s = 500$ is the number of descriptors in

the description, m is the number of reference images or classes ($m = 3$).

We proceed to the implementation of the proposed classification approach based on the values of the parameters θ in the database of three reference images E1, E2, E3 (Fig. 1) and the image E4, transformed by rotation of E1. Note that the representation using KP descriptors provides invariance to the transformations of displacement, rotation and scale of the analyzed object [2].

Fragments of the calculation results for aggregate vectors $\theta^{(1j)}$, $j = 1, 2, 3, 4$, of the form (1) for objects E1, E2, E3, E4 are given in Table 1.

According to formulas (2, 3) with the substitution of indicators E4 we have:

$$D^{(1)} = 0.983 ; ; D^{(2)} = 0.929 ; D^{(3)} = 0.895 .$$

Table 1 – Fragments of component vectors

$$\theta^{(1j)} = \{\theta_i^{(1j)}\}_{i=1}^{512}, j = 1, 2, 3, 4$$

Component number	$\theta^{(11)}$	$\theta^{(12)}$	$\theta^{(13)}$	$\theta^{(14)}$
1	448	397	369	430
2	318	305	394	312
3	410	371	341	395
4	391	365	325	378
5	290	250	385	278
6	306	306	293	305
7	224	162	343	234
...	...	...	...	...
507	213	242	265	220
508	237	247	276	235
509	205	292	303	210
510	145	235	257	154
511	220	299	295	240
512	195	299	298	203



Figure 1 – Reference images with marked KP

We see that the use of classifier (2) leads to the correct recognition of the object E4 as a transformed reference image E1, as its similarity with the first reference image is the greatest.

For the hierarchical features obtained by expression (4) at the size of the fragment $b = 2$ obtained the following values of these indicators: 0,986; 0,934; 0,908.

As you can see, they differ slightly from the values for the full description ($b = 1$), but the data vector is reduced by 2 times, which allows for further reduction of computational volumes.

Confirmation of the fact of statistically significant proximity of E4 to E1, as well as statistically significant difference of E4 from other reference images E2, E3 is obtained using a paired two-sample t-test for averages applied to samples represented by aggregate vectors $\theta^{(1j)} = \{\theta_i^{(1j)}\}_{i=1}^{512}$, $j = 1, 2, 3, 4$, from (1), the results of which are shown in table 2

Table 2 – The results of the application of a paired two-sample t-test

Samples	E1, E4	E2, E4	E3, E4
P-value	0,252	0,002	0,0000000006
Significance	no	yes	yes

In this case, for the hierarchical features obtained by expression (4) at the size $b = 2$ of the fragment, the value of the P-index was obtained: 0.330; 0.022; 0.000001.

Let's move on to consider the vectors aggregated by the second method.

Fragments of the results of the calculation of the components of the vectors $\theta^{(2j)}$, $j = 1, 2, 3, 4$ of the form (5) for objects E1, E2, E3, E4 are given in table 3.

Table 3 – Fragments of component vectors
 $\theta^{(2j)} = \{\theta_v^{(2j)}\}_{v=1}^{500}$, $j = 1, 2, 3, 4$

Component number	$\theta^{(21)}$	$\theta^{(22)}$	$\theta^{(23)}$	$\theta^{(24)}$
1	301	332	248	194
2	272	313	320	273
3	239	311	234	303
4	272	327	344	250
5	302	253	237	137
6	241	318	308	221
7	305	245	298	260
...			...	...
495	200	277	196	253
496	298	269	237	268
497	281	291	268	245
498	213	171	282	267
499	280	196	286	279
500	235	197	268	295

The confirmation of statistically significant proximity of E4 to E1, as well as statistically significant differences of E4 from other reference images E2, E3, obtained using a two-sample t-test for averages with different variances $\theta^{(2j)}$, $j = 1, 2, 3, 4$, are shown in table 4.

Table 4 – The results of the two-sample t-test

Samples	E1, E4	E2, E4	E3, E4
P-value	0.843	0.024	0,0000000005
Significance	no	yes	yes

In order to assess the stability of the studied approach to the application of description variations that may occur under the influence of noise, we performed calculations for descriptions transformed by expression (4), where generalization $b = 2$ for the case occurred by adding pairs of values $\theta^{(2j)}$. This corresponds to the calculation of hierarchical features when performing a merge for pairs of descriptors that appeared next to each other in the description. At the same time, the amount of data is also halved.

