

Possibility Data Analysis with Rough Sets Concept

大阪府立大学工学部

田中英夫 (Hideo Tanaka)

李海寬 (Haekwan Lee)

郭沛俊 (Peijun Guo)

Abstract : This paper is dealing with the upper and lower approximation models for representing the given phenomenon in a fuzzy environment as data analysis. The upper and lower approximation models can be derived from the given data with the possibility and necessity concepts, respectively. Thus, the given phenomenon can be analyzed as two approximation models which represent the upper and lower analyses for the given data. The modalities of the upper and lower models have been illustrated in regression analysis and also in the identification methods of possibility distributions. The comparison of the concepts of possibility data analysis and rough sets is shown clearly. A measure similar to the accuracy measure of rough sets is used to clarify the difference between the data structure and the assumed model.

1. Introduction

As we know, multivariate data analysis is a main tool for analyzing the uncertainty in the real world based on probability theory. Possibility data analysis is an alternative based on possibility distributions. Multivariate data analysis considers the uncertainty as probability phenomena while possibility data analysis considers that as possibility phenomena. Possibility theory based on possibility distributions has been proposed by Zadeh [1] and advanced by Dubois and Prade [2]. There are many applications of possibility theory in various fields. For example, possibility linear regression [3, 4] and possibility portfolio selection [5] have been formulated by using exponential possibility distributions on a multi-dimensional space. The theory of exponential possibility distributions has been proposed by Tanaka and Ishibuchi [6].

Based on possibilities which are more flexible than probabilities, the upper and the lower approximation models can be obtained for representing the given phenomenon. These two models can be derived from the given data with the possibility and necessity concepts, respectively. This modalities of the upper and lower models have been illustrated in interval regression analysis and also in the identification methods of possibility distributions. The upper and lower models are closely connected to the rough sets concept [7]. Pawlak [8] has introduced approximation data analysis based on rough sets. Rough sets can be described as approximate inclusion of sets and provide a systematic framework for the study of the problems arising from imprecise and insufficient knowledge. The comparison of concepts of possibility data analysis and rough sets is shown clearly in this paper. A measure similar to the accuracy measure of rough sets is used to clarify the difference between the data structure and the assumed model.

2. Upper and lower models for interval regression

Interval regression [9] is regarded as the simplest version of possibility regression analysis and thus easily applicable in many uncertain real-life phenomena. If the given output values are intervals, we can formulate two estimation models, i.e., the upper and lower approximation models based on the inclusion relations between the given interval outputs and the estimated intervals.

2.1. Two approximation models

An interval linear system can be written as

$$Y = A_0 + A_1x_1 + \cdots + A_nx_n = A x \quad (1)$$

where $x = (1, x_1, \dots, x_n)^t$ is an input vector, $A = (A_0, \dots, A_n)$ is an interval coefficient vector, and Y is the

corresponding estimated interval. An interval coefficient A_i is denoted as $A_i = (a_i, c_i)$ where a_i is a center and c_i is a radius. Thus, an interval coefficient A_i can also be expressed as

$$A_i = \left\{ u \mid a_i - c_i \leq u \leq a_i + c_i \right\} = [a_i - c_i, a_i + c_i]. \quad (2)$$

By interval arithmetic, the regression model (1) can be expressed as

$$\begin{aligned} Y(x_j) &= (a_0, c_0) + (a_1, c_1)x_{j1} + \cdots + (a_n, c_n)x_{jn} \\ &= (a^t x_j, c^t |x_j|) \end{aligned} \quad (3)$$

where $a = (a_0, \dots, a_n)^t$, $c = (c_0, \dots, c_n)^t$ and $|x_j| = (1, |x_{j1}|, \dots, |x_{jn}|)^t$. Here $a^t x_j$ and $c^t |x_j|$ represent a center and a radius of the estimated interval $Y(x_j)$ respectively.

