

Systematic Study of Gen Profiles Analysis Methods in Disease Classification

Ms. Punam Gulande

Research Scholar,
Department of Electronics Engg VJTI ,Mumbai – 19
mail ID.- psgulande@el.vjti.ac.in,

Dr. R N Awale

Professor
Department of Electrical Engg
VJTI ,Mumbai - 19
mail ID- rnawale@el.vjti.ac.in

Abstract - In today's biomedical research, Detection and Prediction of disease candidate genes from the human genome is a crucial activity. Diseases with the same physical and cycological characteristics of a candidate from genetic have analogous biological characteristics and genes associated with these same diseases inclined to share common functional properties. Therefore diagnose a new disease from a previous disease is possible by applying machine learning methods. An evolution of bioinformatics, disease catagorizing from gene expression data becomes decisive and useful technology for its diagnosis. Thousands of genes and a small number of samples contained in the gene expression data. Therefore gene selection from gene expression data becomes a key step. Different methods are used for gene prediction. This study is focused on improving the accuracy of prediction of diseases using gene expression sequences and is used for diagnosing specific diseases with genomic information and Machine Learning Algorithms.

Keywords – Machine Learning, Disease Gene Identification, Feature Selection, gene expression Profile.

I. INTRODUCTION

A. Basic concepts

Medical interpretation is essential yet and hard task that needs to be done expertly and precisely. The computerized program is much needed to help doctors make better detection and treatments. Depiction of medical awareness, decision making, collection, and reworking to the appropriate model are some of the things the medical system should consider. Medical advances are always supported by data analysis that enhances the ability of medical professionals and establishes a disease treatment process. The ambition of a medical diagnostic program is to assist physicians in determining the level of risk of each patient. Advantages of genetic algorithms and integrated neural networks to predict possibility of heart disease. The genetic algorithm is an efficient algorithm that factitious the principles of natural genes. It finds good and acceptable solutions to problems with immediate acceptance. In most applications, the information that describes the performance of the application you want is contained in the data sets.

When data sets contain information about the system to be built, the neural network promises a solution because it is able to train itself from data sets. Neural networks are flexible models for data analysis that are best suited for handling offline activities. The optimised process of the genetic algorithm is then merged with the learning capabilities of the neural network, and a better precision guessing model is obtained

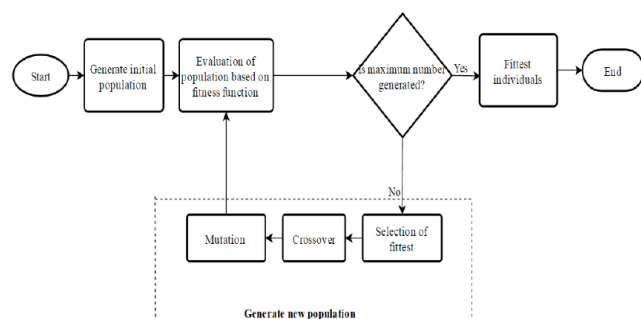


Fig.1.1 Flow chart of Genetic Algorithm

Gene Separation is a research area that presents a new challenge because DNA Microarray contains up to 25000 genetic material. At the same time, these factors may be involved in a number of undesirable factors, and not all of them are essential to the segregation process. Diagnosis, Predictability, and Diagnosis of genetic diseases require a scientific approach. The Soft Computing method helps to facilitate the detection and prediction of these diseases as the data is out of line and incomplete and has intangible parameters. Inactive data sets negatively affect accuracy and performance time. Therefore, selecting a feature is the most important task of selecting the appropriate features that represent the entire database

B. Types of Genetic Disease

Abnormalities in a disease leads to Genetic Diseases in an individual.

Four categories of genetic disorders are listed below.

1. Single gene inheritance
2. Multifactorial inheritance
3. Chromosome abnormalities
4. Mitochondrial inheritance

Single gene inheritance

This type of disorder has diverse shapes of genetic inheritance that includes autosomal dominant inheritance, X-linked inheritance.

If the defective gene is found from either parent then this termed as autosomal dominant inheritance. It is known as X-chromosome inheritance when the defective gene is present on a female or X-chromosome. This may be dominant or recessive. Single gene disorder is also known as a monogenetic disorder.

Multifactorial inheritance

The genetic disorder is caused by environmental factors combined with mutation in multiple genes. Also called complex or polygenic inheritance.

Chromosome abnormalities

The nucleus of each cell hosts the Chromosomes. The same is combination of DNA and Protein and they are the carriers of the genetic material. Any anomalies in chromosome structure or number can result in disease.

Mitochondrial inheritance.

Mitochondria resides in the cytoplasm in plant and animal cells and they are small round or rod-like organelles in appearance and is involved in cellular respiration. Each Mitochondria cell has 5 to 10 circular pieces of DNA and is permanently inherited from a female parent. This type of genetic disorder is usually caused by mutations in the non-

nuclear DNA of Mitochondria.

II. INTRODUCTION TO GENES AND HUMAN DISEASES

The human body is made up of millions of cells and nucleus of each cell hosts the Chromosomes in it and are responsible for passing the genetic information from one generation to the next. Chromosomes are made up of the tightly bonded Deoxyribose-Nucleic Acid commonly called DNA, where information is needed to control protein production. DNA control of when proteins will be produced in bulk. Proteins play a key role in the functioning of the body's cells, tissues, and organs. Genetics are the specific length of DNA that determines the sequence of amino acids used to make proteins. Some proteins are needed for the functioning of all cells. Different types of the same gene are called alleles. Only one genetic form is normally received by an individuals. The effect of alleles when clubbed, leads to differences in appearance such as the colour of their eyes and hair, the shape of their nose, etc. Genetic misconduct is often referred to as mutation. These mutations are responsible for disease outbreaks. According to genetic mutations, diseases are classified as follows:

Chromosomal Diseases: Alteration, losing or doubling of entire or large parts of Chromosomes leads to such diseases. A common example is Down Syndrome. Occurrence of mutation in a gene can lead to stoppage of functioning of the gene. This is termed as Genetic disorder. A sickle-cell anemia is an example of single genetic disease.

