

Applications of Machine Learning in the Industry of Healthcare

Shixuan Tang

Boston University, Boston MA, 02215, United States

billytsx@bu.edu

Abstract. Machine learning, as a branch of Artificial Intelligence, is trying to make computers do identifications, classifications, and predictions as the way humans do, but without human involvement. Machine learning has the ability to deliver quicker and more accurate results than most traditional computer algorithms. As machine learning becomes more established, its applications are widely used. This paper is going to introduce the fundamentals of four traditional machine learning algorithms (DT, RF, SVM, KNN) and one deep learning neural network (DNN). After that, this paper will illustrate how these algorithms function in assisting clinical diagnosis and disease prediction. Final results are provided with actual experiments: DT can help practitioners identify eye diseases patients where the success rate is 92%. RF is used for diagnosing diabetes patients and it is able to achieve as high as 99.7% accuracy. By searching for similar minutiae, SVM can predict Alzheimer's patients 10 years before clinical manifestations appear, and KNN performs an 81.85% prediction accuracy for potential heart disease patients. Besides that, CNN, another form of machine learning, presents a 99% accuracy in predicting Alzheimer's patients and 83% accuracy in predicting heart disease patients.

Keywords: Machine learning; Neural network; Healthcare.

1. Introduction

As one of the greatest inventions of the 20th Century, machine learning [1] has always been the hottest topic since the first introduction by Arthur Samuel in 1959. During the past sixty years, the development and innovation of machine learning have made impressive progress. Machine learning is intended to explore how to manipulate computers to think and behave the way humans do. By achieving this, computers will have the capability to acquire new information and skills without intervention from the outside, then combining those with their existing knowledge structures. This is the reason why we can expect a constant and never-ending improvement in machine learning. That also makes machine learning the core of Artificial Intelligence [2] and the benchmark of high technology. "Let the data speak for itself" is the intention of machine learning. By studying and analyzing the existing dataset, finding out the possible patterns among different samples, and using these patterns to make future identifications and predictions, machine learning is considered to be one of the most powerful weapons for human beings in many areas. Nowadays, machine learning can be found in many different industries all over the world. The applications of machine learning include agriculture [3], bioinformatics [4], economics [5], financial market analysis [6], linguistics [7], fraud detection [8], telecommunication [9], etc.

Healthcare, as one of the most important industries related to human life, has always attracted the attention of the public. Fortunately, the healthcare industry never stops making progress and never lets us down. Between 1959, which is exactly the year the term "machine learning" was coined, and 2020, the life expectancy in the United States has risen from 69.9 years to 78.8 years. According to the CDC (Centers for Disease Control and Prevention), heart disease, cancer, accidents, chronic lower respiratory diseases, cerebrovascular diseases, Alzheimer's disease, diabetes are the leading causes of death in the US.

This paper is going to focus on the applications of machine learning in the industry of healthcare. More specifically, the paper investigates how different machine learning algorithms make impacts on medical diagnosis, including the diagnosis of several top lethal diseases above.

2. Basic Concept of Machine Learning

Before we jump into the discussion of what “machine learning” can do for healthcare, let us take a look at what actually is “machine learning”. Generally, machine learning can be seen as a different genre of Artificial Intelligence. Machine Learning has three categories: supervised machine learning [10], unsupervised machine learning [11], and reinforcement learning [12]. Since supervised machine learning is the most powerful and the most frequently used in the area of healthcare, this paper will mainly focus on how different types of supervised machine learning algorithms show their great influences in assisting medical analysis.

For the model of supervised machine learning, we will train it by using the past data as labeled input. The machine itself will learn from these data and eventually be able to generate the output as prediction and classification on future problems. In order to make sure the output is germane and accurate in different situations, supervised machine learning includes more than 9 different types of algorithms: support vector machines (SVM) [13], linear regression [14], logistic regression [15], decision trees [16], random forest [17], k-nearest neighbors’ algorithm (KNN) [18], etc.

Giving a simple example of what specifically supervised machine learning can achieve in real life, people could use a group of fruit images and let the machine understand what an apple or orange or any other fruits look like and what their features are. As a result, when we have a picture of an unknown fruit in the future, the machine will help us to identify which fruit it is and which category it belongs to. Nonetheless, the process will be way more complicated and dynamic in medical analysis, but the idea is the same.

