

A Study of Various Clustering Algorithms Used for Radar Signal Sorting

Fathi Elnour Hammed¹; Ahmed Abdalla Ali²; Ahmed Awad³
Department of Electrical and Computer Engineering, College of Engineering
Karary University- Khartoum- Sudan

Abstract:- Radar signal sorting plays a significant part in radar countermeasure technology and reconnaissance systems. By means of radar signal sorting, several radars and their parameters in the battlefield are precisely recognized and placed in the radar records for subsequent positioning and jamming processing. The basic sorting methods cannot fulfill the sorting process with accurate and efficient results. Therefore, in this paper, we conduct a study on two main classes clustering techniques, the first is hierarchical based clustering, and the second one is partition based clustering, which have different characteristics into groups of pulses and they can sort and handle large number of radar sequence with high precision and accuracy. The numerical simulations studied and compared both methods under different perspectives to clarify a new directions based on the insightful investigations of these sorting techniques.

Keywords:- Electronic Warfare (EW); Electronic Support Measure (ESM); Radar Signal Sorting; Clustering.

I. INTRODUCTION

Due to recent increase in sophistication of weapons' system and the great progress in signal processing techniques, the ability to find the position of enemy equipment and implement effective countermeasures to minimize hostile threats and maximize the successful of our own weapons are absolutely essential [1]. This is the primary objective of the electronic warfare (EW), which was initiated among the Second World War. EW takes many forms, such as detecting of the hostile emitters and degrading their

performances, etc. There are three parts of EW, that is, electronic support measures (ESM), electronic countermeasures (ECM) and electronic counter-countermeasures (ECCM)[2-4].

Generally, the ESM system consists of three parts: receiver, processor, and identifier [5]. Firstly, the ESM receiver converts the interleaved signal into digital form by using a unit called pulse analyzer, and the parameters of each pulse are involved in small files called pulse descriptor words (PDWs). Generally, PDWs consist of time of arrival (ToA), pulse amplitude (PA), angle of arrival (AoA), pulse width (PW), and radio frequency (RF). Therefore, one or more PDW parameters must be used to accomplish the sorting process. Secondly, the main processing part sorts the interleaved signals into different groups, such that, the pulses of each radar cannot be placed into more than one group. Finally, the sorted pulses are entered the emitter table in order to update the previous information in the table, and also to identify whether the new pulses will be associated with new emitters or not. Figure 1, shows the ESM received the interleaved radar signals, which were emanated from multiple emitters. It is evident that the processor is used to associate each pulse with its emitter.

It should be notable that the words sorting and deinterleaving are almost interchangeable. However, the ToA is often used for deinterleaving, while the PDWs parameters are often used for sorting. It would be preferable to demonstrating the general performances of the PDWs parameters prior to sinking more deeply into sorting process.

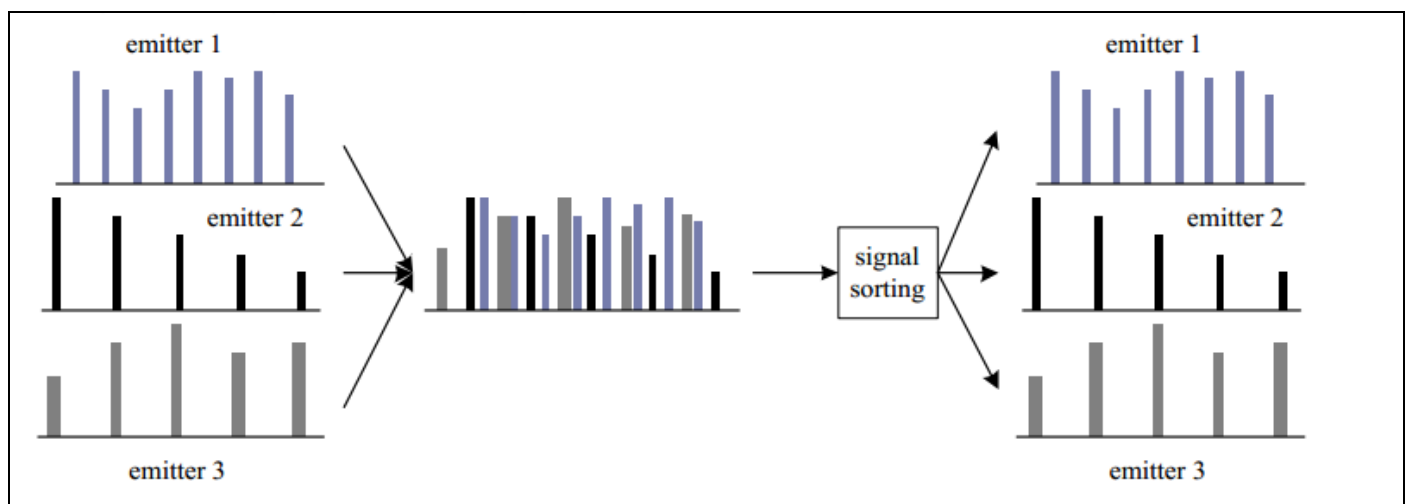


Fig 1 General Demonstration of Radar Signal Sorting.

II. FUNDAMENTAL OF CLUSTERING

Due to modern radar technologies and ample application of radar signals, the basic ToA deinterleaving methods cannot fulfill the deinterleaving process with accurate and efficient results. Clustering is an unsupervised technique, which can be defined as an explorative method that organizes data samples (objects) based on the similarities between the objects into groups of samples. From this definition, we can deploy clustering techniques in order to group large number of interleaved pulses into meaningful partitions, such that, the pulses in each partition are associated with a unique emitter. Implementing clustering technique is more flexible for modern radar signals such as staggered, hopping, and jitter signals. Generally, clustering techniques can be classified into three main classes, the first is hierarchical based clustering, the second one is partition based clustering [6], and the last one density based clustering.

In partition based clustering technique, the data samples are partitioned into a pre-defined number of clusters [7]. The main idea of partition based clustering technique is to minimize the cost function “objective function” based on the measured distance between clusters and prototypes. Commonly, partition based clustering is further classified into two main classes, that is, hard (crisp) clustering and soft (fuzzy) clustering [8]. In hard clustering each data sample belongs to only one cluster, while in soft clustering, each data

sample belongs to all clusters with a specific degree of membership. It should be notable that fuzzy clustering is quite computational than hard clustering; however, it is more suitable and accurate. Partition clustering is simple and particularly appropriated on spherical clusters. On the other hand, partition clustering suffers from the bad initiation of the clusters, which leads to wrong clustering results. Besides, it needs foreknowledge of the number of clusters, which is difficult issue in real time processing.

Hierarchical cluster analysis is an important technique in the data mining field. It has been widely used in various fields, such as pattern recognition, data analysis, and biological studies. The dynamic distance clustering (DDC) algorithm is dynamic, that is, the result of clustering has the dynamic class centers and unfixed the number of classes which depends on the input data. These are important to radar signal sorting.

A. Clustering Concept.

There is no specific notion of cluster, since there are many clustering types and each method defines the concept of cluster in different way. Consider Fig.1. It is clear that, in (a) treated all the entities as one cluster, while in (b) considered the entities as two clusters, and (c) treated the entities as three clusters. Therefore, the definition of cluster is changed from one type to other. Before we get started the clustering type, some concepts must be known.

