

Quadrees for Partial Dependencies in Data

Alexander Vinogradov^{1[0009-0007-0728-9985]}

¹Federal Research Center 'Computer Science and Control', Russian Academy of Sciences,
119333, Moscow, Russia
vngrccas@mail.ru

Abstract. A new approach to the problem of efficient use of available a priori knowledge for deep analysis of complicated data is presented in case when several hypotheses are studied simultaneously that bound, in different and competing ways, some known or expected partial laws in the process of forming values of an aggregated numerical observable. A comparison of the adequacy of competing hypotheses is performed through a top-down refinement of the cumulative measure of difference in bit cells of the feature space. Considerations for organizing a user-friendly research interface are presented, including the use of quadtree techniques, which display the behavior of the measure in various bit layers of the sample on the monitor screen. The goal is to explore ways to leverage applied competencies, primarily imprecise and incomplete ones, as well as informal experience, subjective assumptions, and expert guesses. The CLT and the concept of generalized precedent (GP) as a multidimensional analogue of the Hough transform are used as instruments.

Keywords: Generalized Precedent, Partial Dependency, Stochastic Value.

1 Introduction

The problem of effectively leveraging a priori knowledge is relevant for many applied areas where tasks arise involving analyzing and predicting the behavior of complex data samples. Often, a complex sample arises in today's saturated information environment and is studied "from scratch," and, as consequence, the value of the resulting conclusions is limited by the volume of data analyzed. At the same time, for some recorded object characteristics, many data behavior patterns are discovered earlier, verified through years, and now known to the expert in the applied field. For these reasons, a wide variety of schemes, approaches, models, and algorithms developed based on them are used worldwide, all designed to utilize useful a priori information to the maximum extent [1-9].

This paper explores an approach to the problem based on the concept of generalized precedents (GP). V.V. Ryazanov proposed an approach with GP in order to utilize some properties of a simple description of clusters of a certain shape, found in a sample or on an image. The first successful example of this approach was proposed earlier by Yu.I. Zhuravlev in the form of a logical regularity representing a specific class in a sample $X \subset R^N$ in the form of a DNF, as a logical sum of conjunctions of

predicates for a new object's compliance with threshold constraints on the principal axes in R^N . The main advantage was the computational efficiency of threshold comparisons. Moreover, the decision rule was easily interpreted in generally accepted terms, understandable even to any unconcerned specialist [10], [11].

V.V. Ryazanov developed this idea and proposed considering a hyper-parallelepiped in R^N , a geometric image of a conjunction, as a new object, GP, defined by a point in the $2N$ -dimensional space Y of the parallelepiped's boundaries. It was immediately clear that the transition to a parametric space brings many of the advantages inherent in the classical Hough transform in the IP domain. It also turned out that many other methods of clustering and representing samples by sets of simple shapes can also serve as primary spatial differentiation. Each simple cluster, as a geometric object, is also represented by a point in the parametric space of kind Y [12].

As in the Hough scheme, the task of data analysis can be reduced to searching for point peaks in the density of the secondary distribution $\rho(y)$ on the space Y . In more advanced settings, the parametric space Y itself comes to the fore, the set of parameters y can receive a special meaning and be composed of few derivative quantities constructed in a complex way from primary features. In particular, the study of the behavior of the distribution $\rho(y) = \rho(y, X')$ on certain fragments of the sample $X' \subset X \subset R^N$ may help to identify previously unknown patterns that can control the gathering objects into dense clusters in R^N [13-18].

This paper proposes the next step in developing the approach initiated in the works of V.V. Ryazanov, where the role of primary differentiation operators is played by sets of a priori known partial numerical dependencies $f_m(x)$, $m = 1, 2, \dots, M$, expected by expert or reliably verified earlier as a result of studying many samples of the type $X \subset R^N$, with a similar set of features $x \in R^N$. In what follows, modeling and forecasting of the behavior of some aggregated observable quantity $Z(x)$ is studied, the values of which arise as a result of an interaction of dependencies of the type $f_m(x)$. It is assumed that the model includes a set of control parameters $a \in A$, governing the interaction, and further in the paper, various vectors a of the control parameters act as GPs. Several hypotheses $H_i(a, f)$, $i = 1, 2, \dots, I$, describing possible variants of the interaction of dependencies $f_m(x)$ are considered in parallel. The task is to select a model H_i and values of its control parameters that would ensure the smallest deviations between the simulated values $H_i(a, f(x))$ and the actual observed values $Z(x)$. Friendly interface design options are proposed that allow the best use of the researcher's applied competencies for these purposes. Positional data representation, and in particular quadtrees, are used to visually compare the adequacy of the $H_i(a, f(x))$ model across different fragments X' of the original sample, $X' \subset X \subset R^N$.

