



IJEAST

INTERNATIONAL JOURNAL
OF ENGINEERING APPLIED SCIENCE
AND TECHNOLOGY



VOLUME : 1 ISSUE : 6 Print / Issue Publication Date: 09-Oct-2016



ISSN : 2455-2143



Indexed In



WWW.IJEAST.COM

editor@ijeast.com



DIMENSIONALITY REDUCTION - A STUDY

Dr. S.Vijayarani

Assistant Professor,
Department of Computer Science,
School of Computer Science and Engineering,
Bharathiar University, Coimbatore

Ms. S. Maria Sylviaa

M.Phil Research Scholar,
Department of Computer Science,
School of Computer Science and Engineering,
Bharathiar University, Coimbatore

Abstract: Data mining is the task of exploring new pattern from the huge database. Dimensionality reduction or dimension reduction (DR) is the measure of condensing the number of attributes under required condition. It is of two types: they are feature extraction (FE) and feature selection (FS). FE alters the high dimensional space into fewer dimensions whereas, FS captures a subset of compatible variables or attributes. Many numbers of algorithms are used for dimensionality reduction. Nowadays DR is used in number of applications for getting efficient results. The main objective of this paper is to provide the complete study about the dimensionality reduction concepts,, Basic methods, Techniques and applications.

Keywords: Dimensionality reduction, Feature selection, Feature extraction, Methods, Techniques, Applications.

I. INTRODUCTION

Dimensionality reduction is an important task in machine learning which facilitates clustering, classification, compression, and visualization of high-dimensional data by mitigating undesired properties of high-dimensional spaces [1]. The problem with the high dimensional data is that, it uses number of features for performing various techniques like classification, clustering, association rule, etc. For these techniques some of the features are not important, hence dimensionality reduction is used in order to reduce the features of the data without the loss of any data. Dimensionality reduction can be beneficial not only improving the computational efficiency but also it improves the accuracy of the analysis [2]. Dimension reduction has a long history as a method for data visualization and for extracting low dimensional features from higher dimensions [3]. There are two different techniques in dimensionality reduction: they are feature extraction and feature reduction. DR has been used for number of purposes like projection of data, compression and noise removal. Removing uninformative or dis-informative data columns might indeed help the algorithm to find more

general classification regions and rules to achieve better performances on new data [5]. Statistical and machine reasoning methods faced a formidable problems when dealing with such high dimensional data and normally the number of input variables is reduced before a data mining algorithm can be successfully applied [4]. The motivation of DR is to analyze the diminished set of features which predict the result. In machine learning process the time for training the dataset expands the number of features.

The remaining portion is organized as follows: Section 2 tells about the various methods used for dimensionality reduction, Section 3 defines about the techniques used in dimensionality reduction, Section 4 illustrates the dimensionality reduction algorithms and Section 5 defines about the various fields where the reduction is used. Conclusion is given in section 6.

II. METHODS OF DIMENSIONALITY REDUCTION

There are seven different methods that have been applied in the data analytical space. They are illustrated in the Figure1. In most of the times the huge amount of data in data analytical process does not work well. So dimensionality reduction is applied to the large data. The methods of dimensionality reduction define about the reduction of data elements illustrated in Figure.1.

