

Problems of identification in information systems

UDC 004.9322:004.93'1

doi: <https://doi.org/10.20998/2522-9052.2025.1.01>

Volodymyr Gorokhovatskyi, Yurii Chmutov, Iryna Tvoroshenko, Oleg Kobylin

Kharkiv National University of Radio Electronics, Kharkiv, Ukraine

REDUCING COMPUTATIONAL COSTS BY COMPRESSING THE STRUCTURAL DESCRIPTION IN IMAGE CLASSIFICATION METHODS

Abstract. The research of the article is focused on ways to reduce the amount of analyzed data when applying image classification methods in computer vision systems. **The aim of this work** is to develop approaches to reduce the dimensionality of the vector description of the etalon base using metric granulation, which reduces computational costs and speeds up the classification process while maintaining a sufficient level of accuracy. **Methods used:** keypoint descriptors, metric data granulation apparatus, image classification and processing theory, data structures, software modeling. **Results:** the formalism of granular representation was developed; experimental modeling was carried out using five-level granulation, which reduced the time spent tenfold while maintaining high classification accuracy. In the comparative aspect, we studied ways to reduce the volume of vector descriptions based on data discarding, and researched the effect of the granularity level on the accuracy and classification time. **The practical significance of the work** is to improve the performance of image classification structural methods by implementing granularity and data discarding schemes, which provides much faster data processing without significant loss of classification performance.

Keywords: classification speed; data rejection; description reduction; granularity; image classification; keypoint descriptors.

Introduction. Literature review

Today, one of the most significant problems in computer vision systems is to ensure the required performance of their applied functioning, in particular, the accuracy and speed of data processing. Solving this problem is especially relevant for visual object classification tasks, where a significant number of prototype classes are used, the images of which are represented as a set of vectors. Traditionally, such tasks belong to the field of Big Data and require a volumetric linear search in a multidimensional data space [1–3].

One of the most widespread in applied applications is the continuous-group approach to recognition based on a deterministic model of the formation of the analyzed images [3]. According to this model, an unbounded class of images is generated as a result of a priori unknown topological transformations over some fixed etalon realization [3–6].

Modern classification methods use structural image analysis, where the image of a visual object is represented in the form of a set of keypoint descriptors [2, 4, 5]. The keypoint descriptors are formed by special filters – detectors that also form the coordinates of the keypoints. The keypoint descriptor is represented as a binary vector [4, 6].

In structural methods, the classification process is carried out according to the “bag of words” model; the class of an object is determined by the result of voting for the set of descriptors of the recognized object. Structural methods have certain advantages over modern neural networks, as they are based on direct alignment with the description or parameters of the etalons and do not require long-term training [2, 7, 8]. In addition, when applying structural methods, it is possible to quickly change the composition of the etalon database. Another

significant advantage of implementing structural methods is the ability to make decisions on individual parts or details of the description of a visual object. The functionality of the keypoint descriptors apparatus has the property of invariance to geometric transformations of objects in the field of view, which is necessary in real conditions [4–6].

The issue of reducing the computational cost of classification is urgent for modern developers [7–11]. One of the ways is to build hierarchical data retrieval methods based on clustering, hashing, or evaluation [12]. There are also ways to reduce the dimensionality of the structural description by filtering out descriptors according to the criterion of information content [2] or uniqueness [8]. For example, the informativeness criterion is based on the value of the difference in the distance of the descriptor to its own and other etalons in the classification database. However, the key factor for effective classification is still the group metric correlation between descriptors.

Improving the performance of structural methods can also be achieved by reducing the number of descriptors in the description of the etalons. Any reduction in the number of descriptors is important, as it makes it possible to avoid combinatorial aspects that are difficult to implement in the task of comparing large feature sets.

Based on this, the application of the principles of data granularity can be considered a promising direction [7–13]. Granular computing is a paradigm of information processing based on formal mathematical constructions – information grains. The foundation of the human-centered computing paradigm in soft computing and computational intelligence with the introduction of the fuzzy set apparatus belongs to L. Zadeh [13]. This approach is aimed at reducing the impact of uncertainty

and data complexity. Granulation is the process of obtaining aggregated data that is meaningful in some sense, it organizes complex data for applied decision-making [10, 14].

Information granules are important components of the data processing process, as they can emphasize the characteristics and relationships between the objects of aggregation. In particular, they can be synthesized by following the indistinguishability criterion, which states that components should be grouped together if they demonstrate sufficient functional similarity, proximity, or similarity. Each granule is intended to represent homogeneous semantic information. Different degrees of “granularity” can be used to describe the available data, which contributes to the creation of hierarchical systems with multiple levels of granularity [2, 14, 15].