2 Problem Statement

Let X , $X \subset R^N$, be a sample of digitized applied data, and $f_m(x)$, $m = 1, 2, \dots, M$, be a set of a priori known partial numerical dependencies. Several model hypotheses $H_i(a, f(x))$, $i = 1, 2, \dots, I$, are investigated, corresponding to possible variants of the interaction of dependencies $f_m(x)$ while generating the value of aggregated quanti-

ty $Z(x)$. The problem consists of finding, for each hypothesis H_i , the optimal vector $a_i^* \in A$ that ensures the smallest differences between the values of $H_i(a, f(x))$ and $Z(x)$ over the maximum number of objects in the sample X . We assume that the index i , corresponding to the minimum, supports the hypothesis $H_i(a_i^*, f(x))$ that most adequately describes the way in which the parameters a_i function in the interaction of dependencies $f_m(x)$.

Important features of the formulation under consideration should be noted.

No restrictions on the type of digital model $H_i(a, f(x))$ are imposed a priori. All responsibility is shifted to the shoulders of the researcher proposing this hypothesis.

Similarly, no functional restrictions are imposed on the term 'partial numerical dependency $f_m(x)$ ', $m = 1, 2, \dots, M$. This role can be played by both explicitly defined functional relationships and various constraints, the researcher's assumptions based on their extensive experience, etc.

For different hypothesis $H_i(a, f(x))$, the substantive meaning of the control parameters, as well as their number, may vary. However, in what follows, only the numerical values of the parameters will be important, since the search for the optimal vector $a_i^* \in A$ is performed according to a uniform scheme in a space A .

The modulus $|Z(x) - H_i(a, f(x))|$, $i = 1, 2, \dots, I$, can serve as a natural measure $\mu_i(a, x)$ of the difference between the simulated and actual values of the aggregated quantity $Z(x)$ for any $x \in X$. If the search for $a_i^* \in A$ is performed on a restricted discrete grid $\mathcal{U}$ of vectors $a \in \mathcal{U} \subset A$, then for each $x \in X$ it is easy to find the absolute minimum $a_i(x) \in \mathcal{U}$ of the measure $\mu_i(a, x)$. The main assumption is that the manifestation of many extreme points $a_i(x)$ in the form of a compact cluster C_i , $C_i \subset \mathcal{U}$, while moving x through the sample, will be the evidence of the adequacy of the studied model $H_i(a, f(x))$. Some center $a_i^* \in A$, chosen for this cluster in one way or another, will represent the averaged parameters of the interaction model, that are valid for involved share of the sample X .

A clear illustration of the parallels with the Hough transform here is a variant of this transform designed to detect contour polygons in an image. The set of dependencies $f_m(x)$ here corresponds to operators for detecting corners and line segments, while each set $a \in A$ acts as a variant of logical conditions for combining elements into figures of a particular shape, for example, triangles. This means that for some i , the image is characterized by the presence of the largest number of objects (triangles) of a specific shape, described in the hypothesis H_i by the parameters a_i^* , while for other sets a_j^* , $j \neq i$, hypotheses like H_j assumed triangles of a different shape, or rectangles, or some other shapes that turned out to be less frequently represented in the given image.

In this study, the most adequate hypothesis $H_i(a_i^*, f(x))$ about the interaction of the dependencies $f_m(x)$, which provides the minimum sum of differences between the values of the quantity $Z(x)$ and its model, plays the role of this optimal triangle of the most frequent shape.

3 Correcting Hypotheses on Bit Cells

The choice of the cluster center $\mathbf{a}_i^* \in \mathbf{A}$ is associated with several intermediate problems related to properties such as the cluster's cardinality, its shape, etc. At the same time, the simplicity of defining the measure $\mu(\mathbf{a}, \mathbf{x})$ allows us to apply the CLT here and construct a solution in which many of these problems are automatically resolved.

