

Review

About Granular Rough Computing—Overview of Decision System Approximation Techniques and Future Perspectives

Piotr Artiemjew 

Faculty of Mathematics and Computer Science, University of Warmia and Mazury in Olsztyn,
10-710 Olsztyn, Poland; artem@matman.uwm.edu.pl

Received: 2 February 2020; Accepted: 27 March 2020; Published: 29 March 2020



Abstract: Granular computing techniques are a huge discipline in which the basic component is to operate on groups of similar objects according to a fixed similarity measure. The first references to the granular computing can be seen in the works of Zadeh in fuzzy set theory. Granular computing allows for a very natural modelling of the world. It is very likely that the human brain, while solving problems, performs granular calculations on data collected from the senses. The researchers of this paradigm have proven the unlimited possibilities of granular computing. Among other things, they are used in the processes of classification, regression, missing values handling, for feature selection, and as mechanisms of data approximation. It is impossible to quote all methods based on granular computing—we can only discuss a selected group of techniques. In the article, we have presented a review of recently developed granulation techniques belonging to the family of approximation algorithms founded by Polkowski—in the framework of rough set theory. Starting from the basic Polkowski’s standard granulation, we have described further developed by us concept dependent, layered, and epsilon variants, and our recent homogeneous granulation. We are presenting simple numerical examples and samples of research results. The effectiveness of these methods in terms of decision system size reduction and maintenance of the internal knowledge from the original data are presented. The reduction in the number of objects in our techniques while maintaining classification efficiency reaches 90 percent—for standard granulation with usage of a kNN classifier (we achieve similar efficiency for the concept-dependent technique for the Naive Bayes classifier). The largest reduction achieved in the number of exhaustive set of rules at the efficiency level to the original data are 99 percent—it is for concept-dependent granulation. In homogeneous variants, the reduction is less than 60 percent, but the advantage of these techniques is that it is not necessary to look for optimal granulation parameters, which are selected dynamically. We also describe potential directions of development of granular computing techniques by prism of described methods.

Keywords: rough sets; granular rough computing; granulation techniques; classification

1. Introduction

Granular computing is dedicated to work on data in the form of grouped, similar information vectors. The idea was introduced by Lotfi Zadeh [1,2]. Granulation is an integral part of the fuzzy set theory by the very definition of fuzzy set, where inverse values of fuzzy membership functions are the basic forms of granules. Shortly after Lotfi Zadeh proposed the idea of granular computing, the granules were introduced in terms of rough set theory [3] by T.Y. Lin, L. Polkowski, and A. Skowron. In this theory, granules are defined as classes of indiscernibility relations. Interesting research on more flexible granules based on blocks was conducted by (Grzymala–Busse) (see the LEM2 algorithm), and templates by (H.S. Nguyen), used in classification processes. The granules based on rough

inclusions were introduced by (Polkowski and Skowron [4]), based on tolerance or similarity relations, and, more generally, binary relations by (T.Y. Lin [5], Y. Y. Yao [6–8]). In the context of rough mereology was proposed by (L. Polkowski and A. Skowron), in approximation spaces by (A. Skowron and J. Stepaniuk [9,10]), and finally in logic for approximate reasoning by (L. Polkowski, M. Semeniuk-Polkowska [11], and Qing Liu [12]). Of course, many other authors are conducting considerations on groups of similar objects, which is simply the most natural way of modeling problems; it is impossible to name them all. Let us quote a few very interesting works on various research topics from recent years on granular computing [13–18]. Additionally, interesting research on the field of granular computation with the use of neural network techniques can be found in the works [19–21].

We have developed our methods in terms of granular rough computing paradigm—the internal part of rough sets theory [3]. The computations are based on granules, the groups of objects collected together by fixed similarity measure or metrics. Theoretical background and the framework of discussed methods were proposed by Polkowski in [22–24]—the idea of data approximation using rough inclusions. The basic idea was to create the r -indiscernible groups of objects (objects indiscernible in fixed degree) around each training sample, cover the original training decision system using selected granules and create the granular reflection of training data using granules from the covering in the final step. This particular technique is called standard granulation and was proposed in [24]. The initial work was extended later in many variants and contexts—see [25,26], Polkowski [27,28], Polkowski and Artiemjew [29,30]. These methods, among others, have found application in classification processes [31], data approximations [30], missing values absorption [26,29], and, in the recent work, these were used as a key component of the new Ensemble model—see [32].

In the review, we are focusing on decision system size reduction and maintaining the internal knowledge at the same time. Despite the fact that the granulation of the decision systems in a pessimistic case has a square complexity, it is possible to apply classical techniques of transferring methods to big data for the purpose mentioned. In the article, we have described standard granulation [24], concept-dependent [25], layered [25] and homogeneous granulation [33]—designed for symbolic data, and exemplary variants developed for numerical one—with descriptors indiscernibility ratio–epsilon granulation [33,34].

The rest of the paper has the following content. In Section 2, there is a detailed description of granulation techniques with toy examples. In Section 3, we present the experimental part for a kNN classifier. In Section 4, we have additional results for the SVM and Naive Bayes classifier. In Section 5, we write about possible future developments of these techniques, and we conclude the paper in Section 6.

2. Granulation Techniques

Our methods are based on rough inclusions. Introduction to rough inclusions in the framework of rough mereology is available in Polkowski [22,35]; a detailed, extensive discussion can be found in Polkowski [23]. We refer the reader for a very precise theoretical introduction, but, in the paper, we include the details that allow for understanding its content.

In Polkowski’s granulation procedure, we can distinguish three basic steps.

2.0.1. First Step—Granulation

We begin with computation of granules around each training object using a selected method.

2.0.2. Second Step—The Process of Covering

The training decision system is covered by selected granules.

2.0.3. Third Step—Building the Granular Reflections

The granular reflection of original training decision system is derived from the granules selected in step 2.

We start with detailed description of the basic method—see [24].

2.1. Standard Granulation

Let us consider the decision system (U, A, d) , where U is the universe of objects, A the set of conditional attributes, $d \notin A$ is the decision attribute, and r_{gran} granulation radius from the set $\{0, \frac{1}{|A|}, \frac{2}{|A|}, \dots, 1\}$.

The standard rough inclusion μ , for $u, v \in U$ and for selected r_{gran} is defined as

$$\mu(v, u, r_{gran}) \Leftrightarrow \frac{|IND(u, v)|}{|A|} \geq r_{gran}, \text{ where } IND(u, v) = \{a \in A : a(u) = a(v)\}, \quad (1)$$

For each object $u \in U$, and selected r_{gran} , we compute the *standard granule* $g_{r_{gran}}(u)$ as follows:

$$g_{r_{gran}}(u) \text{ is } \{v \in U : \mu(v, u, r_{gran})\}. \quad (2)$$

In the next step, we use a selected strategy to cover the training decision system U by computed granules—the random choice is the simplest among the most effective studied in [30]. All studied methods are available in [30] (pages 105–220).

In addition, in the last step, granular reflection of training decision set is computed with the use of the Majority Voting procedure. The ties are resolved randomly. In the next section, we show the toy example of the method. To present toy examples, we used the same system from Table 1.

Table 1. Exemplary decision system (U, A, d) by J. R. Quinlan [36].

Day	Outlook	Temperature	Humidity	Wind	Play.golf
u_1	Sunny	Hot	High	Weak	No
u_2	Sunny	Hot	High	Strong	No
u_3	Overcast	Hot	High	Weak	Yes
u_4	Rainy	Mild	High	Weak	Yes
u_5	Rainy	Cool	Normal	Weak	Yes
u_6	Rainy	Cool	Normal	Strong	No
u_7	Overcast	Cool	Normal	Strong	Yes
u_8	Sunny	Mild	High	Weak	No
u_9	Sunny	Cool	Normal	Weak	Yes
u_{10}	Rainy	Mild	Normal	Weak	Yes
u_{11}	Sunny	Mild	Normal	Strong	Yes
u_{12}	Overcast	Mild	High	Strong	Yes
u_{13}	Overcast	Hot	Normal	Weak	Yes
u_{14}	Rainy	Mild	High	Strong	No

Toy Example

For a given training decision system from Table 1, the granulation radius $r_{gran} \in \{0, .25, .5, .75, 1\}$, the steps of the standard granulation are as follows.

In case of $r_{gran} = 0$, each single granule is equal U because objects are treated as indiscernible even if they are completely different. In addition, we expected only one object as the granular reflection of the training data.

The second boundary case is $r_{gran} = 1$; each granule contains only their central object or duplicates because the objects are indiscernible.

Now, allow us to show how the standard granulation works for radius $r_{gran} = 0.5$.

Assuming that $g_{r_{gran}}(u_i) = \{u_j \in U_{trn} : \frac{|IND(u_i, u_j)|}{|A|} \geq r_{gran}\}$
 $IND(u_i, u_j) = \{a \in A; a(u_i) = a(u_j)\}$, U_{trn} is the universe of training objects,
 and $|X|$ is the cardinality of set

The sample standard granules with a 0.5 radius, derived from decision systems from Table 1 look as follows,

$$\begin{aligned} g_{0.5}(u_1) &= \{u_1, u_2, u_3, u_4, u_8, u_9, u_{13}\}, g_{0.5}(u_2) = \{u_1, u_2, u_3, u_8, u_{11}, u_{12}, u_{14}\} \\ g_{0.5}(u_3) &= \{u_1, u_2, u_3, u_4, u_8, u_{12}, u_{13}\}, g_{0.5}(u_4) = \{u_1, u_3, u_4, u_5, u_8, u_{10}, u_{12}, u_{14}\} \\ g_{0.5}(u_5) &= \{u_4, u_5, u_6, u_7, u_9, u_{10}, u_{13}\}, g_{0.5}(u_6) = \{u_5, u_6, u_7, u_9, u_{10}, u_{11}, u_{14}\} \\ g_{0.5}(u_7) &= \{u_5, u_6, u_7, u_9, u_{11}, u_{12}, u_{13}\}, g_{0.5}(u_8) = \{u_1, u_2, u_3, u_4, u_8, u_9, u_{10}, u_{11}, u_{12}, u_{14}\} \\ g_{0.5}(u_9) &= \{u_1, u_5, u_6, u_7, u_8, u_9, u_{10}, u_{11}, u_{13}\}, g_{0.5}(u_{10}) = \{u_4, u_5, u_6, u_8, u_9, u_{10}, u_{11}, u_{13}, u_{14}\} \\ g_{0.5}(u_{11}) &= \{u_2, u_6, u_7, u_8, u_9, u_{10}, u_{11}, u_{12}, u_{14}\}, g_{0.5}(u_{12}) = \{u_2, u_3, u_4, u_7, u_8, u_{11}, u_{12}, u_{14}\} \\ g_{0.5}(u_{13}) &= \{u_1, u_3, u_5, u_7, u_9, u_{10}, u_{13}\}, g_{0.5}(u_{14}) = \{u_2, u_4, u_6, u_8, u_{10}, u_{11}, u_{12}, u_{14}\} \end{aligned}$$

The process of granulation can be tuned with help from the triangular part of *granular indiscernibility matrix* $[c_{ij}]_{(i,j=1)|U|}$, where

$$c_{ij} = \begin{cases} 1, & \text{if } \frac{|IND(u_i, u_j)|}{|A|} \geq r_{gran}, i < j \\ 0, & \text{otherwise} \end{cases}$$

This matrix for $r_{gran} = 0.5$ is in Table 2.