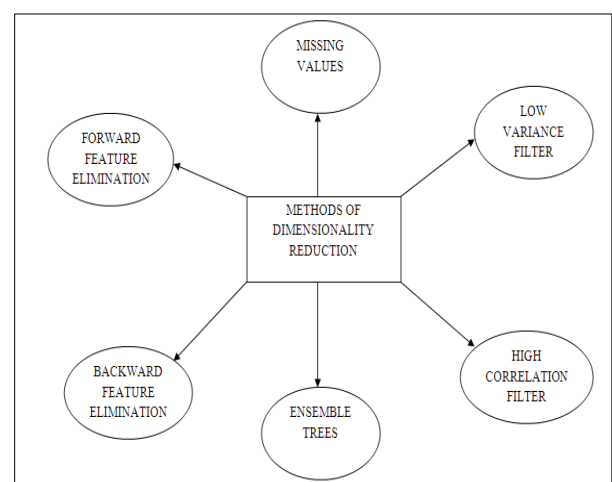




Fig. 1. Methods of Dimensionality Reduction

A. Missing values

The data column with enormous amount of missing values is calculated and then they are removed. The numbers of missing values are calculated using either statistical node or group node. The statistical nodes are used for the analysis of data and the group node is used in the techniques of DR. The missing values larger than the threshold are collected and then they are removed. If the threshold value found is larger than the original values then the reduction will be aggressive. The missing values are calculated using the formula which is shown in Figure.2. In KNIME tool the missing value node calculates the ratio of the missing values is calculated by number of missing values divided by the total number of rows.

$$\text{RATIO OF MISSING VALUES} = \frac{\text{NUMBER OF MISSING VALUES}}{\text{TOTAL NUMBER OF ROWS}}$$

Fig. 2. Formulae for Missing Values

B. Low Variance Filter

The variance of the data is calculated to find the number of information about the data column. In the limit case where the column cells assume a constant value, the variance would be 0 and the column would be of no help in the discrimination of different groups of data [5]. The data column lower than the threshold values of the data are filtered and then they are removed. The filter is applied to the data to find the low variance. First normalization is done to all the columns to find the lower variance.

C. High Correlation Filter

The feature which are given as input are often correlated i.e. the features are dependent to one another and they too have same information. A data column with values highly correlated to those of another data column is not going to add very much new information to the existing pool of input features [5]. To reduce the correlated column, the correlation is measured using the linear correlation node between couple of columns. If any highly correlated column is present between the pair the particular data is removed. Correlation matrix is taken as input to the correlation filter. Filtering highly correlated data columns requires uniform data ranges again which can be obtained with a Normalize node [5].

D. Ensemble Trees

Ensemble trees are also known as random forest. For effective classification feature, selection is

done. One approach to dimensionality reduction is to generate a large and carefully constructed set of trees against a target attribute and then use each attribute's usage statistics to find the most informative subset of features [5]. The score is calculated using level of the candidates and it defines about the relative attributes which are most predictive. A shallow tree has been generated for ensembling and here each tree is trained on the fraction of all attributes. If the particular attribute is selected as the best split one then the informative features are retained.

E. Backward Feature Elimination

The Backward Feature Elimination loop performs dimensionality reduction against a particular machine learning algorithm [5]. It is a simple iterative method and in each method the selected classification algorithm is performed for number of input features. Then one input feature is removed and the model is trained for (n-1) input features for number of times. The input feature with large number of error rate, after the removal of feature is recognized and they are eliminated. Then they are repeated for number of iterations like n-2, n-3, etc to find the larger error rate. The Backward Feature Elimination Filter finally visualizes the number of features that are kept at each iteration and the corresponding error rate [5]. This method can only be applied to the small number of dataset.

F. Forward Feature Elimination

Similarly to the Backward Feature Elimination approach, a Forward Feature Construction loop builds a number of pre-selected classifiers using an incremental number of input features [5]. The forward feature loop opens with one feature and other feature is added to it for every iterations. Both forward and backward are very expensive and computationally high. Running the optimization loop the best cutoffs in terms of lowest number of columns and best accuracy were determined for each one of the six dimensionality reduction methods and for the best performing model [6].

III. TECHNIQUES IN DIMENSIONALITY REDUCTION

There are two different major techniques used in dimensionality reduction. They are feature selection and feature extraction

A. Feature Selection

In machine learning and statistics, feature selection (FS), also known as variable selection or attribute selection or variable subset selection and it is the process of selecting a subset of relevant features (variables, predictors) for use in model construction [7]. This technique has three different reasons: They are Interpretation of model to make them efficient

and simple for the users, Less training time and Reduction of variance for strengthened generalization. Feature selection is applied to the domains with larger features and chooses the feature according to the objective function. A feature selection algorithm can be seen as the combination of a search technique for proposing new feature subsets along with an evaluation measure which scores the different feature subsets [7]. Every subset of the feature is tested separately to minimize the error or noise rate. Feature selection is classified into three different classes: they are embedded, filter and wrapper method. The classes of feature selection are shown in Figure.3.

