

Prediction Of Heart Stroke Diseases Using Machine Learning Technique Based Electromyographic Data

G. Sasikala¹, G. Roja² and D. Radhika³

¹Assistant Professor, Computer Science and Engineering, Department of computer science and Engineering, Vivekanandha College of Engineering for Women. Email: sasikalag@vcew.ac.in

²PG Scholar, Department of Computer Science & Engineering, Vivekanandha College of Engineering for Women (Autonomous), Tiruchengode-637205. Email: 19201006@vcew.ac.in

³Assistant Professor, Computer Science and Engineering, Department of computer science and Engineering, Vivekanandha college of Engineering for women. Email: radhika@vcew.ac.in

Article History: Received: 5 April 2021; Accepted: 14 May 2021; Published online: 22 June 2021

ABSTRACT

Stroke is a prominent cause of disability in adults and the elderly, resulting in a slew of social and financial issues. Stroke can be fatal if left untreated. Patients with stroke have been found to have aberrant bio-signals in the majority of cases. Individuals can obtain appropriate therapy more rapidly if they are observed and have their bio-signals detected and precisely assessed in real-time. However, most stroke diagnostic and prediction systems rely on image analysis methods such as CT or MRI, which are costly and difficult to employ for real-time diagnosis. In this paper, we developed a stroke prediction system that detects stroke using real-time bio-signals with machine learning techniques. The purpose of this study was to analyse and diagnose electromyography data using machine learning. The data cleaning was carried out in accordance with the inclusion criteria that were specifically designed. Following that, two data sets were created, each containing 575 facial motor nerve conduction study reports and 233 auditory brainstem response reports. The data sets were then subjected to four machine learning algorithms: Reinforcement Learning, linear regression, support vector machine, and logistic regression. Comparisons of accuracy and recall rate among several algorithms show that the Reinforcement Learning algorithm outperforms the other two in both data sets, Congenital Heart Disease (CHD) and International Stroke Trial (IST). Furthermore, for each algorithm, comparisons were made with and without deviation standardization, and the results show that deviation standardization has an effect on accuracy improvement. The experiment employs three classification algorithms: linear regression, logistic regression, SVM (support vector machine), and Reinforcement Learning. As a result, Reinforcement Learning has been demonstrated to be an optimal algorithm for diagnosis. It is also worth noting that feature ranking in order of importance can facilitate clinical diagnosis and has clinical potential in diagnosis and diagnostic assessment.

Keywords: support vector machine, International Stroke Trial, Congenital Heart Disease and machine learning algorithms.

1 INTRODUCTION

Stroke has become the world's main cause of incapacity. By 2030 about 70 million stroke survivors were anticipated to survive and over 200 million stroke adjusted life-years (DALYs) lost annually. Stroke burden was very high in high-income countries, with the rapid development of the welfare economy rapidly rising in low-income and mid-income countries. Categorization of the ischemical stroke subtype requires history, testing, laboratory, electrocardiographic and imaging data to deduce and assign causative, etiological or phenotypical classification to a mechanism. Ischemic stroke may occur due to a wide range of vascular disorders leading to brain thromboembolism. Setting the most likely cause is crucial since it affects both short and long-term prognoses and treatment options, in particular in relation to the prevention of recurring episodes [1-4]. The cause of stroke is important. As a result, it influences the design and the provision of critical information for epidemiological research to define the subtype of ischaemic stroke.

In a wide range of research projects, therefore, the employment of a verified subtype classification system that allows for comparisons of outcomes is vital. The method is used to select a subset of important features (variables, predictors) for use in the creation of a model in machine learning and statistics, also known as variable choices, attribute selections, or variable subset choices. For a number of reasons, feature selection approaches are used: Simplification of models to facilitate the researchers' interpretation, shorter training times to avoid the

dimensionality of a curse, increased generalisation by reducing overfit (formally, reduction of the variance) When using feature-selection techniques, the central premise is that data contains certain features which either are redundant or not relevant and therefore can be deleted without input. Two distinct conceptions are redundant and irrelevant as one of the significant features may be redundant if another of its important features is closely associated with. The techniques for selecting features should be distinct from extracting features. Function extraction develops new features from original functions, while the selection of functions returns a subset of characteristics [5 and 6]. Feature selection techniques are often used in areas where many features and few samples are available (or data points). The analysis of written texts and DNA microarray data, in which many thousands and tens of samples are present, is an archetype for the use of feature selection.