Multifactorial Disorders: occurs as a result of change in the gene information which are often linked to natural causes. An commonly seen example of multiple work disruptions is diabetes.

Mitochondrial disorders: are rare disorders caused by non-chromosomal DNA alterations found within mitochondria. (Mitochondria are the organelles under the cells.) It is found that such disorders can affect any part of the body including the brain and muscles.

Many diseases involve many genes in complex collaborations, in addition to environmental impacts. There are situations when a person may not be native of the disease but may be at high risk of confining to it. This is called Genetic Predisposition. Genetic predisposition to a specific disease due to the occurrence of one or more genetic mutations, and/or a mixture of alleles does not have to be exceptional. The WHO has done far-reaching work on major non-communicable diseases, such as diabetes, cancer, asthma, heart disease and other mental illnesses. In some cases, such as cancer, people are native of a genetic predisposition to certain behaviours or to be unprotected to chemicals. Heart disease tends to manifest itself in a number of different ways in different societies. For example, African societies tend to be biased against heart disease, while South Asians are more probable to have heart attacks. Therefore, familiarity of genetic predisposition to disease and to the lifestyle changes that intensify or diminish disease (i.e., no smoking or drinking) is indispensable.

III. MACHINE LEARNING ALGORITHMS IN MEDICINE

Machine learning is a process on how the data is learnt by the computers. It starts from intersection of statistics, which pursues to understand the interactions from data and computer science, with its specific emphasis on effective computing algorithms. The blend of maths and computer science is largely determined by the specific and unique computational experiments of constructing statistical models from massive data sets. The learning types used by computers are appropriately sub classified into supervised learning and unsupervised learning categories

Various building predictors methods have been proposed such as Classification like Naïve Bayes, KNNs (k-Nearest Neighbours), Clustering, GAs (Genetic Algorithms), RSs (Rough Sets), ANNs (Artificial Neural Networks), and Fuzzy Logic Theory. Genetic algorithm is an algorithm that can be optimised to present the principles of natural genetics. Neural networks are normally an adaptive models that are used for data analysis which are predominantly suitable for handling nonlinear functions. ML is more skilled in predicting and used individually for clinical or imaging purposes. Frequently used machine learning algorithms are Support Vector Machines, Artificial Neural Networks, Logistic Regression, Treebased methods and ensemble methods (Brownlee, 2019). Training sets covering the bulk of all available data are used to primarily develop the model, validation sets are used to estimate the overall performance of the model or to fine-tune its hyper parameters.

Every Machine learning techniques has some merits and demerits. The graphical representation of Bayesian Network has advantages in breaking complex problem in smaller models but is very slow in classifying data. The Genetic Algorithm is

capable of resolving problems related to optimization during classification. The SVM is simple for analysis. All computations are performed in space using kernels because of its edge to be used practically, however selection of kernel function is not simple and its requires more memory space. KNN is easy to implement and can solve multi-class problems. **Decision tree** is easy for interpretation and there is no limitation in handling of data sets of high dimensions.

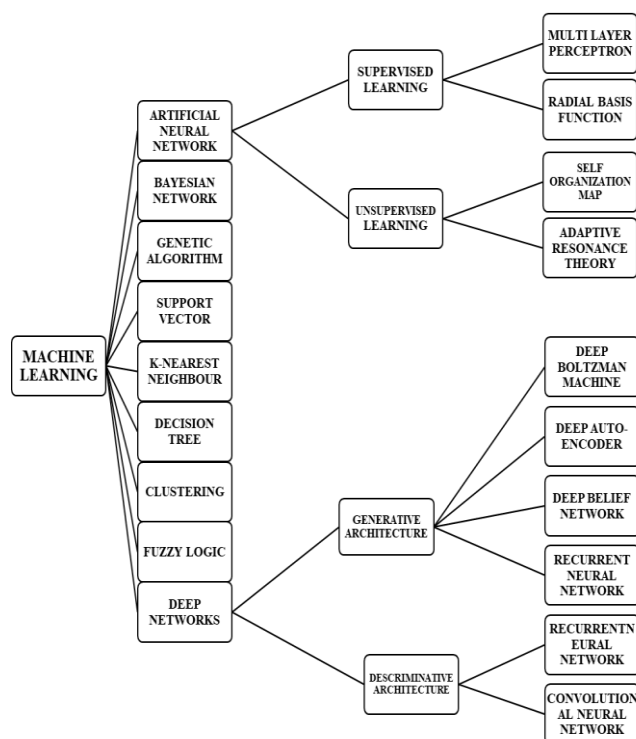


Fig.1.2 Classes of Algorithms of Machine Learning

IV.METHODS FOR DIAGNOSIS OF GENETIC DISEASE AND THEIR LIMITATIONS

A comprehensive experimental trial comprising three major components is required for genetic diagnosis:

1. Physical examination
2. Detailed medical family history
3. Appropriate Clinical and laboratory testing,

Genetic testing can be used for many different purposes, some of which are neonatal tests, carrier tests, prenatal diagnoses, diagnostics / forecasts, forecasts