Since the fine details of the description are also important, they are characterized by a significant number of information granules. As the level of detail decreases, the level of abstraction increases, and the number of granules is expected to decrease. Depending on the level of resolution achieved, atomic units are displayed that describe the feature system in different ways. At the same time, achieving the required level of abstraction based on the data description and the task requires experimental substantiation.

The pyramidal representation is an effective way to implement the granularity principle for integrating image features in space [1, 3, 5, 7]. The pyramidal data structure, which integrates the brightness for spatially close image pixel coordinates, significantly reduces the amount of multi-scale visual information. This granulation method is also appropriate for analyzing and processing sets of descriptors. In general, the concept of data granularity has become widespread in pattern recognition systems [8-15], various aspects of creating an effective feature space in image classifiers are studied [16, 17]. Granularity is used in object-background segmentation methods to improve performance [18, 19]. Effective metrics in the granular data space have been proposed [20]. The results of applied research on the use of compressed description representation have been obtained [21–24].

Methods of clustering, creating hash baskets, or evaluating centers [12] can be considered metric granulation, as they are based on calculating the proximity of data to each other using a metric or data parameters [2, 9]. The implementation of these approaches significantly improves processing speed with an acceptable decrease in classification accuracy. However, achieving the required values of the performance criterion requires a deeper study of all the possibilities and nuances of the applied use of granulation and sifting devices with the need to confirm their objectivity and versatility in the field of multidimensional data mining.

In this article, we study in detail the methods of metric granulation and sifting for image descriptions in the form of a set of keypoint descriptors. The subject of granulation is the detection of pairs of the most similar descriptors. Granulation by similarity (equivalence) involves finding the metrically closest descriptor and

excluding it from further processing. This allows to reduce the amount of computation at times, while retaining the most informative descriptors.

As an alternative, we also consider reducing the description by discarding elements according to some law of order in the full description. Discarding involves excluding, for example, every second element from the data set, which allows you to quickly reduce the number of elements and thus reduces the amount of computation. It is easier to perform, does not require distance estimation, and is fast. However, this method does not take into account the importance and relationship of the elements, which can lead to the loss of important information and a decrease in classification accuracy. Another research approach under investigation is the direct generation of a small description by the keypoint detector.

The research of this paper focuses on the study and experimental implementation of metric granulation and data sifting for image classifiers, analyzing their effectiveness in comparison with the traditional method without granulation. The main goal is to study the impact of the proposed modifications on the speed and accuracy of classification, which is important for improving the performance of computer vision systems in practical applications.

The main contribution of the article is to study the effectiveness of creating a hierarchical description representation with several levels of hierarchy. Another significant contribution is the confirmation of the practical advantages of implementing intelligent analysis by creating granular descriptions compared to trivial sifting schemes and the formation of lower-power descriptions directly by keypoint detectors.

Formalization of the Metric Granulation Method

Let us represent the structural description $Z = \{z_i\}_{i=1}^S$ of a visual object in the form of a set of keypoint descriptors as a collection $F(Z)$ of transformed descriptions f_r , obtained at various levels of hierarchical representation:

$$F(Z) = \cup \{f_r\}_{r=1}^R, s(r) = \text{card } f_r, \quad (1)$$

where $s(r)$ is power of description f_r at the hierarchical level numbered r .

At $r = 0$, we have the lowest level representing the complete description Z . The highest level, numbered R , will be considered as the transformation in the form of a single data vector, for example, a medoid for the set of vectors or the centroid as the mean vector of the elements in the set [4, 21, 24].

The hierarchies f_r will be formed sequentially through a recursive approach, starting from the lowest level (description Z). As a result, we obtain a structured sequence of descriptions

$$Z \rightarrow f_1 \rightarrow f_2 \rightarrow \dots \rightarrow f_R.$$

The aim of constructing the hierarchy is to reduce the dimensionality of the description to enhance the performance of classification by accelerating data processing procedures.

To create a sequential hierarchy of descriptions, we apply a transformation procedure T at each step, which reduces the cardinality of the description when transitioning to the next level

$$f_{r+1} = T\{f_r\}, \quad s(r+1) \ll s(r).$$

We will set the task of forming a higher-level hierarchy until the classifier, utilizing the transformed compressed description $\{f_{r+1}\}$, performs classification effectively with a predetermined accuracy. The main objective of such a construction is to speed up the classification process through the generalization of description data.