Table 2. Triangular indiscernibility matrix for standard granulation ($i < j$), derived from Table 1
 $c_{ij} = 1$, if $\frac{|IND(u_i, u_j)|}{|A|} \geq 0.5$ 0, otherwise.

	u_1	u_2	u_3	u_4	u_5	u_6	u_7	u_8	u_9	u_{10}	u_{11}	u_{12}	u_{13}	u_{14}
u_1	1	1	1	1	0	0	0	1	1	0	0	0	1	0
u_2		1	1	0	0	0	0	1	0	0	1	1	0	1
u_3			1	1	0	0	0	1	0	0	0	1	1	0
u_4				1	1	0	0	1	0	1	0	1	0	1
u_5					1	1	1	0	1	1	0	0	1	0
u_6						1	1	0	1	1	1	0	0	1
u_7							1	0	1	0	1	1	1	0
u_8								1	1	1	1	1	0	1
u_9									1	1	1	0	1	0
u_{10}										1	1	0	1	1
u_{11}											1	1	0	1
u_{12}												1	0	1
u_{13}													1	0
u_{14}														1

Reading the matrix line-wise, we read granules off.

In the next step, we have chosen the random granules to cover the universe of training objects from the Table 1. Our choice is the set.

The U is covered, when, in the set of chosen granules, each object of U appears at least once. The granular reflection of the set Table 1 for the radius 0.5 is in Table 3.

Table 3. Standard granular reflection of the exemplary training system from Table 1, in radius 0.5, 5 attributes, 4 objects; MV is Majority Voting procedure (the most frequent descriptors create a granular reflection).

Day	Outlook	Temperature	Humidity	Wind	Play.golf
$MV(g_{0.5}(u_1))$	Sunny	Hot	High	Weak	Yes
$MV(g_{0.5}(u_4))$	Rainy	Mild	High	Weak	Yes
$MV(g_{0.5}(u_5))$	Rainy	Cool	Normal	Weak	Yes
$MV(g_{0.5}(u_{14}))$	Rainy	Mild	High	Strong	No

Random coverage of training systems is as follows,
 $Cover(U_{trn}) = \{g_{0.5}(u_1), g_{0.5}(u_4), g_{0.5}(u_5), g_{0.5}(u_{14}), \}$

The granular reflection is created by application of majority voting inside selected granules. Ties are resolved randomly.

2.2. Concept Dependent Granulation

A concept-dependent (cd) granule $g_{r_{gran}}^{cd}(u)$ of the radius r_{gran} about u is defined as follows:

$$v \in g_{r_{gran}}^{cd}(u) \text{ if and only if } \mu(v, u, r_{gran}) \text{ and } (d(u) = d(v)) \quad (3)$$

Toy Example

For the decision system from Table 1, we have found concept-dependent granules. For the granulation radius $r_{gran} = 0.25$, the granular concept-dependent indiscernibility matrix (gcdm)—see Table 4—is

$$c_{ij}^{cd} = \begin{cases} 1, & \text{if } \frac{|IND(u_i, u_j)|}{|A|} \geq 0.25, d(u_i) = d(u_j), i < j \\ 0, & \text{otherwise} \end{cases}$$

Table 4. Triangular indiscernibility matrix for concept-dependent granule generation ($i < j$), derived from Table 1.

	u_1	u_2	u_3	u_4	u_5	u_6	u_7	u_8	u_9	u_{10}	u_{11}	u_{12}	u_{13}	u_{14}
u_1	1	1	0	0	0	0	0	1	0	0	0	0	0	1
u_2		1	0	0	0	1	0	1	0	0	0	0	0	1
u_3			1	1	1	0	1	0	1	1	0	1	1	0
u_4				1	1	0	0	0	1	1	1	1	1	0
u_5					1	0	1	0	1	1	1	0	1	0
u_6						1	0	0	0	0	0	0	0	1
u_7							1	0	1	1	1	1	1	0
u_8								1	0	0	0	0	0	1
u_9									1	1	1	0	1	0
u_{10}										1	1	1	1	0
u_{11}											1	1	1	0
u_{12}												1	1	0
u_{13}													1	0
u_{14}														1

hence, the granules in this case are

Assuming that $g_{r_{gran}}(u_i) = \{u_j \in U_{trn} : \frac{|IND(u_i, u_j)|}{|A|} \geq r_{gran}, d(u_i) = d(u_j)\}$

$IND(u_i, u_j) = \{a \in A; a(u_i) = a(u_j)\}$, U_{trn} is the universe of training objects,
and $|X|$ is the cardinality of set

The sample concept-dependent granules with a 0.25 radius, derived from decision systems from Table 1 look as follows,

$$\begin{aligned} g_{0.25}^{cd}(u_1) &= \{u_1, u_2, u_8, u_{14}\}, g_{0.25}^{cd}(u_2) = \{u_1, u_2, u_6, u_8, u_{14}\} \\ g_{0.25}^{cd}(u_3) &= \{u_3, u_4, u_5, u_7, u_9, u_{10}, u_{12}, u_{13}\}, g_{0.25}^{cd}(u_4) = \{u_3, u_4, u_5, u_9, u_{10}, u_{11}, u_{12}, u_{13}\} \\ g_{0.25}^{cd}(u_5) &= \{u_3, u_4, u_5, u_7, u_9, u_{10}, u_{11}, u_{13}\}, g_{0.25}^{cd}(u_6) = \{u_2, u_6, u_{14}\} \\ g_{0.25}^{cd}(u_7) &= \{u_3, u_5, u_7, u_9, u_{10}, u_{11}, u_{12}, u_{13}\}, g_{0.25}^{cd}(u_8) = \{u_1, u_2, u_8, u_{14}\} \\ g_{0.25}^{cd}(u_9) &= \{u_3, u_4, u_5, u_7, u_9, u_{10}, u_{11}, u_{13}\}, g_{0.25}^{cd}(u_{10}) = \{u_3, u_4, u_5, u_7, u_9, u_{10}, u_{11}, u_{12}, u_{13}\} \\ g_{0.25}^{cd}(u_{11}) &= \{u_4, u_5, u_7, u_9, u_{10}, u_{11}, u_{12}, u_{13}\}, g_{0.25}^{cd}(u_{12}) = \{u_3, u_4, u_7, u_{10}, u_{11}, u_{12}, u_{13}\} \\ g_{0.25}^{cd}(u_{13}) &= \{u_3, u_4, u_5, u_7, u_9, u_{10}, u_{11}, u_{12}, u_{13}\}, g_{0.25}^{cd}(u_{14}) = \{u_1, u_2, u_6, u_8, u_{14}\} \end{aligned}$$

Random coverage of training systems is as follows, $Cover(U_{trn}) = \{g_{0.25}^{cd}(u_{13}), g_{0.25}^{cd}(u_{14})\}$

The concept-dependent granular reflection of decision system from Table 1 is in Table 5.

Table 5. Concept-dependent granular reflection of the exemplary training system from Table 1, in radius 0.25, 5 attributes, 2 objects; MV is Majority Voting procedure (the most frequent descriptors create a granular reflection).

Day	Outlook	Temperature	Humidity	Wind	Play.golf
$MV(g_{0.25}^{cd}(u_{13}))$	Overcast	Mild	Normal	Weak	Yes
$MV(g_{0.25}^{cd}(u_{14}))$	Sunny	Hot	High	Strong	No

2.3. Homogeneous Granulation

The homogeneous granules are defined based on standard and concept dependent granules previously defined,

$$g_{r_{gran}}^{homogeneous}(u) = \{v \in U : |g_{r_{gran}}^{cd}(u)| - |g_{r_{gran}}(u)| = 0\}$$

for minimal r_{gran} fulfills the equation

Toy Example

Consider the training decision system from Table 1.

Homogeneous granules for all training objects:

$g_1(u_1) = (u_1)$, $g_{0.75}(u_2) = (u_1, u_2)$, $g_1(u_3) = (u_3)$, $g_1(u_4) = (u_4)$, $g_1(u_5) = (u_5)$, $g_1(u_6) = (u_6)$,
 $g_1(u_7) = (u_7)$, $g_1(u_8) = (u_8)$, $g_{0.75}(u_9) = (u_5, u_9)$, $g_{0.75}(u_{10}) = (u_4, u_5, u_{10})$, $g_{0.75}(u_{11}) = (u_{11})$,
 $g_1(u_{12}) = (u_{12})$, $g_{0.75}(u_{13}) = (u_3, u_{13})$, $g_1(u_{14}) = (u_{14})$.

