

Paper:

Cluster Validity Measures for Network Data

Yukihiro Hamasuna^{*1}, Daiki Kobayashi^{*2}, Ryo Ozaki^{*3}, and Yasunori Endo^{*4}

^{*1}Department of Informatics, School of Science and Engineering, Kindai University

3-4-1 Kowakae, Higashiosaka, Osaka 577-8502, Japan

E-mail: yhama@info.kindai.ac.jp

^{*2}Graduate School of Science and Engineering, Kindai University

3-4-1 Kowakae, Higashiosaka, Osaka 577-8502, Japan

E-mail: 1833340429c@kindai.ac.jp

^{*3}ALBERT Inc.

1-26-2 Nishishinjuku, Shinjuku-ku, Tokyo 163-0515, Japan

E-mail: 1210370107b@gmail.com

^{*4}Faculty of Engineering, Information and Systems, University of Tsukuba

1-1-1 Tennodai, Tsukuba, Ibaraki 305-8573, Japan

E-mail: endo@risk.tsukuba.ac.jp

[Received December 28, 2017; accepted March 13, 2018]

Modularity is one of the evaluation measures for network partitions and is used as the merging criterion in the Louvain method. To construct useful cluster validity measures and clustering methods for network data, network cluster validity measures are proposed based on the traditional indices. The effectiveness of the proposed measures are compared and applied to determine the optimal number of clusters. The network cluster partitions of various network data which are generated from the Polaris dataset are obtained by k -medoids with Dijkstra's algorithm and evaluated by the proposed measures as well as the modularity. Our numerical experiments show that the Dunn's index and the Xie-Beni's index-based measures are effective for network partitions compared to other indices.

Keywords: network clustering, cluster validity measures, modularity, k -medoids

1. Introduction

Network clustering in particular community detection is a recent research topic in many research fields such as social network services, e-commerce, and bioinformatics [1]. A community in network clustering is a group of nodes that are connected strongly to each other than the other nodes. The "community" is considered as the same concept as a "cluster" in the context of clustering. Modularity is a well-known evaluation measure for network partitions [2].

The modularity takes better values when the edges within clusters are dense, and the edges between clusters are sparse. The Louvain method is considered a recent representative method for detecting community structures from network data [3]. The Louvain method merges clusters

one by one by maximizing the modularity and determines the optimal number of clusters by searching the maximum value of the modularity. The modularity is used as not only a merging criterion but also as an evaluation measure in the Louvain method.

In the traditional clustering approach, various cluster validity measures are also used to evaluate cluster partitions and determine the optimal number of clusters [4–9]. These cluster validity measures are constructed by considering geometric features such as compactness and separateness. Compactness means the degree of denseness of the objects in each cluster, while separateness means the degree of distance of the clusters from each other. The modularity and traditional cluster validity measures are assumed to be based on the same concept from the viewpoint of functionality and the construction of measures. However, very few studies have been conducted on the usefulness of traditional cluster validity measures for network partition except for [10]. In addition, several network data did not obtain better partitions by the Louvain method and their optimal number of clusters were determined by the modularity [11]. The novel evaluation measures and network clustering methods are required to obtain better network partitions from massive and complex network data.

Herein, network cluster validity measures are proposed and the effectiveness of these measures are verified through numerical experiments to construct novel evaluation measures and clustering methods for network data. The network cluster validity measures are based on traditional cluster validity measures such as Dunn's index (DI) [4], the sum of the trace of fuzzy covariance matrix (W_{tr}) [5], and Xie-Beni's index (XB) [6]. These measures are re-written to evaluate network partitions. The numerical experiments are conducted with network datasets that are described as a weighted undirected graph. To obtain a network cluster partition, k -medoids, [12] which is known as a variant of k -means [13] is applied to a weighted undi-

rected graph.