V. LITERATURE SURVEY

Author Fei Yuan et Al [1], with the help of machine learning algorithms, probed the genetic information of samples in lung AC and lung SCC. Monte Carlo feature selection (MCFS), a powerful and popularly known feature selection feature, was hired for carrying out the test of gene expression profiles, thus producing a list of prominent and other informative features. Increasing Factor Selection (IFS) and Vector Support Mechanism (SVM) have been used in the list of genetic variants, which can be defined by genetic variation between lung AC and SCC, and to form correct SVM differentiation, separate samples for lung AC and SCC for high performance. Currently for creating the classification rules, educational features have been used for the learning process, which provides a process of clear classification and enhance the understanding of genetic patterns of various lung AC and lung SCC. These new discoveries may divulge probable dissimilarities between lung AC and lung SCC at the transcriptomic level, describing different pathogenic features of the two subtypes of lung cancer. Author Yumin Chen et Al [2] proposes a novel genetic selection approach based on a set neighbour model, which can address real value data while keeping the actual details of the list. It also

addresses the level of entropy under the framework of aggressive sets to deal with the uncertainty and noise of genetic data. The use of this measure could lead to the discovery of subsets of the compact genes. Finally, the genetic selection algorithm was deliberated based on neighbouring granules and entropy measurement. Another study of two genetic data advocates that the anticipated genetic selection is an effective way to increase the plant accuracy. In this paper, they collect genetic data set by a local parameter and develop specific methods of uncertainty in the framework of a neighbouring set. Essentially, integrated entropy is first proposed to assess genetic data set uncertainty and to introduce an integrated genetic selection approach. As a result, some experimental outcomes indicate that the proposed algorithm can identify low-value genes and increase class accuracy. Author Lingyun Gao et Al [3] proposed a Hybrid Genetic Selection method, the Information Gain-Support Vector Machine (IG-SVM). Initially, the Information Gain (IG) was assigned to eliminate non-invasive besides obsolete genes. Subsequently, continuous elimination of unwanted genes was executed using SVM to eradicate noise in databases efficiently. Finally, the genetic educational genes selected by the proposed IG-SVM, act as the input of the LIBSVM separator. IG-SVM presented the maximum truthfulness of performance and high performance as experienced using five sets of cancer data sets on the basis of few selected genes, when compared with other related algorithms. In this study, looking at SVM's high cost of capturing multiple genes, they used a hybrid approach that combines IG and SVM to select educational genes. IG was originally used to select genes to diminish size of the original data. Further airing of unwanted genes was consequently accomplished using SVM. The genes found were finally tested for LIBSVM classification. Author Hernández-Montiel et Al [4] propose Multiple-Filter (MF) with a genetic algorithm (GA) alongwith Tabu Search (TS) in conjunction with Support Vector Machine (SVM) for genetic selection and DNA microarray data classification. The method is considered to identify a subgroup of suitable genes that can segregate the DNA-microarray data more affectively. Initially, the five customary mathematical methods are put in use to select the first genes (Multiple Filter). Later other different suitable subgroups are identified using Wrapper (GA / TS / SVM). A genetic set, comprising of the appropriate genes, is initiated in each mathematical method, by analysing the occurrence of each gene for a different gene. Finally, the most common genes are analysed in the Multiple Wrapper method to find the appropriate final subset. The proposed method is then tested on the fourth data set of DNA microarray. From the test results, it is apparent that the model works considerably superior than other methods described in the literature. This method comprises of 2 basic steps, first the data processing with the help of 5 mathematical filters to mark the preliminary selection from biomedical evidence and then the selection and classification of subgroups. Selection is impacted by the genetic algorithm in conjunction with Taboo Search as a local search method, and the subset feature separation is done with a Vector Support Machine and substantiated using a cross-Validation method. With this hybrid approach, the most suitable genetic variants are explored from 5 public domain sources derived from DNA-microarray technology. This model accomplishes the genetic selection process to acquire a genetic set of correctness with high levels in the four microarray databases. The study by author Sun-Yuan Hsieh et Al [5] improved genetic variation (GEG) to reduce the calculation time. The classified filter fiber filters genes by using the weight of the edge to determine their value, thus being able to compare and accurately classify them. This study attempted to compare the suggested separator with a GEG-based separator using real data and benchmark tests. Another key contribution of the suggested separator is its capability to categorize samples in the suitable class and achieve extracellular samples that are not in the trained class. Otherwise, it reduces the calculation time required for separation. To endorse the efficiency of the proposed method, this study attempted to compare graph-based measurements with GEG-based measurements by performing a set of microarray tests for eight known diseases. Test results show that a weight-based graph separator can reduce the calculation time required for classification, achieve comparative or higher performance in sample handling in the target class library, and separate samples from a non-trained class.

The author V. Elyasigomari et Al [6] established a new hybrid optimization algorithm, COA-HS. This algorithm is a combination of HS algorithms with COA. The results are tested and compared in conjunction with PSO, GA, HS, and COA algorithms. The PSO is an optimization algorithm and is based on the mocking of the migratory birds. In order to detailed the proposed COA-HS algorithm, the details of both the COA and HS methods are described. To address this problem, a two-phase genetic model called MRMR-COA-HS was proposed. Initially, the selection of minimum retrieval and high value (MRMR) is used to identify a set of appropriate genes. Selected genes are then incorporated into a compact set which has the new algorithm, COA-HS, using a vector support device as a separator. This method is applied on the fourth microarray data sets, and its performance is tested by a single break authentication method. Author Wan Li et Al [7], hypothesized that SNPs acting on disease risk in encoding areas and protein communication modules may contribute to the genetic identification of chronic diseases which are of complex nature. Dangerous SNPs were first selected from the SNP data for CHD, HT, and T2D. SNPs in the coding regions which are classified as nonsense and missense, are treated as active SNPs. Subsequently, the regions which are closely associated with each disease were assessed with the help of random permits of SNPs that deal with disease risk. Possible genes associated with these regions were discovered. Due to the path biological similarity between these three diseases, a communication network of a disease-related product is formed based on the possible genes of the disease in these diseases. After a functional analysis of the enrichment of