Randomly selected coverage granules,

$g_{0.75}(u_2) = (u_1, u_2)$, $g_1(u_4) = (u_4)$, $g_1(u_6) = (u_6)$, $g_1(u_7) = (u_7)$, $g_1(u_8) = (u_8)$,
 $g_{0.75}(u_9) = (u_5, u_9)$, $g_{0.75}(u_{10}) = (u_4, u_5, u_{10})$, $g_1(u_{12}) = (u_{12})$, $g_{0.75}(u_{13}) = (u_3, u_{13})$,
 $g_1(u_{14}) = (u_{14})$.

The granular decision system from the above granules is in Table 6.

Table 6. Homogeneous granular decision system formed from covering granules.

Day	Outlook	Temperature	Humidity	Wind	Play Golf
$MV(g_{0.75}(u_2))$	Sunny	Hot	High	Weak	No
$MV(g_1(u_4))$	Rainy	Mild	High	Weak	Yes
$MV(g_1(u_6))$	Rainy	Cool	Normal	Strong	No
$MV(g_1(u_7))$	Overcast	Cool	Normal	Strong	Yes
$MV(g_1(u_8))$	Sunny	Mild	High	Weak	No
$MV(g_{0.75}(u_9))$	Rainy	Cool	Normal	Weak	Yes
$MV(g_{0.75}(u_{10}))$	Rainy	Mild	Normal	Weak	Yes
$MV(g_1(u_{12}))$	Overcast	Mild	High	Strong	Yes
$MV(g_{0.75}(u_{13}))$	Overcast	Hot	High	Weak	Yes
$MV(g_1(u_{14}))$	Rainy	Mild	High	Strong	No

2.4. Layered Granulation

Layered granulation leads to a sequence of granular reflections of decreasing sizes, which stabilizes after a finite number of steps; usually, about five steps are sufficient. Another development that may be stressed here is the heuristic rule for finding the optimal granulation radius giving the highest accuracy.

the optimal granulation radius is located around the value which yields the maximal decrease in size of the granular reflection between the first and the second granulation layers—see [30].

Toy Example

Exemplary multiple granulation of Quinlan's data set [36], see Table 1, for the granulation radius of 0.5 and layers $l_0, l_1, \dots$ runs as follows.

For the decision system from Table 1, granules in the first layer are ($r_{gran} = 0.5$):

$g_{0.5,l_1}^{cd}(u_1) = \{u_1, u_2, u_8\}$, $g_{0.5,l_1}^{cd}(u_2) = \{u_1, u_2, u_8, u_{14}\}$, $g_{0.5,l_1}^{cd}(u_3) = \{u_3, u_4, u_{12}, u_{13}\}$,
 $g_{0.5,l_1}^{cd}(u_4) = \{u_3, u_4, u_5, u_{10}, u_{12}\}$, $g_{0.5,l_1}^{cd}(u_5) = \{u_4, u_5, u_7, u_9, u_{10}, u_{13}\}$, $g_{0.5,l_1}^{cd}(u_6) = \{u_6, u_{14}\}$,
 $g_{0.5,l_1}^{cd}(u_7) = \{u_5, u_7, u_9, u_{11}, u_{12}, u_{13}\}$, $g_{0.5,l_1}^{cd}(u_8) = \{u_1, u_2, u_8, u_{14}\}$,

$$\begin{aligned} g_{0.5,l_1}^{cd}(u_9) &= \{u_5, u_7, u_9, u_{10}, u_{11}, u_{13}\}, g_{0.5,l_1}^{cd}(u_{10}) = \{u_4, u_5, u_9, u_{10}, u_{11}, u_{13}\}, \\ g_{0.5,l_1}^{cd}(u_{11}) &= \{u_7, u_9, u_{10}, u_{11}, u_{12}\}, g_{0.5,l_1}^{cd}(u_{12}) = \{u_3, u_4, u_7, u_{11}, u_{12}\}, \\ g_{0.5,l_1}^{cd}(u_{13}) &= \{u_3, u_5, u_7, u_9, u_{10}, u_{13}\}, g_{0.5,l_1}^{cd}(u_{14}) = \{u_2, u_6, u_8, u_{14}\}. \end{aligned}$$

Covering process of U_{l_0} with usage of order-preserving strategy yields us the covering:

$$\begin{aligned} U_{l_0,Cover} &\leftarrow \emptyset, \\ \text{Step1 } g_{0.5,l_1}^{cd}(u_1) &\rightarrow U_{l_0,Cover}, U_{l_0,Cover} = \{u_1, u_2, u_8\}, \\ \text{Step2 } g_{0.5,l_1}^{cd}(u_2) &\rightarrow U_{l_0,Cover}, U_{l_0,Cover} = \{u_1, u_2, u_8, u_{14}\}, \\ \text{Step3 } g_{0.5,l_1}^{cd}(u_3) &\rightarrow U_{l_0,Cover}, U_{l_0,Cover} = \{u_1, u_2, u_3, u_4, u_8, u_{12}, u_{13}, u_{14}\}, \\ \text{Step4 } g_{0.5,l_1}^{cd}(u_4) &\rightarrow U_{l_0,Cover}, U_{l_0,Cover} = \{u_1, u_2, u_3, u_4, u_5, u_8, u_{10}, u_{12}, u_{13}, u_{14}\}, \\ \text{Step5 } g_{0.5,l_1}^{cd}(u_5) &\rightarrow U_{l_0,Cover}, U_{l_0,Cover} = \{u_1, u_2, u_3, u_4, u_5, u_7, u_8, u_9, u_{10}, u_{12}, u_{13}, u_{14}\}, \\ \text{Step6 } g_{0.5,l_1}^{cd}(u_6) &\rightarrow U_{l_0,Cover}, U_{l_0,Cover} = \{u_1, u_2, u_3, u_4, u_5, u_6, u_7, u_8, u_9, u_{10}, u_{12}, u_{13}, u_{14}\}, \\ \text{Step7 } g_{0.5,l_1}^{cd}(u_7) &\rightarrow U_{l_0,Cover}, U_{l_0,Cover} = U_{l_0}. \end{aligned}$$

The granular reflection of (U_{l_0}, A, d) based on granules from $U_{l_0,Cover}$, with use of Majority Voting, where ties are resolved according to the ordering of granules are shown in Table 7.

Table 7. The decision system (U_{l_1}, A, d) .

Day	Outlook	Temperature	Humidity	Wind	Play Golf
$MV(g_{0.5,l_1}^{cd}(u_1))$	Sunny	Hot	High	Weak	No
$MV(g_{0.5,l_1}^{cd}(u_2))$	Sunny	Hot	High	Weak	No
$MV(g_{0.5,l_1}^{cd}(u_3))$	Overcast	Mild	High	Weak	Yes
$MV(g_{0.5,l_1}^{cd}(u_4))$	Rainy	Mild	High	Weak	Yes
$MV(g_{0.5,l_1}^{cd}(u_5))$	Rainy	Cool	Normal	Weak	Yes
$MV(g_{0.5,l_1}^{cd}(u_6))$	Rainy	Cool	Normal	Strong	No
$MV(g_{0.5,l_1}^{cd}(u_7))$	Overcast	Cool	Normal	Strong	Yes

Exemplary granular reflection formation based on Majority Voting looks as follows. In case, e.g., of the granule $g_{0.5,l_1}^{cd}(u_1)$, we have

$$MV(g_{0.5,l_1}^{cd}(u_1)) = \left\{ \begin{array}{l} \underline{\text{Sunny}} \underline{\text{Hot}} \underline{\text{High}} \underline{\text{Weak}} \\ \underline{\text{Sunny}} \underline{\text{Hot}} \underline{\text{High}} \underline{\text{Strong}} \\ \underline{\text{Sunny}} \underline{\text{Mild}} \underline{\text{High}} \underline{\text{Weak}} \end{array} \right\} = \underline{\text{Sunny}} \underline{\text{Hot}} \underline{\text{High}} \underline{\text{Weak}}$$

Treating all other granules in the same way, we obtain the granular reflection (U_{l_1}, A, d) shown in Table 7.

Granulation performed in the same manner with the granular reflection (U_{l_1}, A, d) from Table 7 yields the granule set in the second layer.

$$\begin{aligned} g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_1))) &= \{MV(g_{0.5,l_1}^{cd}(u_1)), MV(g_{0.5,l_1}^{cd}(u_2))\} \\ g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_2))) &= \{MV(g_{0.5,l_1}^{cd}(u_1)), MV(g_{0.5,l_1}^{cd}(u_2))\} \\ g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_3))) &= \{MV(g_{0.5,l_1}^{cd}(u_3)), MV(g_{0.5,l_1}^{cd}(u_4))\} \end{aligned}$$

$$\begin{aligned} g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_4))) &= \{MV(g_{0.5,l_1}^{cd}(u_3)), MV(g_{0.5,l_1}^{cd}(u_4)), MV(g_{0.5,l_1}^{cd}(u_5))\} \\ g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_5))) &= \{MV(g_{0.5,l_1}^{cd}(u_4)), MV(g_{0.5,l_1}^{cd}(u_5)), MV(g_{0.5,l_1}^{cd}(u_7))\} \\ g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_6))) &= \{MV(g_{0.5,l_1}^{cd}(u_6))\} \\ g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_7))) &= \{MV(g_{0.5,l_1}^{cd}(u_5)), MV(g_{0.5,l_1}^{cd}(u_7))\} \end{aligned}$$

The covering process of $U_{l_1,Cover}$ runs in the following steps:

Step1 $g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_1))) \rightarrow U_{l_1,Cover}$, Step2 $g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_2))) \not\rightarrow U_{l_1,Cover}$,
Step3 $g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_3))) \rightarrow U_{l_1,Cover}$, Step4 $g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_4))) \rightarrow U_{l_1,Cover}$,
Step5 $g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_5))) \rightarrow U_{l_1,Cover}$, Step6 $g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_6))) \rightarrow U_{l_1,Cover}$,
 $U_{l_1,Cover} = U_{l_1}$

Applying Majority Voting to granules in U_{l_1} , we obtain the second granular reflection shown in Table 8.

Table 8. The decision system (U_{l_2}, A, d) , $temp_1 = MV(g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_1))))$, $temp_2 = MV(g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_3))))$, $temp_3 = MV(g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_4))))$, $temp_4 = MV(g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_5))))$, $temp_5 = MV(g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_6))))$.