As a result, a huge number of samples are almost impossible to achieve to meet the data quantity required for deep learning. Taking this into account, thorough study approaches have hindered research and clinical applications [7]. Only by picking appropriate characteristics manually can standard machine-learning algorithms produce high-precise findings even on the basis of limited data sets. Consequently, in the field of Classical Chinese Medicine (TCM), traditional machine-learning techniques are frequently utilized [8]. These investigations have played an important part in examining the guidelines on medical classification. Secondly, IA-related research has rarely been reported on EMG facial and head data, in particular on the FMNC and ABR data, with significant research potential [9].

2 LITERATURE SURVEY

This study examined the application of machine learning approaches in ischemical stroke patients to forecast long-term results. This was a retrospective study using a forward-looking cohort to instruct acute ischemical stroke patients. Advantageous results were characterized as the modified score of 0, 1 or 2 Rankin Scale at 3 months. 3 machine-learning models have been created and their predictability has been examined (deep neural network, enhancement and logistic regression).

We also assessed them against the Lausanne Acute Stroke Registry and Analysis (ASTRAL) score to measure the accuracy of the machine learning models. Results—The study comprised 2604 patients in total, and there were favorable results for 2043 patients (78 percent). The area under the curve was substantially higher for the deep neural network model than the ASTRAL (0,888 compared to 0,839; P) [10]

This work assesses the performance of an AMLwvf differentiator in a machine learning classification from different RCC subtypes using whole tumor slices in the CT data. Material and Methods: In this retrospective analysis, 171 pathologically confirmed renal masses were gathered from a single institution. In three phases including precontrast phases (PCP), corticomedullary (CMP) and reprographic (NP) phases, the texture characteristics were removed from entire-tumor images.

A remedial method of eliminating AMLwvf from all RCC (all-RCC), clear cell RCC (ccRCC), and nonccRCC features was used for a support vector with fivefold cross-validation method (SVM-RFECV), synthetic minority oversampling technique (SMOTE). The performances of the classifiers based on three-phase and single-phase images were compared with each other and morphological interpretations. The results were: the best performance attained in separating AMLwvf from the whole RCC, ccRCC, and non-ccRCC was a machine learning classification. [11].

This work was conducted in conjunction with machine learning approaches to leverage natural language treatment of electronic health records (EHR). Methods: We evaluated unstructured text-based EHR data including neurological progression notes and neuroradiology reports using natural-language processing among IS patients on TOAST subtype observational registries adjudicated by board certified vascular neurologists. We have completed different selection approaches for feature selection to lower the high dimensionality of the features and the 5-times cross-validation to test and minimize fitting.

Using several methods of machine learning and calculation of Kappa values for agreement of each manual adjudication methodology. We next tested the best method blindly against a 50 cases subset. Results: The best machine-based classification achieved a kappa of 25 by radiological reports alone, 57 by progress records alone and 557 by combining data compared to the manual categorization; [12]

These radiological biomarkers have recently been replicated with Deep Learning. We studied Deep Learning techniques for construction models to directly predict good reperfusion (EVT) and good functional results utilizing images of CT angiography, in place of the reproduction of these biomarkers. In this work. These models need no picture annotation and may be calculated quickly.

To compare the Deep Learning model with the model of Machine Learning with classical biomarkers for radiological picture. We have examined topologies of the residual neural network (ResNet), adapted them to weight initialization using the Structured Receptive Fields (RFNN) as well as auto-encoders (AE). We also integrated visualization models for providing insight into the decision-making process in the network. With 1301 patients, we have applied the MR CLEAN data set methods [13].

The purpose of the study was to make available to the public for the purpose of making available for use individual patient data from the International Stroke Trial, one of the biggest randomized acute stroke studies ever performed. Methods: Data on randomised, early outcome point (14-day following randomization, or previous discharge) factors have been retrieved and presented as analyzable data for each randomized patient. Methods: The data for 19 435 acute stroke patients, full follow-up by 99%, are included as part of the IST data set.

More than 26.4% of patients had a history of aging exceeding 80 years. There was limited background stroke and no thrombolytic treatment was received by any patient. Conclusions: The IST data set contains major data sources that might be used to organize further trials, calculate samples and perform new subsequent analyzes. During planning studies of elderly patients in resource deprived settings, data may be valuable due to the distribution of age and nature of the treatment background [14].