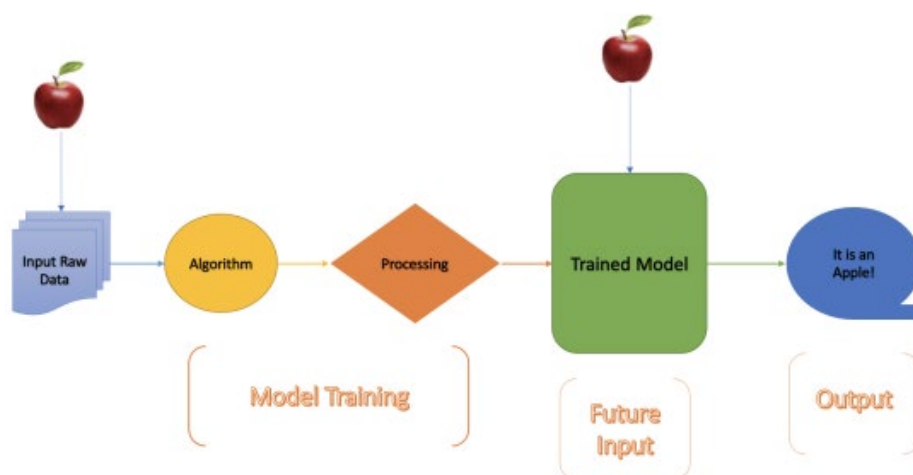


Fig.1 Making Future Prediction by Machine Learning

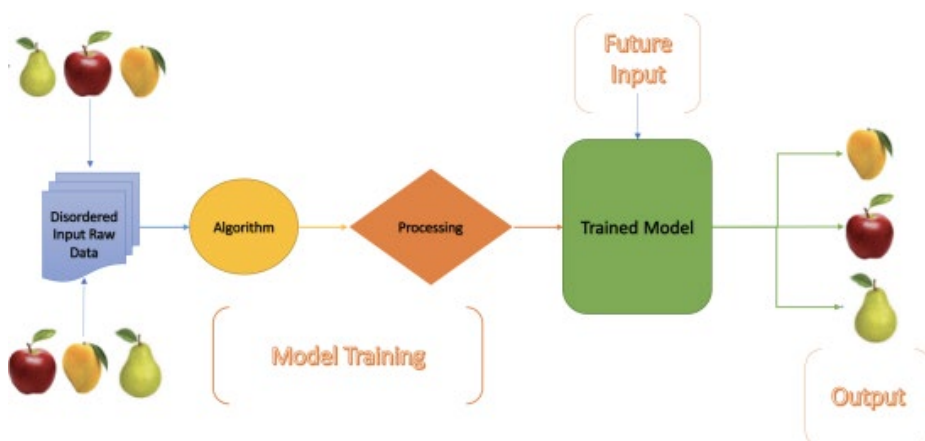


Fig.2 Making Future Classification by Machine Learning

3. Diagnosis

Now let us begin our discussion on what exactly could a supervised learning machine algorithm realize in the industry of healthcare. Diagnosis is a long and complex process that involves multiple considerations and requires tons of decisions. A tiny mistake during this process may cause an inconceivable and irreversible consequence to the patient. In the United States, misdiagnosis can destroy the life of around 12 million patients per year. In addition to that, around 40,000 to 80,000 patients each year in the U.S lose their life due to medical diagnostic errors. According to the National Academies of Sciences, Engineering, and Medicine, we all have experienced different levels of misdiagnosis once in our life, and unfortunately, up to now, there is no way we can prevent the occurrence of misdiagnosis thoroughly. However, there is a silver lining that machine learning can reduce the possibility that misdiagnosis happens to every single one of us again. Researchers suggest that the three major causes of misdiagnosis are: miscommunication and misinterpretation between the doctor and the patient, misjudgment from the doctor, and inaccuracy of the diagnostic testing.

3.1 Decision trees in diagnosing diseases

Decision tree is a decision analysis method to evaluate the project risk and judge its feasibility on the basis of the known probability of occurrence of various situations. The decision tree is one of the most popular and effective algorithms nowadays that supports people making decisions. Just like what the name illustrates, it is a tree-like diagram, and each branch of the tree represents a possible situation. The reason why it is so effective is that it can solve both classification or regression problems, which means, it can cope with both categorical or numerical data. An example image is shown below: If a person wakes up in the morning and is wondering whether he can play golf today or not, the decision tree can help him make the right decision. In order to have a nice golf experience, several conditions must be met. By looking at Figure 3, it is clear and easy to decide whether to play golf or not in different situations.