Assume that input-output data (x_j, Y_j) are given as

$$(x_j, Y_j) = (1, x_{j1}, \dots, x_{jn}, Y_j), j = 1, \dots, p, \quad (4)$$

where x_j is the j -th input vector, Y_j is the corresponding interval output that consists of a center y_j and a radius e_j denoted as $Y_j = (y_j, e_j)$, and p is a data size. From the data set expressed by (4) we consider two approximation models

$$Y^*(x_j) = A_0^* + A_1^* x_{j1} + \cdots + A_n^* x_{jn}, j = 1, \dots, p, \text{ (upper model)} \quad (5)$$

$$Y_*(x_j) = A_{*0} + A_{*1} x_{j1} + \cdots + A_{*n} x_{jn}, j = 1, \dots, p, \text{ (lower model)} \quad (6)$$

where the interval coefficients A_i^* and A_{*i} are denoted as $A_i^* = (a_i^*, c_i^*)$ and $A_{*i} = (a_{*i}, c_{*i})$, respectively. The upper and lower models can be viewed as the possibility and necessity models, respectively. The upper and lower models are defined as follows. The estimated interval $Y^*(x_j)$ by the upper model always includes the observed interval Y_j , whereas the estimated interval $Y_*(x_j)$ by the lower model should be included in the observed interval Y_j . These relations can be expressed as follows:

$$Y^*(x_j) \supseteq Y_j \Leftrightarrow \begin{cases} a^* x_j - c^* |x_j| \leq y_j - e_j \\ y_j + e_j \leq a^* x_j + c^* |x_j| \end{cases} \quad (7)$$

$$Y_*(x_j) \subseteq Y_j \Leftrightarrow \begin{cases} y_j - e_j \leq a_* x_j - c_* |x_j| \\ a_* x_j + c_* |x_j| \leq y_j + e_j \end{cases} \quad (8)$$

$$Y_*(x_j) \subseteq Y_j \subseteq Y^*(x_j). \quad (9)$$

The relations given by (7)-(9) are graphically shown in Figure 1.

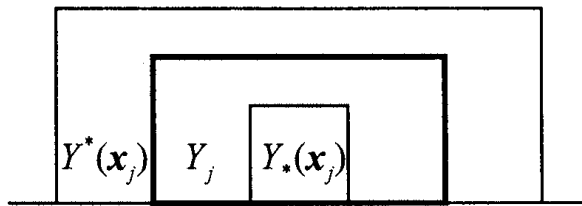


Figure 1. Relations among the upper, lower approximations, and the given interval.

Our main concern is to obtain interval coefficients A_i^* and A_{*i} , $i = 0, \dots, n$ based on the above mentioned mutual relations. The followings are LP problems to obtain two approximation models:

Upper Approximation model;

$$\begin{aligned} \text{Min}_{a^*, c^*} \quad & J^* = \sum_{j=1}^p c^* x_j \\ \text{subject to} \quad & Y^*(x_j) \supseteq Y_j, \quad j = 1, \dots, p, \\ & c_i^* \geq 0, \quad i = 0, \dots, n. \end{aligned} \quad (10)$$

Lower Approximation model;

$$\begin{aligned} \text{Max}_{a_*, c_*} \quad & J_* = \sum_{j=1}^p c_* x_j \\ \text{subject to} \quad & Y_*(x_j) \subseteq Y_j, \quad j = 1, \dots, p, \\ & c_{*i} \geq 0, \quad i = 0, \dots, n. \end{aligned} \quad (11)$$

where the constraint conditions in (10) and (11) refer to (7) and (8), respectively.

To obtain the upper and lower approximation models simultaneously, the following unified LP problem can be considered (see Ishibuchi and Tanaka [9]):

$$\begin{aligned} \text{Min}_{a^*, c^*, a_*, c_*} \quad & \sum_{j=1}^p c^* x_j - \sum_{j=1}^p c_* x_j \\ \text{subject to} \quad & Y^*(x_j) \supseteq Y_j, \\ & Y_*(x_j) \subseteq Y_j, \quad j = 1, \dots, p \\ & a_{*i} + c_{*i} \leq a_i^* + c_i^*, \\ & a_i^* - c_i^* \leq a_{*i} - c_{*i}, \\ & c_i^*, c_{*i} \geq 0, \quad i = 0, \dots, n. \end{aligned} \quad (12)$$

This LP problem is combining (10) and (11) in considering inclusion relations $A_i^* \supseteq A_{*i}$, $i = 0, \dots, n$ between the upper and lower regression coefficients. By adding $A_i^* \supseteq A_{*i}$, $i = 0, \dots, n$ in LP problem (12), we obtain two models satisfying $Y^*(x) \supseteq Y_*(x)$ for any x .