The reminder of this paper is organized as follows: In Section 2, we introduce the notation, k -medoids, cluster validity measures, and the modularity. In Section 3, we propose the network cluster validity measures. In Section 4, we describe the experiments that show the effectiveness of the proposed method. In Section 5, we provide some concluding remarks regarding this research.

2. Preliminaries

A set of objects to be clustered is given and denoted by $X = \{x_1, \dots, x_n\}$, in which x_k ($k = 1, \dots, n$) is an object. In most cases, each object x_k is a vector in the p -dimensional Euclidean space $\mathbb{R}^p$, that is, an object $x_k \in \mathbb{R}^p$. When the network data is assumed, the dissimilarity $d_{ij} \geq 0$ between two objects x_i, x_j is denoted by $d_{ij} = d(x_i, x_j)$ and the dissimilarity matrix is denoted by $D = (d_{ij})_{i=1 \sim n, j=1 \sim n}$. We assume D is symmetric $d_{ij} = d_{ji}$ and $d_{ii} = 0$. A network data (X, D) is assumed to be given, and that objects X and the weight d_{ij} of edges (x_i, x_j) are given. A cluster is denoted by G_i , and a collection of clusters is given by $\mathcal{G} = \{G_1, \dots, G_c\}$. A cluster center of G_i is denoted by $v_i \in \mathbb{R}^p$, and a set of v_i is given by $V = \{v_1, \dots, v_c\}$. The membership degree of x_k belonging to G_i and a partition matrix is denoted as u_{ki} , and $U = (u_{ki})_{1 \leq k \leq n, 1 \leq i \leq c}$.

2.1. k -Medoids

k -medoids is a variant of k -means clustering. The cluster center is used as a cluster representative in k -means [13]. In contrast, an object in each cluster is chosen as a cluster representative in k -medoids [12]. An objective function of k -medoids is as follows:

$$J_{kd}(U, W) = \sum_{i=1}^c \sum_{k=1}^n \sum_{l=1}^n u_{ki} w_{li} r_{kl}. \quad (1)$$

Here, r_{kl} represents a measure of relationship between objects and $W = (w_{li})_{1 \leq l \leq n, 1 \leq i \leq c}$ is a variable called prototype weight. In many cases, r_{kl} is considered as a dissimilarity between objects. An algorithm of k -medoids is based on the alternating optimization with u_{ki} and w_{li} under the constraints on u_{ki} and w_{li} as follows:

$$\mathcal{U}_h = \left\{ (u_{ki}) : u_{ki} \in \{0, 1\}, \sum_{i=1}^c u_{ki} = 1, \forall k \right\}, \quad (2)$$

$$\mathcal{W}_h = \left\{ (w_{li}) : w_{li} \in \{0, 1\}, \sum_{l=1}^n w_{li} = 1, \forall i \right\}. \quad (3)$$

The l -th-object that takes $w_{li} = 1$ is the representative in a cluster. The important feature of k -medoids is that it handles relational data denoted as a table of distances between objects such as network data.

Algorithm 1 k -medoids

KMdd1 Set initial medoids.

KMdd2 Calculate $u_{ki} \in U$ by Eq. (4).

KMdd3 Calculate $w_{li} \in W$ by Eq. (5).

KMdd4 If convergence criterion is satisfied, stop. Otherwise go back to **KMdd2**.

The optimal solutions for u_{ki} and w_{li} are as follows:

$$u_{ki} = \begin{cases} 1 & \left(i = \arg \min_s \sum_{l=1}^n w_{ls} r_{kl} \right) \\ 0 & \text{(otherwise)} \end{cases}, \quad \dots \quad (4)$$

$$w_{li} = \begin{cases} 1 & \left(l = \arg \min_t \sum_{k=1}^n u_{ki} r_{kt} \right) \\ 0 & \text{(otherwise)} \end{cases} \quad \dots \quad (5)$$

The medoid of G_i is denoted in another form as follows:

$$\text{Mdd}(G_i) = \arg \min_{x_k \in G_i} \sum_{x_l \in G_i} d(x_k, x_l). \quad (6)$$

Eqs. (5) and (6) mean the same optimal solution. By considering the optimization problem of J_{kd} , the optimal solution of w_{li} is described in Eq. (5). Further, Eq. (6) is considered by considering k -medoids in the algorithmic procedure. The algorithm of k -medoids is summarized as **Algorithm 1**.