Day	Outlook	Temperature	Humidity	Wind	Play Golf
$temp_1$	Sunny	Hot	High	Weak	No
$temp_2$	Overcast	Mild	High	Weak	Yes
$temp_3$	Rainy	Mild	High	Weak	Yes
$temp_4$	Rainy	Cool	Normal	Weak	Yes
$temp_5$	Rainy	Cool	Normal	Strong	No

The third layer of granulation based on system (U_{l_2}, A, d) from Table 8 is as follows:

$$\begin{aligned} g_{0.5,l_3}^{cd}(MV(g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_1))))) &= \{MV(g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_1)))))\} \\ g_{0.5,l_3}^{cd}(MV(g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_3))))) &= \\ \{MV(g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_3)))), MV(g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_4)))))\} \\ g_{0.5,l_3}^{cd}(MV(g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_4))))) &= \\ = \{MV(g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_3)))), MV(g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_4)))), MV(g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_5)))))\} \\ g_{0.5,l_3}^{cd}(MV(g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_5))))) &= \\ \{MV(g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_4)))), MV(g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_5)))))\} \\ g_{0.5,l_3}^{cd}(MV(g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_6))))) &= \{MV(g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_6)))))\} \end{aligned}$$

Covering process for the third layer is as follows:

Step1 $g_{0.5,l_3}^{cd}(MV(g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_1))))) \rightarrow U_{l_2,Cover}$,
Step2 $g_{0.5,l_3}^{cd}(MV(g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_3))))) \rightarrow U_{l_2,Cover}$,
Step3 $g_{0.5,l_3}^{cd}(MV(g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_4))))) \rightarrow U_{l_2,Cover}$,
Step4 $g_{0.5,l_3}^{cd}(MV(g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_5))))) \not\rightarrow U_{l_2,Cover}$,
Step5 $g_{0.5,l_3}^{cd}(MV(g_{0.5,l_2}^{cd}(MV(g_{0.5,l_1}^{cd}(u_6))))) \rightarrow U_{l_2,Cover}$, $U_{l_2,Cover} = U_{l_2}$

Using Majority voting, we get the third layer of granular reflections shown in Table 9.

Table 9. The decision system (U_{I_3}, A, d) , $temp_1 = MV(g_{0.5,I_3}^{cd}(MV(g_{0.5,I_2}^{cd}(MV(g_{0.5,I_1}^{cd}(u_1))))))$, $temp_2 = MV(g_{0.5,I_3}^{cd}(MV(g_{0.5,I_2}^{cd}(MV(g_{0.5,I_1}^{cd}(u_3))))))$, $temp_3 = MV(g_{0.5,I_3}^{cd}(MV(g_{0.5,I_2}^{cd}(MV(g_{0.5,I_1}^{cd}(u_4))))))$, $temp_4 = MV(g_{0.5,I_3}^{cd}(MV(g_{0.5,I_2}^{cd}(MV(g_{0.5,I_1}^{cd}(u_6))))))$.

Day	Outlook	Temperature	Humidity	Wind	Play Golf
$temp_1$	Sunny	Hot	High	Weak	No
$temp_2$	Overcast	Mild	High	Weak	Yes
$temp_3$	Rainy	Mild	High	Weak	Yes
$temp_4$	Rainy	Cool	Normal	Strong	No

2.5. Epsilon Variants

These methods are designed for numerical data; we can use, for instance, ε -normalized Hamming metric, which, for given ε , is defined as

$$d_{H,\varepsilon}(u, v) = |\{a \in A : \frac{abs(a(u) - a(v))}{max_a - min_a} > \varepsilon\}|, \quad (4)$$

where abs is absolute value,

The methods work analogously to variants for symbolic data; thus, we show only exemplary definition without toy examples.

2.5.1. ε -Modification of the Standard Rough Inclusion

Given a parameter ε valued in the unit interval $[0, 1]$, we define the set

$$IND_\varepsilon(u, v) = \{a \in A : dist(a(u), a(v)) \leq \varepsilon\}, \quad (5)$$

and we set

$$\mu_\varepsilon(v, u, r) \Leftrightarrow \frac{|IND_\varepsilon(u, v)|}{|A|} \geq r \quad (6)$$

Epsilon variant of homogeneous granulation can be defined as follows.

2.6. Epsilon Homogeneous Granulation

The method is defined in the following way:

$$g_{r_u}^{\varepsilon, homogeneous} = \{v \in U : |g_{r_u}^{\varepsilon-cd}| - |g_{r_u}^{\varepsilon}| == 0\}, \text{ for minimal } r_u \text{ fulfills the equation}$$

$$\text{where } g_{r_u}^{\varepsilon, cd}(u) = \{v \in U : \frac{|IND_\varepsilon(u, v)|}{|A|} \leq r_u \text{ AND } d(u) == d(v)\}$$

$$\text{and } g_{r_u}^{\varepsilon}(u) = \{v \in U : \frac{|IND_\varepsilon(u, v)|}{|A|} \leq r_u\}, r_u = \{\frac{0}{|A|}, \frac{1}{|A|}, \dots, \frac{|A|}{|A|}\}$$

$$IND_\varepsilon(u, v) = \{a \in A : \frac{abs(a(u) - a(v))}{max_a - min_a} \leq \varepsilon\}$$

where max_a , min_a are the maximal and minimal attribute values for $a \in A$ in the original data set.

3. A Sample of the Experimental Work Results

In this section, we show the exemplary results for our selected techniques, to show its effectiveness in the context of reducing training data size. For the sake of simplicity, we have chosen the k-NN classifier as a base. We carried out experiments on selected data from the UCI repository [37]—see Table 10. In Tables 11–20 and Figure 1, we have the results for Cross Validation 5 method.

Let us move on to the discussion of selected detailed results, starting from description of the results for the Australian Credit data set. The result for Standard (SG) and Concept-dependent granulation (CDG) is in Table 11, where, in case of SG for radius 0.5, we have reduction in training size of around 90 percent preserving classification accuracy in the range of 84.7 percent. For the CDG variant, we have reduction in training size of about 99.5 percent for radius 0.071, where the exhaustive rule set is reduced in 99.9 percent and accuracy of classification is around 77 percent. The results are comparable, but the concept-dependent variant shows a more stable classification as the radius increases. In case of Homogeneous granulation, see Table 15, we have accuracy equal to 0.835 with a 48 percent reduction of training size. The sample of results for exemplary epsilon variant— ϵ Homogeneous Granulation—is in Table 20, where we have reduction in training size about 50 percent, with accuracy of 0.842. The layered granulation process is visible in Table 16, where the basic method is concept-dependent granulation and the result is similar to a single concept-dependent variant. In the case of Car data set, see Table 12, the concept-dependent variant works best giving accuracy of 0.864, with a reduction in training size of around 73 percent. For a Hepatitis data set, concept-dependent also works best, for radius 0.474, the accuracy is equal to 0.875, with a 90 percent reduction in training size. In addition, finally, the spectacular result is obtained for Heart Disease data set, where with 99 percent reduction in training size, we have obtained for concept-dependent and standard granulation the accuracy 0.8. The results for homogeneous variants are shown in Tables 15 and 20. The best result we have achieved on the tested data are a reduction of 62 percent in the number of objects with full classification efficiency. Allow us to summarize the results obtained in this section. The internal knowledge from the original training decision systems—measured by ability for classification—seems to be preserved in each mentioned case (the accuracy of classification is fully comparable with nil case, without reduction). Both techniques, standard granulation and concept-dependent, prove to be comparable. In the concept-dependent variant, we observe a higher classification stability with an increasing radius. Another advantage of the concept-dependent variant is the creation of granular reflection, which from the smallest radii contain patterns from all decision-making classes. The multiple variant does not produce spectacular results, but, according to our previous research, see [30]—it allows us to look for the optimal granulation radii. Our research shows that the radius for which the reduction of objects between the first and second layer is greatest is close to the optimal one in most tested systems. In this way, the optimum granulation radius can be estimated without classification tests. The last group of tested techniques are recently discovered homogeneous methods, which work dynamically on every data and do not require estimation of optimal parameters. It is obvious that the effectiveness of our methods depends to a large extent on the data under investigation.

We do not plan to present an overview of the effectiveness of the whole range of classification techniques because our aim was to present an example of the effectiveness of approximation methods for decision-making systems. Let us move on to presenting additional test results for selected previously used classifiers.

Table 10. Exemplary decision systems from UCI Machine Learning Repository. Australian credit, Car Evaluation, Heartdisease, and Hepatitis were used in the comparison of standard and concept-dependent granulation with a kNN Classifier. Comparing homogeneous variants with a kNN Classifier, we did not use the car system in the epsilon variant because it is symbolic. We used all four systems to present the effectiveness with the Classifier. To present the effectiveness with the SVM classifier, we used a Wisconsin Diagnostic Breast Cancer system [37].