The obtained P-values for this case were 0.886; 0.034; 0,0000000007. This fact in comparison with the data of tab. 4 indicates the resistance of the considered methods to accidental interference, as the previous classification conclusions are fully confirmed.

5 RESULTS

The main result of this study is the development of images classification models based on the statistical analysis of component sets in the images descriptions and metric means of class selection. The proposed variants of data analysis models are based on the degree of similarity between the object and reference images, are workable and provide sufficient classification efficiency. Computational simulation performed on the example with 3 reference images confirmed efficiency of the proposed approach with use of significant data difference statistical criteria. Variants of the generalized features system synthesis that implies the further compression of the description sets and acceleration of classification procedures are also analyzed.

6 DISCUSSION

The synthesis of an aggregate feature system based on KP descriptor set makes possible to build a classifier that works successfully for the real images database.

As you can see, the first P-value (Table 2) for both options $b = 1$, $b = 2$ significantly exceeds the significance level $\alpha = 0.05$ (equal to the probability of error of the first kind), which indicates the absence of statistically significant differences between $\theta^{(11)}$ and $\theta^{(14)}$ (i.e. between the first reference image E1 and E4, which is transformed from E1); other obtained P-values are significantly less than 0.05, confirming a statistically significant difference in pairs between $\theta^{(14)}$, $\theta^{(12)}$ and $\theta^{(13)}$, $\theta^{(14)}$ (ie between E4, E2 and E4, E3).

Note, that for a large sample size (in the example we have $n = 512$), checking the data for compliance with the normal distribution law when using a paired t-test is not mandatory [11].

Note also, that the visual comparison of bar charts, which is a graphical representation of the vectors aggregated by formula (1), is a clear confirmation of the results obtained on the difference of objects (Fig. 2) (similarity between $\theta^{(14)}$ and $\theta^{(11)}$ and significant difference between pairs $\theta^{(14)}, \theta^{(12)}$ and $\theta^{(14)}, \theta^{(13)}$).

We also see that the first P-value is equal to 0.843 (Table 4) and significantly exceeds the level of significance $\alpha = 0.05$, which indicates the absence of statistically significant differences between $\theta^{(21)}$ and $\theta^{(24)}$ (ie between E1 and E4); other P-values, which are significantly less than 0.05, confirm a statistically significant difference in pairs between $\theta^{(24)}, \theta^{(22)}$ and $\theta^{(24)}, \theta^{(23)}$ (ie between E4, E2 and E4, E3).

Similarly to the previous one, the visual comparison of bar diagrams, which are a graphical interpretation of the vectors aggregated by formula (5), confirms the obtained results (Fig. 3).

The proposed classifier construction method allows further generalization in terms of fragment size aggregation that implies reduction of processing time.

CONCLUSIONS

The urgent problem of mathematical support Statistical data analysis is a powerful research tool for intelligent

decision making, machine learning and data science. The study makes possible to assess the applied effectiveness of the key points descriptors apparatus and build on its basis an aggregate features system for the effective visual objects classifier implementation. Our research has shown that the available information in the form of a bit representation of the descriptors is sufficient for statistical separation of data for different visual objects. Analysis of pairs and other blocks of bits reduces the processing time.

The scientific novelty of the study is the development of a method of image classification based on an integrated statistical features system for structural description, confirmation of the effectiveness of the method and the importance of the created features classification system in the image database.

The practical significance of the work is to confirm the efficiency and effectiveness of the proposed methods on the examples of real images descriptions.

Prospects for the study are related to the further enhancement and application of the developed classifiers in large-scale visual data bases.

ACKNOWLEDGEMENTS

The work was performed within the framework of the state budget research of Kharkiv National University of Radio Electronics “Deep hybrid systems of computational intelligence for data flow analysis and their rapid learning” (№ ДР0119U001403).