We define T as the data granulation procedure based on the model of forming a single vector for the next level, derived from the analysis of the pair of most similar vectors at the previous level. This action is applied in the process of data clustering [14, 15].

The introduced recursion allows us to apply the same procedure, which establishes the rule for transitioning to the next level. This procedure produces one element at the upper level based on a group of elements at the lower level. This is a model of pyramid-recursive structures. In fact, the procedure T performs the partitioning of the set of descriptors into pairs of equivalent elements. The proposed approach implements a “bottom-up” analysis strategy, utilizing the merging of lower-level elements when calculating the elements of the next level in the hierarchy.

By applying the transformation T to the description of the etalon E_k , where the number of descriptors is s , the transformation is defined as

$$T(E_k) \rightarrow E_k^*, \quad \text{card}(E_k^*) = s(k) \ll s, \quad (2)$$

where the result of this transformation is the filtered description E_k^* , which represents a subset $E_k^* \subset E_k$ of components from E_k , selected based on a specific granulation criterion.

In practice, it is proposed to reduce the cardinality of the description by approximately 1.5 to 2 times at each processing stage. As a result of the implementation of (2), the base $E = \cup E_k$ acquires a compressed representation as $E^* = \cup E_k^*$.

Options for Granulation Methods

The criterion in (2) is based on defining the next-level element through the pair of closest elements at the lower level. Thus, the concept of metric proximity underpins the formation of the granule for the next level. This method essentially divides the data into pairs of similar elements, replacing each pair with a single resulting element.

The principle of metric proximity analysis between descriptors can serve as a foundation for creating various filtering (reduction) methods. Let's consider specific variants of building the transformation T based on granulation, which relies on the metric ρ for descriptors.

Clearly, next-level granules can be formed using any arbitrary logical operation.

1. Granulation by equivalence threshold.

For each etalon descriptor $e_d(k) \in E_k$, the closest descriptor $e_*(k)$ in the set E_k is determined based on the

metric ρ . This corresponds to the model of linear metric search

$$e_*(k) = \arg \min_{v \neq d} \rho(e_v(k), e_d(k)). \quad (3)$$

After identifying the descriptor with the minimum metric, the condition of equivalence is checked: $\rho(e_d(k), e_*(k)) \leq \delta_\rho$, where δ_ρ is a predefined threshold parameter for the metric value that approximately establishes the equivalence of two descriptors.

In the space of multidimensional vectors, the value of δ_ρ is often set to 25% of the maximum metric value, though it can be adjusted depending on the task and input data [14, 15].

If equivalence by the threshold δ_ρ in (3) is not established, the first argument is added to E_k^* , this element is excluded from further processing, and the second continues to participate in the analysis.

It's important to note that the order of descriptors in the description and the value of the threshold δ_ρ directly affect the granulation results.

In fact, the model (3) for determining the equivalence of two vectors is supplemented by the predicate P :

$$P(z_1, z_2) = 1 \mid \rho(z_1, z_2) \leq \delta_\rho. \quad (4)$$

The number of resulting vectors in this variant depends on the threshold δ_ρ .

An alternative approach to implementing (3) involves identifying the pair of closest vectors without applying the equivalence threshold (4). In this case, the size of the description is reduced by approximately half from one level to the next.

If the nearest descriptor is chosen based on a threshold, the reduction is less than half, but it preserves a defined distance limit between newly created granules, maintaining a degree of separability in the transformed data. A simpler version of the analysis involves searching with a set equivalence threshold until the first identical element is found.

2. Granulation with filtering by minimum distance.

Involves calculating a distance matrix between all descriptors of the description E_k [3]. The filtering method begins with the smallest distance. A selected descriptor is added to the set E_k^* , and the corresponding row and column of this descriptor are removed from the matrix. Then, the descriptor with the next smallest distance is analyzed. In this way, a fixed number of descriptors are gradually added to the set E_k^* . This method is slightly more complex than the previous one, as it may involve sorting by distance.

However, it is not dependent on the order of the descriptors and does not require adapting the threshold δ_ρ in the data.

As a result, features from lower levels are combined into groups and treated as a single feature in subsequent steps. Compared to clustering, this granulation procedure progressively creates a modified data space governed by the equivalence parameter.

As an alternative to the previously proposed granulation methods, we can consider description reduction by discarding elements at a fixed interval. For

example, we can apply a selective exclusion method, where every second descriptor is removed from the set regardless of its characteristics.

This approach is much simpler and does not require analyzing distances between descriptors. The method can be useful as a first stage in data processing to quickly reduce the volume of input data before applying more complex and precise analysis methods.