Name	Attr No.	Obj No.	Class No.
Australian-credit	15	690	2
Car Evaluation	7	1728	4
Heartdisease	14	270	2
Hepatitis	20	155	2
Wisconsin Diagnostic Breast Cancer	32	569	2

Table 11. Exemplary result for Standard vs. Concept-Dependent Granulation—5 times Cross Validation 5; Australian Credit data set; r_{gran} = Granulation radius, AccSG = Accuracy of classification for Standard Granulation, AccCDG—Accuracy for Concept-Dependent Granulation, SizeSG = Granular decision system size for Standard Granulation, SizeCDG = Granular decision system size for Concept-Dependent Granulation.

r_{gran}	AccSG	SizeSG	AccCDG	SizeCDG
0.071428	0.444928	2.36	0.773	2.64
0.142857	0.444928	5.12	0.779	3.92
0.214286	0.821739	4.76	0.786	5.36
0.285714	0.84058	4.8	0.804	9.12
0.357143	0.768116	9.4	0.813	16.12
0.428571	0.775362	24.2	0.828	32.44
0.5	0.847826	51.2	0.845	71.64
0.571429	0.818841	133.4	0.838	157.96
0.642857	0.833333	297	0.845	318.96
0.714286	0.811594	455.2	0.854	468.16
0.785714	0.855072	533.2	0.858	535.84
0.857143	0.826087	546.4	0.861	547.2
0.928571	0.826087	547.8	0.863	548.8
1	0.826087	552	0.861	552

Table 12. Exemplary result for Standard vs. Concept-Dependent Granulation—5 times Cross Validation 5; Car Evaluation data set; r_{gran} = Granulation radius, AccSG = Accuracy of classification for Standard Granulation, AccCDG—Accuracy for Concept-Dependent Granulation, SizeSG = Granular decision system size for Standard Granulation, SizeCDG = Granular decision system size for Concept-Dependent Granulation.

r_{gran}	AccSG	SizeSG	AccCDG	SizeCDG
0.167	0.388988	8.08	0.396	8.32
0.333	0.456468	17.16	0.539	16.96
0.500	0.495127	38.84	0.681	38.2
0.667	0.546064	106.24	0.804	107.04
0.833	0.611924	368.76	0.864	371.64
1.000	0.359964	1382.4	0.944	1382.4

Table 13. Exemplary result for Standard vs. Concept-Dependent Granulation—5 times Cross Validation 5; Heart Disease data set; r_{gran} = Granulation radius, AccSG = Accuracy of classification for Standard Granulation, AccCDG—Accuracy for Concept-Dependent Granulation, SizeSG = Granular decision system size for Standard Granulation, SizeCDG = Granular decision system size for Concept-Dependent Granulation.

r_{gran}	AccSG	SizeSG	AccCDG	SizeCDG
0.0769231	0.555556	1.2	0.804	2.2
0.153846	0.444444	2.4	0.798	2.96
0.230769	0.555556	3.2	0.799	5.12
0.307692	0.777778	6.2	0.803	8.84
0.384615	0.759259	11	<u>0.819</u>	<u>16.76</u>
0.461538	<u>0.833333</u>	<u>27</u>	0.819	34.08
0.538462	0.814815	58.4	0.824	71.68
0.615385	0.814815	118	0.817	126.56
0.692308	0.796296	177.8	0.827	180.92
0.769231	0.814815	209.8	0.822	210
0.846154	0.814815	216	0.826	216
0.923077	0.814815	216	0.826	216
1	0.814815	216	0.826	216

Table 14. Exemplary result for Standard vs. Concept-Dependent Granulation—5 times Cross Validation 5; Hepatitis data set; r_{gran} = Granulation radius, AccSG = Accuracy of classification for Standard Granulation, AccCDG—Accuracy for Concept-Dependent Granulation, SizeSG = Granular decision system size for Standard Granulation, SizeCDG = Granular decision system size for Concept-Dependent Granulation.

r_{gran}	AccSG	SizeSG	AccCDG	SizeCDG
0.053	0.807742	2	0.803	2
0.105	0.807742	2	0.803	2
0.158	0.807742	2	0.803	2.04
0.211	0.807742	2.12	0.804	2.28
0.263	0.809032	2.72	0.806	2.68
0.316	0.811612	3.48	0.814	3.68
0.368	0.812902	5.2	0.83	5.24
0.421	0.832258	7.16	<u>0.854</u>	<u>7.56</u>
0.474	<u>0.847742</u>	<u>11.28</u>	<u>0.875</u>	11.6
0.526	0.815484	18.56	0.876	18.88
0.579	0.812902	29.8	0.881	31.08
0.632	0.832259	46.36	0.893	46.4
0.684	0.83871	69.6	0.877	69.64
0.737	0.83871	90.08	0.888	89.64
0.789	0.854194	109.68	0.892	109.8
0.842	0.854194	116.96	0.892	116.8
0.895	0.854194	121	0.895	121
0.947	0.854194	121.96	0.895	122
1.000	0.854194	124	0.895	124

Table 15. Exemplary result for Homogeneous Granulation—5 times Cross Validation 5; k – NN classifier; D_1 = Australian-credit, D_2 = Car Evaluation, D_5 = Heartdisease, D_6 = Hepatitis data set; Acc = average accuracy, GS = granular decision system size, TRN_size = training set size, TRN_red = reduction in object number in training size, Radii_range = spectrum of radii.

Data Set	Acc	GS	TRN_size	TRN_red	Radii_range
D_1	0.835	286.52	552	48.1%	$r \geq 0.5$
D_2	0.797	728.5	1382	47.3%	$r \geq 0.667$
D_5	0.833	120.5	216	44.2%	$r \geq 0.461$
D_6	0.88	46.16	124	62.8%	$r \geq 0.579$

Table 16. CV-5; Result of experiments for multi-layer c-d granulation with use of kNN classifier; data set Australian credit; r_{gran} = Granulation radius, Acc = Average accuracy for the considered layer, GranSize The mean size of granular decision system for the considered layer.

r_{gran}	Layer1		Layer2		Layer3		Layer4	
	Acc	GranSize	Acc	GranSize	Acc	GranSize	Acc	GranSize
0	0.768	2	0.768	2	0.768	2	0.768	2
0.071	0.772	2	0.772	2	0.772	2	0.772	2
0.143	0.696	2.6	0.774	2	0.774	2	0.774	2
0.214	0.781	5.6	0.775	2	0.775	2	0.775	2
0.286	0.8	6.8	0.797	2	0.797	2	0.797	2
0.357	0.813	16.4	0.78	2	0.78	2	0.78	2
0.429	0.838	29.6	0.704	3.6	0.67	2.2	0.67	2.2
0.5	0.843	68.6	0.729	15.4	0.37	7.4	0.37	7.4
0.571	0.851	154.8	0.799	70.6	0.69	47.4	0.628	43.2
0.643	0.854	313.2	0.841	245.6	0.806	228.8	0.781	225.6
0.714	0.852	468.2	0.854	444.8	0.855	440	0.857	438.6
0.786	0.858	535.6	0.858	535.4	0.858	535.4	0.858	535.4
0.857	0.854	547.4	0.854	547.4	0.854	547.4	0.854	547.4
0.929	0.864	548.8	0.864	548.8	0.864	548.8	0.864	548.8
1	0.855	552	0.855	552	0.855	552	0.855	552

Table 17. CV-5; Result of experiments for multi-layer c-d granulation with use of kNN classifier; data set Car evaluation; r_{gran} = Granulation radius, Acc = Average accuracy for the considered layer, TRNsize The mean size of granular decision system for the considered layer.

r_{gran}	Layer1		Layer2		Layer3		Layer4	
	Acc	GranSize	Acc	GranSize	Acc	GranSize	Acc	GranSize
0	0.315	4	0.315	4	0.315	4	0.315	4
0.167	0.395	8.6	0.296	4	0.296	4	0.296	4
0.333	0.484	16.4	0.351	6.2	0.326	4.6	0.326	4.6
0.5	0.668	44	0.477	16.2	0.374	9.4	0.296	7
0.667	0.811	102.8	0.723	47.4	0.632	29.8	0.601	25.4
0.833	0.865	370	0.841	199.8	0.832	147.2	0.833	137
1	0.944	1382.4	0.944	1382.4	0.944	1382.4	0.944	1382.4

Table 18. CV-5; Result of experiments for multi-layer c-d granulation with use of kNN classifier; data set Heart disease; r_{gran} = Granulation radius, Acc = Average accuracy for the considered layer, TRNsize The mean size of granular decision system for the considered layer.

r_{gran}	Layer1		Layer2		Layer3		Layer4	
	Acc	GranSize	Acc	GranSize	Acc	GranSize	Acc	GranSize
0	0.811	2	0.811	2	0.811	2	0.811	2
0.077	0.793	2	0.793	2	0.793	2	0.793	2
0.154	0.811	3	0.811	2	0.811	2	0.811	2
0.231	0.796	3.2	0.759	2	0.759	2	0.759	2
0.308	0.804	6.8	0.781	2	0.781	2	0.781	2
0.385	0.807	17	0.763	2.2	0.763	2	0.763	2
0.462	<u>0.833</u>	<u>35.6</u>	0.737	6.6	0.681	4	0.693	3.8
0.538	0.83	69.8	0.778	34.2	0.678	24.6	0.63	23
0.615	0.807	129.4	0.781	100.8	0.667	92.6	0.652	91.4
0.692	0.807	180.2	0.8	172.6	0.804	171	0.804	170.8
0.769	0.83	211	<u>0.826</u>	<u>210.2</u>	<u>0.826</u>	<u>210</u>	0.826	210
0.846	0.83	216	0.83	216	0.83	216	0.83	216
0.923	0.833	216	0.833	216	0.833	216	0.833	216
1	0.837	216	0.837	216	0.837	216	0.837	216

Table 19. CV-5; Result of experiments for multi-layer c-d granulation with use of kNN classifier; data set Hepatitis; r_{gran} = Granulation radius, Acc = Average accuracy for the considered layer, TRNsize The mean size of granular decision system for the considered layer.

r_{gran}	Layer1		Layer2		Layer3		Layer4	
	Acc	GranSize	Acc	GranSize	Acc	GranSize	Acc	GranSize
0	0.8	2	0.8	2	0.8	2	0.8	2
0.053	0.806	2	0.806	2	0.806	2	0.806	2
0.105	0.813	2	0.813	2	0.813	2	0.813	2
0.158	0.826	2	0.826	2	0.826	2	0.826	2
0.211	0.826	2	0.826	2	0.826	2	0.826	2
0.263	0.813	3	0.813	2	0.813	2	0.813	2
0.316	0.806	2.8	0.806	2	0.806	2	0.806	2
0.368	0.819	7.2	0.819	2	0.819	2	0.819	2
0.421	0.832	6.8	0.806	2	0.806	2	0.806	2
0.474	<u>0.871</u>	<u>12.4</u>	0.8	2.2	0.8	2	0.8	2
0.526	0.877	20.2	0.794	4.8	0.703	2.8	0.703	2.8
0.579	0.865	32.2	0.658	10.6	0.652	7.4	0.652	7.4
0.632	0.884	49.6	0.806	27	0.703	22.4	0.69	21.8
0.684	0.89	67	<u>0.865</u>	<u>54.6</u>	<u>0.865</u>	<u>52.6</u>	0.845	52.2
0.737	0.89	88.4	0.877	79	0.871	77.8	0.871	77.6
0.789	0.91	108.6	0.91	104.4	0.91	103.8	0.91	103.8
0.842	0.903	117.4	0.903	114.6	0.903	114.6	0.903	114.6
0.895	0.89	121	0.89	120.2	0.89	120.2	0.89	120.2
0.947	0.916	122	0.916	122	0.916	122	0.916	122
1	0.89	124	0.89	124	0.89	124	0.89	124

Table 20. Exemplary result for Epsilon Homogeneous Granulation (ϵ – HGS)—5 times Cross Validation 5; k – NN classifier; D_1 = Australian-credit, D_3 = Heartdisease, D_4 = Hepatitis data set; Acc = average accuracy of classification, HGS_size = granular decision system size, TRN_size = training set size, HGS_TRN_red = reduction in object number in training set, HG_r_range = spectrum of radii.