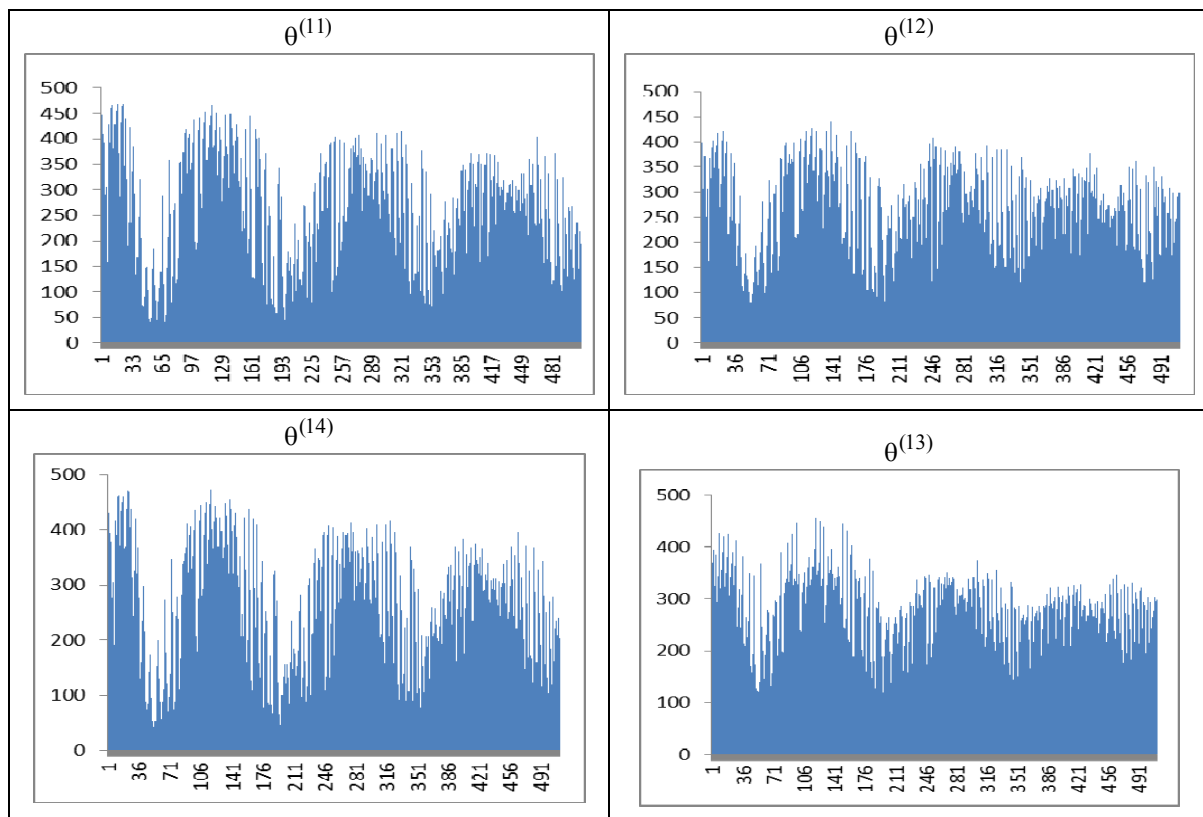


Figure 2 – Graphical representation of vectors $\theta^{(1j)} = \{\theta_i^{(1j)}\}_{i=1}^{512}, j = 1, 2, 3, 4$