However, this approach has several drawbacks. Important descriptors that may be critical for the accuracy of subsequent processing and data analysis could be lost. The elimination of elements is mechanical, without considering their significance or role in the overall structure of the set. In the case of uneven distribution of important and unimportant descriptors within the initial set, it might happen that after discarding every second element, either only important or only unimportant descriptors remain. For complex tasks where high precision and consideration of relationships between descriptors are required, this method is not effective. It does not provide deep analysis and may lead to inaccurate results.

Classification Schemes and Criteria

Note that the classification process needs to be slightly modified compared to the traditional method. It is necessary to consider that the description of the object remains complete, while the descriptions of the etalon models are reduced, i.e., $\text{card } Z \gg \text{card } E_k$ [21]. We will also assume that the etalon models E_k have approximately equivalent descriptive capacities.

Let's introduce the vector $\{h_i\}_{i=1}^N$ – an accumulator of votes for a system with N classes. The classifier is built in two stages [2, 21].

In the first stage, the elements z_i of the analyzed description $Z = \{z_i\}_{i=1}^S$ of the object receive (or do not receive) a vote for a specific etalon model, meaning they find correspondences in the etalon model E_k^* by verifying the truth of the predicate

$$Q(z_i, E_k^*) = 1 \mid \mid \rho_* = \min_{k, v=1, \dots, s(k)} \rho(z_i, e_v(k)) \& (\rho_* \leq \delta_\rho), \quad (5)$$

where $s(k) = \text{card } E_k^*$ – the capacity of the reduced etalon description, $e_v(k) \in E_k^*$.

If the predicate $Q = 1$, then we increment $h_k = h_k + 1$ the number of votes for the class with number k .

Thus, each element $z_i \in Z$ is classified.

In fact, the analysis of model (5) is carried out on the full set $E^* = \cup E_k^*$ of the reduced description base. In this version, the s elements of the object distribute their votes across the system of N classes, and as a result of the analysis, we have $\sum_{k=1}^N h_k \leq s$.

Based on the processing of the full description Z of the object, we compute the vector $\{h_i\}_{i=1}^N$, which contains the distribution of the object's votes across the system of classes.

In the second stage, we determine the object's class as the argument of the maximum

$$Z \rightarrow E_k \mid (k = \arg \max_{i=1, \dots, N} h_i) \& (h_k \geq \delta_h), \quad (6)$$

where δ_h is a threshold for the minimum acceptable number of votes.

If the condition $h_k \geq \delta_h$ is not met, the object's class is not established (classification refusal, there is no basis to assign the object to any of the classes).

The value of δ_h is determined experimentally for a given etalon base E . Typically, it is chosen as the minimum number of votes h_m (with tolerance) necessary for confident classification of a test sample based on the etalon models. From a general theoretical perspective, δ_h represents a compromise in the task of distinguishing etalon descriptors from those of the vast remainder of other images that may enter the classification system [3, 20].

For practical application, we recommend selecting the threshold according to the model $\delta_h = 0.95h_m$, which corresponds to a permissible deviation of 5% below the minimum h_m .

The effectiveness of the classification will be evaluated using the precision criterion pr , which reflects the ratio of the number g of test images with the correct class assignment (or identification as an image outside the class base) to the total number G of experiments [21, 24, 27]:

$$pr = g/G. \quad (7)$$

Another important criterion is the reliability of the classification decision. Reliability β characterizes the confidence level in making the decision and is evaluated by a criterion calculated as the ratio of the local (second-highest) maximum h_{m2} among the votes $\{h_i\}_{i=1}^N$ to the main maximum h_{m1} :

$$\beta = h_{m2}/h_{m1}. \quad (8)$$

The closer the value of (8) is to zero, the better the classification reliability.

Experimental Results and Discussion

The main aim of our experiment was to validate the efficiency and performance of the proposed granulation methods in the classification task, comparing them with traditional (full description) and simplified (element rejection) approaches.

For our research, we chose the OpenCV library [24, 28]. Using the ORB (Oriented FAST and Rotated BRIEF) algorithm [6] to detect and describe keypoints in images, descriptor sets for keypoints were determined for the test image set. The full descriptions of the etalon and test images included 500 ORB descriptors, with possible minor deviations.

The test set in the experiment comprised 72 images, including five etalon images of football club logos (classification base, class alphabet, Fig. 1), three out-of-base images (Fig. 2), and images derived from these eight using geometric transformations such as scaling (10% decrease or increase), rotation (30 degrees left or right), or a combination of scaling and rotation. As can be seen, the second and third out-of-base images (Fig. 2) have significant visual similarities with the etalon images. Fig. 3 shows examples of test images with keypoint coordinates and under the effect of geometric transformations.