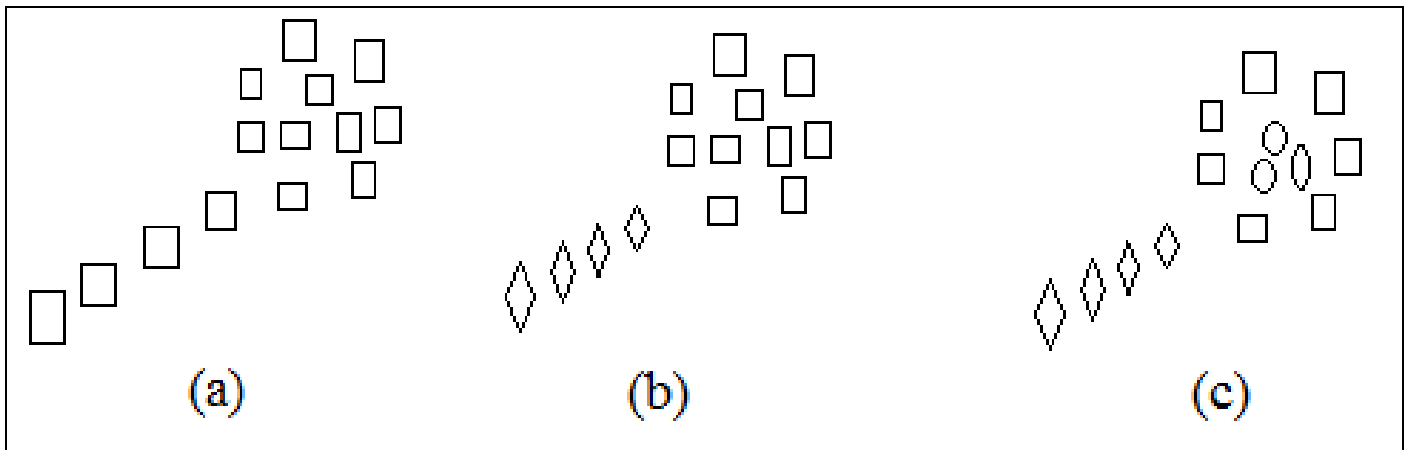


Fig 2 (a) One Cluster. (b) Two Clusters. (c) Three Clusters

➤ Preprocessing Step:

Let the received signal $Y = \{y_j\}$, $j = 1, \dots, n$ consists of n pulses, each pulse has N dimensional space, i.e. $g_i \subseteq \mathbb{R}^N$. Mathematically, this signal can be represented in matrix form as follows.

$$Y = \begin{pmatrix} y_{11} & y_{12} & \cdots & y_{1n} \\ y_{21} & y_{22} & \cdots & y_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ y_{N1} & y_{N2} & \cdots & y_{Nn} \end{pmatrix}_{N \times n} \quad (1)$$

Generally, before we start any type of clustering, some kind of preprocessing have to be done. The preprocessing equation can be represented as

$$\tilde{y}_{ij} = \frac{y_{ij} - \min(y_i)}{\max(y_i) - \min(y_i)} \quad (2)$$

Where $i = 1, \dots, N$, and $j = 1, \dots, n$. The above equation is also called normalization equation.

➤ Distance Measure;

Clustering is technique for grouping samples with similar attributes. In order to achieve this goal, some distance measure must define whether these samples are similar or

not. There are a lot of types of distance measure; however, the most famous one is Euclidean distance. Euclidean distance, d , between $y_1 = (y_{11}, y_{21}, \dots, y_{N1})^T$ and $y_2 = (y_{12}, y_{22}, \dots, y_{N2})^T$ is defined as follows

$$d = \sqrt{(y_{11} - y_{12})^2 + (y_{21} - y_{22})^2 + \dots + (y_{N1} - y_{N2})^2} \quad (3)$$

➤ Global Minima and Local Minima:

The Global minima can be defined as the minimum point that corresponds to the smallest value of the all-error function, while local minima are the minimum values of error compared with the nearby errors. Fig.2 delineated the concept of local minima and Global minima.

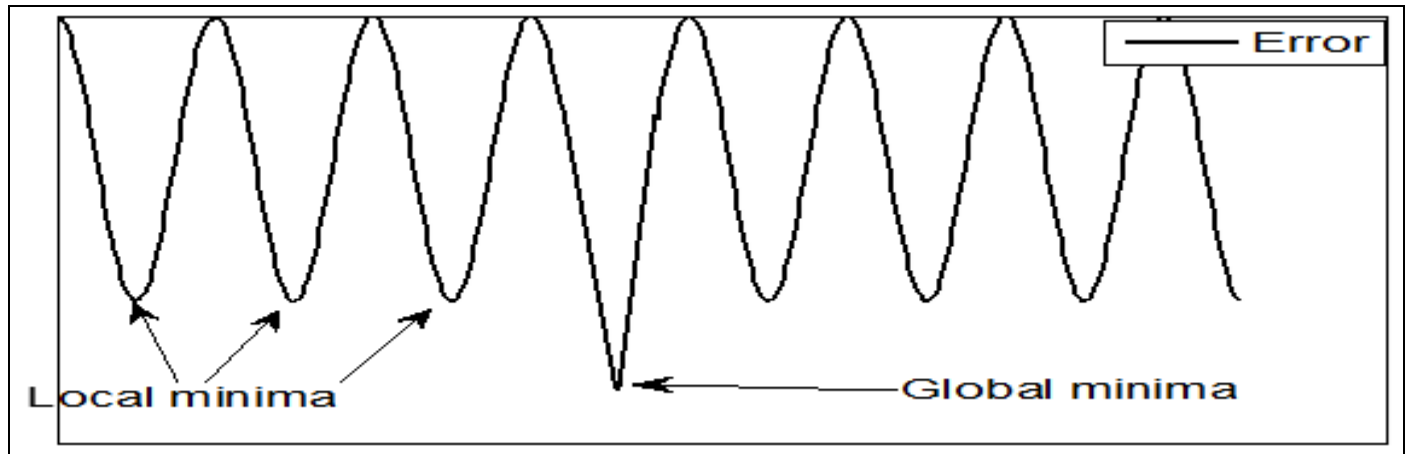


Fig 3 Global and Local Minima

III. THE PROPOSED CLUSTERING ALGORITHMS

A. Hierarchical Clustering Methods

Hierarchical based clustering technique is a method which creates a hierarchy's structure of clusters. Its main concept is that; objects are more likely to be linked with close data sample rather than farther one. In general Hierarchical clustering has been categorized into two types, i.e. Agglomerative clustering and Divisive clustering [8,9]. Agglomerative clustering first and foremost, by considers each data sample as a cluster and then the algorithm merges the elements into larger clusters based on their distances. At the end, these clusters are merged to form one cluster. While

divisive clustering considers all data samples as a unique cluster, and then the algorithm splits the cluster until each data sample represents a cluster on its own. Both Agglomerative clustering and Divisive clustering are represented by cluster tree or "dendrogram". Fig.3 shows the difference between Agglomerative clustering and Divisive clustering. The main merits of hierarchical are that, it does not need prior knowledge of clusters' number, and there is no effect of initialization. However, the drawbacks of hierarchical clustering are that, it only deals with local neighbors, it cannot incorporate information about clusters shape and size, and it associated with static algorithm, thus each data sample belonging to a cluster at the initial steps cannot belong to another cluster during the final steps.