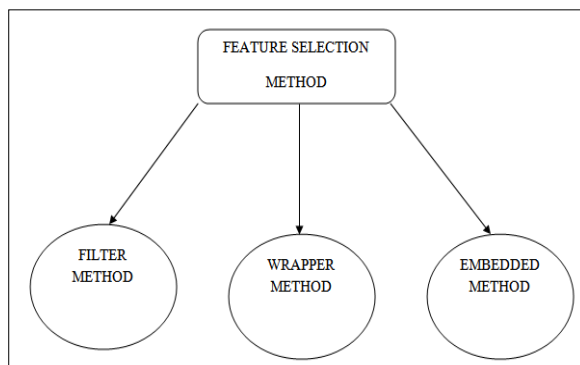


Fig.3. Feature Selection Methods

A.1 Filter Method

In filter method, the variables are selected based on the model used. They are based only on general features like the correlation with the variable to predict [7]. Filter method contains the minimum interesting variables. Other variables will be used for either classification or regression models. These methods work efficiently under time consumption. It is only used as a preprocessing method and they choose the redundant variable because they do not matter the relationship between the variable. The flow of this class is proposed in Figure4. In this all the features are taken and from that the best subsets are chosen and the algorithms for further tasks are applied to it to find the performance.

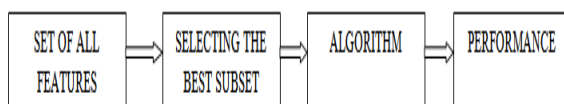


Fig. 4. Filter Method

A.2 Wrapper Method

Unlike the filter approach wrapper method calculates the subset of variables to encounter the interaction between the variable. There are two disadvantages of wrapper method, they are

- The number of observations are inadequate when the over fittings are increased.
- Larger the variable, the computation time is increased.

Here the subset feature is generated using the search method and the performance is calculated. Wrapper method uses search algorithm for selecting the subset. Alternative search-based techniques are based on targeted projection pursuit which finds low-dimensional projections of the data that score highly: the features that have the largest projections in the lower-dimensional space are then selected [7]. The flow of wrapper method is illustrated in Figure5.

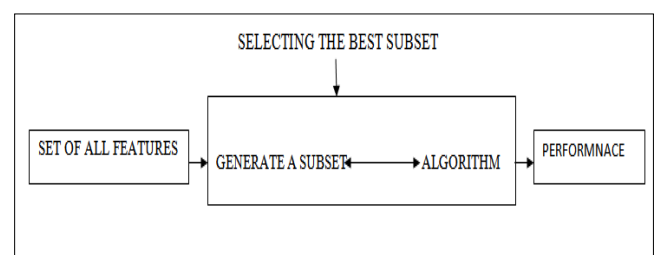


Fig.5. Wrapper Method

Search Approaches

- Exhaustive search
- Best first search
- Simulated annealing
- Genetic algorithm
- Greedy forward selection
- Greedy backward selection
- Scatter search
- Variable neighborhood search

A.3 Embedded Method

Recently, embedded methods have been proposed to reduce the classification of learning [7]. To form the embedded class, the advantages of filter and wrapper classes are used. Many embedded feature selection methods have been introduced during the last few years unifying theoretical framework has not been developed to date[8]. The learning algorithm takes their own variable selection algorithm. The flow is shown in Figure6. This method is varied from other feature selection method

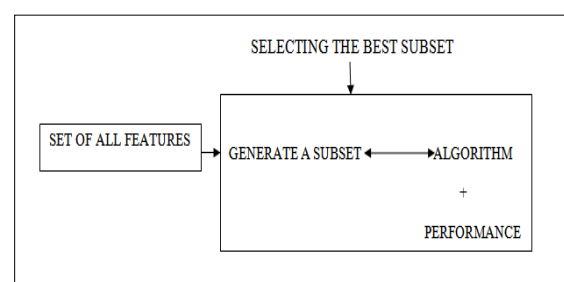




Fig.6. Embedded Method

The parameterized function has been shown in Figure7 which defines the function F to find the vector α , let X be the variable R defines the relation between the variable X and the vector α . Embedded method takes the set of all features as the input and selects the best subset by generating the subset to find the performance of the method. In this method the features are added or removed iteratively to find the subset.