As known, the probability density for the sum of $|X|$ identically distributed random variables quickly normalizes in the neighborhood of the mean, as $|X|$ increases. In this case, the standard deviation σ decreases at the rate of $\sqrt{|X|}^{-1}$, whatever the original distribution.

As such variables, the values of the difference measure $\mu(\mathbf{a}, \mathbf{x})$ on all vectors $\mathbf{x} \in X$ can be considered. Each of these variables obeys one and the same distribution of the empirical density of the sample $X \subset R^N$. The corresponding cumulative sum $\mu_i(\mathbf{a}, X) = \sum_{\mathbf{x} \in X} \mu_i(\mathbf{a}, \mathbf{x})$ will be represented on the grid $\mathcal{A}$ by the distribution of the summarized differences between $\mathbf{Z}(\mathbf{x})$ and $H_i(\mathbf{a}, \mathbf{f}(\mathbf{x}))$, which are averaged, smoothed, and stable. The mean of the distribution is the best candidate for the role of $\mathbf{a}_i^* \in \mathbf{A}$.

Unfortunately, the value of $\mu_i(\mathbf{a}_i^*, X)$ does not guarantee zero or even a small total difference between the values of $\mathbf{Z}(\mathbf{x})$ and $H_i(\mathbf{a}_i^*, \mathbf{f}(\mathbf{x}))$. This value only indicates the best accuracy of reconstructing the interaction of dependencies $f_m(\mathbf{x})$ that the hypothesis H_i is capable of providing on the full sample X in principle. If this accuracy is low, it is necessary to find ways to adjust the internal construction of the numerical model $H_i(\mathbf{a}, \mathbf{f}(\mathbf{x}))$ itself.

The main goal of this study is to develop a toolkit that allows for the best use of the experience of an expert in the applied field, and even their intuitive guesses. If successful, the latter may then be transformed into new knowledge. Indeed, in the presence of the measure $\mu_i(\mathbf{a}, X)$, both the reverse transition to the original sample and the verification of the adequacy of the measure $\mu_i(\mathbf{a}, X')$ for its various parts X' , $X' \subset X \subset R^N$ become possible. As a rule, an expert in the applied field who has formulated a particular hypothesis $H_i(\mathbf{a}, \mathbf{f}(\mathbf{x}))$ also has knowledge of the specifics of the data in various parts of the original sample. Having knowledge of the adequacy of the model $\mu_i(\mathbf{a}, X')$ on various fragments of the sample $X' \subset X \subset R^N$ and of the degrees of their negative contribution to the value of the difference measure $\mu_i(\mathbf{a}, X')$, the expert can change the place and the sense of participation of individual components of the vector $\mathbf{a}$ in the model $H_i(\mathbf{a}, \mathbf{f}(\mathbf{x}))$, in particular, remove unnecessary or add additional control parameters.

Below, for the sake of convenience and clarity for the expert, a regular method of dividing the sample X into bit cells is used, in each of which the study of the behavior of the measure $\mu_i(\mathbf{a}, X)$ can be performed separately and in one and the same way. It is assumed that the components of the vectors $\mathbf{x} \in X$ are positive, specified in binary code, and in any of the components of the vector $\mathbf{x}$, only the K elder binary digits are used [19-20]. Each vector $\mathbf{x}$ is then represented by the binary matrix $[\chi_{nk}]$, $n=1, 2, \dots, N$, $k=1, 2, \dots, K$, $\chi_{nk}=0, 1$. Thus, the addressing is specified of the location of objects of the sample X in bit cells of the positive quadrant of the space R^N . A special case of this type of addressing is, in fact, quadtree. If, for example, for one reason or another, an

expert is interested in projecting a sample X onto a monitor screen in features (x_α, x_β) , $\alpha, \beta = 1, 2, \dots, N$, fragmented to the level $K=2$, then a possible type of projection could be the quadtree shown in Figure 1. As X , some multilevel GIS image was chosen, RGB components of vectors $\mathbf{x}$ are represented in the quadtree for two spatial coordinates, hidden features are implied, and any two of them can be used in similar manner, too.