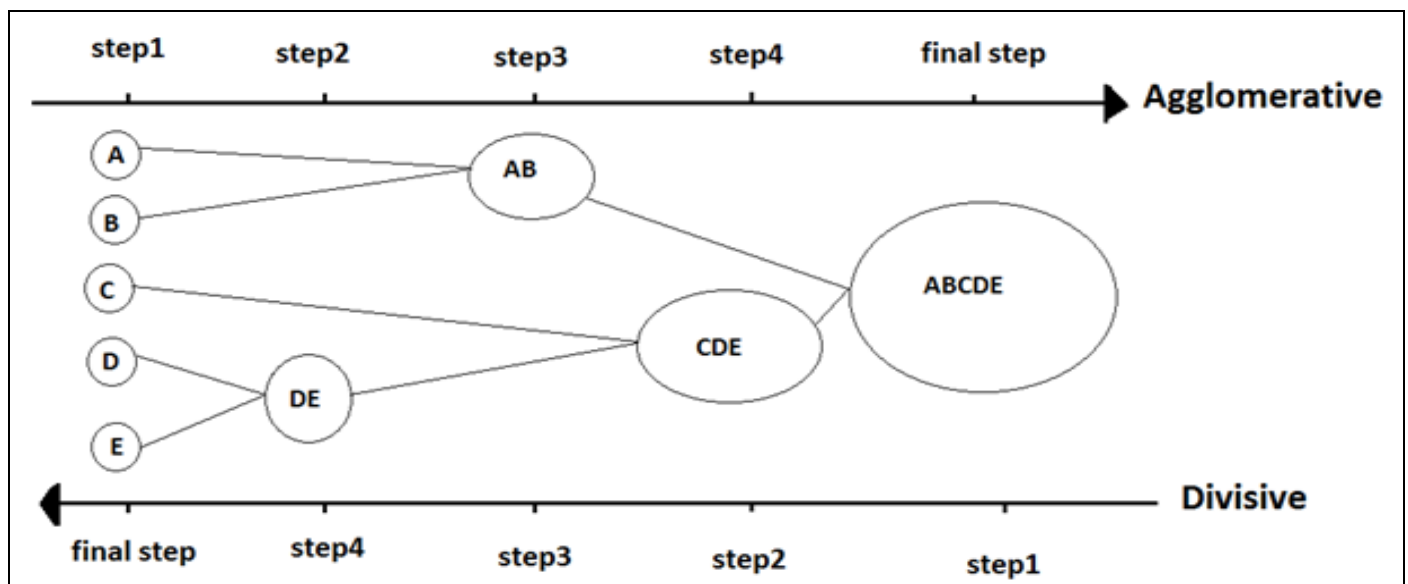


Fig 4 The difference between Agglomerative clustering and Divisive clustering

➤ *There are Three Commonly used Strategies to Calculate the Distance in Hierarchical Agglomerative Clustering:*

• *Single Linkage Method*[6].

It also called neighbor joining, minimum method, or the nearest neighbor method. In this method the role of combination is based on “the shortest distance”, that is, distance between two clusters is the nearest distance between any object in the first cluster and any object in the second cluster. Suppose that A and B are two clusters in 2-D space, $\mathbb{R}^2$, then the Single Linkage distance, d_{SL} , is defined as.

$$d_{SL}(A, B) = \min_{a \in A, b \in B} d(a, b) \quad (4)$$

• *Complete Linkage Method* [7]:

It also called the maximum method or the furthest neighbor method. Its role of combination is based on “maximum distance”, that is, distance between two clusters is the farthest distance between any object in the first cluster and any object in the second cluster, i.e

$$d_{CL}(A, B) = \max_{a \in A, b \in B} d(a, b) \quad (5)$$

Where d_{CL} is the Complete Linkage distance.

• *Average Linkage Method* [8]:

It also called minimum variance method. The distance between two clusters is the distance between their centers (mean value), i.e.

$$d_{AL}(A, B) = d(\mu_A, \mu_B) \quad (6)$$

Where d_{AL} is the Average Linkage distance, and

$$\mu_A = \frac{\sum_{a \in A} a}{|A|}, \quad \mu_B = \frac{\sum_{b \in B} b}{|B|} \quad (7)$$

B. Partition Based Clustering].

In partition based clustering technique, the data samples are partitioned into a pre-defined number of clusters. The main concept Partition based clustering technique is to minimize the cost function or objective function based on the distance measured between clusters and prototypes. Generally, partition based clustering are further classified into two main classes, that is, hard (crisp) clustering [10], and soft (fuzzy) clustering [11]. In hard clustering each data sample belongs to only one cluster, while in soft clustering, each data sample belongs to all clusters with a specific degree of membership. Fuzzy clustering methods are more computational than hard clustering methods. Moreover, to cluster the data samples, fuzzy clustering methods are more suitable than hard clustering methods.

➤ *Hard Clustering*

• *k-means Clustering Algorithm*.

The well-known hard clustering is k-means algorithm [12,13], which has been used in many fields such as, pattern recognition, image processing, data mining, etc. k-means is very simple algorithm, and it is used to cluster our data samples based on given number of clusters. The steps of the algorithm can be given as follows:.

- Initialize the algorithm by setting k centers randomly.
- Assign each sample (object) to the nearest center.
- Modify the centers by calculating the average samples in each cluster.

Repeat step (3) and (4) until the number of samples remain constant in each cluster or some stopping criteria satisfied Figure.5 shows the final clustering result of some data sample. We set the number of clusters to three. There are three types of data samples, which are represented by circles, and center, are denoted by ($\bullet$). In case A, the centers are located in the center of each cluster, and the algorithm converged into correct results. In case B, one center is located between two circles, while two centers are located in one circle, and the algorithm converged into false results. In case C, the algorithm converged into two clusters, and hence one cluster is empty.

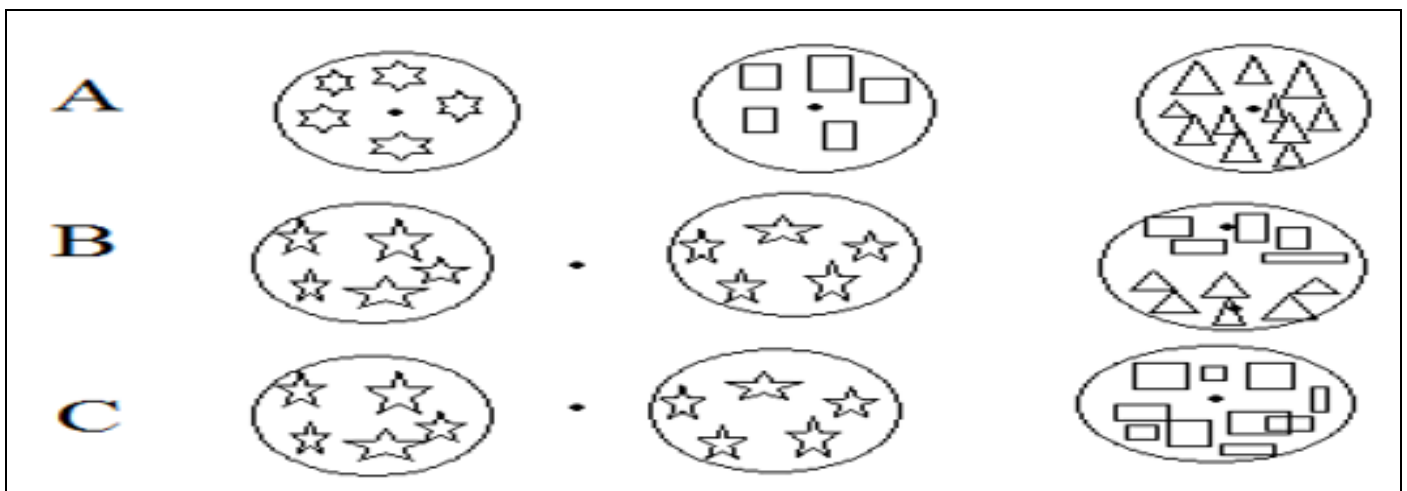


Fig 5 Three k-means for one Data set.

Therefore, it clear that k-means algorithm is very sensitive to the initial centers. Thus it may not converge or yield to the accurate results (global solution). Besides, they need foreknowledge of the number of clusters, which is difficult issue in real time processing.