$$F: R(\alpha, X) \rightarrow f(\alpha, X)$$

Fig.7. Formulae for parameterized function

B. Feature Reduction

In feature reduction, the data of high dimensional space are transformed into lower dimensional space. Transforming the data is done in linear method or non linear method. The main linear technique for dimensionality reduction performs a linear mapping of the data to a lower-dimensional space in such a way that the variance of the data in the low-dimensional representation is maximized. In order to reduce the features the correlation coefficient is calculated by finding the Eigen vectors and Eigen values. In this model the relevant features for the class is selected by eliminating the redundant features. The feature reduction is done with four aspects which are defined in Figure.10.

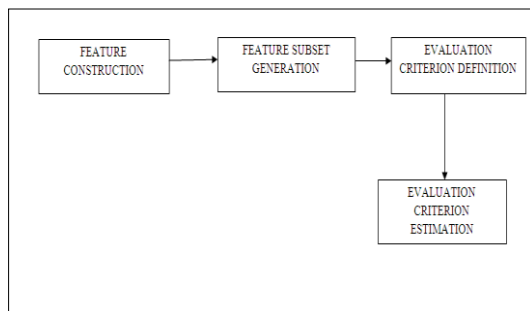


Fig. 8. Aspects of Feature Extraction

First, the features are taken/constructed as input and then the subset of the features are extracted from the original features in which the redundant features are eliminated and the similar features are extracted using the algorithms such as PCA, SVD, LDA, etc. Then the filters use criteria not involving any learning machine for example a relevance index based on correlation coefficients or test statistics. Then the filters are used to estimate the feature value.

IV. ALGORITHM

Many different algorithms are used for performing in dimensionality reduction process.

A. PRINCIPAL COMPONENT ANALYSIS (PCA)

Principal component analysis (PCA) is a statistical procedure that uses an orthogonal transformation to convert a set of observations of possibly correlated variables into a set of values of linearly uncorrelated variables called principal components [10]. The count of principal component is lower or equal to the original features. The Eigen vector, Eigen value and covariance matrix is found in order to perform PCA.

$$Z = \sum_{i=1}^n a_i X_i$$

In which X_i is the i^{th} variables where $i=1, 2, 3, 4, p$ is the linear coefficient of Z which is denoted in matrix and places as $a_1^T a_1 = 1$. Here, a_1 vector maximizes the variance and here Z is called the first principal component. The matrix $C = Cov(X)$ is the covariance matrix of X. The Eigen value λ_i is calculated by determining the equation.

$$\det(C - \lambda I)$$

The Eigen vector is defined to be the column of the matrix X. $C = A D A^T$. Where

$$D = \begin{bmatrix} \lambda_1 & 0 & 0 \\ 0 & \lambda_2 & 0 \\ 0 & 0 & \lambda_3 \end{bmatrix}$$

STEP 1: X Create $N \times d$ data matrix with one row vector x_n per data point
 STEP 2: X subtract mean from each row vector x_n in X
 STEP 3: $\Sigma \leftarrow$ Covariance matrix of X
 STEP 4: Find eigenvectors and eigenvalues of Σ
 STEP 5: Principal Component $\leftarrow$ the M eigenvector with largest eigenvalues.

Fig.9. PCA Algorithm

PCA algorithm is applied to large dataset for efficient result. The algorithm is shown in the figure.9. In this algorithm first takes the high dimensional data matrix as input which then calculates the mean for the data matrix and then it is subtracted with each row vector of the data matrix. Then the covariance is found to find the Eigen vector and Eigen value in which the highest



Eigen value column is taken as output and that features are extracted. The principal components are orthogonal because they are the eigenvectors of the covariance matrix which is symmetric and it is sensitive to the relative scaling of the original variables [10]. PCA is used as a tool to compose predictive model. The outcome of this algorithm is defined as factor score.

