

TECHNIQUES FOR PREDICTING LIVER DISEASE USING MACHINE LEARNING

Sayed Taha Ali¹, Roshan Kumar Gouda¹, MD Sohail Khan¹, and Guriya Kumari¹

¹NSHM Knowledge Campus, Durgapur

Abstract: The liver, being a crucial organ in the body, plays a primary role in various functions such as bile production, bile and bilirubin excretion, protein and carbohydrate metabolism regulation, enzyme activation, glycogen and vitamin storage, synthesis of plasma proteins, and clotting factor regulation. However, the liver is susceptible to adverse effects from factors like alcohol consumption, painkiller usage, dietary habits, and other unconventional practices. Presently, the identification of liver-related diseases relies mainly to analyze liver function blood test reports and scan results, which is both time-consuming and expensive. To streamline the diagnosis of liver diseases, various data mining algorithms are employed. Utilizing a larger dataset enhances the accuracy of predictions. To address issues related to local storage limitations in healthcare centers, cloud storage is implemented, given the substantial volume of documents generated in healthcare settings. The problem is framed as a binary classification issue, utilizing ten features for each data point and a label indicating whether the patient is afflicted with liver disease or not.

Keywords: liver disease, machine learning, binary classification, naive bayes, KNN, random forest.

1. INTRODUCTION

Currently, liver-related diseases have become prevalent due to the rise in sedentary lifestyles and physical inactivity. While the severity is relatively manageable in rural areas, liver diseases are highly common in urban settings, particularly in large cities. Liver disease is a leading cause of millions of annual fatalities, with viral hepatitis alone claiming the lives of 1.34 million individuals each year. Diagnosing liver problems early poses a challenge, as even partial damage may not immediately affect its normal functioning. The timely identification of liver issues is crucial for improving patient survival rates. Notably, Indians exhibit higher rates of liver failure.

Machine Learning: The study of machine learning (ML) revolves around comprehending and refining “learning” methodologies—strategies that utilize data to enhance performance in specific tasks. Regarded as an integral component of AI, machine learning algorithms autonomously construct models by utilizing sample data, known as training data, enabling them to make predictions or assessments without explicit programming. These algorithms find applications in diverse fields such as computer vision, speech recognition, email filtering, medicine, and agriculture, addressing tasks where traditional algorithms prove either impractical or unattainable. A segment of machine learning shares a close association with

computational statistics, concentrating on the use of computers for predictive purposes. However, it's important to note that not all machine learning falls under the category of statistical learning. The exploration of mathematical optimization adds methods, theories, and application domains to the realm of machine learning. Another connected field, data mining, focuses on exploratory data analysis through unsupervised learning. Certain machine learning implementations replicate the functions of a biological brain, employing data and neural networks. When applied to address business challenges, machine learning is frequently termed predictive analytics.

Supervised Learning: In supervised learning, algorithms create a mathematical model by utilizing a dataset that contains both input values and their corresponding desired outputs. This dataset, known as training data, is composed of various training examples. Each training example includes one or more inputs and the desired output, often referred to as a supervisory signal. Through iterative optimization of an objective function, supervised learning algorithms develop a function capable of predicting outputs for new inputs. An optimal function enables the algorithm to accurately determine outputs for inputs not present in the training data. Algorithms which leads to the improvement of the accuracy of its outputs or predictions over time can be considered to have learned to perform the given task.

Unsupervised Learning: In case of unsupervised learning algorithms, a collection of data is taken which consists solely of inputs and use them to identify patterns within the data, such as clustering or

grouping of data points. In this case test data which is unlabeled is used to teach the algorithms.

Reinforcement Learning: The subset of machine learning, centering on the decision-making process of intelligent agents within an environment to maximize cumulative rewards. Reinforcement learning is considered as one of the fundamental paradigms in machine learning, apart from supervised learning and unsupervised learning mechanisms.

2. LITERATURE SURVEY

This section delves into research concerning the application of AI to diagnose liver diseases within the timeframe from 2017 to 2019. A number of studies have leveraged the Indian Liver Patient Dataset (ILPD) obtained from the UCI machine repository [1]. This dataset encompasses 416 records of patients afflicted with liver conditions and 167 records of individuals without liver ailments, sourced from the North East region of Andhra Pradesh, India. The “Dataset” column functions as a class label, categorizing individuals into those with liver diseases and those without. The dataset comprises of patient records of 142 female patient, 441 male patients records. Patients aged 90 or older are uniformly labeled as being of age “90.”

Comprising 416 unique instances, each characterized by 10 attributes, the dataset also includes an additional attribute indicating the presence and severity of the disease.