Results	D_1	D_3	D_4
Acc	0.842	0.831	0.87
HGS_size	274.52	109.4	46.2
TRN_size	552	216	124
HG_TRN_red	50.3%	49.4%	62.7%
HG_r_range	$r_u \geq 0.65$	$r_u \geq 0.615$	$r_u \geq 0.579$

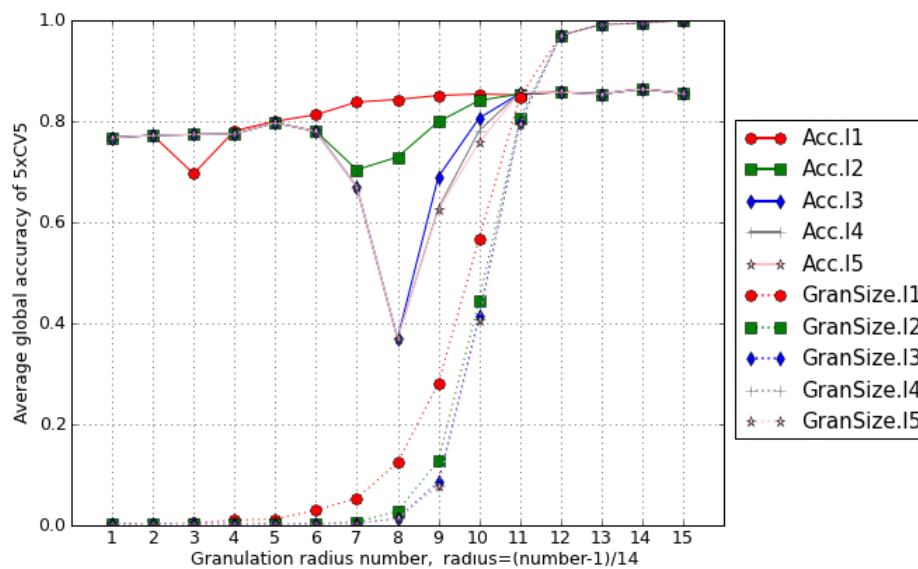


Figure 1. Visualization of results for Australian credit.

4. Application of Selected Other Classifiers on Granular Data

In our previous research, we checked the performance of the tens of classifiers; each variant examined matched well with the granular data. Some of the most interesting results were obtained for the Naive Bayes classifier (see the results in Chapter 7 of [30]), the SVM technique [38], and Deep Learning [39]. Examples of results are presented in this section.

In Figure 2, we have the accuracy of the classification of the granular data using the SVM method with an RBF kernel. We use the ϵ concept-dependent granulation—see Section 2.5. It is the result for Wisconsin Diagnostic Breast Cancer data set (see [37]) 569 objects and 32 attributes. Analyzing Figures 2 and 3, we see that the level of accuracy of the classification is reasonable with a considerable percentage of the size reduction of granular systems.

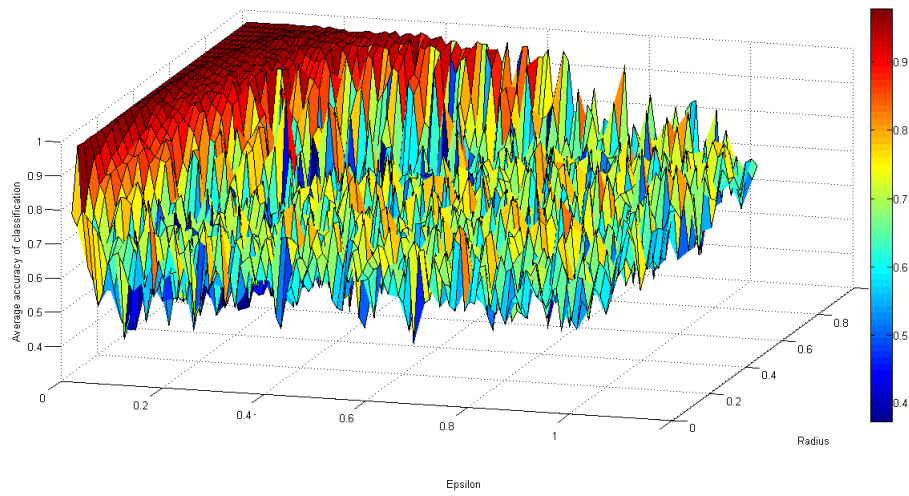


Figure 2. Results of classification accuracy for SVM with RBF kernel, $5 \times CV5$ test, ε concept-dependent granulation; Wisconsin Diagnostic Breast Cancer data set; Epsilon = is descriptors indiscernibility ratio, Radius = granulation radius.

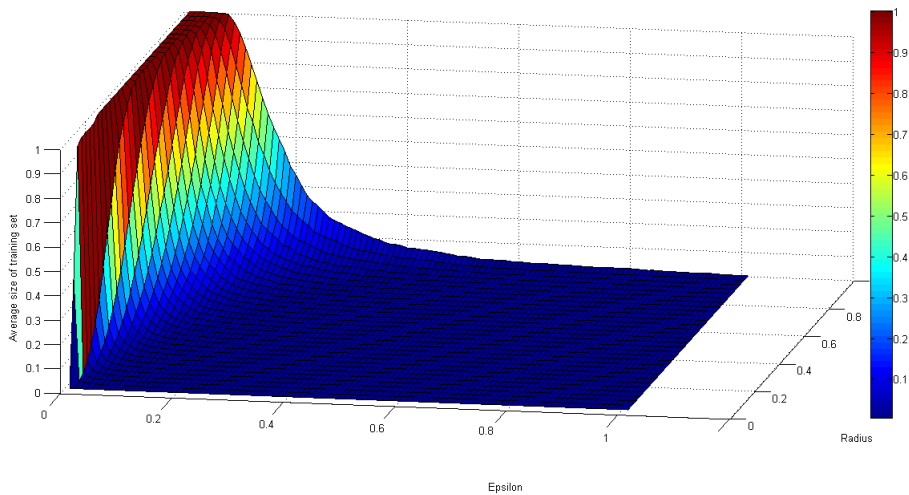


Figure 3. Percentage size of granulated data, $5 \times CV5$ test, ε concept-dependent granulation; Wisconsin Diagnostic Breast Cancer data set; Epsilon = is descriptors indiscernibility ratio, Radius = granulation radius.

Considering four variants of classification for the Naive Bayes classifier (for which the parameters determining the classification are as follows):

1. $Param_{d=d_i}^{V1} = \sum_{m=1}^n P(b_m = a_m(v) | d = d_i).$
2. $Param_{d=d_i}^{V2} = P(d = d_i) * \sum_{m=1}^n P(b_m = a_m(v) | d = d_i).$
3. $Param_{d=d_i}^{V3} = \prod_{m=1}^n P(b_m = a_m(v) | d = d_i).$
4. $Param_{d=d_i}^{V4} = P(d = d_i) * \prod_{m=1}^n P(b_m = a_m(v) | d = d_i).$

The results showing the effectiveness of the Naive Bayes classifier can be found in Tables 21–24 (the details can be found in [30]). The most spectacular approximation is for the 0.428571 radius, where, with an Australian credit data set, accuracy of classification is 0.852, and the average number of objects is reduced by about 94 percent.

Table 21. $5 \times CV-5$; The result of experiments for four variants of the Naive Bayes classifier; data set Australian credit; concept dependent granulation; r_{gran} = Granulation radius; nil = result for data without missing values; Acc = Accuracy of classification; GranSize = The size of data set after granulation in the fixed r .

r_{gran}	Acc				GranSize			
	V1	V2	V3	V4	V1	V2	V3	V4
0.0714286	0.789	0.703	0.813	0.788	2.32	2.32	2.52	2.4
0.142857	0.788	0.682	0.812	0.76	3.4	3.84	3.52	3.76
0.214286	0.789	0.707	0.79	0.759	5.2	5.4	5.16	5.32
0.285714	0.806	0.738	0.656	0.628	8.8	9.08	8.56	9.36
0.357143	0.827	0.727	0.692	0.707	16.64	15.16	16.32	16.12
0.428571	0.853	0.772	0.717	0.745	32.84	30.72	32.28	31.28
0.5	0.85	0.814	0.749	0.732	71.56	70.76	71	69.68
0.571429	0.852	0.77	0.725	0.721	157	158.36	157.16	155.92
0.642857	0.857	0.764	0.734	0.732	319	320.4	317.8	318.08
0.714286	0.843	0.83	0.732	0.737	468.56	468.44	467.88	468.28
0.785714	0.843	0.813	0.732	0.739	536.28	536.24	536	536.04
0.857143	0.843	0.799	0.73	0.739	547.36	547.16	547.16	547.28
0.928571	0.843	0.8	0.73	0.739	548.92	548.76	548.72	548.8
1	0.843	0.799	0.729	0.739	552	552	552	552

Table 22. $5 \times CV-5$; The result of experiments for four variants of the Naive Bayes classifier; data set Car evaluation; concept dependent granulation; r_{gran} = Granulation radius; nil = result for data without missing values; Acc = Accuracy of classification; GranSize = The size of data set after granulation in the fixed r .