Fig. 1. Quadtree for a sample X shown in spatial coordinates (x_α, x_β) . Color range reflects the degree a of expert's interest to the data in different cells. Here $(00,00)$ and $(11,11)$ are the cells of the least interest. Scalar a serves as control parameter of certain dependencies $f_m(\mathbf{x})$ related to hidden features of the water body infrastructure.

The expert does not necessarily need to place the entire sample X in the quadtree. The expert can select for analysis a particular data slice where only some of the bit indices are allowed. For example, for a feature x_γ , $\gamma = 1, 2, \dots, N$, (or for several such features), the expert can optionally place in the quadtree cells, being constructed for the dimensions (x_α, x_β) , only points of those vectors $\mathbf{x} \in X$, where certain bits $x_{\gamma k}$ for some specific indices k , $k = 1, 2, \dots, K$, have a fixed value: $\chi_{nk} = 0$ or $\chi_{nk} = 1$. This opportunity can utilize the expert's experience, visually observing the influence of known features of the selected data fragment X' , $X' \subset X \subset R^N$, on the local behavior of the measure $\mu_i(\mathbf{a}, X')$, and hints to the expert a way to improve the hypothesis $H_i(\mathbf{a}, f(\mathbf{x}))$ under consideration.

4 Sequential Refinement of the Difference Measure

So far, only the problem $\mu_i(\mathbf{a}_i^*, X) = \min_{\mathbf{a}} \mu_i(\mathbf{a}, X)$ has been solved, indicating the accuracy achievable for the $H_i(\mathbf{a}, f(\mathbf{x}))$ model on the entire sample X . A similar problem can also be solved for each bit cell X^χ of the quadtree, where χ is the bit address of the cell. Knowing the optimal values of the total measure $\mu_i(\mathbf{a}, X^\chi)$ in different cells, the

expert can make informed modifications to the model. In particular, the components of the control parameter vector $\mathbf{a}_i$ then become functions dependent on the bit index of a cell: $\mathbf{a}_i = \mathbf{a}_i(\chi)$. This confidence is supported by the following simple lemma of three means, proven by different authors in many variations:

Lemma 1.

Let $X = X^1 \cup X^2$, $X^1 \neq X^2$, $a_i^1 = \frac{1}{|X^1|} \min_{\mathbf{a}} \mu_i(\mathbf{a}, X^1)$, $a_i^2 = \frac{1}{|X^2|} \min_{\mathbf{a}} \mu_i(\mathbf{a}, X^2)$, $a_i^* = \frac{1}{|X|} \min_{\mathbf{a}} \mu_i(\mathbf{a}, X)$. Then $a_i^1 \leq a_i^*$ or $a_i^2 \leq a_i^*$.

For example, if X^1, X^2 are bit cells of the same level, and for the i -th hypothesis the inequality $a_i^1 < a_i^*$ is valid, then the expert has grounds for a more detailed study of the applicability of the set of control parameters $\mathbf{a}_i^1$, extreme in the problem $\min_{\mathbf{a}} \mu_i(\mathbf{a}, X^1)$, also in other cells of the quadtree. Conversely, if $a_i^2 > a_i^*$, then the corresponding set of control parameters $\mathbf{a}_i^2$ most likely requires modification in this cell, or entire replacement with a more adequate hypothesis itself.

The best case is when, in all cells describing increasingly filled levels of detail k , $k > 0$, all values of type a_i^1, a_i^2 satisfy the conditions $a_i^1 \leq a_i^*$, $a_i^2 \leq a_i^*$, and when k increases from 1 to K , the corresponding vectors of control parameters of type $\mathbf{a}_i^1, \mathbf{a}_i^2$ remain close to $\mathbf{a}_i^*$. This then means that the solution to the problem $\mu_i(\mathbf{a}_i^*, X) = \min_{\mathbf{a}} \mu_i(\mathbf{a}, X)$ was initially successful, and the set $\mathbf{a}_i^*$ does not require correction with increasing detailing the sample X .

However, as k increases, the number of objects in a single bit cell, such as X^1, X^2 , decreases at an average rate not lower 2^{kN} . Certain share of solutions of the form $\mathbf{a}_i^1, \mathbf{a}_i^2$ become increasingly less reliable, some of the cells become empty, and then normalizing factors of the type $|X^1|, |X^2|$ become meaningless.