Now suppose we have data sets $(x_j^O, Y_j^O) = (1, x_{j1}^O, \dots, x_{jm}^O; Y_j^O)$, $j = 1, \dots, p$ which satisfies a following linear system:

$$Y^O(x_j^O) = A_0^O + A_1^O x_{j1}^O + \dots + A_n^O x_{jm}^O, \quad j = 1, \dots, p. \quad (13)$$

Then we obtain the following equalities

$$A^O = A^* = A_*, \quad Y^O(x) = Y^*(x) = Y_*(x) \quad (14)$$

by solving the upper approximation model (10) and the lower approximation model (11) (see Tanaka et al. [3]).

Unfortunately, the lower approximation model is not always guaranteed if we fail to assume a proper regression model. In case of no solution for linear regression with the lower approximation model, we can take a following polynomial:

$$Y = F(x) = A_0 + \sum A_i x_i + \sum A_{ij} x_i x_j + \sum A_{ijk} x_i x_j x_k + \dots \quad (15)$$

Since a polynomial such as (15) can represent any function, the center of the lower approximation model $Y_*(x_j)$ can meet the center of the observed output Y_j . Thus, one can obtain an optimal solution in the lower approximation model by increasing the number of terms of the polynomial (15) until a solution is found. The existence of the lower approximation model can be interpreted as the fact that the assumed model is somewhat reliable. If we find a lower model, since there always exists an optimal solution for the upper model, the measure of fitness φ can be introduced as

$$\varphi = \frac{1}{p} \sum_{j=1}^p \frac{c_i^* |x_j|}{c^* |x_j|} \quad (16)$$

where $0 \leq \varphi \leq 1$. This measure of fitness φ indicates how approximately the upper and lower models are assumed to the given data. It is desirable to assume regression models which give higher value of φ . If the given input-output data satisfy a linear system (13), then we can obtain the upper and the lower models which are identical. In this case the measure of fitness φ becomes 1.

Now assuming that an analyst may consider a tolerance limit ω such that $\varphi_Y \geq \omega$, we propose a new algorithm which gives two approximate models for the interval regression analysis.

Algorithm obtaining two approximation models :

Step 1: Take a linear function as regression model: $Y = A_0 + \sum A_i x_i$. (17)

Step 2: Solve the lower approximation model (11). If there is an optimal solution in the model (11), go to Step 4. Otherwise, go to Step 3.

Step 3: Increase the number of terms of the polynomials, i.e., $Y = A_0 + \sum A_i x_i + \sum A_{ij} x_i x_j$. (18)

Go to Step 2.

Step 4: Solve the unified LP problem (12) and calculate the measure of fitness φ of the two models. If $\varphi \geq \omega$, then go to Step 5 (We already have the optimal upper model $Y^*(x)$ and the optimal lower model $Y_*(x)$ satisfying $Y^*(x) \supseteq Y_*(x)$ for any x). Otherwise, go to Step 3.

Step5: End the procedure.

2.2. A numerical example

The data set of crisp inputs and interval outputs is shown in Table 1. The proposed algorithm is applied to obtain two approximate interval regression models under the assumption of $\omega = 0.25$ as a tolerance limit.

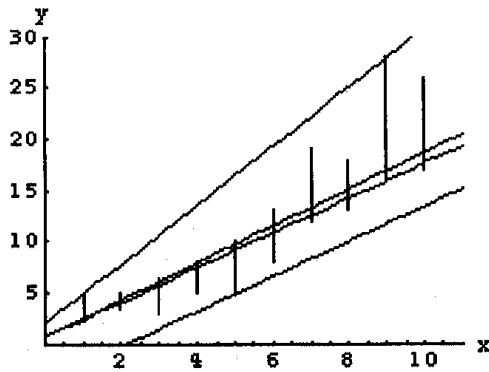
Table 1. Data set.