disease-related network modules, the affected genes were finally identified and authorized by other literature. Author Ozra Nikdelfaz et Al [8], suggests a new way of predicting genes based on Markov's hidden model and genetic sequence. The results of the experiments demonstrated the effectiveness of modeling for students' genetic prediction problems using a single phase classification method, HMM. The predictive effects of HMM-based disease genes studied without coagulation, showing genetic interactions significantly improve the values of clarity, memory, and F1 metrics. In the case of cancer, it is shown that the proposed method outperforms the other methods by almost 12.7%, 11.7%, and 12.4% better than the EPU method (the best method of previous methods) in terms of accuracy, memory, and F1, respectively. Author Dimitrios Klefogiannis et Al [9] proposes YamiPred (Another miRNA predictor), an embedded separation method that chains the effectiveness and strength of SVMs and Gas with the selection of features and limitations for optimization. YamiPred was proficient using human pre-miRNA sequencing. The feature vector includes the thermo dynamical structure, structure, and sequence of state elements, as well as the 10 recently introduced features, recommended in the literature. Test outcomes show that YamiPred exceeds C-state-of-the-art methods and SVM-based predictions in terms of class performance and easiness of classified classifiers. These paybacks are due to the exceptional way to deal with the problematic of class inequality, slower integration, and easy-to-interpret translation of the ratio of direct and negative samples. An additional contribution of YamiPred is the newly presented function of a specific solidity problem, which attains a balanced split function between the many and a few categories and simultaneously forces the algorithm to search for segmentation models effortlessly based on input features and model features. Author Chandra Babu Gokulnath et Al [10] proposes the use of a vector support system (SVM). The proposed function is applied to the genetic algorithm (GA) to select the most key factors for heart disease. GA-SVM test results are equated with various selected algorithms such as Relief, CFS, filtered subset, Info gain, Consistency subset, Chi-squared, One attribute-based, Filtered attribute, Gain ratio and GA. Examination of the receiver detection feature was executed to test the usefulness of the SVM separator. The complete framework is revealed in the MATLAB area with a database received from the Cleveland cardiovascular database.

Author Alfredo Benso, et Al [11] attempt to address their issues by familiarizing a new graph-based data-based algorithm, entitled the Gene Expression Graph (GEG), which is used to signify genetic data. GEGs clearly specify the association between genetically modified and silenced genes in a various conditions (e.g., healthy and relative to disease). Therefore, they are well matched for the analysis of DNA microarrays (cDNA) that provide a distinct point of support for both healthy and diseased models. The conception of using graph-based data modeling to symbolize genetic data, first published by authors in [4] and [5], permits creating of distinctions based on higher comparisons between known class graphs and graphs representing sample samples. Author Bhuvaneswari Amma Al [12], proposed system, benefits of genetic algorithm and neural network collective, to envisage the threat of heart disease. The genetic algorithm is an effective algorithm that impersonators the ideologies of natural genes. It finds good and adequate solutions to problems with instantaneous acceptance. In most applications, the information that defines the performance of the application you want is contained in the data sets. When data sets holding information about the system to be constructed, the neural network provides a solution because it is able to train itself from data sets. Neural networks are flexible models for data analysis that are best suited for handling offline activities. The combination of the optimization process of the genetic algorithm along-with the learning abilities of the neural network, a model with superior precision predicting can be obtained. Author Pradipta Maji et Al [13] presented a innovative genetic selection algorithm, named RelSim, to identify genetic diseases. It incorporates comprehensively pro-file information and protein communication networks. This delivers an effective way to analyse functional similarities between the two distinct genes. The algorithm proposed by the author, chooses a set of genes data, processing the informations of the protein-protein interactions and microarray, by enhancing the similarity and functionality of the selected species. While genetic profiles are normally used for differentiate genetically modified genes, protein-protein communication networks supports to calculate functional similarities among genes. Colon cancer test results show that the algorithm which has been proposed, recognizes genetic cancers with multiple colors equated to prevailing methods of genetic selection. The result announces the performance of a suggested genetic selection algorithm, called RelSim, based on the performance measurement of total performance correlation (MRMFS). The effectiveness of the proposed MRMFS policy is comparable to the implementation of certain other conditions, namely, high compliance (MR), minimum downtime (mRMR), and high-density value (MRMS). The performance of the RelSim algorithm is similar to the specific integrated genetic selection methods. Algorithms are compared to MR + PPIN, MRMS + PPIN, mRMR + PPIN, GenePEN and CLAIM. To evaluate the effectiveness of various genetic screening methods and methods, using a set of microarray data sets, namely GSE10950, GSE25070, GSE4 988, GSE44 861, and GSE11223, associated to colon cancer. Each data set is pre-processed by determining each sample at zero definition and unit variance. Author Narayan Behera et Al [14] proposes a new algorithm called the Evolutionary Clustering Algorithm (ECA). This algorithm has been used to analyze data on stomach cancer. Complete analysis of algorithms in the gastrointestinal cancer database shows that several comparisons containing high-level genes that provide high differentiation of microarray test samples are more in ECA than CMI and kmeans. This proves that ECA finds the

right cancer genes with a higher probability than other algorithms. The author, Rasmita Dash [15], suggested a selection of features based on a selection of features based on compliance and the use of the Pareto method. Here feature selection is done in two steps. The initial selection of feature inclusion as a function of optimization, the genetic selection was performed using a flexible compliance search (AHS GS) search process. Then in step 2 to find the full number of the feature with the reference to the appropriate feature, Pareto's multi-sectoral solution is explored. According to the two standard filtering methods the Pareto optimization process is used and a few of the most suitable features are selected. Considering the advantages of both threats and filtering methods this analysis is a way to wrap the filter selection filter into the most advanced information. Finally, statistical analysis was performed to prove the superiority of AHS GS over two other forms of evolution