B. Singular Value Decomposition (SVD)

In linear model SVD is a type of factorization for the complex model. SVD is most only used in applications like signal processing and statistics. The SVD of an $a \times b$ complex or real matrix M is decomposition of the form $M = U\Sigma V^*$ where U is an $a \times a$ real or complex unitary matrix, Σ is an $a \times b$ rectangular diagonal matrix with non-negative real numbers on the diagonal and V^* (the conjugate transpose of V or simply the transpose of V if V is real) is an $b \times b$ real or complex unitary matrix and the diagonal entries $\Sigma_{i,i}$ of Σ are known as the singular values of M [11]. The a column and b column are called the left and right singular vectors of m . The X is found by multiplying $U^*D^*V_T$. The data set is organized into matrix and the data are normalized. The SVD is calculated by $X=UDV^T$. U is calculated by converting row into column matrix, D is the diagonal matrix of U . The transformation matrix X_T is multiplied on both the sides of U . Then the new coordinated axis is found by applying X_T on X and then the Singular value is decomposed.

STEP 1: Organize the dataset in a matrix X

STEP 2: Normalize the dataset using Z-score

STEP 3: Calculate the singular value decomposition of the data matrix $X=UDV_T$

STEP 4: Calculate the variance using the diagonal element of D

STEP 5: Multiply both sides with the transformation matrix X_T

STEP 6: Find the reduced projected dataset Y in a new coordinate axis by applying X_T on X

Fig.10. SVD Algorithm

Linear Discriminant Analysis (LDA)

Linear discriminant Analysis(LDA) is a generalization of Fisher's linear discriminant, a method used in statistics, pattern recognition and machine learning to find a linear combination of features that characterizes or

separates two or more classes of objects or events[21]. In this the class can be separated by discriminant analysis. Ronald A. Fisher formulated the *Linear Discriminant* in 1936 [22]. The discriminates can be found by coefficient. The discriminant function is used to calculate discriminant. In this, D_i is linear combination, X_k is predicted variable.

$$D_i = b_0 + \sum_{k=1}^p b_k X_k$$

First the matrix is taken as input for different classes then the scatter matrices are computed. The Eigen vectors and Eigen values are calculated for the scatter matrices. The Eigen vectors are sorted by decreasing the Eigen values then the samples are transformed into newer subspace by the formulae $Y=X \times W$ where X is $n \times d$ dimensional matrix is $d \times k$ dimensional matrix.

STEP 1: Compute the d -dimensional mean vectors for the different classes from the dataset.

STEP 2: Compute the scatter matrices (between-class and within-class scatter matrix).

STEP 3: Compute the eigenvectors ($e_1, e_2, \dots, e_d$) and corresponding eigenvalues ($\lambda_1, \lambda_2, \dots, \lambda_d$) for the scatter matrices.

STEP 4: Sort the eigenvectors by decreasing eigenvalues and choose k eigenvectors with the largest eigenvalues to form a $d \times k$ -dimensional matrix W (where every column represents an eigenvector).

STEP 5: Use this $d \times k$ eigenvector matrix to transform the samples onto the new subspace. This can be summarized by the equation $Y = X \times W$ (where X is an $n \times d$ -dimensional matrix; the i th row represents the i th sample, and Y is the transformed $n \times k$ -dimensional matrix with the n samples projected into the new subspace

Fig.11. LDA Algorithm

Multi-Dimensional Scaling (MDS)

Multidimensional scaling (MDS) is a classical approach for DR used to transform high dimensional space into lower space. MDS transforms the dimensional space by calculating the distance. MDS addresses the problem of constructing a configuration of t points in



Euclidean space by using information about the distances between the patterns [15]. To solve the problems of MDS two types of methods have been used. They are metric method and non-metric method. The metric method will produce the original metrics, whereas non-metric method will rank the values of the object. Here, the matrix is taken as input in the proximity and the double centering B is applied to the matrix dataset with the formulae $B = 1/2 JP^{(2)}$ and extract the largest Eigen vector value and the particular largest value is derived from the coordinate matrix.