No.(j)	1	2	3	4	5
Input (x)	1	2	3	4	5
Output(y)	[2.5, 5]	[3.5, 5]	[3, 6.5]	[5, 8]	[5, 10]
No.(j)	6	7	8	9	10
Input (x)	6	7	8	9	10
Output(y)	[8, 13]	[12, 19]	[13, 18]	[16, 28]	[17, 26]

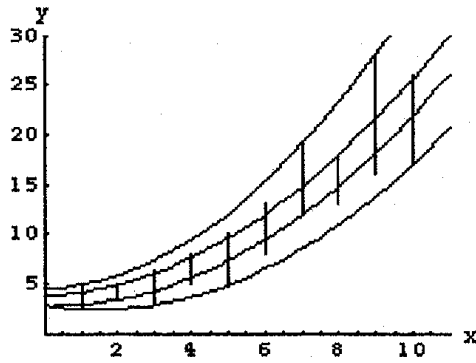
Using a linear regression model, we obtained the following upper and lower models

$$Y^*(x) = (-0.6563, 2.7813) + (2.2813, 0.5938)x \quad (19)$$

$$Y_*(x) = (0.8125, 0) + (1.7422, 0.0547)x \quad (20)$$



(a) $Y = A_0 + A_1 x$ ($\varphi_Y = 0.0458$)



(b) $Y = A_0 + A_1x + A_2x^2$ ($\varphi_T = 0.3281$)

Figure 2. Obtained upper and lower models for the given data set.

which are depicted in Figure 2 (a) where the outer two lines represent the upper model and the inner two lines represent the lower model. But the measure of fitness is $\varphi (=0.0458) < \omega (=0.25)$. Thus we rejected the obtained models. Therefore we increase the number of terms of the polynomials as (18). By solving the unified LP problem (12), we obtained the upper model $Y^*(x)$ and the lower model $Y_*(x)$ denoted as

$$Y^*(x) = (3.7118, 0.8507) - (0.1958, 0.3625)x + (0.2340, 0.0368)x^2 \quad (21)$$

$$Y_*(x) = (3.3417, 0.4806) + (0.0597, 0.1347)x + (0.1972, 0)x^2 \quad (22)$$

which are depicted in Figure 2 (b). The measure of fitness is greater than the tolerance limit, i.e., $\varphi (=0.3281) > \omega (=0.25)$. Thus we accept (21) and (22) as optimal models which satisfy $Y^*(x) \supseteq Y_*(x)$ for any x .

3. Identification methods of upper and lower possibility distributions

Let us begin with the given data (x_j, h_j) , ($j=1, \dots, p$) where $x_j = (x_{j1}, \dots, x_{jm})$ is a vector and h_j is an associated possibilistic grade given by expert knowledge. Assume that h_j , ($j=1, \dots, p$) are expressed by a possibility distribution A defined as

$$\Pi_A(x_j) = \exp \{-(x_j - a)' D_A^{-1} (x_j - a)\} = (a, D_A)_e, \quad (23)$$

where $a = (a_1, \dots, a_n)$ is a center vector and D_A is a symmetric positive definite matrix, denoted as $D_A > 0$.

Given the data, the problem is to determine an exponential possibility distribution of (23), i.e., a center vector a and a symmetric positive definite matrix D_A . According to two different viewpoints, the upper and lower possibility distributions are introduced in this section. The upper and lower possibility distributions denoted as Π_u and Π_l , respectively should be assumed to satisfy the inequality $\Pi_u(x) \geq \Pi_l(x)$, $\forall x$ with considering some similarities between our proposed methods and rough sets.

The center vector a can be approximately estimated as

$$a = x_j \quad (24)$$

whose $h_j = \max_{k=1, \dots, p} h_k$ and the associated possibility degree h_j is revised to be 1. Taking the transformation $y_j = x_j - a$, a possibility distribution with the zero center vector is as follows:

$$\Pi_A(y_j) = \exp(-y_j' D_A^{-1} y_j) = (0, D_A)_e, \quad (25)$$

where D_A is denoted as D_u and D_l corresponding to Π_u and Π_l , respectively.