Fig. 1. Images of the logos from the etalons database



Fig. 2. Images outside the etalon database



Fig. 3. An example of test images with keypoint coordinates

A classification error was recorded in cases where a transformed etalon image was classified as belonging to the “wrong” class or when a transformed out-of-base image was assigned to one of the existing classes.

To perform granulation using the equivalence threshold (4), a threshold of $\delta p = 37.5\%$ of the maximum metric (256) was established at each level. The granulation was performed several times, and based on the modeling results; it was decided to stop at the fifth level of the hierarchy, where the level of accuracy

begins to decrease. Using the threshold analysis, the number of elements was reduced by less than half at each stage.

Table 1 presents data demonstrating the number of descriptors remaining in the base after granulation at each level. It also includes the number of votes for the correct class and the evaluation of the classification time for a single etalon image at different levels of granulation. The results in Table 1 were obtained for the descriptions in the existing etalon base.

Table 1 – Classification results using granulation

Granulation level	Time, s	Number of base descriptors	Number of votes per class
Full description	29.45	2492	500, 500, 495, 497, 500
I	15.19	1279	411, 409, 435, 451, 441
II	10.05	830	304, 331, 404, 417, 368
III	7.27	586	242, 267, 390, 373, 306
IV	5.55	449	212, 237, 368, 318, 256
V	4.48	376	213, 217, 343, 267, 227

As can be observed from Table 1, the higher the level of granulation, the fewer descriptors remain in the reduced descriptions of the etalons, and the classification time proportionally decreases.

Despite the significant reduction in the number of informative elements at the 5th level of granulation (six and a half times compared to the full description!), the classification accuracy is $pr = 1$, meaning it remains at the highest level, and the maximum number of votes always corresponds to the “correct” class.

Specifically, the values of the vote maxima $\{h_{ij}\}_{i=1}^N$ for one of the reduced description variants were 175, 148, 164, 172, 192, which means that, according to this data, the threshold δ_h equals $\delta_h = 0.95 \times 148 = 141$ votes. Now, the decision to assign a class is made if the number of votes received is 141 or higher. In the case of images outside the database, the number of votes must be less than 141.

Unlike granulation using an equivalence threshold, the approach of discarding elements with a fixed step always reduces the number of elements by a specific factor. Specifically, the number of descriptors decreased at each level of the hierarchy as 2492, 1245, 622, 309, 154, 75.

The modeling conducted on a test set of 72 images with geometric transformations showed the following results.

For experimental evaluation, the granulation results at the 5th level were selected, where the number of descriptors in the description of each etalon was reduced to 15, which is 33 times less compared to the volume of the full description. The alternatives included filtering methods and the traditional method with full descriptions of the etalons.

In comparative terms, we analyzed the accuracy and reliability of these reduction **schemes**:

1. The detector ORB independently selects only 15 (not 500!) descriptors for each image from the etalon database.
2. Filtering by selecting every 33rd descriptor in the etalon description.
3. Performing granulation at the 5th level with a threshold.
4. Filtering by discarding every other descriptor 5 times until 15 descriptors are left in each description.
5. Performing granulation at the 5th level without using a threshold.
6. Selecting the first 15 descriptors from each etalon description.

7. Classification using full descriptions without reduction (traditional method).

The modeling results are presented in Table 2.

Analyzing the data in Table 2, we see that the traditional method (scheme 7) based on the full description demonstrates the best accuracy and reliability scores. However, the rest of the schemes (1...6), which aim to speed up classification, require approximately 33 times less processing time but have lower performance metrics.

Granulation with a threshold (scheme 3) showed slightly better results in both criteria compared to granulation without a threshold (scheme 5).

Table 2 – Accuracy and reliability indicators for different schemes

Scheme	Accuracy, pr	Credibility, β
1	0.63	0.57
2	0.95	0.63
3	0.75	0.48
4	0.85	0.62
5	0.63	0.60
6	0.50	0.72
7	1.00	0.20

This can be explained by the introduction of additional analysis based on the minimum distance value.

Scheme 6, which uses the first 15 descriptors from the etalon descriptions, has the worst performance in both accuracy and reliability. Schemes 2 and 4 showed high accuracy but with a low reliability coefficient. This emphasizes that the effectiveness of the reduction methods (schemes 2, 4, 6) significantly depends on the data sequence and the reduction scheme used. This may require additional multiple iterations to achieve the necessary metrics.