2.2. Cluster Validity Measures

Cluster validity measures are used to evaluate cluster partitions and determine the optimal number of clusters [7, 9].

2.2.1. Dunn's Index

Dunn's index (DI) [4] is constructed by the cluster compactness $\text{dia}(G_l)$, and the separateness $\text{dis}(G_i, G_j)$ and is described as follows:

$$DI = \frac{\min_{1 \leq i, j \leq c, i \neq j} \text{dis}(G_i, G_j)}{\max_{1 \leq l \leq c} \text{dia}(G_l)}, \quad (7)$$

$$\text{dis}(G_i, G_j) = \min_{x \in G_i, y \in G_j} d(x, y),$$

$$\text{dia}(G_l) = \max_{x, y \in G_l} d(x, y).$$

$d(x, y)$ is the dissimilarity between objects x and y . In addition, $\text{dis}(G_i, G_j)$ indicates the minimum dissimilarity between two clusters and $\text{dia}(G_l)$ is the maximum dissimilarity within a cluster. If DI is large, the cluster partition is considered to be good.

2.2.2. Trace of Fuzzy Covariance Matrix

Gath-Geva's index [5] is based on a fuzzy covariance matrix F_i . The sum of the traces of F_i is considered as

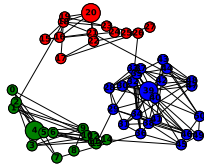


Fig. 1. Data 1 (25%–1%).

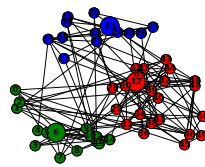


Fig. 2. Data 2 (25%–5%).

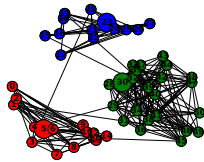


Fig. 3. Data 3 (50%–1%).

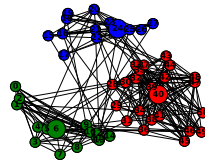


Fig. 4. Data 3 (50%–5%).

sists three clusters of 51 objects, and has two attributes in the original form. Each cluster has 13, 15, and 23 objects. The edges within the same clusters are generated by 25% or 50% of the whole, and the edges between clusters are also generated 1% or 5% of the whole. The edge weights are calculated as the squared L_2 -distance from each attribute in the original form. Four conditions are considered in each combination. For each condition, 100 patterns of network datasets are generated randomly. For each generated network data, k -medoids is executed 100 times with different initial values.

Figures 1–4 are illustrative examples of the network data, which are classified into adequate clusters. In these figures, the big circle means the medoid in each cluster obtained by the k -medoids. Each network data is assumed to be a complete graph by Dijkstra's algorithm, and is clustered by k -medoids. By applying Dijkstra's algorithm, the dissimilarity between nodes that are not connected directly are obtained and used in the k -medoids procedures. When the minimal value of the objective function is obtained from the 100 trials, the value of measures are calculated. Subsequently, the average and standard derivation of 100 network dataset patterns are shown in **Figs. 5–12**. These procedures are executed with the number of clusters from two to six.

4.2. Experimental Results

We show the experimental results to demonstrate the effectiveness of the cluster validity measures and the modularity. **Table 1** shows the summary of the evaluation by each measure. In this table, “o” means that the measure estimates the optimal number of clusters, while “—” means that the measure can not estimate the optimal number of clusters. DI and NXB evaluate the optimal number

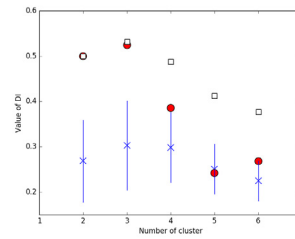


Fig. 5. Result of DI with Data 1 (25%–1%).