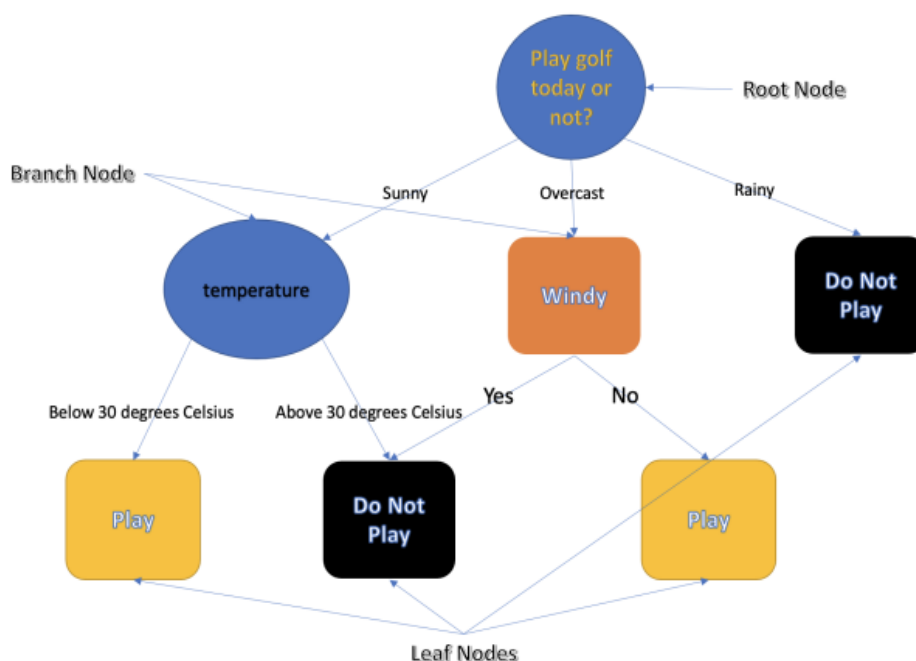


Fig.3 Example of using DT to decide whether playing golf

However, how do we know the decision tree we created for each situation is the best? Due to the fact that there are so many possible elements corresponding to influence the decision, and it is not realistic and necessary to include all of them, so which specific classifiers or filters should we choose as branch nodes (decision nodes) when building the decision tree? As to dispel misgivings, “entropy” is a good tool for us to solve the issue. Entropy is a crucial term to rely on when building the proper decision tree model. By definition, entropy measures the randomness and uncertainty of the dataset;

if the dataset is homogeneous, entropy equals zero. Its formula is illustrated below: each possible outcome X_i corresponds to its occurrence probability P_i .

$$H = - \sum_{i=1}^k P_i \times \log_2(P_i)$$

The lower the entropy is, the better the classifier is. Once we calculated the entropy of each possible scenario, we could just choose the ones with lower entropy.

This design has already been conducted as an experiment by [20]. They collect 400 patients from 18-year-old to 70-year-old where 80% of them are randomly chosen to be the training dataset and the rest 20% will be served as test data. Combined with a trained neural network, they use the decision tree model as a classifier for seven different eye diseases including glaucoma which is mentioned before. They generate as many as 22 input variables that could be served as the decision nodes, as shown in Table 1. Rather than mix categorical and numerical data, their extraction rule for the decision tree is simply an “If and then” rule. Figure 4 illustrates how their decision tree model functions and how it can eventually be used to diagnose eye diseases. It’s worth noting that their system is able to perform as high as 92% successful diagnosis.