3.1. Upper and lower distributions by the integrated model

The upper and lower distributions are used to reflect two kinds of distributions from the upper and lower directions.

Let us identify the upper and lower possibility distributions with the following assumptions whose meanings are illustrated in Figure 3.

- 1) Maximize $\Pi_l(y_1) \times \cdots \times \Pi_l(y_p)$ subject to $\Pi_l(y_j) \leq h_j$, $j = 1, \dots, p$. (Lower distribution)
- 2) Minimize $\Pi_u(y_1) \times \cdots \times \Pi_u(y_p)$ subject to $\Pi_u(y_j) \geq h_j$, $j = 1, \dots, p$. (Upper distribution)
- 3) $\Pi_u(y) \geq \Pi_l(y)$. (Relation of upper and lower distributions)

The upper and lower distributions can be simultaneously obtained by the following optimization problem.

$$\begin{aligned}
 & \underset{D_u, D_l}{Min} \quad \sum_{j=1}^p y_j' D_l^{-1} y_j - \sum_{j=1}^p y_j' D_u^{-1} y_j \\
 & \text{subject to} \quad y_j' D_u^{-1} y_j \leq -\ln h_j, \\
 & \quad y_j' D_l^{-1} y_j \geq -\ln h_j, \quad j = 1, \dots, p, \\
 & \quad D_u - D_l \geq 0, \\
 & \quad D_l \geq 0.
 \end{aligned} \tag{26}$$

It is obvious that (26) is a nonlinear optimization problem which is difficult to be solved.

To cope with this difficulty, we will use principle component analysis (PCA) to rotate the given data by the linear transformation T to obtain a positive definite matrix easily. The columns of T are characteristic vectors of matrix $\Sigma = [\sigma_{ij}]$ of given data, where σ_{ij} is defined by the following formulation.

$$\sigma_{ij} = \left\{ \sum_{k=1}^p (x_{ki} - a_i)(x_{kj} - a_j) h_k \right\} / \sum_{k=1}^p h_k. \tag{27}$$

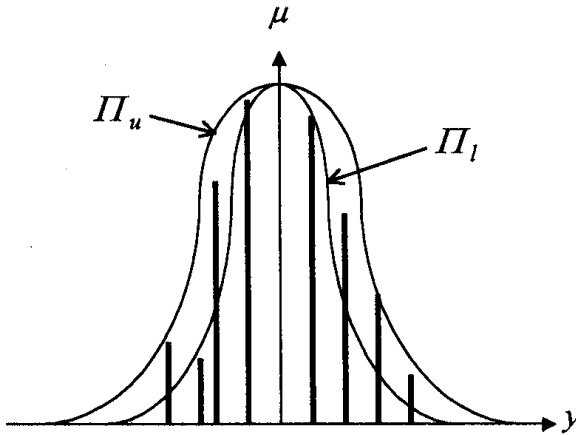


Figure 3. Graphic explanation of upper and lower distributions.

Using the linear transformation, the data $\{y_j\}$ can be transformed into $\{z_j = T'y_j\}$. Then we have

$$\Pi_A(z_j) = \exp(-z_j' T' D_A^{-1} T z_j). \tag{28}$$

According to the feature of PCA, $T' D_A^{-1} T$ is assumed to be a diagonal matrix as follows:

$$T' D_A^{-1} T = C_A = \begin{bmatrix} c_1 & & 0 \\ & \ddots & \\ 0 & & c_n \end{bmatrix}. \tag{29}$$