Scheme 1 is fully based on the ORB descriptor properties, and its performance does not guarantee high classification efficiency.

Conclusions

The traditional method is based on a voting procedure and essentially implements a nearest neighbor method on a set of descriptors for the entire base of etalon descriptions.

Granulation and feature reduction in etalon descriptions reduce the search scope, which generally increases classification speed in proportion to the degree of reduction.

Based on the research, it can be concluded that the choice of the most effective scheme for compressing descriptions in classification structural methods directly depends on the defined set of etalons and the allowable diversity of input images. Some applied tasks can be effectively implemented by direct feature reduction of a full description, possibly by introducing a variety of methods.

More complex tasks require the use of computationally expensive metric granulation, based on principles of intelligent analysis, both with and without a threshold.

At the same time, the effectiveness of the feature reduction method varies widely and depends not only on the sequence of descriptors in the description but also on the processing scheme, thus requiring multiple iterations of reduction options.

On the other hand, granulation methods create a reduced description in a single pass using a fixed procedure during the preliminary data analysis stage, and the classification time does not increase as a result.

Acknowledgements

The results of this research were obtained under the international research project INITIATE under the grant No.

101136775-HORIZON-WIDERA-2023-ACCESS-03.

REFERENCES

1. Tymchyshyn, R., Volkov, O., Gospodarchuk, O. and Bogachuk, Yu. (2018), "Modern Approaches to Computer Vision", *Control systems and computers*, vol. 6, pp. 46–73, doi: <https://doi.org/10.15407/usim.2018.06.046>
2. Gorokhovatskyi, V., Tvoroshenko, I., Yakovleva, O., Hudáková, M. and Gorokhovatskyi, O. (2024), "Application a committee of Kohonen neural networks to training of image classifier based on description of descriptors set," *IEEE Access*, vol. 12, pp. 73376–73385, doi: <https://doi.org/10.1109/ACCESS.2024.3404371>
3. Putyatin, E. P. and Averin, S. I. (1990), *Image processing in robotics*, Mashinostroeniye, 1990, 320 p., available at: https://scholar.google.com.ua/citations?view_op=view_citation&hl=uk&user=dfWDBoAAAAJ&citation_for_view=dfWDBoAAAAJ:isC4tDSrTZIC
4. Gadetska, S. V., Gorokhovatskyi, V. O., Stiahlyk, N. I. and Vlasenko, N. V. (2022), "Statistical data analysis tools in image classification methods based on the description as a set of binary descriptors of key points", *Radio Electronics, Computer Science, Control*, no. 4, pp. 58–68, Jan. 2022, doi: <https://doi.org/10.15588/1607-3274-2021-4-6>
5. Lowe, D. G. (2004), "Distinctive image features from scale-invariant keypoints", *International Journal of Computer Vision*, vol. 60 (2), doi: <https://doi.org/10.1023/B:VISI.0000029664.99615.94>
6. Rublee, E., Rabaud, V., Konolige, K. and Bradski, G. (2011), "ORB: An efficient alternative to SIFT or SURF", *Proc. Int. Conf. Comput. Vis.*, Nov. 2011, Barcelona, Spain, pp. 2564–2571, doi: <https://doi.org/10.1109/ICCV.2011.6126544>
7. Crowley, J. and Riff, O. (2003), "Fast computation of scale normalized Gaussian receptive fields", *Proc. Scale-Space'03*, Isle of Skye, Scotland, *Springer Lecture Notes in Computer Science*, 2695, doi: https://doi.org/10.1007/3-540-44935-3_41
8. Li, H., Xie, F., Zhou, J. and Liu, J. (2024), "Object Detection, Segmentation and Categorization in Artificial Intelligence", *Electronics*, vol. 13, no. 2650, doi: <https://doi.org/10.3390/electronics13132650>
9. Daradkeh, Y.I., Gorokhovatskyi, V., Tvoroshenko, I. and Zeghid, M. (2022), "Cluster representation of the structural description of images for effective classification", *Computers, Materials & Continua*, vol. 73 (3), pp. 6069–6084, doi: <https://doi.org/10.32604/cmc.2022.030254>
10. Baldini, L., Martino, A. and Rizzi, A. (2019), "Stochastic information granules extraction for graph embedding and classification", *Proc. of the 11th Int. Joint Conf. on Comp. Intelligence*, doi: <https://doi.org/10.5220/0008149403910402>
11. Petrovska, I., Kuchuk, H. and Mozhaiev, M. (2022), "Features of the distribution of computing resources in cloud systems", *2022 IEEE 4th KhPI Week on Advanced Technology*, KhPI Week 2022 - Conference Proceedings, 03-07 October 2022, Code 183771, doi: <https://doi.org/10.1109/KhPIWeek57572.2022.9916459>
12. Gorokhovatskyi, V., Gadetska, S. and Stiahlyk, N. (2023), "Accelerating Image Classification based on a Model for Estimating Descriptor-to-Class Distance", *International Journal of Computing*, vol. 22(4), pp. 485–492, doi: <https://doi.org/10.47839/ijc.22.4.3355>
13. Zadeh, L. A. (1997), "Toward a theory of fuzzy information granulation and its centrality in human reasoning and fuzzy logic", *Fuzzy sets and systems*, vol. 90 (2), pp. 111–127, doi: [https://doi.org/10.1016/S0165-0114\(97\)00077-8](https://doi.org/10.1016/S0165-0114(97)00077-8)
14. Fang, Y. and Liu, L. (2022), "Angular quantization online hashing for image retrieval", *IEEE Access*, vol. 10, pp. 72577–72589, doi: <https://doi.org/10.1109/ACCESS.2021.3095367>
15. Flach, P. (2012), "Model ensembles", *The Art and Science of Algorithms That Make Sense of Data*, New York, NY, USA: Cambridge Univ. Press, pp. 330–342, doi: <https://doi.org/10.1017/CBO9780511973000>
16. Aggarwal, C.C. (2023), *Neural networks and deep learning*, Textbook, Springer International Publishing, 529 p., doi: <https://doi.org/10.1007/978-3-031-29642-0>
17. Xiong, H. and Li, Z. (2014), *Data Clustering: Algorithms and Application*, 1st ed., CRC Press, Boca Raton, 652 p., doi: <https://doi.org/10.1201/9781315373515>
18. Zhang X., Yu, F. X., Karaman, S. and Chang, S.-F. (2017), "Learning discriminative and transformation covariant local feature detectors", *IEEE Conf. on Computer Vision and Pattern Recognition*, CVPR, Honolulu, HI, USA, pp. 4923–4931, doi: <https://doi.org/10.1109/CVPR.2017.523>
19. Butenkov, S. (2004), "Granular Computing in Image Processing and Understanding", *Proc. of IASTED International Conf. On AI and Applications «ALA 2004»*, Innsbruck (Austria), February 10-14, 2004, doi: <https://doi.org/10.1016/j.procs.2017.01.111>
20. Vorobel, R. A. (2012), *Logarithmic image processing*, Kyiv, Naukova Dumka, 232 p., available at: https://scholar.google.com/citations?view_op=view_citation&hl=uk&user=vEtLugEAAAAJ&citation_for_view=vEtLugEAAAAJ:qjMakFHDy7sC