- **Kernel K-means.**

Given a number of k prototypes, the K-means clustering algorithm looks for clusters.

$$J(\{P_l\}_{l=1}^k) = \sum_{l=1}^k \sum_{u_i \in P_l} \|u_i - m_l\|^2, \text{ where } m_l = \sum_{u_i \in P_l} u_i / |P_l| \quad (8)$$

That minimizes the cost function m_l is the mean of the l th center, and u_i is defined before. One of the main drawbacks of the K-means algorithm is an inability to find

optimal clusters that are nonlinearly separable in the input space. To overcome this drawback, we use Kernel method as follows. The objective function of the Kernel K-means can be represented as;

$$J(\{P_l\}_{l=1}^k) = \sum_{l=1}^k \sum_{u_i \in P_l} \|\Phi(u_i) - m_l\|^2, \text{ where } m_l = \sum_{u_i \in P_l} \Phi(u_i) / |P_l| \quad (9)$$

The norm Euclidean distance $\|\Phi(u_i) - m_l\|^2$ may be represented as;

$$= \Phi(u_i) \cdot \Phi(u_i) - \frac{2 \sum_{u_j \in P_l} \Phi(u_i) \cdot \Phi(u_j)}{|P_l|} + \frac{\sum_{u_j \in P_l} \Phi(u_j) \cdot \Phi(u_j)}{|P_l|^2} \quad (10)$$

Table 1 The Kernel K-means Algorithm

=Kernel K-means

Input:

, , = Kernel Matrix, number of the prototypes, and maximum number of iterations, respectively

Output:

=Assignment of the data samples to their clusters.

- 1- Initialize the clusters randomly. set .
- 2- Compute the distance .
- 3- Update the clusters according to the computed distance. .
- 4- If stop, otherwise go to step 2.

➤ **Soft Clustering**

- **Fuzzy C-Means (FCM)**

iffCM algorithm is used to partition n samples to number of known clusters (C). It offers a degree of membership (μ_{ji}) between every sample and cluster , with the cost function is given by

$$J = \sum_{j=1}^C \sum_{i=1}^N (\mu_{ji})^l \|x_i - v_j\|^2 \quad (11)$$

Where $l(l > 1)$ is a fuzzz function constant, $\sum_{i=1}^N \mu_{ji} = 1$, and v_j is the jth clustering center. We can minimize the cost function

$$v_j = \frac{\sum_{i=1}^N (\mu_{ji})^l x_i}{\sum_{i=1}^N (\mu_{ji})^l} \quad (12)$$

$$\mu_{ji} = \frac{(1 / x_i - v_j^2)^{1/l-1}}{\sum_{k=1}^C (1 / x_i - v_k^2)^{1/l-1}} \quad (13)$$

Then FCM algorithm will be calculated by setting the number of clusters, initializing V_j , and reiteratively solving (12), and (13).

- **KFCM.**

KFCM is an improvement of FCM that map the data points from input space into kernel space. The objective function of the Kernel k-means can be represented as

$$J(\{P_l\}_{l=1}^k) = \sum_{l=1}^k \sum_{i=1}^N \beta_{il}^h \|\Phi(u_i) - \Phi(m_l)\|^2 \quad (14)$$

Where $\Phi(m_l)$ and $\|\Phi(u_i) - \Phi(m_l)\|^2$ are defined previously.

IV. EVALUATION OF SORTING ALGORITHMS

In this Part, we are going to compare some sorting algorithms, which introduced before. It is worth nothing that, each of these algorithms might combine different types of clustering in order to enhance the efficiency of the sorting.

➤ *Dynamic Distance Clustering (DDC) and (IDDC) Algorithms:*

DDC belongs to partition clustering, and its aim to sort radar pulses without prior knowledge of radars' number. It is mainly based on the minimum distance principle. The steps of this algorithm are summarized as follows:

- Choose any sample, usually the first one, to be the first cluster center, z_1
- Compute the Euclidean distances between all samples and the obtain cluster center, and then select the sample, which corresponding to the maximum distance as the next cluster center, z_2 .
- Calculate the threshold based on

$$\lambda = \theta \|z_1 - z_2\| \quad (15)$$

$$\theta_i = \frac{\alpha_i \max(y_i)}{\max(y_i)(1 - \alpha_i) - \min(y_i)} \quad (16)$$

$$\theta = \left(\sum_i \theta_i^2 \right)^{0.5} \quad (17)$$

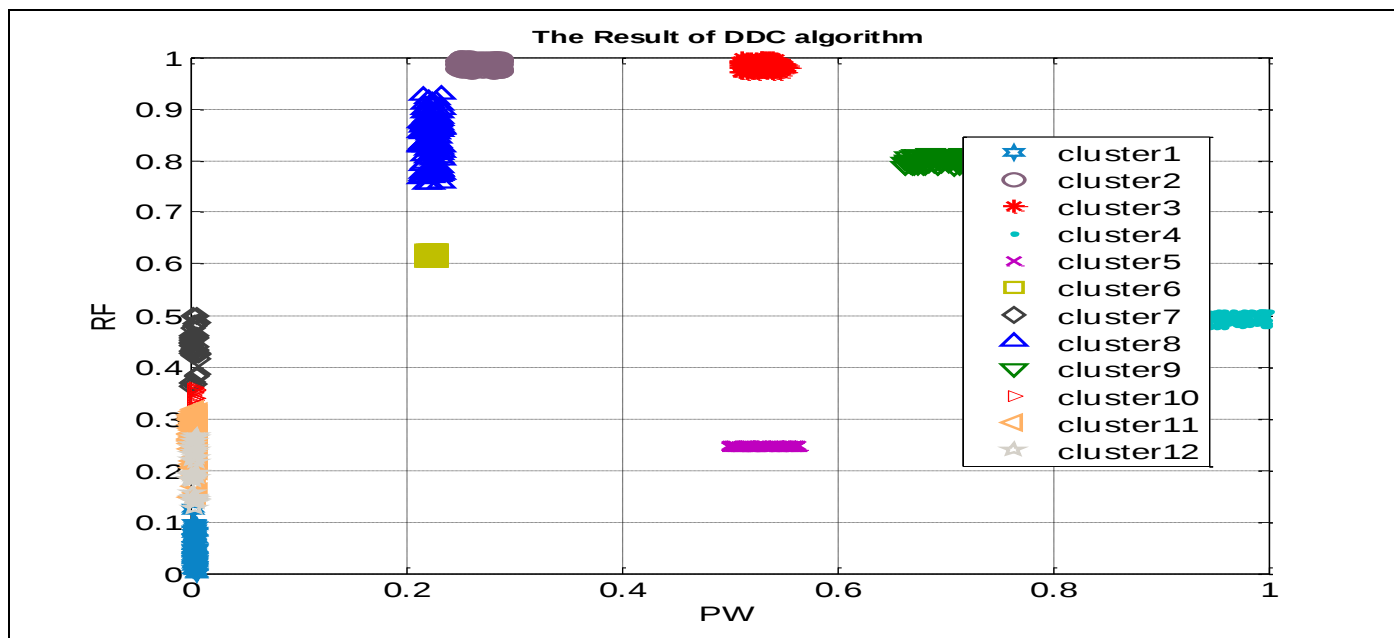
- Compute the minimum distance between all samples and the obtained centers, and then if the maximum distance, $D_{\max}$, satisfy $D_{\max} > \lambda$, then the corresponding sample will be considered as a next cluster center.
- Continue the previous step until $D_{\max} > \lambda$, and then all centers will be estimated.
- Set each sample with the nearest center.
- Combine the adjacent clusters which have too less pulses (samples).

In order to verify the validity of sorting algorithm, nine radars' pulse signals which are mixed according to the TOA are simulated. Considering radar signals from the same direction, we use four-dimensional clustering. The clustering parameters include RF, PW, modulation slope k for linear frequency modulation (LFM) and bit rate R for phase shift keying (PSK). Taking into account measurement error inevitably, random quantities are added to the parameters in simulation. Set the bias of RF to 2% or less, and set the bias of PW, k and R to 10% or less. The variation range of stable PRI and staggered PRI is 1% to 3% of the mean PRI value, and the variation range of jittered PRI is 5% to 10% of the mean value. The preset parameters of the nine radars are shown in Table 2.