Denote C_A as C_u and C_l for the upper and lower possibility distribution cases, respectively. The corresponding diagonal elements in C_u and C_l are denoted as c_{ui} and c_{li} , ($i = 1, \dots, n$), respectively. The integrated model can be rewritten as follows:

$$\begin{aligned} \underset{C_u, C_l}{Min} \quad & \sum_{j=1}^p z_j' C_l z_j - \sum_{j=1}^p z_j' C_u z_j \\ \text{subject to} \quad & z_j' C_u z_j \leq -\ln h_j, \\ & z_j' C_l z_j \geq -\ln h_j, \quad j = 1, \dots, p, \\ & c_{ui} \geq \varepsilon, \\ & c_{li} \geq c_{ui}, \quad i = 1, \dots, n, \end{aligned} \quad (30)$$

where $c_{li} \geq c_{ui} \geq \varepsilon > 0$, ($i = 1, \dots, n$) make D_u and D_l positive definite, and $D_u - D_l$ semi-positive definite. Thus, we have

$$\begin{aligned} D_u &= T C_u^{-1} T', \\ D_l &= T C_l^{-1} T'. \end{aligned} \quad (31)$$

Similar to regression analysis, we can define the measure of fitness η as

$$\eta = \frac{1}{p} \sum_{j=1}^p \frac{\Pi_l(x_j)}{\Pi_u(x_j)}. \quad (32)$$

3.2. A numerical example

The data in the possibility portfolio problem are given in the following table.

Table 2. Return rate on two securities and possibility degrees.

hi	year	#1 Am.T	#2 A.T.&T.
0.2	1937(1)	-0.305	-0.173
0.241	1938(2)	0.513	0.098
0.282	1939(3)	0.055	0.2
0.324	1940(4)	-0.126	0.03
0.365	1941(5)	-0.28	-0.183
0.406	1942(6)	-0.003	0.067
0.447	1943(7)	0.428	0.3
0.488	1944(8)	0.192	0.103
0.529	1945(9)	0.446	0.216
0.571	1946(10)	-0.088	-0.046
0.612	1947(11)	-0.127	-0.071
0.653	1948(12)	-0.015	0.056
0.694	1949(13)	0.305	0.038
0.735	1950(14)	-0.096	0.089
0.776	1951(15)	0.016	0.09
0.818	1952(16)	0.128	0.083
0.859	1953(17)	-0.01	0.035
0.9	1954(18)	0.154	0.176

From the proposed approach explained in Section 3.1, we obtained

$$\begin{aligned}
a &= (0.154, 0.176), \\
D_u &= \begin{bmatrix} 0.2665 & 0.0972 \\ 0.0972 & 0.1689 \end{bmatrix}, \\
D_l &= \begin{bmatrix} 0.0313 & 0.0165 \\ 0.0165 & 0.0148 \end{bmatrix}.
\end{aligned} \tag{33}$$

Using the formulation (30) and (31), we obtained the two possibility distributions as shown in Figure 4 where the outer ellipse is the upper possibility distribution and the inner ellipse is the lower one for $h = 0.5$, respectively. From (33), we obtained $\eta = 0.226$.

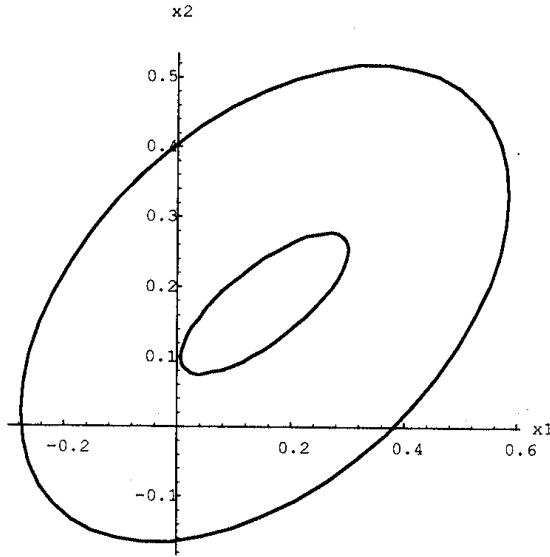


Figure 4. The upper and lower possibility distributions.