3 PROPOSED SYSTEM

In below figure 1 shows that the flow diagram of proposed method. Initially the given two dataset is given to the pre-processing methods. The data is pre-process and feature scaling by using standardization method. Then select the features from the pre-processed data, which make the classifier as better prediction and classification purpose. Then the different machine learning classifier model to predict the heart stroke significantly.

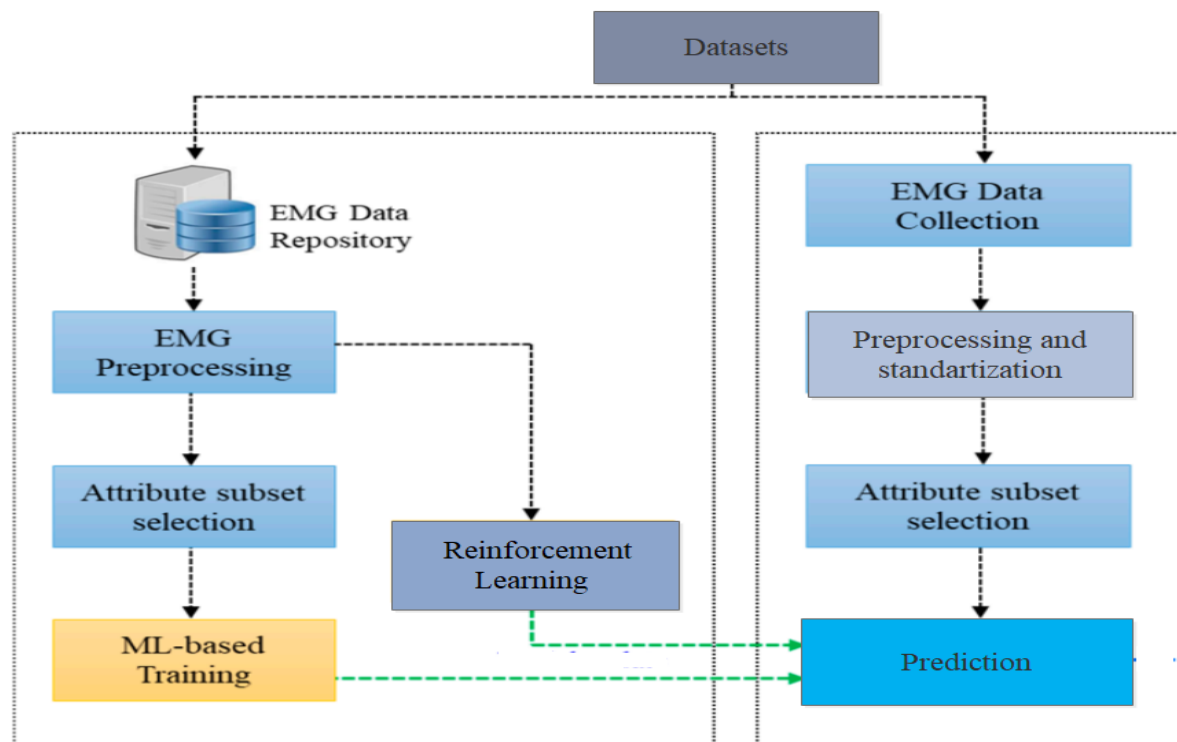


Figure 1: proposed methodology

4 DATA COLLECTION PREPROCESSING

Firstly, under the Confidentiality Agreement and the Authority clearance, 2352 EMG medical records received from TCM Sichuan Provincial Hospital are collected. As indicated earlier, even in a well-known hospital it is almost impossible to gather a huge amount of data. These more than 2300 reports have been collected for 10 months already. The measuring equipment employed in this work is a potential inspector MEB-9200K, developed by Nihon Kohden Corporation in Japan which was electromyographically revealed. Electricity is consumed at 430W. This unit is strongly extensible and leads to the position in the production fields of EMG. This equipment can be used for up to eight independent tests simultaneously. This equipment MEB-9200K is used for many years in the TCM Sichuan Provincial Hospital.