Author Jessa Bekker [16], proposed an valuation of the current state of the art in PU learning. It upraises seven key research questions that have arisen in the field and offers an outline of how the sector has tried to resolve them. Author Kyungsook Han et Al [17], has established a fresh method that recognizes a strong regulatory collaboration between a group of genes and a single gene. It produces an authoritative network of genetic control relations of several types, which have not been deliberated in prevailing ways. This paper presents a way to identify and visualize mutating genetic rules and experimental results with accurate genetic information from the Yeast Genome Project. The author, Banu Dost et Al [18], proposed a quick way to identify the thickest pieces in a composite graph described in genes (or, research sets). If performance-related genetic collections were all tense, and are the only type of compression, the aggregation graph must look like a combination of excluded genes. Errors and other variations which are biological in nature, will supplement extra edges and remove certain true edges. To simplify exposure, they view these additional and missing layers as invalid data (false positives (FP), and negatives (FN), respectively), or they may include a real biological object. The resultant graph is thus a broad graph with "thick" characters that comprise functional groups. Author Yukyee Leung et Al [19] recommends a multi-filter (MFMW) hybrid model method. In the proposed MFMW hybrid model, for selection of genes, the use of multiple filters is considered and the clubbed to provide a supplementary integrated set of genes. The usage of multiple filters with other different filter metrics endorses that useful biomarkers are not established in the initial phase of the filter. The usage of multiple wrappers is proposed to support the trustworthiness of the split by creating consistency between several different sections. Due of this, there exist some agreement between the various different dividers in the wrapper month, of which genes should be selected. Therefore, the last selected genes can be measured as more advanced alongwith a combination of traits equal to a few metals, so they are better matched as biomarkers. Since the MFMW model previously includes many other filter elements and metals, it is not recommended to try a different filter combination to wrap an appropriate combination that provides the highest correctness of separation. They show that the proposed MFMW model delivers projected data that can be compared or superior than the prevailing positive results which is achieved using all SFSW methods.

Author Bo Li et Al [20], a proposed algorithm that creates vector modification and retrieval model results in improving the accuracy of the actual LLE recognition, which enables samples with different class labels to be well classified and makes similar class approaches closer. Under such circumstances, embedding can be interpreted and sent to any location. To get the right translation and standardization, the margin condition is greatly increased (MMMC) introduced to automatically determine the correct change. Currently, in LLDE, class data is also used to check for correct translation vectors. Author Jiaqi Zhang et Al [21] developed an anonymous discovery of a Positive and Unlabeled (PU) learning problem wherein only labeled (normal) data and non-labeled (normal and abnormal) data, for unidentified detector learning are required. Since label-less data is not very consistent, , we present a PU novel reading method that can deal with a situation where a set of non-labeled data is generally in good condition. They start with usage of a straightforward model for removing the most reliable negative situations, followed by the learning process itself to add negative and promising reliable conditions at different speeds based on the previously positive phase. Classifiers in the process of independent study are measured in combination based on the estimated error number to construct the final division. A detailed examination of the six real data sets and one performance data displays that our methods have better resulted under different conditions when compared to existing methods. First they used a compatible model with different functions for the cost of a good and unstructured set to detect reliable negative conditions. Reliable good and bad conditions are gradually released at a consistent pace depending on the predictable results of the current phase. A good section ahead is estimated to adjust the learning speeds of various good and bad classes. To minimize the error of the previous positive rating of the category, the remaining unlabeled cases will be added to the dividing section in the final step, rather than separated from the previous positive rating section. The first step is removing of the most unreliable N set from the unlabeled set U. With standard PU learning modes, the positive data is the target acquisition and non-label conditions are negative data. Therefore, traditional methods often view unregistered data as erroneous to form the first division. However, due to data detection issues, label cases that are normal data and malformations are still too low for a label less label. Therefore, standard PU learning methods that view non-labeled data as bad data cannot be used directly for problems of anonymous detection. Motivation often increases the performance of individual isolation particularly

when separation is more sensitive to minor deviations of the training set. Each separator that comes out of the self-study process is considered a weak individual separator here. Reinforcement is accepted with a limited rating of a series of dividers. Author Sucheta Chauhan et Al [22] introduces a predictive approach to detecting unfavourable behaviour on Deep LSTM neural networks. The advantage of the proposed method is to incorporate minimal or no data processing, it does not demand hand-coded features but functions directly on the green signal. It does not require prior information of abnormalities. LSTMs are more profitable than traditional RNNs in activities that include longer delays among events. In order to extract information posted during intervals between these events, the network requirement is there for some operation. LSTM can resolve such indirect tasks and, by learning to measure precise time intervals, uses LSTM cells with high connectivity that permits them to explore their prevailing internal states. Author Qiuwen Chen et Al [23] introduces anonymous recognition and acquisition (AnRAD), an independent framework for the acquisition of heterogeneous multicore platforms, which provides unattended real-time learning and stimulates inaccurate data streams in real-time data streams. It is driven by the idea that the processing of human data is very similar and that learning comes from organizations within the elements derived from unlabeled data. Those combined elements create a common expectation and any unexpected patterns will create a surprise. Author Minlong Lin et Al [24] introduces a two-pronged approach, firstly it should integrate samples and training processes, and then should be capable of dealing straight with multiclass difficulties. Multilayer perceptron's (MLPs) are having strong capabilities towards learning complex boundary parameters and therefore they can be used directly in multiclass problems. The sampling process can be fully incorporated into the training process, by using a sequential training model. Therefore, MLP was accepted as the elementary model and was developed as a data selection approach to address the issues of multiclass inequality. For further clarification, a heuristic novel is proposed to determine if the training model (e.g., used to refresh MLP instruments) should be studied in the training period. The heuristic decision-making is based on the existing state of the MLP and therefore is independent of the initial setting of the sample rate. Further, as MLP status changes as learning advances, an unreadable configuration in one era may be useful in the next. Using consecutive training mode, no configuration will be discarded throughout the duration of the training process, and thus the loss of information caused by the pre-sampling process can be avoided. Since such a system can be viewed as a vigorous selection of training models for each term, the proposed method is called the conversion method (DyS) for MLPs.