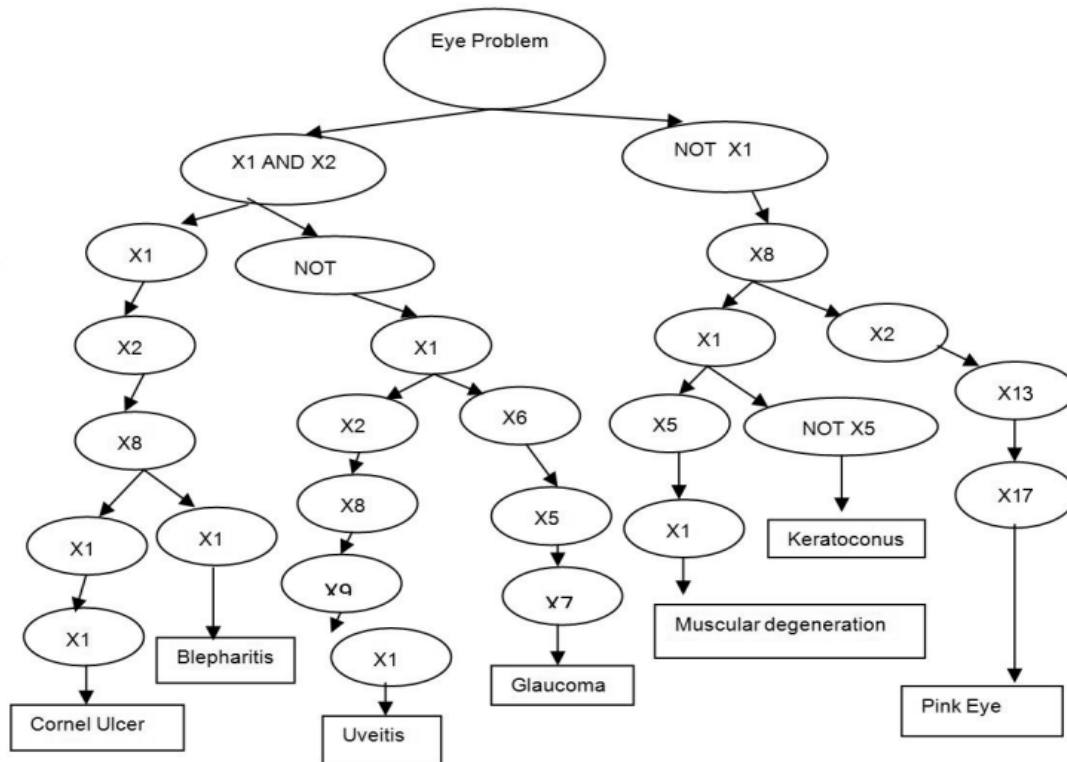


Fig.4 Diagnosing Eye Problem by DT from [20]

Another approach in eye disease diagnosis using decision trees is exhibited by [21]. Their primary target focuses on the diagnosis of dry eye disease (DED) since the lack of significant biomarkers of it may induce ophthalmologists to overlook it. They collect samples of 600 patients from both urban and rural areas. Their decision nodes include Schirmer test, tear film breakup time, corneal fluorescein staining, lissamine green staining, and ocular surface disease index. By combining all the results from these tests, they build a decision tree model which can diagnose whether a patient has DED or not. Their decision tree model is illustrated in Figure 5 and the model has a success rate of 86.1%.

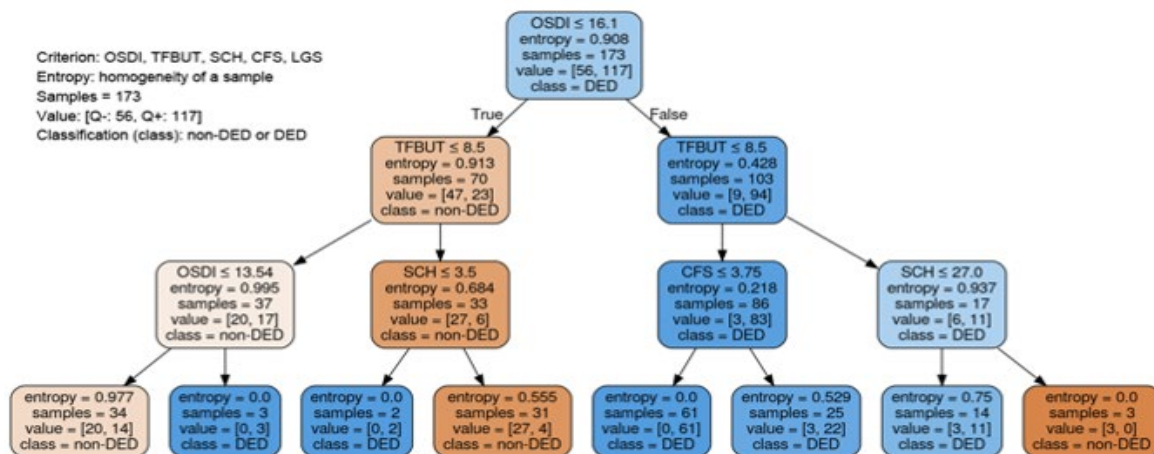


Fig.5 Diagnosing Dry Eye Disease by DT from [21]

3.2 Random forest in diagnosing diseases

Random forest is a combination of multiple decision trees. The final decision made by random forest depends on the majority votes from the decision trees [Figure 6]. Each decision tree is built by the same process mentioned in the previous section. However, the difference is, when a training dataset is available, each decision tree will only randomly receive part of the dataset as training inputs and generate its own system, so different decision trees will have different training inputs. Unlike a single decision tree which may become sensitive to the small variation in the dataset, random forests with multiple decision trees reduce the risk of overfitting and increase the accuracy of classification. For example, if there is an outlier in the dataset, it may force a single decision tree to misclassify a future new input, however, since the random forest considers the outputs from all the decision trees, the mistakes made by one won't influence the final result.