21. Daradkeh, Y.I., Gorokhovatskyi, V., Tvoroshenko, I., and Zeghid, M. (2024), "Improving the effectiveness of image classification structural methods by compressing the description according to the information content criterion", *Computers, Materials & Continua*, vol. 80, no. 2, pp. 3085–3106, doi: <https://doi.org/10.32604/cmc.2024.051709>
22. Yakovleva, O., Kovtunenkov, A., Liubchenko, V., Honcharenko, V. and Kobylin, O. (2023), "Face Detection for Video Surveillance-based Security System (COLINS-2023)", *CEUR Workshop Proceedings*, vol. 3403, pp. 69–86, available at: <https://www.sytoss.com/blog/face-detection-for-video-surveillance-based-security-system>
23. Ullah, Z., Qi, L., Pires, E.J.S., Reis, A. and Nunes, R.R. (2024), "A systematic review of computer vision techniques for quality control in end-of-line visual inspection of antenna parts. Computers", *Materials & Continua*, vol. 80(2), pp. 2387–2421, doi: <https://doi.org/10.32604/cmc.2024.047572>
24. Gorokhovatskyi, O., and Yakovleva, O. (2024), "Medoids as a packing of orb image descriptors", *Advanced Information Systems*, vol. 8(2), pp. 5–11, doi: <https://doi.org/10.20998/2522-9052.2024.2.01>
25. Kuchuk, H., Kovalenko, A., Ibrahim, B.F. and Ruban, I. (2019), "Adaptive compression method for video information", *International Journal of Advanced Trends in Computer Science and Engineering*, vol. 8(1), pp. 66-69, doi: <http://dx.doi.org/10.30534/ijatcse/2019/1181.22019>
26. Svyrydov, A., Kuchuk, H. and Tsiapa, O. (2018), "Improving efficiency of image recognition process: Approach and case study", *Proceedings of 2018 IEEE 9th International Conference on Dependable Systems, Services and Technologies, DESSERT 2018*, pp. 593–597, doi: <https://doi.org/10.1109/DESSERT.2018.8409201>
27. Gorokhovatskyi, V. and Vlasenko, N. (2021), "The image description reduction in the set of descriptors on informativeness metric criteria base", *Advanced Information Systems*, vol. 5 (4), pp. 10–16, doi: <https://doi.org/10.20998/2522-9052.2021.4.02>
28. (2024), *OpenCV*, available at: <https://docs.opencv.org/>