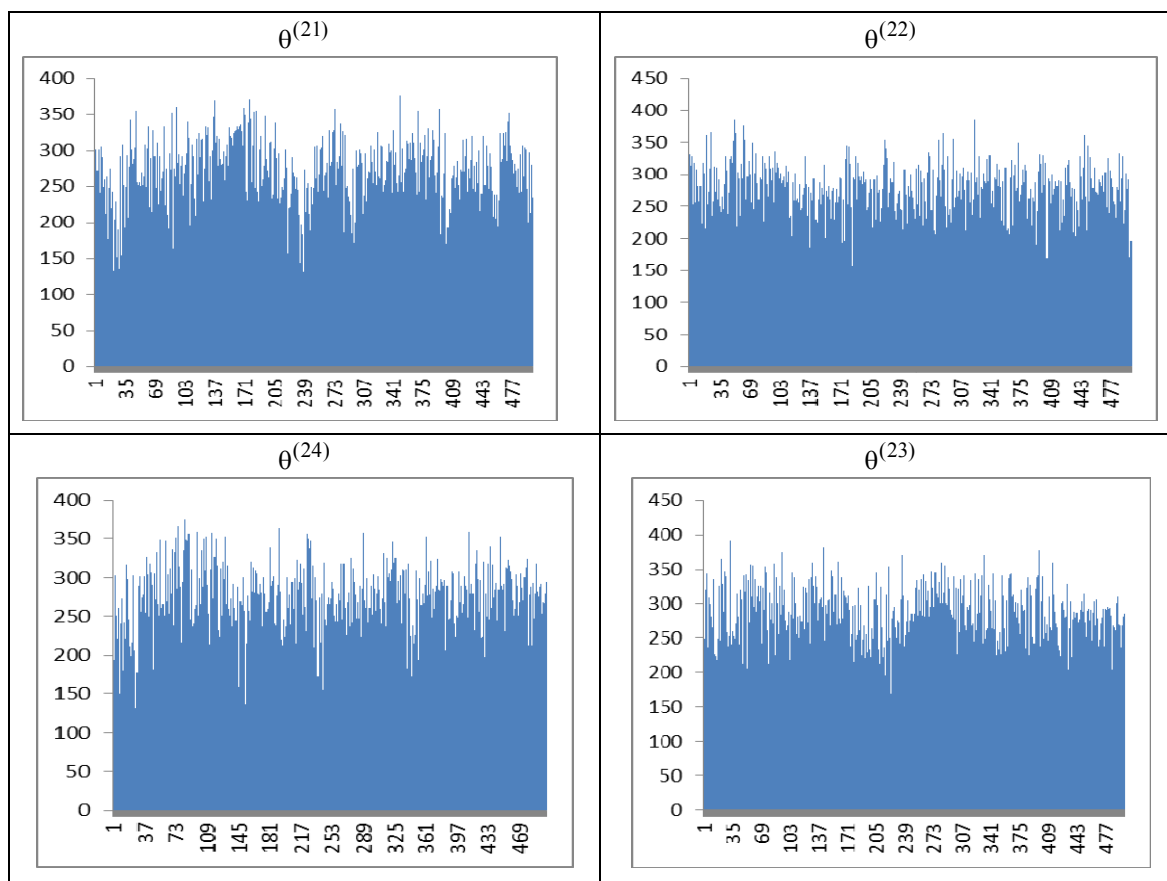


Figure 3 – Graphical representation of vectors $\theta^{(2j)} = \{\theta_v^{(2j)}\}_{v=1}^{500}, j=1,2,3,4$

REFERENCES

1. Flach P. Machine learning. The Art and Science of Algorithms that Make Sense of Data. New York, NY, USA, Cambridge University Press, 2012. DOI: 10.1017/CBO9780511973000
2. Gorokhovatskyi V., Gadetska S. Statistical processing and data mining in structural image classification methods. Kharkiv, Ukraine, FOP Panov A.N., 2020, 128 p.
3. Gorokhovatsky V. O., Gadetska S. V., Stiahlyk N. I. Vivchennja statistichnih vlastivostej modeli blochnogo podannja dlja mnozhini deskriptoriv ključovih točok zobrazen', *Radio Electronics, Computer Science, Control*. – 2019, No. 2, pp. 100–107. DOI: 10.15588/1607-3274-2019-2-11.
4. Gorokhovatsky V., Gadetska S. Determination of Relevance of Visual Object Images by Application of Statistical Analysis of Regarding Fragment Representation of their Descriptions, *Telecommunications and Radio Engineering*, 2019, No. 78 (3), pp. 211–220. – DOI: 10.1615/TelecomRadEng.v78.i3.20
5. Daradkeh Y. I., Gorokhovatskyi V., Tvoroshenko I., Gadetska S., and Al-Dhaifallah M. Methods of Classification of Images on the Basis of the Values of Statistical Distributions for the Composition of Structural Description Components, *IEEE Access*, 2021, No. 9, pp. 92964–92973. DOI: 10.1109/ACCESS.2021.3093457
6. Gorokhovatskyi V., Gadetska S., Stiahlyk N. Image structural classification technologies based on statistical analysis of descriptions in the form of bit descriptor set, *In CEUR Workshop Proceedings: Computer Modeling and Intelligent Systems (CMIS-2020)*, 2608, pp. 1027–1039. Available online: <http://ceur-ws.org/Vol-2608/>
7. Oliinyk A., Subbotin S., Lovkin V. et al. A The System of Criteria for Feature Informativeness Estimation in Pattern Recognition, *Radio Electronics, Computer Science, Control*, 2017, No. 4, pp. 85–96. DOI: 10.15588/1607-3274-2017-4-10
8. Peters J. F. Foundations of computer vision: Computational Geometry, Visual Image Structures and Object Shape Detection. Cham, Switzerland, Springer International Publisher, 2017, 417 p. DOI: 10.1007/978-3-319-52483-2
9. Hietanen A. and et al. A comparison of feature detectors and descriptors for object class matching, *Neurocomputing*, 2016 No. 184, pp. 3–12. DOI: 10.1016/j.neucom.2015.08.106
10. Gadetska S. V., Gorokhovatsky V. O. Statistical Measures for Computation of the Image Relevance of Visual Objects in the Structural Image Classification Methods, *Telecommunications and Radio Engineering*, 2018, Vol. 77 (12), pp. 1041–1053. DOI: 10.1615/TelecomRadEng.v77.i12.30
11. Kobzar' A. I. Prikladnaja matematicheskaja statistika. Dlja inzhenerov i nauchnyh rabotnikov, uchebnoe posobie, 2-e izd. Moscow, FIZMATLIT, 2012, 816 p.
12. Berk K., Kjejri P. Analiz dannyh s pomoshh'ju Microsoft Excel, Per. s angl. Moscow, Izdatel'skij dom «Vil'jams», 2005, 560 p.
13. Leutenegger S., Chli M., Roland Y. Siegwart. BRISK, Binary Robust Invariant Scalable Keypoints, *Computer Vision (ICCV)*, 2011, pp. 2548–2555.
14. Shapiro L., Stockman J. Computer vision. Moscow, Russia, Binom, Knowledge Laboratory, 2013, 752 p.