STEP 1: Setup the matrix of squared proximities $P^{(2)} = [p^2]$

STEP 2: Apply the double centering $B = 1/2 JP^{(2)}$ using the matrix $J = I - n^{-1}11$, Where n is number of objects.

STEP 3: Extract the largest positive Eigen value $\lambda_1, \dots, \lambda_m$ of B and the corresponding m Eigen vectors $e_1, \dots, e_m$

STEP 4: A m-dimensional spatial configuration of n objects is derived from the coordinate matrix

$X = E_m \lambda_m^{1/2}$ where E_m is the matrix of m eigenvectors and λ_m is the diagonal matrix of m eigenvalues of B respectively

Fig.12. MDS Algorithm

V. DIMENSIONALITY REDUCTION IN VARIOUS FIELD

Dimensionality reduction is used in numerous field is as follows.

A. Data Mining

Dimensionality reduction is applied to high dimensional data in order to reduce the features. Features are the attributes of the data which is reduced by certain condition. In practice researchers and practitioners interchangeably use dimension, feature, variable, and attribute [15]. The reduced features are then used to perform data mining techniques. Removal of irrelevant data increases the accuracy of the data. Data dimensionality reduction (DDR) has become an important aspect of data mining since human experts and corporate managers are able to make better use of lower dimensional data compared to high dimensional ones [16]. When preprocessing of the data is done for dimension reduction, it must focus on describing the dataset using minimum number of attributes but it must give the

performance comparable to that of original dataset containing all the attributes [17]. The techniques have to be considered which performing retention like dependency and significance.

B. Image Processing

Pattern recognition technology in image is an essential branch in information processing, financial perspective and management perspective. It has broad applications and exploring prospects, it has already become a key research projecting pattern recognition and artificial intelligence [12]. Dimensionality reduction like Feature extraction is done to the high dimensional image to distort the image and the recognition rate of the image is reduced. Over-fitting in the training model of the image is recovered by feature extraction. Linear dimension reduction method performs dimensional reduction to high dimensional data through performance target looking for the linear transformation matrix, but since light, expression, posture and factors that make the high dimensional image that have obvious characteristics of nonlinear manifold [13]. Due to curse of dimensionality, high dimensional data leads to lower result. Class label context coming from the spatial configuration of images provides an additional source of classification information and therefore taking this contextual information into account can be beneficial [14]. There are different problems arise in image analysis like

- High dimensionality reduction problem in which it has dimension ranging from hundreds to thousands of factors and they are repeated in different point of time and space.
- Soft dimensionality reduction problem in which the data are not much high and have direct interpretation.
- Visualization problem is one of the techniques to represent the data. There are several techniques. When the time variable is added to the image data there arise two different types of problems like statistical dimensionality reduction and time dependent dimensionality reduction.

C. Network Logs

All networks and systems can be vulnerable to different types of intrusions [18]. When evaluating textual information the features and space/area grows larger and will be inadequate. Diffusion map is a manifold learning method that maps high-dimensional data to a low dimensional diffusion space [18]. The manifold learning method performs the Eigen decomposition of the network data in the matrix form. The textual information is converted into feature space for the reduction. The network



data are numerous because of mapping the textual form into the matrix for reduction. To find the coordinate of the system only few variables are needed to describe the data. The amount of log lines that needs to be inspected is reduced. Dimensionality reduction means reducing the complexity of data while preserving some quantity of interest [20]. This is useful for system administrators trying to identify intrusions. The number of interesting log lines is low compared to the total number of lines in the log file [18]. Component analysis methods takes the network data as input. The initial weight which is attached to the network is very high and fine tuning i.e. multi layer encode network is done to reduce the dimensions.

used. Further work is done by performing various dimensionality reduction algorithms and the techniques like clustering, classification, etc can be performed for comparative analysis