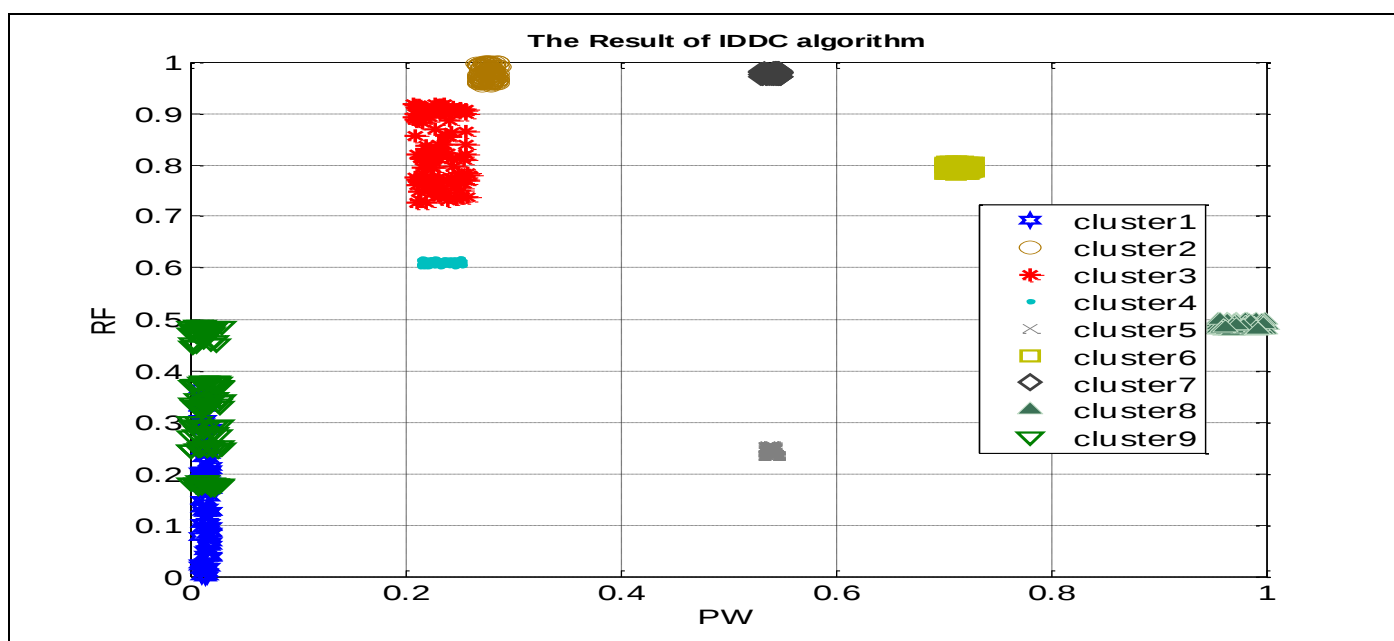
To improve the DDC method, we combine the last four steps together into one step, that is, the cluster centers are estimated with their corresponding samples, simultaneously. Table 2 shows nine radar signals are simulated at clustered with DDC and IDDC algorithm and the result depicted in Fig.6(a), and (b), respectively.

Table 2 The DDC Method Radar Parameters

No	Type	RF (MHz)	PW (μS)	PRI(μS)	K	R
1	LFM	2200-2800 A(32)	40	3400(F)	50	0
2	PSK	3400-3700 A(16)	65	2000(F)	0	2
3	Monopulse	3200	65	3400(J)	0	0
4	PSK	2600	100	1100,1200(S)	0	2.5
5	Monopulse	3500	120	900,850,1200(S)	0	0
6	LFM	3800	100	2500(F)	150	0
7	PSK	3800	70	38(J)	0	5
8	Monopulse	3000	150	4000(F)	0	0
9	LFM	2400-3000 A(16)	40	1300,1100, 1600(S)	100	0



(a)



(b)

Fig 6 (a) DDC Clustering Results. (b) IDDC Clustering Results.

It is obvious to see that, the DDC estimated twelve clusters instead of nine, whereas IDDC obtained the correct number of clusters, which indicate the superiority of IDDC.

➤ *Tolerance Threshold Clustering (TTC) algorithm.*

The concept of tolerance in this method means the allowable level in each parameter that can correspond to some cluster [10]. More simplicity, consider Fig.7 which consist of one cluster and some noise pulses. Suppose the center of the cluster is denoted by c_1 . Then the tolerance

become $\left(c_1 \pm \frac{\Delta}{2} RF, c_1 \pm \frac{\Delta}{2} PW \right)$.

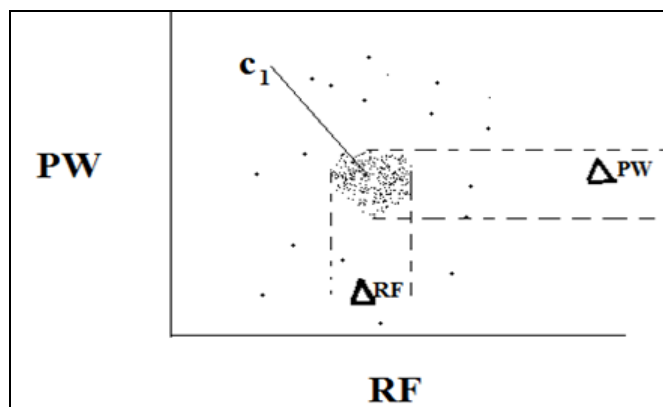


Fig 7 The Concept of Tolerance in Clustering

Therefore, if we have good estimate of this tolerance (threshold), then we can obtains our cluster. The algorithm is work as follows:

- Select arbitrary point as the first center, c_j , usually the first one.
- Compute the distance between c_j and the other pulses.

$$\begin{cases} \text{if } d(c_j, y_i) \leq Th & y_i \in c_j \\ \text{Else} & y_i \notin c_j \end{cases} \quad (18)$$

- Update the center when a new sample belong to it as

$$c_j = avg(c_j + y_i) \quad (19)$$

- Discard all samples which belong to the current cluster
- If the remaining number of sample larger than 5, go to first step. Else, break the algorithm.

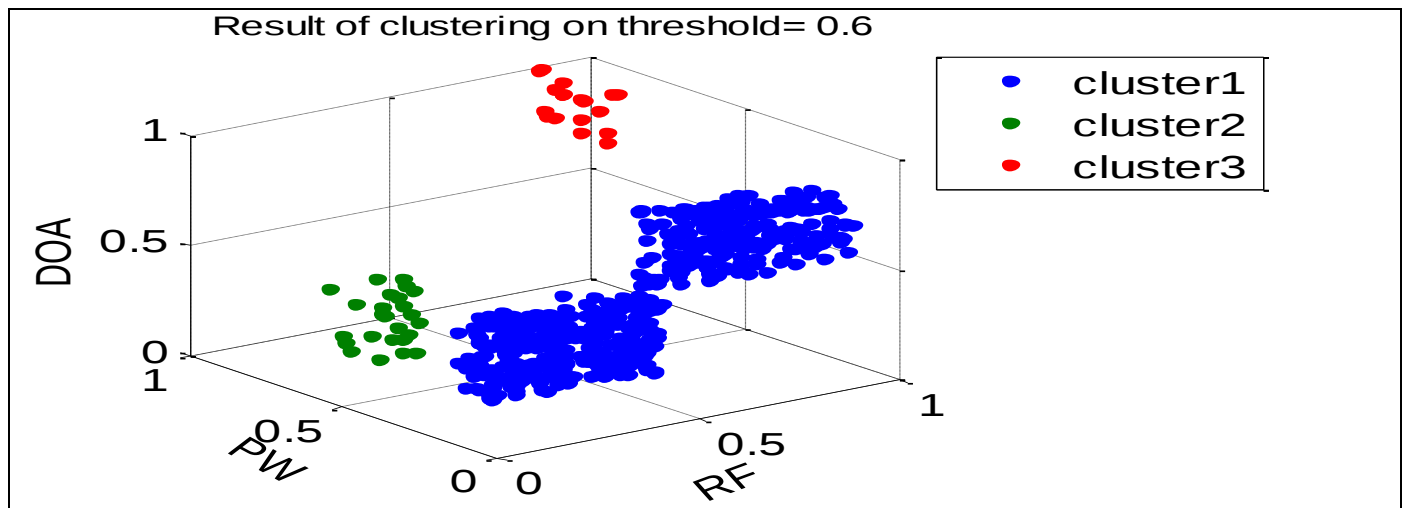
Let consider the following radar pulses which distributed as shown in Table 3.