15. Prohorenok N. A. OpenCV i Java. Obrabotka izobrazhenij i komp'yuternoe zrenie. SPb., BHV-Peterburg, 2018, 320 p.
16. Adelson E. H., Anderson C., Bergen J., Burt P., Ogdan J. Pyramid methods in image processing, RCA Engineer, Vol. 29(6), pp. 33–41. [Elektronnij resurs]. Rezhim dostupu: http://persci.mit.edu/pub_pdfs/RCA84.pdf
17. Fazal Malik, Baharum Baharudin Analysis of distance metrics in content-based image retrieval using statistical quantized histogram texture features in the DCT domain. [Elektronnij resurs]. Rezhim dostupu: <https://www.sciencedirect.com/science/article/pii/S1319157812000444>.
18. Muja M., Lowe D. G. Fast Matching of Binary Features, *Conference on Computer and Robot Vision*, 2012, pp. 404–410. DOI: 10.1109/CRV.2012.60.
19. Kim S., Kweon I.-S. Biologically motivated perceptual feature: Generalized robust invariant feature, *In Asian Conf. of Comp. Vision (ACCV-06)*, 2006, pp. 305–314. DOI: 10.1007/11612704_31
20. Grosse K., Manoharan P., Papernot N., Backes M., and McDaniel P. On the (statistical) detection of adversarial examples. arXiv preprint arXiv:1702.06280, 2017.
21. Open CV Open Source Computer Vision, <https://docs.opencv.org/master/index.html>.
22. Wei P., Ball J., and Anderson D. Fusion of an ensemble of augmented image detectors for robust object detection, *Sensors*, 2018, No. 18 (3), P. 894. DOI: 10.3390/s18030894
23. Gorokhovatskiy V. A. Compression of Descriptions in the Structural Image Recognition, *Telecommunications and Radio Engineering*, 2011, Vol. 70, No 15, pp. 1363–1371. DOI: 10.1615/TelecomRadEng.v70.i15.60
24. Tvoroshenko I. S., Kramarenko O. O. Software determination of the optimal route by geoinformation technologies, *Radio Electronics, Computer Science, Control*, 2019, No. 3, pp. 131–142. DOI: 10.15588/1607-3274-2019-3-15
25. Gayathiri P., Punithavalli M. Partial Fingerprint Recognition of Feature Extraction and Improving Accelerated KAZE Feature Matching Algorithm, *International Journal of Innovative Technology and Exploring Engineering (IJITEE)*, 2019, Vol. 8, Issue 10, pp. 3685–3690 DOI: 10.35940/ijitee.I9653.0881019
26. Oliynyk A., Subbotin S. A. A stochastic approach for association rule extraction, *Pattern Recognition and Image Analysis*, 2016, Vol. 26, No. 2, pp. 419–426. DOI: 10.1134/S1054661816020139

Received 06.07.2021.
Accepted 21.10.2021.