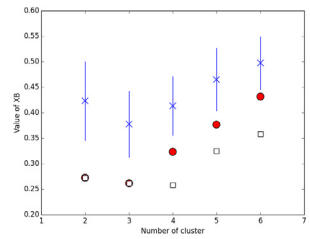


Fig. 6. Result of NXB with Data 1 (25%–1%).

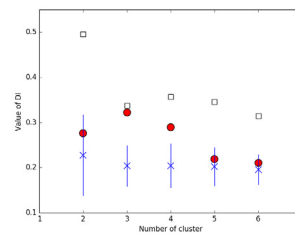


Fig. 7. Result of DI with Data 2 (25%–5%).

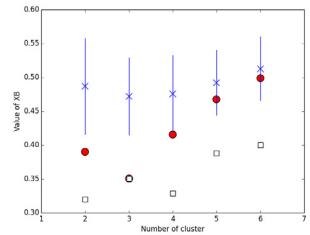


Fig. 8. Result of NXB with Data 2 (25%–5%).

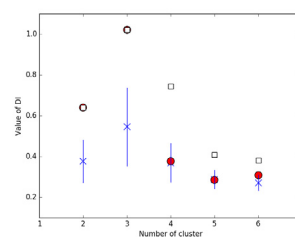


Fig. 9. Result of DI with Data 3 (50%–1%).

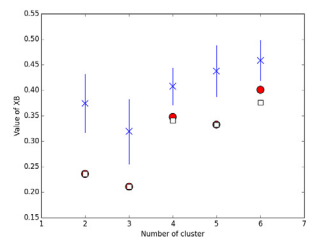


Fig. 10. Result of NXB with Data 3 (50%–1%).

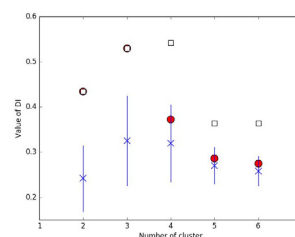


Fig. 11. Result of DI with Data 4 (50%–5%).

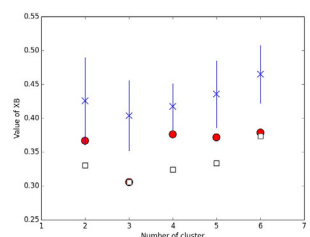


Fig. 12. Result of NXB with Data 4 (50%–5%).

of clusters in data 1, 3, and 4. However, DI and NXB can not evaluate the optimal number of clusters in data 2. NW_{tr} and modularity can not evaluate the optimal number of clusters for all data. The value of NW_{tr} and modularity monotonically increase/decrease for all data.

Figures. 5–12 are the results of DI and NXB with each condition. The results of NW_{tr} and modularity are omitted here because these measures can not obtain the optimal number of clusters by a monotonic increase/decrease. “×” means the average of 100 datasets and the vertical

Table 1. Result of network cluster validity measures with each condition.

	Data 1 (25%–1%)	Data 2 (25%–5%)	Data 3 (50%–1%)	Data 4 (50%–5%)
DI	○	–	○	○
NW _{tr}	–	–	–	–
NXB	○	○	○	○
Q	–	–	–	–

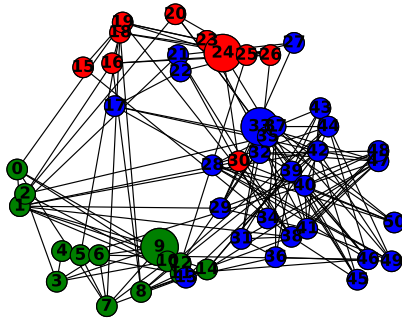


Fig. 13. An example of Data 2 (25%–5%).