Table 3 The Radar Parameters

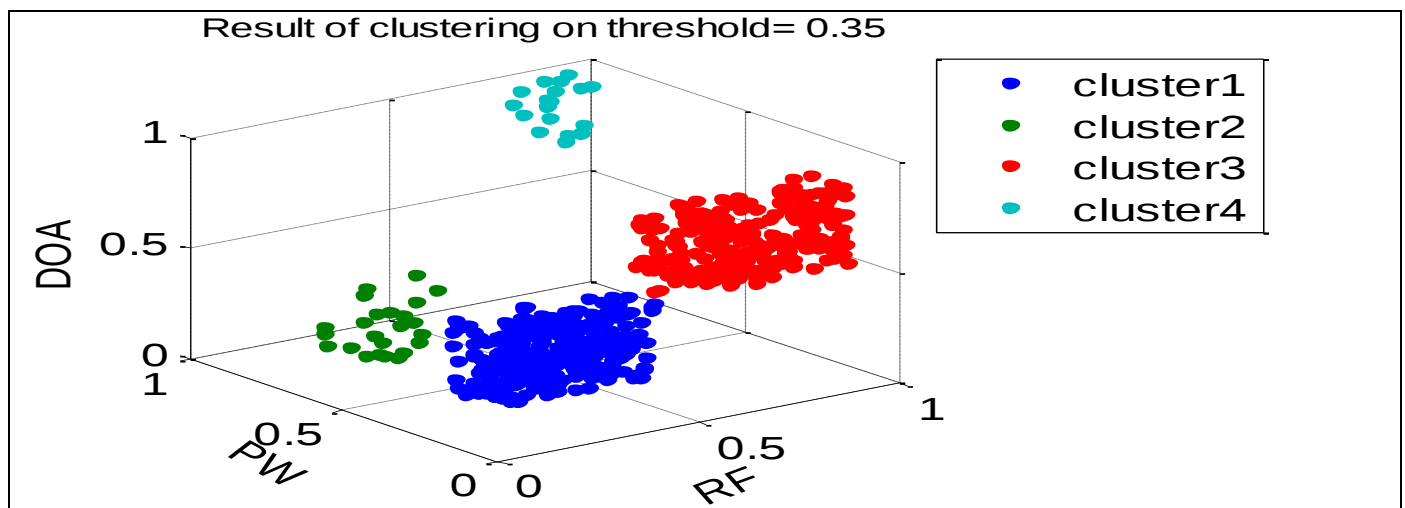
radar	PRI (μ s)	RF (GHz)	PW	DoA	Tot
1	3~4	8.9~9.3	0.3~0.4	38~41	130
2	30~60	9.1~9.3	0.8~0.5	36~39	26
3	4~5	9.4~9.8	0.4~0.5	40~43	130
4	40~70	9.6~9.8	0.9~1	42~45	18

Fig.7 shows the results of clustering using TTC algorithm in different threshold to indicate the importance of the threshold criteria. In addition, the selection of the first

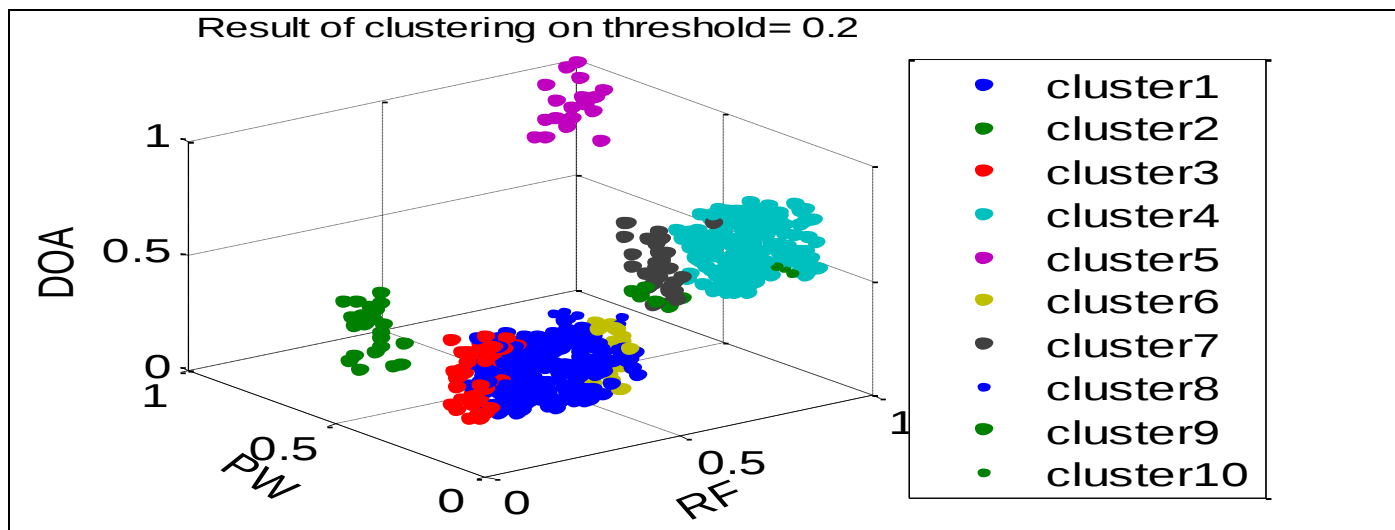
class center has little effect on the entire performance of both approaches. It indicates that DDC and IDDC are not sensitive to the sequence of input data.



(a)



(b)



(c)
Fig 8 The Clustering Results for Different Threshold

In Fig.8(a) the threshold is very large, i.e Th=0.6, and hence two clusters are merged into one cluster. As the threshold decreases the performances of clustering become well Th=0.35 as depicted in Fig.7 (b). In Fig.7 (c) the threshold is very small, and thus large numbers of clusters are obtained. Accordingly, if the threshold is very big, then all samples will be clustered into one cluster, while when the threshold is very small, then every sample may be considered as a unique cluster.

➤ *SVC and k-means Algorithm [8].*

SVC method is used to cluster data points with nonlinear boundaries in data space. Its main idea is to map

data samples from low dimension feature space to high dimension feature space by nonlinear transformation. The common nonlinear transformations are shown in Table 4.

Let, $u_i \subseteq U$, $i=1, \dots, N$ be a radar pulse chain consisting of N pulses with $U \in \mathfrak{R}^d$, where d is the features' dimension. Applying nonlinear transformation Φ from U to some high dimension space, the clusters take a far better form. We look for the smallest sphere which comprises almost all the data samples.

Table 4 The Common Nonlinear Transformation

Gaussian Kernel	$K(u_i, u_j) = \exp(-h \ u_i - u_j\ ^2)$
Polynomial Kernel	$K(u_i, u_j) = (u_i \cdot u_j + c)^2$
Sigmoid Kernel	$K(u_i, u_j) = \tanh(c(u_i \cdot u_j) + \theta)^2$

The enclosing sphere of radius R which contains all data samples can be represented by

$$\|\Phi(u_i - a)\|^2 \leq R^2 \quad (20)$$

Where $\|\cdot\|^2$ is Euclidean norm distance and a is the sphere's center. The soft constraint can be obtained by adding slack variables η ($\eta \geq 0$).

$$\|\Phi(u_i - a)\|^2 \leq R^2 + \eta_j \quad (21)$$

To solve these constraints we apply the Largangian,

$$L = R^2 - \sum_j \left(R^2 + \eta_j - \|\Phi(u_i - a)\|^2 \right) \alpha_j - \sum_j \eta_j \gamma_j + C \sum_j \eta_j \quad (22)$$

Where $\alpha_j \geq 0, \gamma_j \geq 0, \alpha_j \geq 0, \gamma_j \geq 0$ are the Largangian multipliers, C is a constant, and $C \sum_j \eta_j$ is a penalty factor. Minimizing L with respect to R, and a respectively lead to

$$\sum_j \alpha_j = 1 \quad (23)$$

$$a = \sum_j \alpha_j \Phi(u_i) \quad (24)$$

$$\alpha_j = C - \gamma_j \quad (25)$$

The Karush-Kuhn-Tucker (KKT) conditions are

$$\left(R^2 + \eta_j - \|\Phi(u_i - a)\|^2\right) \alpha_j = 0 \quad (26)$$

$$\eta_j \gamma_j = 0 \quad (27)$$

The first equation of KKT implies that, the image $\Phi(u_i)$ with $\alpha_j > 0$ and $\eta_j > 0$ locates out of the sphere. While the second condition of KKT implies $\gamma_j = 0$. So, if $\alpha_j = C$, then the image $\Phi(u_i)$ locates outside the boundary of the sphere, which is known as outliers. If $0 < \alpha_j < C$, then the image $\Phi(u_i)$ locates on the boundary surface of the sphere, known as Support Vector (SV). The reminder points locate inside the sphere, i.e. $\alpha_j = 0$.