УДК 004.932.2:004.93*1

ЗАСОБИ СТАТИСТИЧНОГО АНАЛІЗУ ДАНИХ У МЕТОДАХ КЛАСИФІКАЦІЇ ЗОБРАЖЕНЬ НА ПІДСТАВІ ОПИСУ ЯК МНОЖИНИ БІНАРНИХ ДЕСКРИПТОРІВ КЛЮЧОВИХ ТОЧОК

Гадецька С. В. – канд. фіз.-мат. наук, доцент, доцент кафедри вищої математики, Харківський національний автомобільно-дорожній університет, Харків, Україна.

Гороховатський В. О. – д-р техн. наук, професор, професор кафедри інформатики, Харківський національний університет радіоелектроніки, Харків, Україна.

Стяглик Н. І. – канд. пед. наук, завідувач кафедри інформаційних технологій та математичного моделювання, Навчально-науковий інститут «Каразінський банківський інститут» Харківського національного університету ім. В.Н. Каразіна, Харків, Україна.

Власенко Н. В. – канд. техн. наук, ст. викладач кафедри інформатики та комп'ютерної техніки, Харківський національний економічний університет ім. С. Кузнеця, Харків, Україна.

АНОТАЦІЯ

Актуальність. Сучасні системи комп'ютерного зору потребують дієвих класифікаційних рішень на підґрунті вивчення природи оброблюваних даних. Статистичні розподіли на цей час є досконалим засобом подання та аналізу візуальних даних у системах розпізнавання образів. Якщо опис розпізнаваного об'єкту представлено множиною векторів, статистичний апарат стає фундаментальним для прийняття класифікаційного рішення. Вивчення розподілів даних у складі системи блоків для дескрипторів ключових точок показали свою результативність у аспекті забезпечення потрібних показників якості класифікації та швидкодії оброблення. Виникає необхідність поглибленого вивчення статистичних властивостей для множини дескрипторів у аспекті головного фактору – розрізнення багатовимірних даних задля класифікації. Особливе значення набуває ця задача при побудові нових ефективних просторів ознак, наприклад, шляхом агрегування множини дескрипторів за їх складовими компонентами, в тому числі за окремими бітами. Для цього природним є напрацьоване використання апарату статистичних критеріїв, призначених для порівняння параметрів розподілу досліджуваних вибірок. Незважаючи на широке застосування і прикладну результативність апарату дескрипторів для класифікації зображень, до цих пір залишається не дослідженим статистичне підґрунтя цих методів при впровадженні їх у агрегованих системах ознак візуальних даних і вибір ефективних засобів для оцінювання їх дієвості для розрізнення реальних зображень у прикладних базах даних.

Мета роботи. Розроблення ефективного за швидкодією методу результативної класифікації зображень шляхом впровадження агрегованих статистичних ознак для складу компонентів опису.

Метод. Запропоновано метричний класифікатор зображень на основі агрегації ознак для множини дескрипторів опису із використанням статистичних критеріїв щодо оцінювання значущості класифікаційного рішення.

Результати. Здійснено синтез методу класифікації на підставі впровадження агрегованих статистичних ознак для множини дескрипторів опису зображення. Підтверджено працездатність і ефективність розробленого класифікатора. На прикладах застосування варіантів методу для системи ознак реальних зображень експериментально оцінена його результативність.

Висновки. Проведене дослідження дає можливість оцінити прикладну ефективність застосування апарату дескрипторів ключових точок зображення і побудови на його основі агрегованої системи ознак для результативного здійснення класифі-

кації візуальних об'єктів. Наше дослідження показало, що наявної інформації у вигляді бітового подання дескрипторів опису достатньо для значущого статистичного розрізнення описів візуальних об'єктів. Аналіз пар і інших блоків для бітів дескрипторів дає перспективну можливість скорочення часу оброблення.

Наукову новизну дослідження складає розроблення методу класифікації зображень на підставі системи інтегрованих статистичних ознак для структурного опису, підтвердження результативності методу та значущості створеної системи ознак при класифікації у межах бази зображень.

Практична значущість роботи полягає у підтвердженні працездатності та результативності запропонованих методів на прикладах дескрипторних описів реальних зображень.