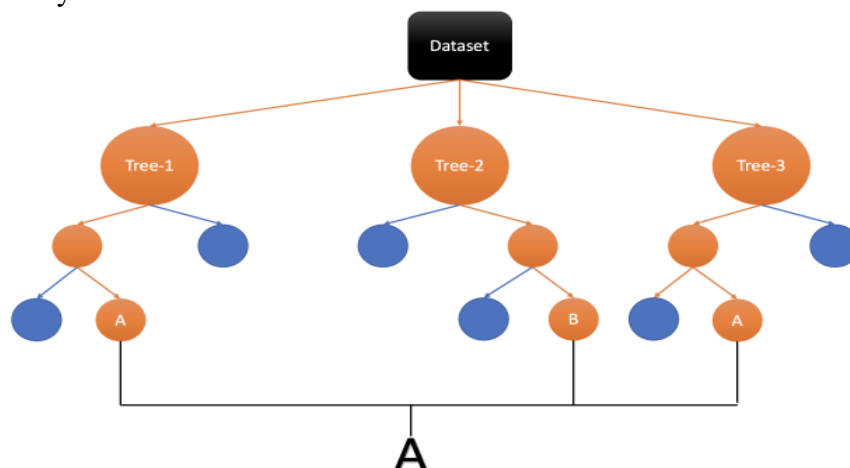


Fig.6 Random Forest Illustration

Diabetes, one of the leading causes of death in the world, causes around eighty thousand Americans to die each year. According to ADA (American Diabetes Association), around 34.2 million Americans, which is 10.5% of the entire American population, have diabetes. Moreover, For Americans whose age is above 65, 26.8% of them have diabetes. The most terrible thing is out of these 34.2 million people, 7.2 million people still remain undiagnosed. Diabetes is a chronic disease induced when people's blood glucose levels are too high. Insulin is a hormone that can transport blood glucose into the cells to give people energy. Diabetes occurs when the pancreas does not produce enough insulin or the body cannot utilize the insulin effectively. The consequences led by

diabetes include blindness, kidney failure, heart attacks, strokes, lower limb amputations, and death. Diabetes can be classified into two types. Type one diabetes is when patients' pancreas cannot produce insulin, and type two diabetes is when patients' bodies cannot utilize the insulin very well. Unfortunately, most type one diabetes is congenital.

[22] shows a practical experiment of such an idea that uses the random forest as a classifier to diagnose diabetes mellitus. They collect the information of 768 patients from the "Pima Indians Diabetes Database." 60% of them are the training dataset and the rest 40% will be the test dataset for the built random forest model. Eight attributes include age, body mass index and 2-Hour serum insulin are chosen to be the parameters for each decision tree to make the future classification. Figure 7 shows that their random forest classifier approach provides a 99.7% accuracy when classifying diabetes patients. This achievement indicates the random forest has the potential to be a very effective and decisive algorithm that will make a huge impact on future clinical diabetes diagnoses.

[24] combines different machine learning algorithms involving random forest into a new hybrid model to check if this can provide better performance. They use the same dataset with [22] and choose the identical eight parameters. In the end, their new hybrid model achieves an 81.89% accuracy in diagnosing diabetes patients.

Despite the high accuracy that random forest is able to perform in diagnosing diabetes patients, it still has some drawbacks. First of all, the problem dealing with building an optimal random forest system is the number of trees in the forest. Which number is the most suitable and how to choose is the question. There is no doubt that more trees can give more accurate results, thus, this brings us to the next problem of random forest. One limitation of using random forest clinically and widely so far is that a random forest with a large number of decision trees can take a very long time to operate. The training part is easy to finish, but the clinical classification process for each future patient might be too slow to wait. However, if the extremely precise result isn't the priority of using random forest clinically, but the time is, we could put fewer trees in the forest which will reduce the classification period. Random forest is a double-edged sword. How to find the balance between the two sides and how to make this trade-off are the key steps for future practitioners.

4. Prediction

Assisting in diagnosing diseases is a huge step for machine learning, but assisting in predicting and preventing the occurrence of diseases will bring machine learning to another level.

4.1 Support-vector machines in predicting diseases

Support-vector machine is a binary classification algorithm that allows us to classify subjects into two categories. Just to make the point unambiguous, the support vector machine is also able to show its capability doing multiple classifications, but the best battlefield for it will no doubt be a binary classification. During the set-up of a support-vector machine model, past data of the two different categories will be used as input to create a hyperplane (decision boundary) that best separates them from each other. Hyperplane exists as a line in a one-dimensional situation and some two-dimensional situations, such as identifying gender. However, in a multi-dimensional situation, where the hyperplane won't be able to serve as a line anymore. The Kernel function will be used to transform the one-dimensional input into multi-dimensional output, thus we are able to draw the different formations of the hyperplane to separate the observations. As indicated above from the example of fruits images, when a new unknown data is sent into the built model, the model will automatically classify it into one side of the hyperplane based on its characteristics. As shown in the several graphs below, blue dots and red dots represent two different categories and the x-axis and y-axis represent different features that are used to measure these two categories, such as height and weight in the example of identifying gender.

Alzheimer's, one of the most prevalent diseases among the elder population, will cause damage to the brain which will result in loss of memory and eventually lose the ability to live independently.

Alzheimer's is the most common form of dementia. Up to now, more than six million people in the U.S are experiencing different levels of Alzheimer's, and unfortunately, no treatments are available so far to cure this disease. Even more depressing, the inducement of Alzheimer's still remains unclear. One more tough fact about Alzheimer's is when symptoms or even just early signs are detected, it is already too late to turn back. Therefore, prophylactic measures are the only way to alleviate or even prevent the harm from Alzheimer's. Hopefully, Alzheimer's will be eliminated in its embryonic stage. However, giving all the elderly people around the world an appropriate prophylactic measure is not quite practical since it would cause a huge amount of human and social capital. For that reason, narrowing the range and only aiming for potential targets will be the most effective and realistic way to solve the issue. The question now, of course, is how to narrow the range and find these potential targets, hence, the special property of the support vector machine makes us wonder whether it can help us achieve this.