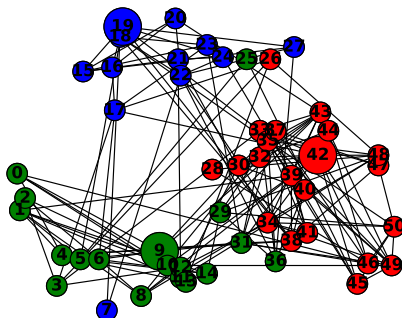


Fig. 14. Another example of Data 2 (25%–5%).

bar means the standard derivation. “○” means the result when the value of the objective function is minimized in 100 datasets with 100 different initial values. “□” means the best value of measures in 100 datasets. **Figs. 5, 6, 9–12** show that DI and NXB obtain better results for optimal cluster number $c = 3$ on average, i.e., the result when the objective function is minimal, or the best value of the measures. However, the average of DI fails to estimate the optimal number of clusters, as shown in **Fig. 7**. In addition, NXB also fails to estimate the optimal number of clusters, as shown in **Fig. 8**.

4.3. Discussions

First, we show the overview of the indices. **Table 1** shows that DI and NXB are suitable for the network clus-

ter validity measures. However, NW_{tr} and Q can not obtain the optimal number of clusters. **Figs. 13** and **14** are illustrative examples of data 2. If many edges are present between clusters, the dissimilarity between objects in the other cluster becomes smaller. This causes adverse effects to the cluster partitions.

The features of the measures are summarized as follows:

- DI obtains better results except for data 2. The medoids in each cluster is not considered in DI. If a few edges are present between clusters, the minimum distance between clusters is large. In such cases, DI obtains better results than others, as shown in **Figs. 5** and **9**.
- NXB obtains better results for all datasets. If many edges are present between clusters, the minimum distance between medoids is small. In such cases, NXB can not obtains better results, as shown in **Figs. 8** and **12**.
- NW_{tr} does not obtain better results for all datasets. The value of NW_{tr} monotonically decreases for the crisp partition as the number of clusters increases. This feature of W_{tr} for crisp partition is trivial. The fuzzy membership degree is an effective extension to overcome this problem. The fuzzification of NW_{tr} is required to evaluate the network partition and to estimate the optimal number of clusters from that sense.
- Modularity does not obtain better results for all datasets. The modularity value changes monotonically for network partitions obtained by k -medoids. Modularity is considered the suitable evaluation measure when it is used in the Louvain method.

These experiments show that DI and NXB are suitable for the network cluster validity measures. The advantage of the proposed network cluster validity measures is their usefulness for complex structures by introducing fuzzification and kernelization. For example, kernelized XB and W_{tr} have been proposed to improve the performance of traditional cluster validity measures [9].

In addition, traditional k -medoids and cluster validity measures have been actively studied until now. By extending the traditional approach to network clustering, various extensions could be proposed as novel cluster validity measures for network partitions. In addition, diffusion kernel [15] is a well-known kernel method to calculate the similarity between nodes for an unweighted

graph. As a natural extension of the kernel method to network clustering and cluster validity measures, the kernel method is one of the useful methods.

5. Conclusions

Herein, network cluster validity measures that are based on traditional indices were studied and the effectiveness of these measures were verified through numerical experiments. The numerical experiments showed that DI and NXB are effective for network partitions compared to other indices.

In future works, we will consider novel cluster validity measures including fuzzification and kernelization according to numerical experiments. Next, we will construct new clustering methods for network data based on the proposed indices. Moreover, we will conduct numerical experiments with large network data to demonstrate the effectiveness of our proposed methods.

Several network data exist, such as unweighted, directed, and bipartite. To handle various network structures, the fuzzy clustering approach [8, 14], kernel method [9, 15], probabilistic model, and other approaches are required. In addition, the relation among cluster validity measures, modularity, kernel methods, and the probabilistic should be discussed.

Acknowledgements

This work was partly supported by JSPS KAKENHI Grant Numbers JP16K16128 and JP17K00332.