It is further assumed that, as k increases, the optimal set of control parameters from the highest enclosing cell of the bit structure is automatically transferred to the younger empty cells. Figure 2 shows the evolution of this solution with increasing k in the case where the bit structure of the sample is represented in projection onto a single specific dimension x_n , which, for one reason or another, is of particular interest to the researcher.

A natural assumption is that the expert should be especially attentive to those hypotheses that change optimal sets of the form $\mathbf{a}_i^k$ less frequently in diminishing bit cells, as k increases from 1 to K . In this case, the graph of the function $\mathbf{a}_i(\chi)$, represented for the full sample X by the constant vector $\mathbf{a}_i^*$, is transformed into a certain sequence of discontinuous hypersurfaces $S_k \subset \mathcal{H}$. Since the values of the total difference measure $\mu_i(\mathbf{a}, X^k)$ for lower-order cells become increasingly less reliable, the expert faces two interrelated problems. Among the hypotheses H_i , it is necessary to select the one that ensures minimal transformations for S_k , and this property is preserved through the maximum levels of detail $k, k=1,2,\dots,K$.

All the above-mentioned granularity features are depicted in Figure 2.

For Hypothesis 1, the difference measure for the full sample is initially high and barely decreases with granularity.

For Hypothesis 2, the difference measure decreases at all three levels of detail, but at the lowest level, this is achieved via unacceptable changes in the set of control parameters.

For Hypothesis 3, the sample is divided into two subregions, where two distinct vectors of control parameters consistently emerge.

Hypothesis 4 yields the best solution: the optimal vector of control parameters changes little with detailing, and the difference measure is initially low.

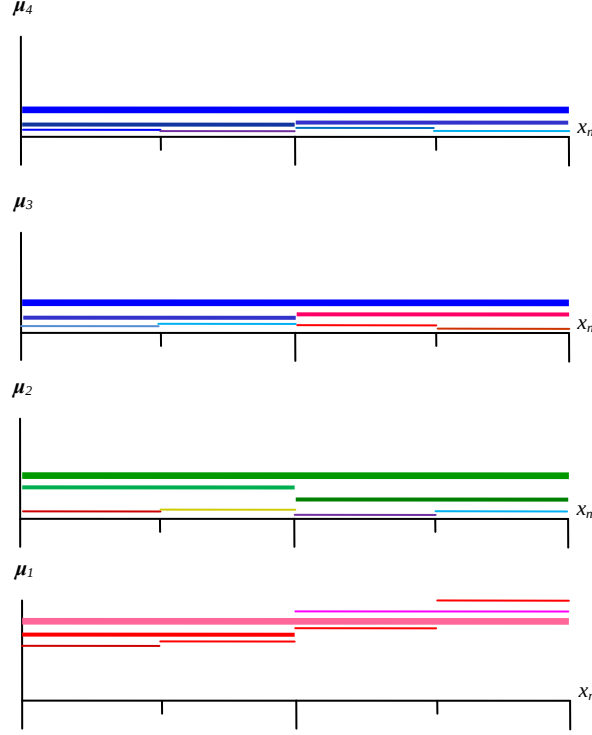


Fig. 2. Modeled 3-level refinement of the normalized measure $\mu_i = \frac{1}{|X^X|} \mu_i(\mathbf{a}_i^*, X^X)$ for 4 hypotheses H_i , $I=4$. RGB-components depict possible 3-dimensional vector of control parameters $\mathbf{a} \in \mathcal{A} \subset \mathbb{R}^3$

It is beyond the scope of this article to detail the various techniques and considerations an expert might use when selecting the best hypothesis. For example, they might consider not only absolute but also local minima of the measure $\mu_i(\mathbf{a}, X^X)$ in cells, the combination of which would ensure minimal transformations. Conversely, this approach may reveal disjoint data regions requiring significantly different optimal sets of the form $\mathbf{a}_i^X$ or different hypotheses $H_i(\mathbf{a}, f)$, $i=1,2,\dots,I$, that turn out the best for disjoint regions. In all cases, Lemma 1 guarantees that, as the level of detail k increases, the corresponding hypersurface S_k provides an increasingly accurate reconstruction of the interaction of dependencies $f_m(\mathbf{x})$, although local solutions become increasingly less well-founded.