Author Guo Haixiang et Al [25] aims to provide a comprehensive detail for the measurement of measured data including its strategies and applications. At the technical level, introducing common ways of dealing with unequal learning and suggesting a common framework in which each algorithm can be set. This framework behaves as a model of integrated data mining and involves processing, classification, and evaluation. At the operational level, 162 papers were reviewed that have attempted to construct specific systems to address occasional pattern acquisition problems and increase tax rates for unequal learning application domain. The existing documentation illustrates a extensive range of research fields from medical to industry to management. Author Divya Krishnani et Al [26] analyze and evaluates the use of different machine learning algorithms in the diagnosis of heart disease by combining all the features in the database to develop differentiation models. Random Forest, K-Nearest Neighbors, and Decision Tree classification models of coronary heart disease predictions are made. They also use K-fold Cross-Validation in algorithms. The results show that these types can accurately predict the risk of heart disease, concluding that a person is more likely to suffer from CHD over a period of 10 years. Author Oscar Barquero-Pe' rez et Al [27] studied experimental interactions (CI) with HR and HRT. Instead of using a signal ratio, HRT was observed and measured by insertion of TS into individual VPC computer. For carrying out this purpose, data were obtained and recorded from patients with normal cardiac output, who received electrophysiological studies (EPS). The clinical protocol was explicitly considered for CI to be administered and controlled with a systemic heart rate protocol. HR was administered using isoproterenol. The electrophysiological protocol is consisting of two sub-studies, one for analysing the relationship between HR and HRT, and the other for analysing the collective effect of HR and CI on HRT. Analysis of this interaction was carried out using Holter recording from patients after the episode of AMI, so that EPS data, acquired in good times, could be equated with data from ambulance monitoring where very high noise levels were noticed. Appropriate retrospective models were used to replicate the previous results and to explain the collective effect of CI and HR on HRT. Author Juan Wang et Al [28] investigated the possibility of automatic and accurate detection of BAC in mammograms by in-depth study methods to support the long-term goal of building an automatic BAC detector that could be used to provide risk indications for coronary artery disease. They create a problem like pixel-wise, patch-based two-classification problems. That is, for any conceivable pixel, use the surrounding image patch. Image dots are provided by a deep CNN trained to distinguish whether the middle pixel belongs to the BAC category or not. Compared to other machine learning and other algorithms, deep CNN can automatically extract features while achieving better segmentation performance. It has been shown that the auto-read features of deep CNNs often surpass the hand-drawn features of human experts. Author Zahra Ghafouri et Al [29] proposes a simultaneous action that detects the appropriate hyper parameter setting and removes suspected abnormalities instead of changing the OCSVM objective function. Numerical tests show that our system overcomes high-level statistics

and unregulated methods for measuring OCSVM hyper parameters. Author Qinbao Song et Al [30], based on the MST process, suggested the fast clustering-based Selection algorithm (FAST). The said algorithm is performed in two steps. Initially the features are collected together with the use of graph-theoretic interaction methods. Then the most representative element closely linked to the target categories are selected in each group for formation of the final set of features. The elements, normally, in the various assemblies are self-contained, a FAST-based accumulation approach has a better chance of constructing a subset of valuable and self-regulating features. The FAST algorithm of the planned selection feature subset was verified on 35 widely obtainable images, microarrays, and transcript data sets. Test results show that, the proposed algorithm not only reduces the number of features but also increases the performance of the four known different types of configurations, when compared with the other five different types of subset selection algorithms. The Internet of Things (IoT) [31-34] assisted healthcare services received significant attentions since from decade especially disease detection and monitoring.

VI. DISCUSSION

Main Findings

Research in this domain down with different types of cancer like lung, breast, colon and prostate. Main steps are found in literature is feature selection and classification using machine learning algorithm on selected genes. Some algorithms are used for different data sets and accuracy is compared with the performance of the different filters and classifiers. Based on the result optimum methods for filtration with some algorithms on the specific genes are detected. The below table shows the accuracy of classifiers used for the mentioned disease.

It is observed that when Genetic Algorithm combines with SVM classifier and Neural Network gives better results on heart disease. Accuracy found in [9] is 93.21% on miRNA datasets and when Cleveland database is used then it is 88.34% with some selected features[10]. When some database is tested with GA and NN, then accuracy is 97.14% [12]. GA used with Evolutionary Clustering Algorithm on gastric cancer, the accuracy was seen as 88.34% [14]. Anomaly detection of ECG signals with Supervised Machine Learning Algorithm used for classifier namely RF, KNN and Decisions Tree, then the accuracy was found as 96.80%, 92.8% and 92.45% respectively.