VI. REFERENCE

1. Laurens van der Maaten “An Introduction to Dimensionality Reduction Using Matlab”, MICC, Maastricht University.
2. G.N.Ramadevi and K.Usharani- “Study On Dimensionality Reduction Techniques And Applications” - 2- Vol 04, Special Issue01; 2013 -Publications Of Problems & Application In Engineering Research – Paper
3. Christopher J. C. Burges “Dimension Reduction: A Guided Tour-Foundations and TrendsR_ in Machine Learning” Vol. 2, No. 4 (2009) 275–365 _c 2010 DOI: 10.1561/22000000002.
4. C.O.S. Sorzan, J.Vargas and A. Pascual Montano “A survey of dimensionality reduction techniques”
5. Seven Techniques for Dimensionality Reduction-Open for innovative KNIME.
6. <http://www.kdnuggets.com/2015/05/7-methods-data-dimensionality-reduction.html>.
7. https://en.wikipedia.org/wiki/Feature_selection.
8. Thomas Navin Lal1, Olivier Chapelle, Jason Weston, and André Elisseeff - Learning with Local and Global Consistency-Max Planck Institute for Biological Cybernetics, Tubingen, Germany.
9. https://en.wikipedia.org/wiki/Dimensionality_reduction#Feature_extraction
10. https://en.wikipedia.org/wiki/Principal_component_analysis
11. https://en.wikipedia.org/wiki/Singular_value_decomposition
12. Dong Li “An Improved Algorithm for Nonlinear Dimensionality Reduction in Image Processing” The 2012 2nd International Conference on Circuits,

S.No	RESEARCH PAPER NAME	DATASET USED	TECHNIQUES USED		BEST TECHNIQUE \$ NAME
			EXISTING	PROPOSED	
1.	Study On Dimensionality Reduction Techniques And Applications[2]	Swiss roll dataset	PCA,LDA,LSI,CCA,PLS	-	Labeled data-PCA, Unlabeled data-LDA
2.	Learning with Local and Global Consistency[8]	Handwritten digits, newsgroups dataset	Normal SVM, Cluster Kernel , SLLGC	LLGC	LLGC
3.	An Improved Algorithm for Nonlinear Dimensionality Reduction in Image Processing[12]	ORL face database, CMUPIE face library	LDA,LLE	LLE+LDA	LLE+LDA
4.	A Review On Linear And Non-Linear Dimensionality Reduction Techniques[24]	Five different image databases (African people,Beach,Building, Bus, Dinosaurs)	LDA,PCA,ICA	-	ICA
5.	Dimensionality Reduction Framework for Detecting Anomalies from Network Logs[18]	Real life web service data	PCA, Diffusion map	-	PCA
6.	Dimensionality Reduction Of Multimodal Labeled Data By Local Fisher Discriminant Analysis[25]	Banana, breast cancer, flare-solar, german, heart, ring norm, thyroid titanic, waveform_USPS	PCA,LDI,LFDA, FDA	-	LFDA
7.	Analysis Of Unsupervised Dimensionality Reduction Techniques[26]	Medline,Cranfield, CACM and CISI datasets	SVD,NMF,ICA, FKM	-	FKM
8.	Principle Component Analysis and Partial Least Squares: Two Dimension Reduction Techniques for Regression[27]	BOP(business of business owners policies)data	PCA,PLS	-	PLS

Table 1. Comparative analysis on dimensionality reduction

VI. CONCLUSION

The main objective of this paper is to provide the overview of the dimensionality reduction. Before performing the reduction the estimation should be made for dimensions. Dimensionality reduction has been used in numerous fields for further techniques. It is used for providing the better result for the data analysis. The guidance should be given to the users handling dimensionality reduction in their data analytics. This paper discussed about the dimensionality reduction methods, techniques used, algorithms, fields where the reduction has been