r_{gran}	Acc				GranSize			
	V1	V2	V3	V4	V1	V2	V3	V4
0.166667	0.315	0.653	0.092	0.369	8.12	8.48	7.72	8.52
0.333333	0.357	0.723	0.044	0.118	17.96	17.44	17.36	17.4
0.5	0.383	0.715	0.077	0.32	38.96	38.52	36.72	38.84
0.666667	0.403	0.7	0.108	0.382	105.28	106.12	106.84	107.32
0.833333	0.436	0.7	0.06	0.328	368.88	369.08	369.28	374.68
1	0.451	0.7	0.052	0.196	1382.4	1382.4	1382.4	1382.4

Table 23. $5 \times CV-5$; The result of experiments for four variants of the Naive Bayes classifier; data set Heart disease; Concept dependent granulation; r_{gran} = Granulation radius; nil = result for data without missing values; Acc = Accuracy of classification; GranSize = The size of data set after granulation in the fixed r .

r_{gran}	Acc				GranSize			
	V1	V2	V3	V4	V1	V2	V3	V4
0.0769231	0.801	0.774	0.785	0.793	2.04	2.2	2.12	2.16
0.153846	0.802	0.752	0.773	0.781	2.68	3.08	2.96	2.88
0.230769	0.807	0.736	0.731	0.758	4.56	4.96	4.72	4.56
0.307692	0.802	0.784	0.722	0.735	9.2	8.28	8.52	9
0.384615	0.824	0.806	0.79	0.79	16.6	16.04	16.48	16.72
0.461538	0.823	0.824	0.763	0.753	34.84	34.64	34.36	35.32
0.538462	0.841	0.814	0.722	0.709	69.44	70.2	69.44	70.32
0.615385	0.827	0.814	0.696	0.707	127.24	127.2	126.76	127.8
0.692308	0.83	0.821	0.73	0.727	181.36	181.28	181.28	180.28
0.769231	0.83	0.796	0.738	0.737	210.56	210.12	210.24	210.36
0.846154	0.829	0.776	0.739	0.739	216	216	216	216
0.923077	0.829	0.776	0.739	0.739	216	216	216	216
1	0.829	0.776	0.739	0.739	216	216	216	216

Table 24. $5 \times CV$ -5; The result of experiments for four variants of the Naive Bayes classifier; data set Hepatitis; Concept dependent granulation; r_{gran} = Granulation radius; nil = result for data without missing values; Acc = Accuracy of classification; GranSize = The size of data set after granulation in the fixed r .

r_{gran}	Acc				GranSize			
	V1	V2	V3	V4	V1	V2	V3	V4
0.0526316	<u>0.839</u>	0.828	0.846	0.821	<u>2</u>	2	2	2
0.105263	0.839	0.828	0.846	0.821	2	2	2	2
0.157895	0.827	0.831	0.846	0.821	2.2	2	2	2
0.210526	0.825	0.831	0.826	0.835	2.4	2.12	2.2	2.36
0.263158	0.826	0.841	0.791	0.844	2.6	2.68	2.68	2.76
0.315789	0.813	0.822	0.76	0.859	3.52	3.36	3.88	3.52
0.368421	0.822	0.836	0.693	0.855	5.32	5.08	4.96	4.68
0.421053	0.827	0.817	0.639	0.823	7.48	7.56	7.4	6.88
0.473684	0.868	0.827	0.761	0.84	11.64	12.16	11.44	11.72
0.526316	<u>0.876</u>	0.806	0.804	0.885	<u>18.28</u>	19.28	18.48	18.12
0.578947	0.871	0.796	0.8	0.863	31.36	30.84	29.8	30.68
0.631579	0.866	0.794	0.766	0.883	46.68	46.4	45.84	47.48
0.684211	0.857	0.794	0.804	0.871	70.28	70.04	69.4	70.2
0.736842	0.852	0.794	0.813	0.879	89.32	90.2	89.6	90.72
0.789474	0.855	0.794	0.83	0.886	109.28	110	109.88	110
0.842105	0.845	0.794	0.843	0.879	116.92	117.04	116.8	117.12
0.894737	0.845	0.794	0.844	0.876	121	121	121	121
0.947368	0.843	0.794	0.845	0.876	122	121.92	121.96	121.96
1	0.841	0.794	0.845	0.876	124	124	124	124

In Table 25, we have presented an example of the result of a deep neural network on the granulated data—see [39]. It turns out that it learns the internal knowledge of decision-making systems and maintains a high level of classification effectiveness. In Table 25 and Figure 4, we have the result for Australian Credit data set, for radius 0.66, with a reduction of 40 percent, and classification efficiency is around 84 percent.

Table 25. Results for Australian Credit dataset (mean from 10 experiments). Classification based on learning of deep neural networks—see [39].

Gran_rad	No_of_gran_objects	Percentage_of_objects	Time_to_learn	Accuracy
	Mean	Mean	Mean	Mean
0.0667	2.0	0.4149	0.36664	0.5646
0.1333	2.0	0.4149	0.3607	0.5337
0.2000	3.4	0.7054	0.3691	0.5423
0.2667	5.1	1.0581	0.3685	0.5154
0.3333	8.2	1.7012	0.3696	0.5192
0.4000	16.0	3.3195	0.3778	0.5577
0.4667	31.6	6.5560	0.3777	0.6236
0.5333	65.3	13.5477	0.3916	0.7764
0.6000	145.3	30.1452	0.4287	0.8125
0.6667	<u>283.8</u>	58.8797	0.7464	<u>0.8399</u>
0.7333	412.9	85.6639	0.8210	0.8534
0.8000	468.8	97.2614	0.8585	0.8587
0.8667	477.9	99.1494	0.8532	0.8553
0.9333	479.3	99.4398	0.8817	0.8553
1.0000	482.0	100.0000	0.8995	0.8562

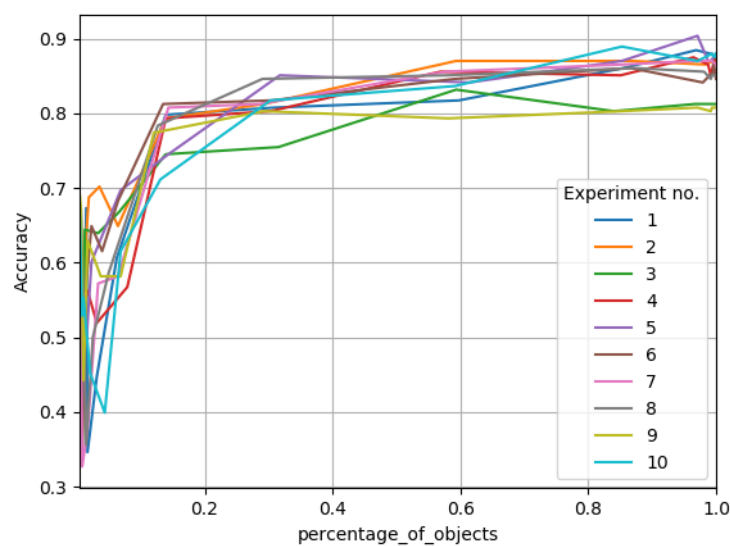


Figure 4. Visualization of classification efficiency for ten learning cycles of the neural network taking into account the percentage reduction of objects.

The additional experimental results presented were to show that our granular techniques are compatible with various classification methods. In the next section, we discuss the potential directions of development of granular computing methods, through the prism of the possibilities of our own methods.

5. Future Directions in Granular Computing Paradigm

Granular computing techniques will undoubtedly play a key role in building artificial intelligence because intelligent handling of data are based on analyzing its similarity and abstracting from the vast amount of information available in the environment. One of the problems to be solved is the ability to use real-time granular computing techniques on large data. The only barrier against using these methods is scalability problem. To deal with possible scalability problems, the following methods can be considered: Data sampling method and creation of model based on samples; Decomposition method, to use the algorithms on the split data and work on them separately; the streaming computing method, incremental data processing; the mass parallel computing technique on the computer cluster, with the use of classic ways to compute in parallel, like MPI implementation (Message Passing Interface); and mass parallel computing methods based on future technologies like quantum calculations. Without a doubt, deep neural networks is one of the promising fields for using granular computing. New methods of data preprocessing can be expected to emerge, before feeding it into deep neural networks. In particular, we mean the use of granular computing in the convolutionary and pooling part of the convolutional neural networks. The granular structures of the granular computing paradigm can intuitively be used to build such new network architectures at a time when we have no clear limit of creating neural network structures. Modeling the world using granular computing is a very natural process for us, which will undoubtedly play a crucial role in the development of future technologies.

6. Conclusions

In this work, we offer a review of selected recently developed granular computing techniques dedicated to the approximation of decision systems (from the family of methods proposed by Polkowski in [22,24]). That is, techniques which, among other things, aim at reducing the size of data while maintaining their classification efficiency. A very important family of techniques is dedicated to speeding up decision-making processes. Our approximation techniques reduce the size of decision systems significantly maintaining the internal knowledge at the same time, which was proven

in many experimental works. In our research, the main problem for standard, concept dependent, and layered methods is the need to estimate the optimal granulation radius searching among all possible ones. The problem has been partially solved for these methods—in the previous works, we have developed heuristics for searching optimal parameters by a double granulation technique (see [30]). In our last technique, homogeneous granulation, this problem does not apply because parameters are automatically set in the process of approximation. Our last method seems to be an important discovery, as it is immediately applicable, without the need to estimate the parameters, and it turns out to work very well in all the contexts we have studied. Particularly noteworthy is its application in the new technique of boosting classification—Ensemble of Random Granular Reflections [32]. To sum up our work, the presented granulation techniques allow for reducing the number of exhaustive set of rules by up to 99 percent while maintaining classification efficiency at the level obtained on the original unreduced data. Such efficiency was obtained, for example, for the concept-dependent technique using the kNN classifier. On the other hand, our methods achieve a reduction in the number of objects to more than 90 percent while maintaining classification efficiency on the original data. We have achieved such results, for example, for standard granulation with the kNN classification and concept-dependent granulation using the Naive Bayes classifier. As the closest directions of research on the development of our knowledge granulation methods, we can point out the work on hybrids with deep neural network learning and Random Forests technique. Another direction of work is the application in the process of convolution and pooling for the convolutionary neural networks and development of our proposed Ensemble model based on random granular reflections of decision systems. In conclusion of this review, we may add that, without any doubt, real-time granular computing methods will play an important role in creating artificial intelligence. Therefore, it is worthwhile to develop methods for the approximation of decision systems in order to invest in research into this prospective paradigm of knowledge.