Received (Надійшла) 20.10.2024

Accepted for publication (Прийнята до друку) 22.01.2025

ABOUT THE AUTHORS / ВІДОМОСТІ ПРО АВТОРІВ

Гороховатський Володимир Олексійович – доктор технічних наук, професор, професор кафедри інформатики, Харківський національний університет радіоелектроніки, Харків, Україна;

Volodymyr Gorokhovatskyi – Doctor of Technical Sciences, Professor, Professor of Informatics Department, Kharkiv National University of Radio Electronics, Kharkiv, Ukraine;

e-mail: gorokhovatsky.vl@gmail.com; ORCID Author ID: <http://orcid.org/0000-0002-7839-6223>;

Scopus ID: <https://www.scopus.com/authid/detail.uri?authorId=6506997369>.

Чмутов Юрій Вадимович – аспірант, Харківський національний університет радіоелектроніки, Харків, Україна;

Yurii Chmutov – PhD student, Kharkiv National University of Radio Electronics, Kharkiv, Ukraine;

e-mail: yurii.chmutov@nure.ua; ORCID Author ID: <http://orcid.org/0000-0002-6164-6126>.

Творошенко Ірина Сергіївна – кандидат технічних наук, доцент, доцент кафедри інформатики, Харківський національний університет радіоелектроніки, Харків, Україна;

Iryna Tvoroshenko – Candidate of Technical Sciences, Associate Professor, Associate Professor of Informatics Department, Kharkiv National University of Radio Electronics, Kharkiv, Ukraine;

e-mail: iryina.tvoroshenko@nure.ua; ORCID Author ID: <https://orcid.org/0000-0002-7184-8143>;

Scopus ID: <https://www.scopus.com/authid/detail.uri?authorId=57211556518>.

Кобилін Олег Анатолійович – кандидат технічних наук, доцент, завідувач кафедри інформатики, Харківський національний університет радіоелектроніки, Харків, Україна;

Oleg Kobylin – Candidate of Technical Sciences, Associate Professor, Head of Informatics Department, Kharkiv National University of Radio Electronics, Kharkiv, Ukraine;

e-mail: oleg.kobylin@nure.ua; ORCID Author ID: <https://orcid.org/0000-0003-0834-0475>;

Scopus ID: <https://www.scopus.com/authid/detail.uri?authorId=56845919400>.

Зниження обсягу обчислювальних затрат шляхом стиснення структурного опису у методах класифікації зображень

В. О. Гороховатський, Ю. В. Чмутов, І. С. Творошенко, О. А. Кобилін

Анотація. Дослідження статті зосереджені на способах скорочення обсягу аналізованих даних при застосуванні методів класифікації зображень у системах комп'ютерного зору. **Мета** – розвинути підходи щодо скорочення розмірності векторного опису бази еталонів за допомогою метричної грануляції, що знижує обчислювальні витрати та прискорює процес класифікації при збереженні достатнього рівня точності. **Застосовувані методи:** дескриптори ключових точок, апарат метричної грануляції даних, теорія класифікації та обробки зображень, структури даних, програмне моделювання. **Отримані результати:** розроблено формалізм гранульованого подання, здійснено експериментальне моделювання з використанням п'ятирівневої грануляції, що скоротило часові витрати у десятки разів при збереженні високої точності класифікації. У порівняльному аспекті вивчено способи скорочення обсягу векторних описів на підставі відкидання даних, досліджено вплив рівня грануляції на точність та час класифікації. **Практична значущість роботи** полягає у підвищенні продуктивності структурних методів класифікації зображень шляхом впровадження грануляції та схем відкидання даних, що забезпечує набагато швидшу обробку даних без суттєвої втрати результативності класифікації.

Ключові слова: швидкодія класифікації; відкидання даних; скорочення опису; грануляція; класифікація зображень; дескриптори ключових точок.