Study	Classifier used	Cancer Type	Accuracy (%)
[1]	SVM	Lung (AC) & Lung (SCC)	96.7%
[3]	SVM	Lung Colon Prostate	100% 90.32% 96.08%
[4]	SVM	Colon Lung	96.7% 97.8%
[6]	SVM	Prostate Colon	100% 100%

Data Base

Datasets used in most of the existing literature are of microarray gene expression profiles (GEP). Few of these are extracted by combining GENECARD, OMIM, Gene Expression Omnibus (GEO) and few are downloaded from National Center for Biotechnology Information (NCBI). Data sets and a human gene annotation were also downloaded from the GO database. The dataset collected for heart disease are from Cleveland heart disease database

Limitation

Preliminary survey shows that the analysis of the methods proposed has the major limitation of having the consistent and data on large scale. Given the size of the datasets which are small for testing purpose, the methods which are proposed in these studies need to be further confirmed in larger datasets, which can be more real-time and will help to obtain a better training model. Also, an emphasis on refining the preprocessing further to give veracious outcomes

Performance of many of the algorithms can be improved by using Genetic Algorithms and Principal Component Analysis. This is in order to reduce the dimension of the dataset and is used to predict the risk of any of the diseases. In combination of other techniques of feature selection can be used for getting better classification accuracy and by minimizing the number of genes and also by setting the different parameters. A review of other rare event detection techniques and its applications, not endorsed by the author.

VII. CONCLUSION

The main advantage of genetic algorithms is due to its flexibility and sturdiness, as a universal search process. Genetic algorithms are very useful in collaboration with neural networks. They can deal with non-linear problems and inseparable activities. By looking at genetic profiles as a sequence of observations, HMM is the best model for identifying a gene and increasing predictive accuracy. Genetic identification from Micro-array data, Genetic profile combination, and PPI Networks provides better results.

Most of the algorithms used the SVM Classifier which plays an important role in improving efficiency. SVM works compared to other methods, in the diagnosis of heart disease. GA-SVM includes Tabu Search used for genetic selection and DNA Microarray Data editing. The GEG-based separator is useful in reducing calculation time compared to other methods. The selection of the MRMR feature for appropriate genetic selection and integrated with the COA-HS SVM wrapper algorithm provides the best results. The risk of coronary heart disease can be determined based on mammograms through in-depth study methods.

Genetic algorithms are very useful in collaboration with neural networks. They can deal with non-linear problems and inseparable activities. Many algorithms and techniques can be used to predict a gene in different data sets. An algorithm or strategy needs to be developed in terms of diagnostic accuracy.

The integration of alternative selection strategies and factor classification is required for the identification and identification of disease-related genes.

REFERENCES

- [1] Fei Yuan, Lin Lu, Quan Zou in “Analysis of gene expression profiles of lung cancer subtypes with machine learning algorithms “,BBA - Molecular Basis of Disease 1866 (2020) 165822.
- [2] Yumin Chen, Zunjun Zhang , Jianzhong Zheng, Ying Ma, Yu Xue in “ Gene selection for tumor classification using neighborhood rough sets and entropy measures “, in Journal of Biomedical Informatics 67 (2017) 59–68.
- [3] Lingyun Gao , Mingquan Ye , Xiaojie Lu , Daobin Huang in “Hybrid Method Based on Information Gain and Support Vector Machine for Gene Selection in Cancer Classification” in Genomics Proteomics Bioinformatics 15 (2017) 389–395.
- [4] Hernández-Montiel Alberto Luis, Bonilla-Huerta Edmundo, Morales-Caporal Roberto and Guevara-García Antonio José in “ Selection and Classification of Gene Expression Data Using a MF-GA-TS-SVM Approach” in Springer International Publishing Switzerland 2014.
- [5] Sun-Yuan Hsieh and Yu-Chun Chou in “A Faster cDNA Microarray Gene Expression Data Classifier for Diagnosing Diseases” in IEEE/ACM Transactions on Computational Biology and Bioinformatics.
- [6] V. Elyasigomari, D.A. Lee, H.R.C. Screen, M.H. Shaheed in “Development of a two-stage gene selection method that incorporates a novel hybrid approach using the cuckoo optimization algorithm and harmony search for cancer classification” in Journal of Biomedical Informatics 67 (2017).
- [7] Wan Li, Lina Zhu, Hao Huang, Yuehan He, Junjie Lv, Weimin Li, Lina Chen, Weiming He in “Identification of susceptible genes for complex chronic diseases based on disease risk functional SNPs and interaction networks” in Journal of Biomedical Informatics 74 (2017).
- [8] Ozra Nikdelfaz, Saeed Jalili, “ Disease genes prediction by HMM based PU-learning using gene expression profiles, “ in Journal of Biomedical Informatics 81 (2018).
- [9] Dimitrios Klefogiannis, Konstantinos Theofilatos, Spiros Likothanassis, and Seferina Mavroudi in “ YamiPred: A novel evolutionary method for predicting pre-miRNAs and selecting relevant features”, In IEEE Transactions On Computational Biology And Bioinformatics.
- [10] Chandra Babu Gokulnath1, S. P. Shantharajah1 in “ An optimized feature selection based on genetic approach and support vector machine for heart disease” Published in Springer Science + Business Media, LLC, part of Springer Nature 2018.
- [11] Alfredo Benso, Stefano Di Carlo, and Gianfranco Politano in “ A cDNA Microarray Gene Expression Data Classifier for Clinical Diagnostics Based on Graph Theory” in IEEE/ACM Transactions On Computational Biology And Bioinformatics 2011.