5 EXAMINATION ITEMS SELECTION

Given a huge number of EMG exam items, this study includes only sections of these things. In particular, there are two categories of EMG exam data. Due to the lesser data dimension but bigger amounts, F-MNCS and ABR have been picked. The F-MNCS test is typically employed as a major auxiliary approach for diagnosing facial paralysis in the clinical evaluation. In the meanwhile, the ABR test is commonly used in patients to evaluate tinnitus. Both test items are carried out in the head. Certain patients would have both exam items. Below are the qualities of each item. The facial engine nerve conduction study(F-MNCS) is measured by marginal jaw stimulation and the mental mental muscle's evoked EMG. The common value of F-MNCS is 48.8 ± 3.68 m/s.

The F-MNCS test can reflect the forecast of face paralysis as well as the degree of nerve damage such as Bell's paralysis (already referred to as idiopathic palsy) and of other nerve-related ailments. Therefore, F-MNCS research and applications are popular worldwide. There are numerous fields of data for the typical F-MNCS test. Measuring equipment from various manufacturers would have little impact on fields of data. 19 data fields are collected in the initial reports.

6 DATA STANDARDIZATION

For four machine-learning methods, including reinforcement, linear regression, SVM and logistic regression, both data sets stated above are utilised. In this work, before using the algorithms, the data standardization is implemented. In addition, in scenarios with and without standardisation the performance of four algorithms is compared. Comparisons are performed on the precision and retrieval rate between four algorithms. The standardization of the data is the basis for machine study. In order to evaluate an indicator both the dimension and the value of an indicator would make a huge impact. The outcomes of the data analysis would be altered without data processing. Common standardization processes include standardization of decimal calibration, standard deviations and standardization of deviations. The approach of standardized decimal calibration is to map $[-1,1]$ attribute values by shifting decimal values. The major purpose of this strategy is to avoid the units' influence.

7 ML ALGORITHM FOR CLASSIFICATION

7.1 LINEAR REGRESSION MODEL

The linear regression approach describes the link between the data and a straight line that can be expressed by the function as accurately as feasible. It is possible to express the function of the linear regression model. Since the efficacy of data fitting varies with the kind of the linear model, the model selection is important. A function that reflects the difference between the linear regression model and the actual data is used to select a more precise model to describe the linear connection between data.

It is similar to that of L2 or Euclidean distance to mean the cost function given above. Since the lower cost function value is the closer relationship and the better results.

7.2 LOGISTIC REGRESSION MODEL

The logistic regression model is a common nonlinear probabilistic regression model and can only be either 0 or 1. Probability of Y similar to 1 is indicated by the supposition that p are indie variables $X = [x_1, x_2, \dots, x_p]$, $p = P(y = 1|X)$. The Way of the Way.

7.3 SVM

SVM is one of the classic categories (Support Vector Machine). SVM has been widely used and extended to the extreme by Guang-Bin Huang et al. in their studies. There are, however, several SVM weaknesses. For example, when the sample is large and the kernel function has a large dimension, the calculation becomes significant; SVM is very sensitive to missing data; kernel features do not have a universal standard for non-linear issues. Due to the foregoing shortcomings, alternative methods eventually substituted SVM. Kernel function, the gamma kernel function and the C error penalty coefficient are key parameters of SVM. The kernel functions are the Gaussian, the polynomial, the sigmoid and the linear kernel functions. The higher the C-coefficient of error penalty the greater the exactness of the exercise sample. The generalization capacity of the model would nonetheless be less.

7.4 REINFORCEMENT LEARNING MODEL

The paradigm for reinforcement education is built on the idea of an ensemble technique of learning bagging. For ultimate prediction, many small decision trees are combined into a complex forest. Currently, ID3, C4.5 and CART are the major decision tree algorithms. First, the ID3 is considered to be the most traditional algorithm. Attributes depending on the gain of data at any node in the tree are selected at their core. Secondly, improving the C4.5 algorithm over ID3 depends instead of gaining information on the selection of the node property. C4.5 can therefore solve the problem of the continuous attributes which ID3 cannot solve. The CART technique for building reinforcement training is employed as a decision tree algorithm in this project. The distinction between the previous three methods is that the CART decision tree selects the best segmentation function based on the Gini index instead of gaining information. Every decision tree branch is binary in the meantime.