MRI, as known as “magnetic resonance imaging”, is a medical imaging technique used in computer-generated radiology to form pictures of the organs and tissues in the human body. This technology has the ability to create both a panoramic and detailed view of our brain. Based on the illustrations of SVM above, the question that comes up next is can we use the embryonic stage MRI brain images of all Alzheimer's patients as input, to let the machine discover what are the common features among them. Once the model is being successfully created, the next biggest challenge will be are we able to use this as a prediction model in which one side of the hyperplane is labeled as “potential Alzheimer's patients”, and the other side is “no signs are detected”. In the future application stage, if the model detects similarity between the new image and the training input, the new image will be classified as “potential Alzheimer's patients”. The most effective part of this model is it has the capability to detect the minutiae which humans cannot even notice normally. As long as this idea could be realized, the usage of this model in clinical practice will be helping doctors determine which patient will eventually develop Alzheimer's disease and which one may not. Therefore, it answers the initial question which is how to narrow the range and find the potential targets.

4.2 K-nearest neighbors algorithm(KNN) in predicting diseases

Just like what the support vector machine can accomplish, KNN is another classification algorithm belonging to supervised machine learning. However, unlike SVM which could only identify binary distributed patterns, KNN is able to classify undistributed patterns into multiple categories. Even though SVM is pretty good against the threat of outliers while KNN is quite sensitive to them, this won't prevent KNN from being one of the most powerful classification algorithms. The approach a KNN classifier uses for identification is called features similarity. Different from SVM in which the hyperplane is generated to classify new data points, KNN classifies a new data point based on how its neighbors are classified. “K” in “KNN” stands for how many nearest points will be selected in the classification process. Among these “K” nearest points, the majority side will welcome the new data point to join them. An example image of KNN classification is shown in Figure 11: the red dot is considered as the new data point. When the “K” is chosen to be 3, among the three nearest points around the red dot, two out of three are blue, so the new data point will be classified as class B. However, if the “K” equals 5, in this case, three out of five closest points around the red dot are yellow, thus the new data point will go to class A. As you can see, how to choose the right “K” will be the key step in the entire ‘KNN’ classification process; it is being called “parameter tuning”: The better the “K” is, the more accurate the classification will be. Unfortunately, no well-defined methods are available to find the optimal “K” value so far. In most situations, the “K” value is chosen randomly at the beginning, then narrows down to the possible optimal one.

Due to the fact that KNN is going to perform the best when dealing with continuous features as input, heart rate, blood pressure, LDL (low-density lipoprotein cholesterol), triglycerides, or any other measurements from different types of heart disease patients could be selected as the training input. Besides, such measurements of some healthy people could also be chosen as a reference in order to make sure the model won't misclassify a completely robust person into a heart disease group. Once

the model is being successfully built, it could be applied in an annual physical examination for everyone. As long as there are similarities between a person's test results with the training dataset, the model will alarm the doctor or the inspector to notice that this person might be a potential heart disease patient. Afterward, more comprehensive examinations will be carried out and future prophylactic measures will be implemented.

KNN is facing the same type of problems that SVM faces, but even more severe. Finding all the possible measurement data of different types of heart disease patients is hard but not impossible, however, the real difficulty is how to store them. Unlike SVM where the training samples will only be used to create the hyperplane, the training samples of KNN need to be stored permanently in order to cope with different situations. For the mission of storing all the possible training data, a very large memory is required. Again, if the model is desired to serve the entire human population, the total amount of training data is inconceivable, and no existing hardware so far is available to handle such a tremendous task. Furthermore, the computational cost would be extremely expensive if the model is used widely. For every single new data, the prediction needs to calculate all the distance between the data point and each individual training data in order to find the nearest ones. Such a computation assignment requires a very long period of time and a very strong software platform. Ideally, using KNN as an effective and decisive tool to discover potential heart disease patients and classify them into each type is possible, but before all the barriers are cleared away, it still has a long way to go to be applied clinically.

4.3 Convolutional neural network(CNN) in predicting diseases

CNN is not a typical supervised machine learning algorithm, but generally, it is trained using a supervised method. More specifically, CNN is a class of neural networks that belong to deep learning [30], and deep learning is a subset of machine learning. Similar to those traditional supervised machine learning algorithms illustrated before, the CNN is also a popular and powerful tool for image analysis. The first CNN called "LeNet"[24] was introduced in 1988 by Yann LeCun, who is currently the director of Facebook's AI researcher group. The applications of CNN nowadays are widely used: facial recognition [35].