Here some considerations are discussed why the proposed toolkit is primarily oriented toward incorporating various incomplete, imprecise, and heuristic knowledge

into the analysis, often discovered intuitively, including comparison of hypothesis, in particular, with the help of visualization. Indeed, if some of the used dependencies $f_m(\mathbf{x})$ have exact polynomial approximations of the form $f_t(\mathbf{x})=0$, $t=1,2,\dots,T$, while the remaining dependencies are imprecise and not universally defined, then the sample X is located in R^N on an algebraic hypersurface of codimension T , making it easier to work with. If, moreover, the corresponding system of T equations has a unique solution with respect to some T initial features, then the main problem of using a set of a priori dependencies returns to the original formulation, since all these T features obtain representations in terms of remaining ones, and now the problem may be solved in the feature space of reduced dimensionality R^{N-T} .

Of course, it wouldn't be very difficult to develop an application for this task that would implement all the steps described above in a digital environment. However, it's hardly feasible to take into account many of the expert's informal considerations and hypotheses in advance and integrate them into the interface of this hypothetical application. Moreover, many of them will be rejected during experimentation. From this perspective, directly using the expert's knowledge and intuition remains a working option.

5 Conclusion

The paper presents a new approach to leveraging the a priori knowledge and experience of specialists accumulated in an applied field. Particular attention is given to poorly formalized aspects of experience. It is shown that precise a priori dependencies, reliably established and specified digitally, can be taken into account and excluded from analysis in advance. The problem of comparing competing hypotheses about the role of partial a priori dependencies in generating the values of an aggregated quantity is investigated. The numerical model of each hypothesis is loaded with a set of control parameters included in the model to govern the interaction of dependencies. The space of control parameters plays the role of the Hough space, whose points act as GPs. Among all of the points, the extremums of the cumulative measure, which calculates differences between modeled and observed values of the aggregated quantity in the bit cells of the feature space, are identified. Proposals are formulated for organizing a user-friendly interface to leverage the informal knowledge and practical experience of specialists, including the use of a quadtree technique that displays various layers of estimates of the difference measure in bit cells on a monitor screen. The paper utilizes a lemma that guarantees increasing (non-decreasing) accuracy of estimates despite granularity of the measure in smaller bit cells, but the increasing randomness of the sample fraction that forms this local estimate is a problem. An interface is proposed for selecting the hypothesis that provides the minimum of the measure along with increasing granularity. It is noted that the search for the best hypothesis is easily automated, but the development of the corresponding software will rely heavily on the same a priori knowledge that is directly used in this case. Therefore, the approach presented in the paper is a viable alternative.