Using the above relations, we can write the Lagrangian formula in wolf dual form as;

$$W = \sum_j \Phi(u_i)^2 \alpha_j - \sum_{i,j} \alpha_i \alpha_j \Phi(u_i) \Phi(u_j) \quad (28)$$

Where $0 < \alpha_j < C$. In this section we used Gaussian Kernel which is defined in Table 4-4, thus the Lagrangian W can be written as;

$$W = \sum_j K(u_j, u_j) \alpha_j - \sum_{i,j} \alpha_i \alpha_j K(u_i, u_j) \quad (29)$$

The distance from each image's sample u to sphere's center a can be represented as

$$R^2 = \|\Phi(u_i - a)\|^2 \quad (30)$$

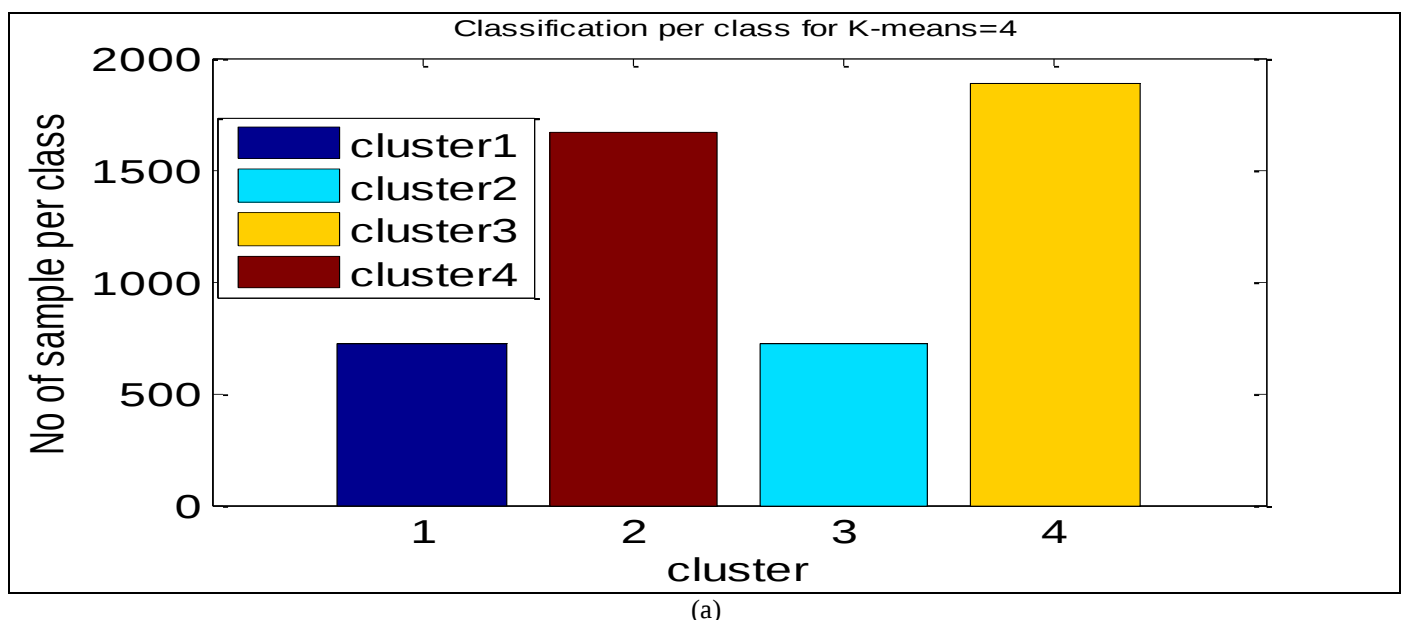
The radius of the sphere is $\{R = R(u_i) | u_i \text{ is support vector}\}$. The tightness of the boundaries is controlled by two parameters, h and C , while the number of outliers are governed by C .

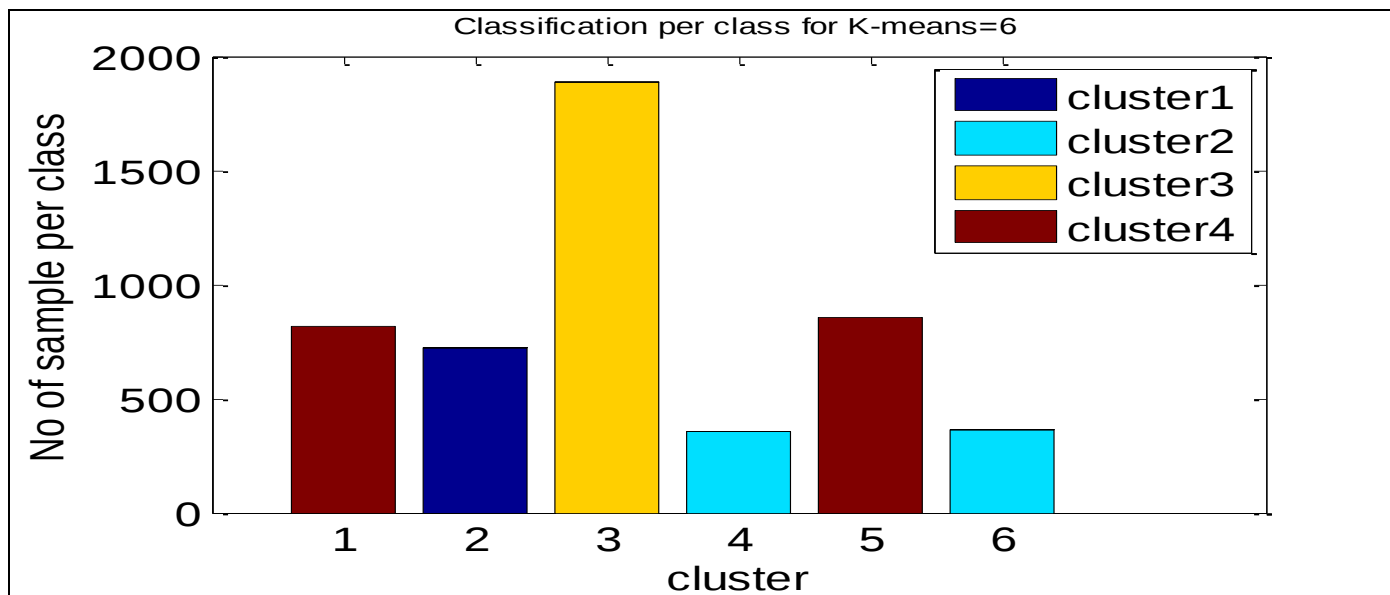
To distinguish the pair of data samples are belonged to same cluster or not, the following steps should be considered. First, connect the path between the pair of points in the feature space. Next, divide the line path to segment of points, z . Finally, if $R(z) \leq R$, then the pair of points belong to same cluster. Otherwise, they belong to different clusters.