References

- [1] M. Newman, "Networks: An Introduction," Oxford University Press, 2010.
- [2] M. E. J. Newman, "Modularity and community structure in networks," PNAS, Vol.103, No.23, pp. 8577-8582, 2006.
- [3] V. D. Blondel et al., "Fast unfolding of communities in large networks," J. of Statistical Mechanics: Theory and Experiment, p. P10008, 2008.
- [4] J. C. Dunn, "Well separated clusters and optimal fuzzy partitions," J. of Cybernetics, Vol.4, pp. 95-104, 1974.
- [5] I. Gath and A. B. Geva, "Unsupervised optimal fuzzy clustering," IEEE Trans. on Pattern Analysis and Machine Intelligence, Vol.11, No.7, pp. 773-780, 1989.
- [6] X. L. Xie and G. Beni, "A validity measure for fuzzy clustering," IEEE Trans. on Pattern Analysis and Machine Intelligence, Vol.13, No.8, pp. 841-847, 1991.
- [7] W. Wang and Y. Zhang, "On fuzzy cluster validity indices," Fuzzy Sets and Systems, Vol.158, No.19, pp. 2095-2117, 2007.
- [8] S. Miyamoto, H. Ichihashi, and K. Honda, "Algorithms for Fuzzy Clustering," Springer, 2008.
- [9] W. Hashimoto, T. Nakamura, and S. Miyamoto, "Comparison and evaluation of different cluster validity measures including their kernelization," J. Adv. Comput. Intell. Intell. Inform., Vol.13, No.3, pp. 204-209, 2009.
- [10] I. J. Sledge et al., "Relational Generalizations of Cluster Validity Indices," IEEE Trans. on Fuzzy Systems, Vol.18, No.4, pp. 771-786, 2010.
- [11] Y. Hamasuna, R. Ozaki, and Y. Endo, "A study on cluster validity measures for clustering network data," Joint 17th World Congress of Int. Fuzzy Systems Association and 9th Int. Conf. on Soft Computing and Intelligent Systems (IFSA-SCIS2017), #89, 2017.
- [12] L. Kaufman and P. J. Rousseeuw, "Finding Groups in Data: An Introduction to Cluster Analysis," Wiley, 1990.
- [13] A. K. Jain, "Data clustering: 50 years beyond K-means," Pattern Recognition Letters, Vol.31, No.8, pp. 651-666, 2010.
- [14] J. C. Bezdek, "Pattern Recognition with Fuzzy Objective Function Algorithms," Plenum Press, 1981.
- [15] R. I. Kondor and J. Lafferty, "Diffusion kernels on graphs and other discrete structures," Proc. of the 19th Int. Conf. on Machine Learning (ICML'02), pp. 315-322, 2002.



Name:

Yukihiro Hamasuna

Affiliation:

Lecturer, Department of Informatics,
School of Science and Engineering,
Kindai University

Address:

3-4-1 Kowakae, Higashiosaka, Osaka 577-8502, Japan

Brief Biographical History:

2009-2010 Research Fellow of the Japan Society for the Promotion of Science

2011-2014 Assistant Professor, Kindai University

2014- Lecturer, Kindai University

Main Works:

- Y. Hamasuna, N. Kinoshita, and Y. Endo, "Comparison of Cluster Validity Measures Based x -means," J. Adv. Comput. Intell. Intell. Inform., Vol.20, No.5, pp. 845-853, 2016.
- Y. Hamasuna and Y. Endo, "On Sequential Cluster Extraction Based on L_1 -Regularized Possibilistic c -Means," J. Adv. Comput. Intell. Intell. Inform., Vol.19, No.5, pp. 655-661, 2015.
- Y. Hamasuna and Y. Endo, "On Semi-supervised Fuzzy c -Means Clustering for Data with Clusterwise Tolerance by Opposite Criteria," Soft Computing, Vol.17, No.1, pp. 71-81, 2013.
- Y. Hamasuna, Y. Endo, and S. Miyamoto, "On Tolerant Fuzzy c -Means Clustering and Tolerant Possibilistic Clustering," Soft Computing, Vol.14, No.5, pp. 487-494, 2010.