- [12]. Bhuvaneswari Amma N.G.in “ Cardiovascular Disease Prediction System using Genetic Algorithm and Neural Network”.
- [13]. Pradipta Maji, Ekta Shah, Sushmita Paul, “ RelSim: An integrated method to identify disease genes using gene expression profiles and PPIN based similarity measure “ in Information Sciences 384 (2017) 110–125.
- [14] Narayan Behera, Shruti Sinha, Rakesh Gupta, Ann Geoncy, Nevenka Dimitrova, and Javed Mazher in “” Analysis of gene expression data by evolutionary clustering algorithm” in 2017 International Conference on Information Technology.
- [15] Rasmita Dash in “ An Adaptive Harmony Search Approach for Gene Selection and Classification of High Dimensional Medical Data” in Journal of King Saud University – Computer and Information Sciences xxx (2018).
- [16] Jessa Bekker , Jesse Davis in “Learning from positive and unlabeled data: a survey” in part of Springer Nature 2020.
- [17] Kyungsook Han and Jeonghoon Lee in “GeneNetFinder2: Improved Inference of Dynamic Gene Regulatory Relations with Multiple Regulators” in IEEE/ACM Transactions On Computational Biology And Bioinformatics 2016.
- [18] Banu Dost, Chunlei Wu, Andrew Su, and Vineet Bafna in “TCLUST: A Fast Method for Clustering Genome-Scale Expression Data” in IEEE/ACM Transactions On Computational Biology And Bioinformatics 2011.
- [19] Yukye Leung and Yeungsam Hung in “A Multiple-Filter- multiple-Wrapper Approach to Gene Selection and Microarray Data Classification” in IEEE/ACM Transactions On Computational Biology And Bioinformatics 2011.
- [20] Bo Li , Chun-HouZheng, De-ShuangHuang , LeiZhang , KyungsookHan in “Gene expression data classification using locally linear discriminant embedding” in Computers in Biology and Medicine 40 (2010).
- [21] Jiaqi Zhang, Zhenzhen Wang, Jingjing Meng, Yap-Peng Tan, Senior Member, IEEE and Junsong Yuan, Senior Member, IEEE in “ Boosting Positive and Unlabeled Learning for Anomaly Detection with Multi-features” in IEEE Transactions on Multimedia, Citation information: DOI 10.1109/TMM.2018.2871421.
- [22] Sucheta Chauhan and Lovekesh Vig in “ Anomaly Detection in ECG Time signals via Deep Long Short-Term Memory” in 2015 IEEE.
- [23] Qiuwen Chen, Ryan Luley, Qing Wu, Richard W. Linderman, and Qinru Qiu, in “AnRAD: A Neuromorphic Anomaly Detection Framework for Massive Concurrent Data Streams” in IEEE Transactions On Neural Networks And Learning Systems.
- [24] Minlong Lin, Ke Tang, and Xin Yao in “Dynamic Sampling Approach to Training Neural Networks for Multiclass Imbalance Classification,” in IEEE Transactions On Neural Networks And Learning Systems, Vol. 24, No. 4, April 2013.
- [25] Guo Haixiang, Li Yijing , Jennifer Shang, Gu Mingyun, Huang Yuanyue , Gong Bing in “ Learning from class-imbalanced data: Review of methods and applications” in Expert Systems With Applications 73 (2017).
- [26] Divya Krishnani, Anjali Kumari, Akash Dewangan, Aditya Singh, Nenavath Srinivas Naik in “Prediction of Coronary Heart Disease using Supervised Machine Learning Algorithms” in 2019 IEEE Region 10 Conference (TENCON 2019).
- [27]. Oscar Barquero-Pe´ rez, Carlos Figuera, Rebeca Goya-Esteban, Inmaculada Mora-Jim´enez, Francisco Javier Gimeno-Blanes, Pablo Laguna, Senior Member, IEEE, Juan Pablo Mart´inez, Eduardo Gil, Leif S’ ornmo, Senior Member, IEEE, Arcadi Garc´ia-Alberola, and Jose´ Luis Rojo-A´ lvarez, Senior Member, IEEE in “On the Influence of Heart Rate and Coupling Interval Prematurity on Heart Rate Turbulence”, IEEE Transactions On Biomedical Engineering, Vol. 64, No. 2, February 2017.
- [28]. Juan Wang, Huanjun Ding, Fatemeh Azamian Bidgoli, Brian Zhou, Carlos Iribarren, Sabee Molloy, and Pierre Baldi “ Detecting Cardiovascular Disease from Mammograms With Deep Learning”, in IEEE Transactions On Medical Imaging, Vol. 36, No. 5, May 2017.
- [29] Zahra Ghafoori , Sarah M. Erfani, Sutharshan Rajasegarar, James C. Bezdek, Shanika Karunasekera, and Christopher Leckie in “Efficient Unsupervised Parameter Estimation for One-Class Support Vector Machines” in IEEE Transactions On Neural Networks And Learning Systems.
- [30] Qinbao Song, Jingjie Ni, and Guangtao Wang in “A Fast Clustering-Based Feature Subset Selection Algorithm for High-Dimensional Data” In IEEE Transactions On Knowledge And Data Engineering 2013.
- [31] Mahajan, H.B., Badarla, A. & Junnarkar, A.A. (2020). CL-IoT: cross-layer Internet of Things protocol for intelligent manufacturing of smart farming. J Ambient Intell Human Comput. <https://doi.org/10.1007/s12652-020-02502-0>.
- [32] Mahajan, H.B., & Badarla, A. (2018). Application of Internet of Things for Smart Precision Farming: Solutions and Challenges. International Journal of Advanced Science and Technology, Vol. Dec. 2018, PP. 37-45.
- [33] Mahajan, H.B., & Badarla, A. (2019). Experimental Analysis of Recent Clustering Algorithms for Wireless Sensor Network: Application of IoT based Smart Precision Farming. Jour of Adv Research in Dynamical & Control Systems, Vol. 11, No. 9. 10.5373/JARDCS/V11I9/20193162.
- [34] Mahajan, H.B., & Badarla, A. (2020). Detecting HTTP Vulnerabilities in IoT-based Precision Farming Connected with Cloud Environment using Artificial Intelligence. International Journal of Advanced Science and Technology, Vol. 29, No. 3, pp. 214 - 226.