8 PERFORMANCE COMPARISON

The small amount of data, on the one hand, leads to inadequate model training. As a result, ABR data training's reinforcement learning approach isn't particularly dependable. The increase in the data set size can be projected to encourage model accuracy without standardization and to reduce the influence on the algorithm accuracy and functional values of the deviation standardization. On the other hand, the dimension would have a big impact on the tiny data set model. The classical machine learning technique in EMG data is therefore better than the DL method. The first is that the random values on the PC are not actual or ideal random values, but pseudo-random values or approximation values that could lead to over-concentrations of certain data segmentation traits, which would therefore reduce accuracy.

The second reason is that only 233 data are accessible, which results in inadequate training for the strengthening learning model, in the small amount of ABR examination data. Like the previous F-MNCS investigation, the ABR examination data feature extraction process is conducted to evaluate whether standardization deviations will affect the correlation of the data. At the same time, gains in reinforcement learning, SVM and logistical regression have been enormous. Since, however, the amount of data to make the algorithm work efficiently is not sufficient, the precision of linear regression is rather small. In short, because of the low amount of data, ABR data may overfit, even if the usual machine learning algorithms are used without large amounts of data.

9 EXPERIMENTAL SETUP AND RESULT

Firstly we ran a non-playback test on the trained linear regression model with the F-MNCS test system after a large amount of parameter tweaking for the linear regression model. The label is considered abnormal in tables of data set and all other but abnormal data fields are separate variable.

9.1 PERFORMANCE METRICS

The performance of the proposed system is verified by using different parametric measure as accuracy, sensitivity and specificity, which are expressed in the following equations correspondingly.

$$SE = \frac{tp}{tp + fn} \quad (1)$$

$$SP = \frac{tn}{tn + fp} \quad (2)$$

$$AC = \frac{tp + tn}{tp + fp + tn + fn} \quad (3)$$

$$FM = \frac{tp}{tp + 1/2(fp + fn)} \quad (4)$$

Where,

TP represent as the true positive rate, that is, the correctly classified Disease

FP is identified false positive, that is, classified the Disease incorrectly

FN is identified as the false negative, that is, the classified the Disease incorrectly

TN is identified as true negative, that is, classified the Disease correctly.

Table 1: performance evaluation of different classification technique for Congenital Heart Disease (CHD)

Methodology	Accuracy	Precision	Recall	F-measure	Execution time
Linear regression	0.741	0.785	0.695	0.654	0.52
Logistic regression	0.713	0.801	0.698	0.661	0.34
Svm (support vector machine)	0.699	0.791	0.711	0.713	0.23
Reinforcement Learning	0.805	0.813	0.724	0.743	0.18