КЛЮЧОВІ СЛОВА: комп'ютерний зір, ключова точка, дескриптор, агрегація даних, статистичний розподіл, значущість класифікаційного рішення, швидкодія оброблення.

УДК 004.932.2:004.93*1

СРЕДСТВА СТАТИСТИЧЕСКОГО АНАЛИЗА ДАННЫХ В МЕТОДАХ КЛАССИФИКАЦИИ ИЗОБРАЖЕНИЙ НА ОСНОВании ОПИСАНИЯ КАК МНОЖЕСТВА БИНАРНЫХ ДЕСКРИПТОРОВ КЛЮЧЕВЫХ ТОЧЕК

Гадецкая С. В. – канд. физ.-мат. наук, доцент, доцент кафедры высшей математики, Харьковский национальный автомобильно-дорожный университет, Харьков, Украина.

Гороховатский В. А. – д-р техн. наук, профессор, профессор кафедры информатики, Харьковский национальный университет радиоэлектроники, Харьков, Украина.

Стяглик Н. И. – канд. пед. наук, заведующая кафедрой информационных технологий и математического моделирования, учебно-научный институт «Каразинский банковский институт» Харьковского национального университета им. В.Н. Каразина, Харьков, Украина.

Власенко Н. В. – канд. техн. наук, ст. преподаватель кафедры информатики и компьютерной техники, Харьковский национальный экономический университет им. С. Кузнеца, Харьков, Украина.

АННОТАЦИЯ

Актуальность. Современные системы компьютерного зрения требуют действенных классификационных решений на основе изучения природы обрабатываемых данных. Статистические распределения в настоящее время являются совершенным средством представления и анализа визуальных данных в системах распознавания образов. Если описание распознаваемого объекта представлено множеством векторов, статистический аппарат становится фундаментальным для принятия классификационного решения. Изучение распределений данных в составе системы блоков для дескрипторов ключевых точек показали свою результативность в аспекте обеспечения требуемых показателей качества классификации и быстродействия обработки. Возникает необходимость углубленного изучения статистических свойств для множества дескрипторов в аспекте главного фактора – различения многомерных данных для классификации. Особое значение приобретает эта задача при построении новых эффективных пространств признаков, например, путем агрегирования множества дескрипторов по их составляющим компонентам, в том числе по отдельным битам. Для этого естественным является наработанное использование аппарата статистических критериев, предназначенных для сравнения параметров распределения исследуемых выборок. Несмотря на широкое применение и прикладную результативность аппарата дескрипторов для классификации изображений, до сих пор остается не исследованной статистическая основа этих методов при внедрении их в агрегированных системах признаков визуальных данных и выбор эффективных средств для оценки их действенности при различении реальных изображений в прикладных базах данных.

Цель работы. Разработка эффективного по быстродействию метода результативной классификации изображений путем внедрения агрегированных статистических признаков для состава компонентов описания.

Метод. Предложено метрический классификатор изображений на основе агрегации признаков для множества дескрипторов описания с использованием статистических критериев оценки значимости классификационного решения.

Результаты. Осуществлен синтез метода классификации на основании внедрения агрегированных статистических признаков для множества дескрипторов описания изображения. Подтверждено работоспособность и эффективность разработанного классификатора. На примерах применения вариантов метода для системы признаков реальных изображений экспериментально оценена его результативность.

Выводы. Проведенное исследование дает возможность оценить прикладную эффективность применения аппарата дескрипторов ключевых точек изображения и построения на его основе агрегированной системы признаков для результативного осуществления классификации визуальных объектов. Наше исследование показало, что имеющейся информации в виде битового представления дескрипторов описания достаточно для значимого статистического различия описаний визуальных объектов. Анализ пар и других блоков для битов дескрипторов дает перспективную возможность сокращения времени обработки.

Научную новизну исследования составляет разработка метода классификации изображений на основе системы интегрированных статистических признаков для структурного описания, подтверждение результативности метода и значимости созданной системы признаков при классификации в пределах базы изображений.

Практическая значимость работы заключается в подтверждении работоспособности и результативности предложенных методов на примерах дескрипторных описаний реальных изображений.

КЛЮЧЕВЫЕ СЛОВА: компьютерное зрение, ключевая точка, дескриптор, агрегация данных, статистическое распределение, значимость классификационного решения, быстродействие обработки.