The main idea of SVC & K-means is to partition our data into small parts, in order to decrease the time computation of SVC. Then we apply K-means into each small segment. In order to illustrate the performance of SVC & k-means, four radar signals are simulated as shown in Table 5.

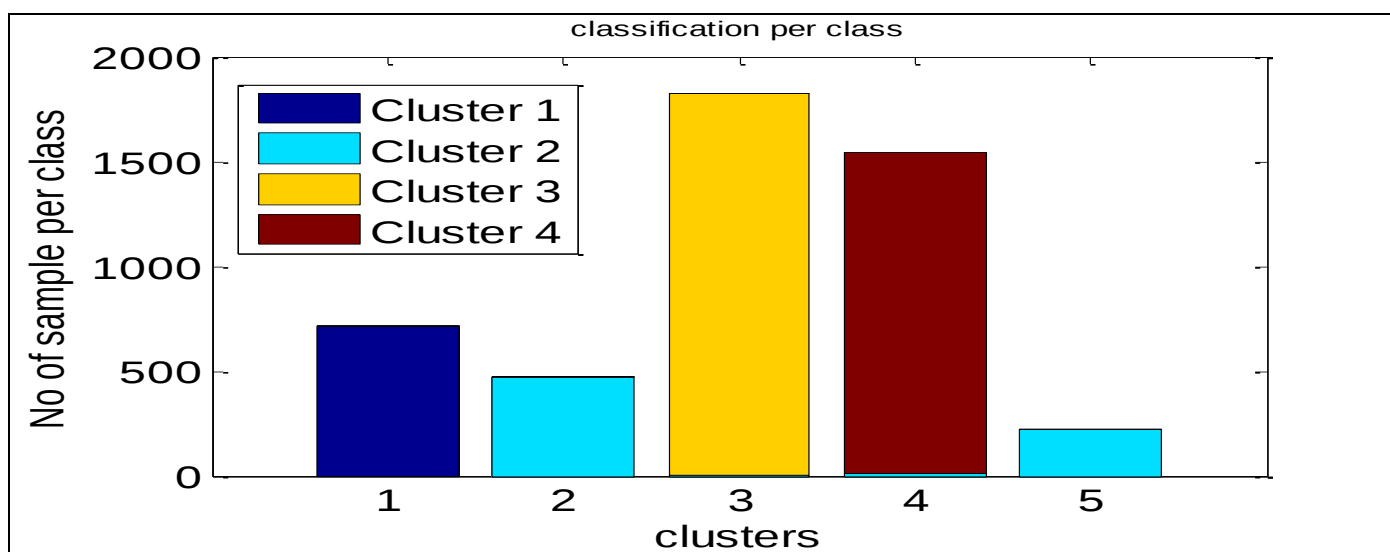
Table 5 The Radars' Parameters

Radar	RF(GHz)	PW(μ s)	DOA	No.of pulses
1	2.08~2.25	1.2~1.3	48~50	824
2	2.75~2.85	1~1.1	60~65	823
3	2.25~2.35	1.2~1.25	68~70	2149
4	2.22~2.75	1.3~1.4	56~60	1891





(b)



(c)

Fig 9 (a) Three clusters Clustering result with K-means=4. (b) Clustering result with K-means=6. (c) Clustering result for with SVC and K-means.

In Fig. 9(a) we cluster the received signal using k-means, and we set k-means signal to 4 clusters. It is clear detecting that, all radar signals are clustered well. Whereas, in Fig. 9(b) we set the number of cluster to be 6, resulting that, the algorithm converged into six clusters. Since k-means is largely depend on two factors, that is, the number of clusters and the position of centers. Finally, these information conducted by using SVC and hence we clustered the signal using k-means. The method worked well; however, this may lead to false estimation, since small part of signal cannot always reflect the actual number of clusters as depicted in Fig 9(c).

V. CONCLUSION

In this paper we presented a comparison between hierarchical and partition clustering methods in order to obtain the amenable algorithm for sorting radar signal. It is clear that all compared algorithms have good performances when the initial centers are inherent in the problem. Additionally, we observe the hard algorithms are highly sensitive to the initial values as compared with soft algorithms. The initiation, missing pulses, and noise pulses all are considered for the comparison. Next, a new Fuzzy c-means validity index is proposed and the simulation results show the efficiency of this method in time consuming and resistance to noise. Additionally, an improved density clustering algorithm for sorting radar signal is proposed to improve the time processing of the previous version. It is also compared with the other benchmark density clustering algorithms.

REFERENCES

- [1]. Ahmed Abdalla Ali. Mohammed Ramadan. Yongjian Liao . Shijie Zhou , “An adaptive filtering algorithm in pulse-Doppler radar for counteracting range-velocity jamming”, International Journal of Electronics, 109(10),pp 1695–1713, 2022.
- [2]. Ahmed Abdalla Ali., Wang, W. Q., Yuan, Z., Mohamed, S., & Bin, T. “Subarray-based FDA radar to counteract deceptive ECM signals”, EURASIP Journal on Advances in Signal Processing, 104, 1–11,2016.
- [3]. Ahmed Abdalla Ali., Yuan, Z., & Bin, B. “ECCM schemes in netted radar system based on temporal pulse diversity”. Journal of Systems Engineering and Electronics, 27(5), 1001–1009, 2016.
- [4]. Ahmed Abdalla Ali., Shokrallah, A. M. G., Yuan, Z. H. A. O., Ying, X., & Bin, T. A. N. G. “ Deceptive jamming suppression in multistatic radar based on coherent clustering”, Journal of Systems Engineering and Electronics, 29(2), 269–277, April 2018.
- [5]. Rokach, Lior, and Oded Maimon. "Clustering methods." Data mining and knowledge discovery handbook. Springer US, 2005. 321-352.
- [6]. Murtagh, F. A survey of recent advances in hierarchical clustering algorithms which use cluster centers. Comput. J. 26 354-359, 1984.
- [7]. Chen, Long, CL Philip Chen, and Mingzhu Lu. "A multiple-kernel fuzzy C-means algorithm for image segmentation." IEEE Transactions on Systems, Man, and Cybernetics, Part B: Cybernetics, 41.5: 1263-1274,2011.
- [8]. He, Ai-Ling, et al. "Multi-Parameter Signal Sorting Algorithm Based on Dynamic Distance Clustering." Journal of Electronic Science And Technology of China7.3: 249-253,2009.
- [9]. D. Xu, M. Xu, H. Wang, F. Feng, L. Tang and M. Gu, "A Real-time Radar Signal Sorting Method and Implementation Based on DSP," 2020 IEEE MTT-S International Wireless Symposium (IWS), Shanghai, China, 2020, pp. 1-3, 2020.
- [10]. Z. Cui, X. Fu, P. Lang, J. Dong, F. Wu and H. Gao, "Radar Signal Sorting Based on Adaptive SOFM and Coyote optimization," 2022 7th International Conference on Signal and Image Processing (ICSIP), Suzhou, China, 2022, pp. 157-161,2022.
- [11]. M. Wan, Y. Zhang, Y. Bai, Y. Sun, Q. Yu and Q. Wang, "A Real-Time Radar Signal Sorting Method Under Bayesian Framework With Dynamic Cluster Merging," in IEEE Sensors Journal, vol. 24, no. 17, pp. 27859-27869, 1 Sept.1, 2024.
- [12]. Z. Zhou, X. Fu, J. Dong and M. Gao, "Radar Signal Sorting With Multiple Self-Attention Coupling Mechanism Based Transformer Network," in IEEE Signal Processing Letters, vol. 31, pp. 1765-1769, 2024.
- [13]. Zhizhong Zhang, Xiaoran Shi, Xinyi Guo, Feng Zhou, "TR-RAGCN-AFF-RESS: A Method for Radar Emitter Signal Sorting", Remote Sensing, vol.16, no.7, pp.1121, 2024.