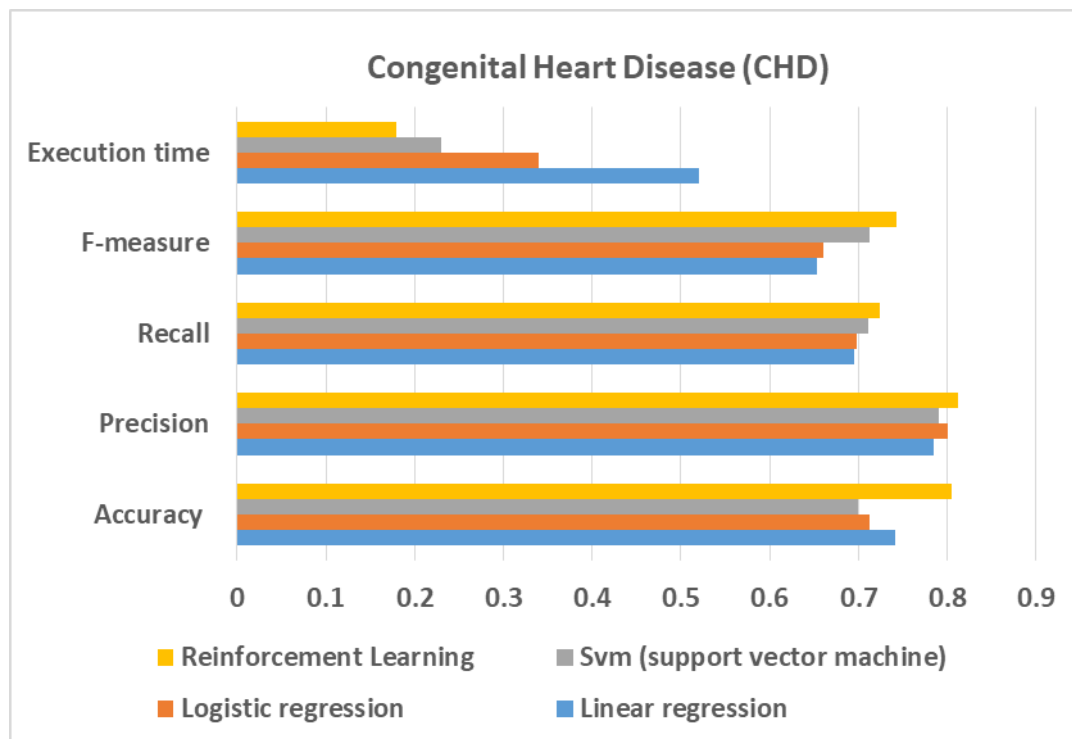


Figure 2: Graphical representation of CHD dataset performance analysis

The table 1 and figure 2 represent that the performance analysis of different classifier under CHD dataset. By this comparison analysis, Reinforcement Learning model achieve the better accuracy and less computation time than another classifier.

Table 1: performance evaluation of different classification technique for International Stroke Trial (IST) dataset.

Methodology	Accuracy	Precision	Recall	F-measure	Execution time
Linear regression	0.798	0.719	0.695	0.654	0.51
Logistic regression	0.723	0.801	0.698	0.651	0.37
Svm (support vector machine)	0.687	0.791	0.711	0.716	0.13
Reinforcement Learning	0.809	0.814	0.726	0.745	0.19

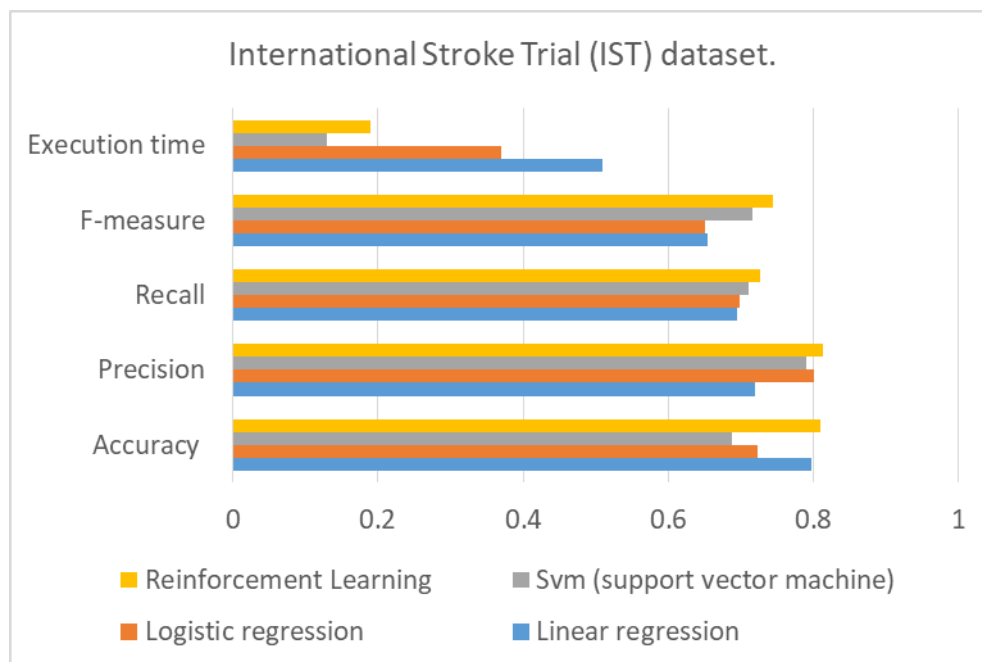


Figure 3: Graphical representation of IST dataset performance analysis

The table 2 and figure 3 represent that the performance analysis of different classifier under IST dataset. By this comparison analysis, Reinforcement Learning model achieve the better accuracy and less computation time than another classifier.