CNN is not only able to predict Alzheimer's, but heart diseases as well. If you can recall, heart disease is one of the leading causes of human death today. One obstacle preventing KNN to fully examine and predict potential heart disease patients is the lack of existing data and strong hardware. Due to this unavoidable challenge which is being called "imbalance" by [36], a convolutional neural network model with good resilience is generated by them to counter such difficulty. They use the data from National Health and Nutritional Examination Survey (NHANES) and they choose two periods which are 1999-2000 and 2015 -2016. The dataset is compiled by around 37,000 CHD (coronary heart disease) patients and 35,000 non-CHD individuals in their sample. As many as 30 continuous variables and 6 categorical variables are chosen to be the parameters for predicting the latent CHD. The list of these variables comprises gender, age, body mass index, white blood cells, red blood cell width, uric acid, albumin, alkaline phosphatase (ALP), lactate dehydrogenase (LDH), etc. Since there are too many variables and it will take so much time to train the model with all of them, they rely on LASSO (least absolute shrinkage and selection operator) to help them find out which several variables are the most suitable for predicting. Parameters, such as aspartate aminotransferase (AST), cholesterol, glycohemoglobin, uric acid, blood glucose levels, and platelet count, are being selected. Figure 7 gives a quick look at what their CNN architecture looks like and its best prediction accuracy is around 83%.

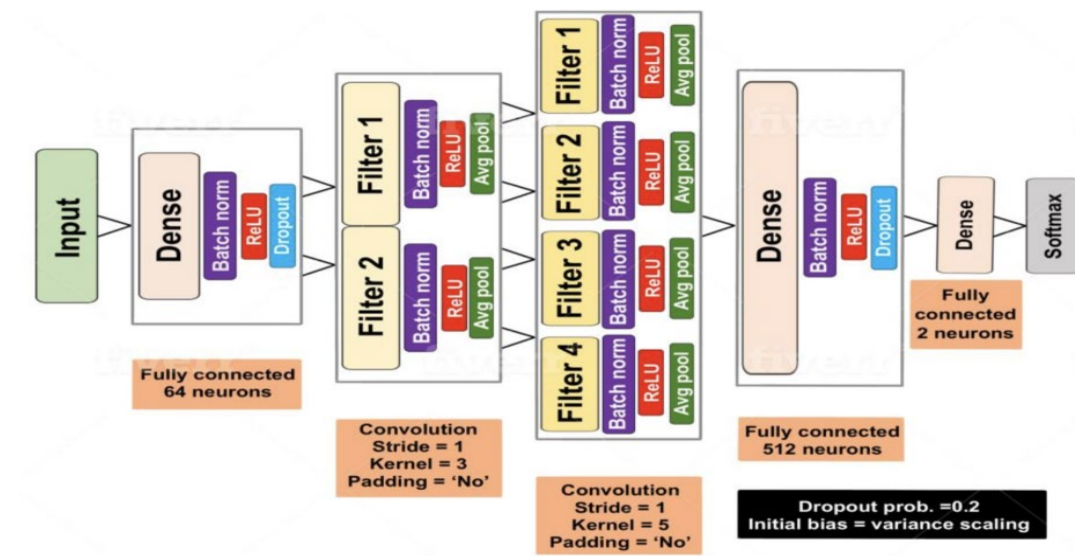


Fig.7 Illustration of the Convolutional Neural Network Architecture from [36]

If we only compare the final result, CNN clearly outperforms both SVM and KNN in terms of predicting Alzheimer's and heart diseases. It is not hard to understand since CNN is a more complicated algorithm and requires a complex build-up process compared to all traditional supervised machine learning algorithms. For future practitioners, more accurate results by using CNN or easier manipulations by using traditional supervised machine learning algorithms is the dilemma they have to face.

5. Conclusion

This paper has illustrated four different types of traditional supervised machine learning algorithms (DT, RF, SVM, KNN) and one deep learning neural network (CNN). By taking the examples from references, this paper has shown the applications of these machine learning algorithms in the current industry of healthcare and their potential applications in the future. More than that, this paper also points out the limitations of using these machine learning algorithms clinically and the suggestions for overcoming these limitations. More specifically, DT and RF could address the misdiagnosis of eye diseases and diabetes. RF can give a higher diagnosis accuracy, but the operation time might be quite long in clinical situations. DT can deliver the result in a short period of time, but it may have some misclassification in similar diseases. SVM and KNN exhibit their promise in assisting practitioners to predict dementia and heart disease even before clinical manifestations appear, while CNN could achieve both. SVM and KNN are easier to implement for clinical practitioners, but the prediction accuracy might be less than CNN. CNN requires a very complex build-up process, but the prediction accuracy will be optimal. This paper could stand as an introduction for future practitioners or a reference for future researchers. Hopefully, this paper can help them find a possible solution to incorporate machine learning with clinical experiments into a high-performance model that will take a decisive role in clinical diagnosis situations. In the future, more and more machine learning algorithms will take the dominant positions in the industry of healthcare, and by then, wish to see no more people suffer from diseases anymore.