4. Similarities between proposed models and rough sets

Let a set $X \subset U$ be given. Then an upper approximation of X in A denoted as $A^*(X)$ means the least definable set containing X , and a lower approximation of X in A denoted as $A_*(X)$ means the greatest definable set contained in X . The upper approximation $A^*(X)$ and the lower approximation $A_*(X)$ can be defined as

$$A^*(X) = \bigcup_{E_i \cap X \neq \emptyset} E_i, \quad A_*(X) = \bigcup_{E_i \subseteq X} E_i, \tag{34}$$

where E_i is the i -th elementary set in A . An accuracy measure of a set X in the approximation space $A = (U, R)$ is defined as

$$\alpha_A(X) = \frac{\text{Card}(A_*(X))}{\text{Card}(A^*(X))} \tag{35}$$

where $\text{Card}(A_*(X))$ is the cardinality of $A_*(X)$. When the classification $C(U) = \{X_1, \dots, X_n\}$ is given, the accuracy of the classification $C(U)$ is defined as

$$\beta_A(U) = \text{Card}(\bigcup A_*(X_j)) / \text{Card}(\bigcup A^*(X_j)) \tag{36}$$

whose concept is used to define the measure of fitness in interval regression analysis and the identification methods of exponential possibility distributions. Furthermore, the upper and lower approximations of X , $A_*(X)$ and $A^*(X)$ are corresponding to the upper and lower approximation models in regression analysis and in the identification methods of possibility distributions. Thus, we can summarize the similarities between our models and rough sets in Table 3.

5. Concluding remarks

In the foregoing research, we are going to formulate portfolio selection problems, DEA (Data Envelopment Analysis), and possibility decision problems with the upper and lower models. It is emphasized that the upper and lower models can really represent possibility phenomena which are a kind of uncertainty.

Table 3. Similarities of the concepts between the proposed models and rough sets.

Interval regression analysis	Possibility distributions	Rough sets
Upper approximation model: $Y^+(x)$	Upper distribution: $\Pi_u(x)$	Upper approximation: $A^+(X)$
Lower approximation model: $Y_-(x)$	Lower distribution: $\Pi_l(x)$	Lower approximation: $A_-(X)$
Spread of $Y^+(x)$: $c^+ x $	Spread of $\Pi_u(x)$: $\sum \Pi_u(x_j)$	Cardinality of $A^+(X)$: $\text{Card}(A^+(X))$
Spread of $Y_-(x)$: $c^- x $	Spread of $\Pi_l(x)$: $\sum \Pi_l(x_j)$	Cardinality of $A_-(X)$: $\text{Card}(A_-(X))$
Inclusion relation: $Y^+(x) \supseteq Y_-(x)$	Inequality relation: $\Pi_u(x) \geq \Pi_l(x)$	Inclusion relation: $A^+(X) \supseteq A_-(X)$
Measure of fitness: φ (The higher, the better.)	Measure of fitness: η (The higher, the better.)	Accuracy measure of classification $C(U)$: $\beta_A(U)$ (The higher, the better.)

References

- [1] L.A. Zadeh, "Fuzzy sets as a basis for a theory of possibility", *Fuzzy Sets and Systems*, vol.1, 1977, pp.3-28.
- [2] D. Dubois, and H. Prade, *Possibility Theory*, Plenum Press, New York, 1988.
- [3] H. Tanaka, I. Hayashi, and J. Watada, "Possibilistic linear regression analysis for fuzzy data", *European J. of Operational Research*, vol.40, 1989, pp. 389-396.
- [4] H. Tanaka and H. Ishibuchi, "Identification of possibilistic linear systems by quadratic membership functions of fuzzy parameters", *Fuzzy sets and Systems*, vol.41, 1991, pp. 145-160.
- [5] H. Tanaka, and P. Guo, "Portfolio selections based on upper and lower exponential possibility distributions", *European J. of Operational Research*, (to appear).
- [6] H. Tanaka, and H. Ishibuchi, "Evidence theory of exponential possibility distributions", *Int. J. of Approximate Reasoning*, vol.8, 1993, pp. 123-140.
- [7] Z. Pawlak, "Rough sets", *Int. J. Information and Computer Sciences*, vol. 11, 1982, pp. 341-356.
- [8] Z. Pawlak, "Rough classification", *Int. J. Man-Machine Studies*, vol. 20, 1984, pp. 469-483.
- [9] H. Ishibuchi, and H. Tanaka, "A unified approach to possibility and necessity regression analysis with interval regression models", *Proc. of 5th IFSA World Congress*, Seoul, Korea, 1993, pp.501-504.