References

1. Xie, Y. et al.: Optimal scheduling method integrating priori expert knowledge for water distribution systems based on historical operational information tensor: A new perspective: Resources, Conservation and Recycling, **219**, Elsevier, (2025). doi.org/10.1016/j.resconrec.2025.108321.
2. Ma, W. Dong, F. Li, Y. Li, B. Zhou, C.: Application of a priori knowledge-enhanced fuzzy clustering to acoustic emission-based damage identification of composite laminates. Applied Acoustics **229** (2025). doi.org/10.1016/j.apacoust.2024.110404.
3. Yang, R. et al.: A large-scale ultra-high-resolution segmentation dataset augmentation framework for photovoltaic panels in photovoltaic power plants based on priori knowledge. Applied Energy **390** (2025). doi.org/10.1016/j.apenergy.2025.125879.
4. Wei, H. et al.: Real-time remote sensing detection framework of the earth's surface anomalies based on a priori knowledge base. International Journal of Applied Earth Observation and Geoinformation **122**, 1-9 (2023). doi.org/10.1016/j.jag.2023.103429.
5. Li, D. Fai, C. Wang, B. Liu, M.: A study of a priori knowledge-assisted multi-scopic metrology for freeform surface measurement. Procedia CIRP **75**, 337-342 (2018). doi.org/10.1016/j.procir.2018.05.002.
6. Mirpulatov, I. Kedrov, A. Illarionova, S.: Improving plot-level growing stock volume estimation using machine learning and remote sensing data fusion. Computer Optics **49**(4), 682-691 (2025).
7. Mandrikova, B.: A method for analyzing complex structured data with elements of machine learning. Computer Optics **46**(3), 506-512 (2022).
8. Duan, H. et al.: Large-area urban TomoSAR method with limited a priori knowledge and a complex deep learning model. International Journal of Applied Earth Observation and Geoinformation **139**, 1-16 (2025). doi.org/10.1016/j.jse.2018.11.035.
9. Wang, Y. Liu, C. Zhu, R. Liu, M. Ding, Z. Zeng, T.: MAda-Net: Model-adaptive deep learning imaging for SAR tomography. IEEE Transactions on Geoscience and Remote Sensing **61**, 1-13 (2023).
10. Zhouravlev, Yu.I. Ryazanov, V.V. Sen'ko, O.V.: Raspoznavanie. Matematicheskiye metody. Programmaya sistema. Prakticheskiye primeneniya. Fazis, Moskva (2006). (Russian)
11. Ryazanov, V. Vinogradov, A.: Analogues of image analysis tools in the problems of finding latent regularities in big applied data. Pattern Recognition and Image Analysis **32**(3), 639-644 (2022). DOI: 10.1134/S105466182203035X.
12. Vinogradov, A.: Searching for Hidden Patterns of Interacting Partial Regularities in Data". In: ICPR 2024 International Workshops and Challenges, LNCS, vol. 15616, pp. 400-408. Springer, Heidelberg (2025). doi.org/10.1007/978-3-031-87663-9_33.
13. Ryazanov, V. Vinogradov, A. Angalt, E.: Weakened generalized precedents in problems of searching for informative aspects of data. In: Proceedings of the X International Conf. on Information Technology and Nanotechnology (ITNT), IEEEExplore, pp. 1-5 (2024). DOI: 10.1109/ITNT60778.2024.
14. Naumov, V. Nelyubina, E. Ryazanov, V. Vinogradov, A.: Analysis and prediction of hydrological series based on generalized precedents. In: Book of abstracts of the 12-th Int. Conf. Intelligent Data Processing (IDP-12), Gaeta, Italy, pp.178-179 (2018).
15. Nelyubina, E. Ryazanov, V. Vinogradov, A.: Transforms of Hough Type in Abstract Feature Space: Generalized Precedents. In: Proceedings of the 12th International Joint Confer-

- ence on Computer Vision, Imaging and Computer Graphics Theory and Applications VISIGRAPP 2017. INSTICC Press, Porto, Portugal, pp. 651–656 (2017).
16. Nelyubina, E. Ryazanov, V. Vinogradov, A.: Analogs of Image Analysis Tools in the Search of Latent Regularities in Applied Data. In: Pattern Recognition, Computer Vision, and Image Processing. ICPR 2022 International Workshops and Challenges, Part II, LNCS, vol. 13644, pp. 529–540. Springer, Heidelberg (2023). <https://doi.org/10.1007/978-3-031-37742-6>.
 17. Ryazanov, V. Vinogradov, A.: Analogues of Image Analysis Tools in the Problems of Finding Latent Regularities in Big Applied Data. Pattern Recognition and Image Analysis **32**(3), 639–644 (2022).
 18. Ryazanov, V. Vinogradov, A.: Podkhody k vyboru sodержatel'nykh aspektov pri analize slozhnykh vyborok». In: Sbornik trudov 21-y Vserossiyskoy konferentsii s mezhdunarodnym uchastiyem «Matematicheskiye metody raspoznavaniya obrazov» (MMRO-2023), str. 6–8 (2023). (Russian)
 19. Vinogradov, A. Laptin, Yu.: Mining Coherent Logical Regularities of Type 2 via Positional Preprocessing. In: Gourevich, I. Niemann, H. Salvetti O. (eds.) VISIGRAPP/IMTA 2013 4th International Workshop on Image Mining Theory and Applications, in conjunction with VISIGRAPP 2013, pp.56–62, INSTICC Press, Setubal (2013).
 20. Vinogradov, A. Laptin, Yu.: Using Bit Representation for Generalized Precedents. In: Proceedings of International Workshop OGRW-9, Koblenz, Germany, December 2014, pp.281–283 (2015).