References

- [1] Alpaydin E. Introduction to machine learning[M]. MIT press, 2020.
- [2] Russell S, Norvig P. Artificial intelligence: a modern approach[J]. 2002.
- [3] Liakos K G, Busato P, Moshou D, et al. Machine learning in agriculture: A review[J]. Sensors, 2018, 18(8): 2674.

- [4] Larranaga P, Calvo B, Santana R, et al. Machine learning in bioinformatics[J]. *Briefings in bioinformatics*, 2006, 7(1): 86-112.
- [5] Mullainathan S, Spiess J. Machine learning: an applied econometric approach[J]. *Journal of Economic Perspectives*, 2017, 31(2): 87-106.
- [6] Henrique B M, Sobreiro V A, Kimura H. Literature review: Machine learning techniques applied to financial market prediction[J]. *Expert Systems with Applications*, 2019, 124: 226-251.
- [7] Lin Z. A Methodological Review of Machine Learning in Applied Linguistics[J]. *English Language Teaching*, 2021, 14(1): 74-85.
- [8] Yee O S, Sagadevan S, Malim N H A H. Credit card fraud detection using machine learning as data mining technique[J]. *Journal of Telecommunication, Electronic and Computer Engineering (JTEC)*, 2018, 10(1-4): 23-27.
- [9] Qureshi S A, Rehman A S, Qamar A M, et al. Telecommunication subscribers' churn prediction model using machine learning[C]//*Eighth international conference on digital information management (ICDIM 2013)*. IEEE, 2013: 131-136
- [10] Kotsiantis S B, Zaharakis I, Pintelas P. Supervised machine learning: A review of classification techniques[J]. *Emerging artificial intelligence applications in computer engineering*, 2007, 160(1): 3-24.
- [11] Gentleman R, Carey V J. *Unsupervised machine learning[M]//Bioconductor case studies*. Springer, New York, NY, 2008: 137-157.
- [12] Sutton R S, Barto A G. *Reinforcement learning: An introduction[M]*. MIT press, 2018.
- [13] Noble W S. What is a support vector machine? [J]. *Nature biotechnology*, 2006, 24(12): 1565-1567.
- [14] Seber G A F, Lee A J. *Linear regression analysis[M]*. John Wiley & Sons, 2012.
- [15] Wright R E. *Logistic regression[J]*. 1995.
- [16] Quinlan J R. Induction of decision trees[J]. *Machine learning*, 1986, 1(1): 81-106.
- [17] Breiman L. Random forests[J]. *Machine learning*, 2001, 45(1): 5-32.
- [18] Guo G, Wang H, Bell D, et al. KNN model-based approach in classification[C]//*OTM Confederated International Conferences" On the Move to Meaningful Internet Systems"*. Springer, Berlin, Heidelberg, 2003: 986-996.
- [19] Neely DC, Bray KJ, Huisinigh CE, Clark ME, McGwin G, Owsley C. Prevalence of Undiagnosed Age-Related Macular Degeneration in Primary Eye Care. *JAMA Ophthalmol*. 2017;135(6):570–575. doi:10.1001/jamaophthalmol.2017.0830
- [20] Kabari, Ledisi & Nwachukwu, Enoch. (2012). Neural Networks and Decision Trees For Eye Diseases Diagnosis. 10.5772/51380.
- [21] Denny Marcos Garcia, Leidiane Adriano, Regina Celia Nucci Pontelli, Eduardo M Rocha; A decision tree approach in dry eye disease: short questionnaire versus signs and symptoms diagnosis based in a household epidemiologic study.. *Invest. Ophthalmol. Vis. Sci*. 2020;61(7):95.
- [22] M., Butwall, and S., Kumar, "A Data Mining Approach for the Diagnosis of Diabetes Mellitus using Random Forest Classifier," *International Journal of Computer Applications*, vol. 120, pp. 0975– 8887, 2015.
- [23] Daghistani, Tahani & Alshammari, Riyad. (2016). Diagnosis of Diabetes by Applying Data Mining Classification Techniques. *International Journal of Advanced Computer Science and Applications*. 7. 10.14569/IJACSA.2016.070747.
- [24] A.K., Dewangan, and P., Agrawal, "Classification of Diabetes Mellitus Using Machine Learning Techniques," *International Journal of Engineering and Applied Sciences*, vol. 2, 2015.
- [25] Battineni G, Chintalapudi N, Amenta F. Machine learning in medicine: Performance calculation of dementia prediction by support vector machines (SVM)[J]. *Informatics in Medicine Unlocked*, 2019, 16: 100200.