Funding: This work has been fully supported by the grant from the Ministry of Science and Higher Education of the Republic of Poland under the project number 23.610.007-000.

Conflicts of Interest: The author declares no conflict of interest. The funders had no role in the design of the study; in the collection, analyses, or interpretation of data; in the writing of the manuscript, or in the decision to publish the results.

References

1. Zadeh, L.A. Fuzzy Sets and Information Granularity. 1979. Available online: <https://digitalassets.lib.berkeley.edu/techreports/ucb/text/ERL-m-79-45.pdf> (accessed on 13 February 2020).
2. Zadeh, L.A. Graduation and granulation are keys to computation with information described in natural language. In Proceedings of the 2006 IEEE International Conference on Granular Computing, Atlanta, GA, USA, 10–12 May 2006.
3. Pawlak, Z. Rough sets. *Int. J. Comput. Inform. Sci.* **1982**, *11*, 341–356. [CrossRef]
4. Skowron, A.; Polkowski, L. Synthesis of decision systems from data tables. In *Rough Sets and Data Mining*; Lin, T.Y., Cercone, N., Eds.; Springer: Boston, MA, USA, 1997; pp. 289–299.
5. Lin, T.Y. Granular computing: examples, intuitions and modeling. In Proceedings of the 2005 IEEE International Conference on Granular Computing, Beijing, China, 25–27 July 2005; Volume 1, pp. 40–44.
6. Yao, Y.Y. Granular computing: Basic issues and possible solutions, In Proceedings of the 5th Joint Conference on Information Sciences, Atlantic, NJ, USA, 27 February 2000; Volume 1, pp. 186–189.
7. Yao, Y. Information Granulation and Approximation in a Decision-Theoretical Model of Rough Sets. In *Rough-Neural Computing*; Pal, S.K., Polkowski, L., Skowron, A., Eds.; Springer: Berlin, Germany, 2004; pp. 491–516.
8. Yao, Y. Perspectives of granular computing. In Proceedings of the 2005 IEEE International Conference on Granular Computing, Beijing, China, 25–27 July 2005; Volume 1, pp. 85–90.
9. Skowron, A.; Stepaniuk, J. Information granules: Towards foundations of granular computing. *Int. J. Intell. Syst.* **2001**, *16*, 57–85. [CrossRef]

10. Skowron, A.; Stepaniuk, J. Information Granules and Rough-Neural Computing. In *Rough-Neural Computing*; Pal, S.K., Polkowski, L., Skowron, A., Eds.; Springer: Berlin, Germany, 2004; pp. 43–84.
11. Polkowski, L.; Semeniuk–Polkowska, M. On rough set logics based on similarity relations. *Fund. Inform.* **2005**, *64*, 379–390.
12. Liu, Q.; Sun, H. Theoretical study of granular computing. In *Rough Sets and Knowledge Technology*; Wang, G.Y., Peters, J.F., Skowron, A., Yao, Y., Eds.; Springer: Berlin, Germany, 2006; Volume 4062, pp. 92–102.
13. Cabrerizo, F.J.; Al-Hmouz, R.; Morfeq, A.; Martínez, M.A.; Pedrycz, W.; Herrera-Viedma, E. Estimating incomplete information in group decision-making: A framework of granular computing. *Appl. Soft Comput.* **2020**, *86*, 105930. [\[CrossRef\]](#)
14. Hryniewicz, O.; Kaczmarek, K. Bayesian analysis of time series using granular computing approach. *Appl. Soft Comput.* **2016**, *47*, 644–652. [\[CrossRef\]](#)
15. Martino, A.; Giuliani, A.; Rizzi, A. (Hyper) Graph Embedding and Classification via Simplicial Complexes. *Algorithms* **2019**, *12*, 223. [\[CrossRef\]](#)
16. Martino, A.; Giuliani, A.; Todde, V.; Bizzarri, M.; Rizzi, A. Metabolic networks classification and knowledge discovery by information granulation. *Comput. Biol. Chem.* **2020**, *84*, 107187. [\[CrossRef\]](#) [\[PubMed\]](#)
17. Pownuk, A.; Kreinovich, V. Granular approach to data processing under probabilistic uncertainty. In *Granular Computing*; Springer: Berlin, Germany, 2019; pp. 1–17.
18. Zhong, C.; Pedrycz, W.; Wang, D.; Li, L.; Li, Z. Granular data imputation: A framework of granular computing. *Appl. Soft Comput.* **2016**, *46*, 307–316. [\[CrossRef\]](#)
19. Leng, J.; Chen, Q.; Mao, N.; Jiang, P. Combining granular computing technique with deep learning for service planning under social manufacturing contexts. *Knowl.-Based Syst.* **2018**, *143*, 295–306. [\[CrossRef\]](#)
20. Ghiasi, B.; Sheikhan, H.; Zeynolabedin, A.; Niksokhan, M.H. Granular computing-neural network model for prediction of longitudinal dispersion coefficients in rivers. *Water Sci. Technol.* **2020**, *80*, 1880–1892. [\[CrossRef\]](#) [\[PubMed\]](#)
21. Capizzi, G.; Lo Sciuto, G.; Napoli, C.; Polap, D.; Woźniak, M. Small Lung Nodules Detection Based on Fuzzy-Logic and Probabilistic Neural Network with Bio-inspired Reinforcement Learning. 2019. Available online: <https://ieeexplore.ieee.org/abstract/document/8895990> (accessed on 13 February 2020).
22. Polkowski, L. Formal granular calculi based on rough inclusions. In Proceedings of the 2005 IEEE Conference on Granular Computing, Beijing, China, 25–27 July 2005; pp. 57–62.
23. Polkowski, L. *Approximate Reasoning by Parts. An Introduction to Rough Mereology*; Springer: Berlin, Germany, 2011.
24. Polkowski, L. A model of granular computing with applications. In Proceedings of the 2006 IEEE Conference on Granular Computing, Atlanta, GA, USA, 10 May 2006; pp. 9–16.
25. Artiemjew, P. Classifiers from Granulated Data Sets: Concept Dependent and Layered Granulation. 2007. Available online: <https://pdfs.semanticscholar.org/e46a/0e41d0833263220680aa1ec7ae9ed3eddbb42.pdf#page=7> (accessed on 13 February 2020).
26. Artiemjew, P.; Ropiak, K.K. On Granular Rough Computing: Handling Missing Values by Means of Homogeneous Granulation. *Computers* **2020**, *9*, 13. [\[CrossRef\]](#)
27. Polkowski, L. Granulation of knowledge in decision systems: The approach based on rough inclusions. The method and its applications. In *Rough Sets and Intelligent Systems Paradigms*; Kryszkiewicz, M., Peters, J.F., Rybinski, H., Skowron, A., Eds.; Springer: Berlin, Germany, 2007; Volume 4585, pp. 69–79.
28. Polkowski, L. Granulation of Knowledge: Similarity Based Approach in Information and Decision Systems. In *Encyclopedia of Complexity and System Sciences*; Meyers, R.A., Ed.; Springer: Berlin, Germany, 2009.
29. Polkowski, L.; Artiemjew, P. On granular rough computing with missing values. In *Rough Sets and Intelligent Systems Paradigms*; Kryszkiewicz, M., Peters, J.F., Rybinski, H., Skowron, A., Eds.; Springer: Berlin, Germany, 2007; Volume 4585, pp. 271–279.
30. Polkowski, L.; Artiemjew, P. *Granular Computing in Decision Approximation - An Application of Rough Mereology*; Springer: Cham, Switzerland, 2015.
31. Polkowski, L.; Artiemjew, P. On granular rough computing: Factoring classifiers through granular structures. In *Rough Sets and Intelligent Systems Paradigms*; Kryszkiewicz, M., Peters, J.F., Rybinski, H., Skowron, A., Eds.; Springer: Berlin, Germany, 2007; Volume 4585, pp. 280–290.
32. Artiemjew, P.; Ropiak, K. A Novel Ensemble Model - The Random Granular Reflections. 2018. Available online: <http://ceur-ws.org/Vol-2240/paper17.pdf> (accessed on 13 February 2020).
33. Ropiak, K.; Artiemjew, P. Homogenous Granulation and Its Epsilon Variant. *Computers* **2019**, *8*, 36. [\[CrossRef\]](#)

34. Artiemjew, P. A Review of the Knowledge Granulation Methods: Discrete vs. Continuous Algorithms. In *Rough Sets and Intelligent Systems - Professor Zdzisław Pawlak in Memoriam*; Skowron, A., Suraj, Z., Eds.; Springer: Berlin, Germany, 2013; Volume 43, pp. 41–59.
35. Polkowski, L. *Rough Sets*; Springer: Berlin, Germany, 2002.
36. Quinlan, J.R. *C4.5: Programs for Machine Learning*; Elsevier: New York, NY, USA, 2004.
37. University of California, Irvine Machine Learning Repository. Available online: <https://archive.ics.uci.edu/ml/index.php> (accessed on 13 February 2020).
38. Szypulski, J.; Artiemjew, P. The Rough Granular Approach to Classifier Synthesis by Means of SVM. In *Rough Sets, Fuzzy Sets, Data Mining, and Granular Computing*; Yao, Y., Hu, Q., Yu, H., Grzymala-Busse, J., Eds.; Springer: Cham, Switzerland, 2015; Volume 9437, pp. 256–263.
39. Ropiak, K.; Artiemjew, P. On a Hybridization of Deep Learning and Rough Set Based Granular Computing. *Algorithms* **2020**, *13*, 63. [[CrossRef](#)]



© 2020 by the authors. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).