- System and Simulation (ICCSC 2012) IPCSIT vol. 46(2012) © (2012) IACSIT Press, Singapore.
13. Sam Roweis, Lawrence Saul, Geoff Hinton. Global coordination of local linear models. *Advances in Neural Information Processing Systems* 2001, 14: 452-456.
 14. Marco Loog, Bram van Ginneken, Robert P.W. Duin "Pattern Recognition" Pattern Recognition Society. Published by Elsevier Ltd. All rights reserved. doi:10.1016/j.patcog.2005.04.011
 15. Manoranjan Dash, Huan Liu "Dimensionality Reduction"
 16. Xiuju Fu, Lipo Wang "Data dimensionality reduction with application to improving classification performance and explaining concept of dataset" - *Int.J.Business intelligence and data mining* vol.1.no.1, 2005.
 17. Mohini D Patil, Dr. Shirish S. Sane "Effective classification after dimension reduction: Comparative study" *International Journal of Scientific and Research Publications*, Volume 4, Issue 7, July 2014 1 ISSN 2250-3153
 18. Tuomo Sipola Antti Juvonen Joel Lehtonen- "Dimensionality Reduction Framework for Detecting Anomalies from Network Logs" *Engineering Intelligent Systems* 20 (1), 87-97- ©2012 CRL Publishing Ltd.
 19. Coifman, R.R., Lafon, S., Lee, A.B., Maggioni, M., Nadler, B., Warner, F., and Zucker, S.W. "Geometric diffusions as a tool for harmonic analysis and structure definition of data" *Diffusion maps In Proceedings of the National Academy of Sciences of the United States of America* volume 102, p. 7426 (2005)
 20. Prakash Balachandran "Dimensionality reduction of network data" SAMS Transition Workshop Monday, June 06, 2011.
 21. https://en.wikipedia.org/wiki/Linear_discriminant_analysis
 22. http://sebastianraschka.com/Articles/2014_python_lda.html
 23. Ali Ghodsi-Dimensionality Reduction A Short Tutorial- http://www.stat.washington.edu/courses/stat539/spring14/Resources/tutorial_nonlinear-dim-red.pdf
 24. Arunasakthi. K, 2kamatchipriya. L- A Review On Linear And Non-Linear Dimensionality Reduction Techniques- *Machine Learning And Applications: An International Journal (MLAIJ)* Vol.1, No.1, September 2014
 25. Masashi Sugiyama- Dimensionality Reduction Of Multimodal Labeled Data By Local Fisher Discriminant Analysis- *Journal Of Machine Learning Research* 8 (2007) 1027-1061
 26. Ch. Aswani Kumar- Analysis of Unsupervised Dimensionality Reduction Techniques- Networks and Information Security Division, School of Information Technology and Engineering, VIT University, UDC 004.423, DOI: 10.2298/csis0902217K
 27. Saikat Maitra and Jun Yan- Principle Component Analysis and Partial Least Squares: Two Dimension Reduction Techniques for Regression- *Casualty Actuarial Society*, 2008 Discussion Paper Program

IJEAST

INTERNATIONAL JOURNAL
OF ENGINEERING APPLIED SCIENCE
AND TECHNOLOGY

ABOUT IJEAST

International Journal of Engineering Applied Science and Technology (IJEAST) is a peer-reviewed, open access journal that publishes high-quality research papers in the field of Engineering, Applied Science and Technology.

IJEAST aims to provide a platform for researchers, academicians, and professionals to share their innovative ideas, research findings, and practical experiences with the global scientific community.

FOCUS AREAS

- Engineering
- Applied Science
- Technology
- Innovation & Development
- Interdisciplinary Studies



PEER REVIEWED

All submissions are rigorously peer reviewed to ensure quality.



OPEN ACCESS

Free and unrestricted access to research for all.



GLOBAL REACH

Connecting researchers and professionals worldwide.



TIMELY PUBLICATION

We ensure a swift and efficient publication process.



For more information, visit our website
www.ijeast.com



INTERNATIONAL JOURNAL
OF ENGINEERING APPLIED SCIENCE
AND TECHNOLOGY

✉ editor@ijeast.com

🌐 www.ijeast.com

📍 India



2455-2143