10 CONCLUSION

Furthermore, data standardization, such as derivation standardization, is a useful way for improving performance, such as accuracy. At the moment, the doctor performs this operation manually based on the waveform's properties. A comprehensive system can be developed by combining the method used in this work with additional research. We contrast three classifier models: linear regression, logistic regression, and SVM with Reinforcement Learning for prediction of heart stroke. For the experimentation study we used Congenital Heart Disease (CHD) and International Stroke Trial (IST). Datasets. According to this comparison analysis, the Reinforcement Learning model outperforms better than the other classifier in terms of accuracy and computation time as 0.805 on CHD dataset and 0.809 on IST dataset. Simultaneously, the results provide a significant reference for clinical data-based diagnosis and the enhancement of medical efficiency.

Reference

- [1]. Turkmen HI, Karsligil ME. Advanced computing solutions for analysis of laryngeal disorders. *Medical & biological engineering & computing*. 2019 Nov;57(11):2535-52.
- [2]. Lapham AC, Bartlett RM. The use of artificial intelligence in the analysis of sports performance: A review of applications in human gait analysis and future directions for sports biomechanics. *Journal of Sports Sciences*. 1995 Jun 1;13(3):229-37.
- [3]. Orjuela-Cañón AD, Ruíz-Olaya AF, Forero L. Deep neural network for EMG signal classification of wrist position: Preliminary results. In *2017 IEEE Latin American Conference on Computational Intelligence (LA-CCI) 2017 Nov 8 (pp. 1-5)*. IEEE.
- [4]. Reaz MB, Hussain MS, Mohd-Yasin F. Techniques of EMG signal analysis: detection, processing, classification and applications. *Biological procedures online*. 2006 Dec;8(1):11-35.
- [5]. Buongiorno D, Cascarano GD, Brunetti A, De Feudis I, Bevilacqua V. A survey on deep learning in electromyographic signal analysis. In *International Conference on Intelligent Computing 2019 Aug 3 (pp. 751-761)*. Springer, Cham.
- [6]. S. F. Weng, J. Reps, J. Kai, J. M. Garibaldi, and N. Qureshi, "Can machine-learning improve cardiovascular risk prediction using routine clinical data?" *PLoS ONE*, vol. 12, no. 4, Apr. 2017, Art. no. e0174944.
- [7]. J. Heo, J. G. Yoon, H. Park, Y. D. Kim, H. S. Nam, and J. H. Heo, "Machine learning-based model for prediction of outcomes in acute stroke," *Stroke*, vol. 50, no. 5, pp. 1263–1265, May 2019.
- [8]. R. Garg, E. Oh, A. Naidech, K. Kording, and S. Prabhakaran, "Automating ischemic stroke subtype classification using machine learning and natural language processing," *J. Stroke Cerebrovascular Diseases*, vol. 28, no. 7, pp. 2045–2051, Jul. 2019.
- [9]. P. Xu, G. Zhao, Z. Kou, G. Fang, and W. Liu, "Classification of cancers based on a comprehensive pathway activity inferred by genes and their interactions," *IEEE Access*, vol. 8, pp. 30515–30521, 2020.
- [10]. J. Bamford, P. Sandercock, M. Dennis, C. Warlow, and J. Burn, "Classification and natural history of clinically identifiable subtypes of cerebral infarction," *Lancet*, vol. 337, no. 8756, pp. 1521–1526, Jun. 1991.
- [11]. R. I. Lindley, C. P. Warlow, J. M. Wardlaw, M. S. Dennis, J. Slattery, and P. A. Sandercock, "Interobserver reliability of a clinical classification of acute cerebral infarction.," *Stroke*, vol. 24, no. 12, pp. 1801–1804, Dec. 1993.
- [12]. P. Amarenco, J. Bogousslavsky, L. R. Caplan, G. A. Donnan, and M. G. Hennerici, "Classification of stroke subtypes," *Cerebrovascular Diseases*, vol. 27, no. 5, pp. 493–501, 2009.
- [13]. S. Ricci, S. Lewis, and P. Sandercock, "Previous use of aspirin and baseline stroke severity: An analysis of 17 850 patients in the international stroke trial," *Stroke*, vol. 37, no. 7, pp. 1737–1740, Jul. 2006.
- [14]. S. F. Weng, J. Reps, J. Kai, J. M. Garibaldi, and N. Qureshi, "Can machine-learning improve cardiovascular risk prediction using routine clinical data?" *PLoS ONE*, vol. 12, no. 4, Apr. 2017, Art. no. e0174944.