Membership in Academic Societies:

- The Japan Society for Fuzzy Theory and Systems (SOFT)
- The Japanese Society for Artificial Intelligence (JSAI)
- The Institute of Electrical and Electronics Engineering (IEEE)



Name:

Daiki Kobayashi

Affiliation:

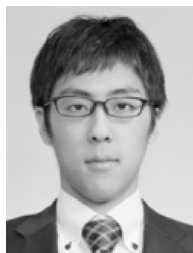
Graduate School of Science and Engineering,
Kindai University

Address:

3-4-1 Kowakae, Higashiosaka, Osaka 577-8502, Japan

Brief Biographical History:

2018- Graduate School of Science and Engineering, Kindai University



Name:
Ryo Ozaki

Affiliation:
ALBERT Inc.

Address:

1-26-2 Nishishinjuku, Shinjuku, Tokyo 163-0515, Japan

Brief Biographical History:

2016-2018 Graduate School of Science and Engineering, Kindai University

2018- ALBERT Inc.

Main Works:

- R. Ozaki, Y. Hamasuna, and Y. Endo, "Agglomerative Hierarchical Clustering Based on Local Optimization for Cluster Validity Measures," Proc. of 2017 IEEE Int. Conf. on Systems, Man, and Cybernetics (IEEE SMC 2017), pp. 1822-1827, 2017.
- R. Ozaki, Y. Hamasuna, and Y. Endo, "A Method of Two-Stage Clustering Based on Cluster Validity Measures," Joint 8th Int. Conf. on Soft Computing and Intelligent Systems and 17th Int. Symp. on Advanced Intelligent Systems (SCIS & ISIS 2016), pp. 410-415, 2016.



Name:
Yasunori Endo

Affiliation:
Professor, Faculty of Engineering, Information and Systems, University of Tsukuba

Address:

1-1-1 Tennodai, Tsukuba, Ibaraki 305-8573, Japan

Brief Biographical History:

1994-1997 Research Assistant, Waseda University

1997-2001 Assistant Professor, Tokai University

2001-2004 Assistant Professor, University of Tsukuba

2004-2013 Associate Professor, University of Tsukuba

2012 Guest Research Scholar, International Institute for Applied Systems Analysis (IIASA)

2013- Professor, University of Tsukuba

Main Works:

- Y. Endo and S. Miyamoto, "Spherical k-Means++ Clustering," The 12th Int. Conf. on Modeling Decisions for Artificial Intelligence (MDAI 2015), LNAI 9321, pp. 103-114, 2015.
- Y. Hamasuna and Y. Endo, "On a Family of New Sequential Hard Clustering," J. Adv. Comput. Intell. Inform., Vol.19, No.6, pp. 759-765, 2015.
- Y. Endo and N. Kinoshita, "Objective-Based Rough *c*-Means Clustering," Int. J. of Intelligent Systems, Vol.28, Issue 9, pp. 907-925, 2013.
- Y. Endo et al., "Fuzzy *c*-Means Clustering for Uncertain Data using Quadratic Penalty-Vector Regularization," J. Adv. Comput. Intell. Inform., Vol.15, No.1, pp. 76-82, 2011.
- Y. Endo et al., "Fuzzy *c*-Means for Data with Tolerance," Proc. 2005 Int. Symp. on Nonlinear Theory and Its Applications, pp. 345-348, 2005.
- Y. Endo, H. Haruyama, and T. Okubo, "On Some Hierarchical Clustering Algorithms Using Kernel Functions," Proc. IEEE Int. Conf. on Fuzzy Systems, #1106, 2004.

Membership in Academic Societies:

- The Japan Society for Fuzzy Theory and Intelligent Informatics (SOFT)
- The Institute of Electronics, Information and Communication Engineers (IEICE)
- The Institute of Electrical and Electronics Engineers (IEEE)