The SPSS Software showcased the highest accuracy, reaching 87.91% [2]. Extracting significant features from the Region of Interest (ROI) in CT scans from 80 patients was achieved using both SFSS and GA (genetic algorithm). To classify fatty and cirrhotic livers techniques Linear Vector Quantization (LVQ) Neural Network, Probabilistic Neural Network (PNN), Back Propagation Neural Network (BPNN) were used, where the performance of PNN is better than LVQ and BPN, achieving an accuracy (recognition rate) of 95% [3]. A study which compares SWE measurements with biopsy scores for estimation of fibrosis in patients suffering from chronic liver disease is considered. The empirical findings have suggested a correlation between SWE estimates of liver stiffness and the severity of fibrosis. This suggests that SWE could be considered as a valuable tool to distinguish patients with different levels of fibrosis [4]

A machine learning model which aims to predict the stage of liver fibrosis, employed the Decision Tree (DT) classifier based on the records of 100 HCV patients. Following data preprocessing, crucial features were identified using ANOVA and BOX plot tests, resulting in an accuracy of 93.7% [5]. Rather than relying on blood markers, a different set of variables were extracted from the medical records of 2060 patients. This included alcoholic cirrhosis, sex, age, viral hepatitis, various types of chronic hepatitis, alcoholic fatty liver disease, other types of fatty liver disease, and hyperlipidemia. The dataset was divided

into 70% for training and 30% for testing, employing four classifiers—ANN, LR, SVM, and DT. Among these, ANN exhibited the most promising results with an AUROC of 0.873 [6]. Meta-learning algorithms, specifically Ad boost, Logit boost, Bagging, and Grading, underwent comparison for the classification of ILPD [7]. It was discerned that grading stood out as the most effective algorithm, demonstrating the highest Correct Classification rate, the minimum incorrect classification rate, and the shortest execution time. The feature selection stage have identified a set of 40 significant features, including age, BMI, and 37 bacterial species. By utilization of these features, the Random Forest (RF) classifier accomplished an AUROC of 0.936 in detection of advanced fibrosis [8].

A research study [9] presented a decision-support system tailored for detecting the stage of hepatitis B, utilizing 513 Radiomics Texture Energy (RTE) images from patients. Age, gender, and 11 RTE features extracted from the images were utilized to train Naive Bayes (NB), Random Forest (RF), k-Nearest Neighbors (KNN), and Support Vector Machine (SVM) classifiers. Among these classifiers, the RF algorithm showcased the highest accuracy at 91.25%. Following this, four machine learning classifiers (logistic regression, ridge regression, decision tree, and AdaBoost) underwent training with these features. Their performance was evaluated to identify the model with the most accurate predictions. The results indicated that ridge regression achieved the highest Area Under the Receiver Operating Characteristic curve (AUROC) values of 0.87 and 0.88 in the training and validation groups, respectively [10].

In a research study, the C5.0, CHAID (Chi-square Automatic Interaction Detector), and CART (Classification and Regression Tree) algorithms were integrated with MLPNN, and their performance was evaluated in classifying ILPD. Among them, MLPNN-C5.0 demonstrated the highest accuracy at 94.12% [11]. In the classification of ILPD, Decision Stump surpassed J48, Logistic Model Tree (LMT), RF, Random Tree, REPTree, and Hoeffding Trees, achieving an accuracy of 70.67% [12]. J48, MLP, Bayesnet, RF, and SVM classifiers were employed to classify the ILPD dataset, pre-processed through min-max normalization. The Kruskal-Wallis and Chi-Square tests identified a robust correlation between age, platelet count, AST (aspartate aminotransferase), and albumin with advanced fibrosis, leading to their inclusion in model creation. The dataset, divided into training (22,690 patients) and testing (16,877) sets through random uniform sampling, underwent classification using the four algorithms with 10-fold cross-validation. The results indicated that the ADT model demonstrated the highest efficiency and AUROC of 84.4% and 0.76, respectively [13].

The proposed approach achieved an ROC of 93.06%, surpassing other technologies. This innovative method involves training the model with genetic algorithms, normalization, feature selection, and 5-fold cross-validation. Approximately ten machine learning algorithms were employed for classification, with

SVM (type C-SVC) using the proposed method demonstrating the highest accuracy at 88.49% [14]. In a study comparing Fuzzy Logic, Fuzzy Neural Networks, and Decision Trees for classifying the BUPA liver dataset, it was concluded that the Neuro-Fuzzy system, specifically Fuzzy NN, attained the highest accuracy of 91% [15]. A Deep Learning Radiomics of Elastography (DLRE) diagnosed liver fibrosis using a CNN classifier. Augmentation of 1330 two-dimensional shear wave elastography (2D-SWE) images was conducted through random transformations to expand the dataset. A DLRE Region of Interest (ROI) with the entire 2D-SWE was selected for classifying 660 images using CNN. Comparative analysis with 2D-SWE, APRI, and FIB-4 revealed that DLRE provided superior classification, boasting an AUROC of 0.97 for cirrhosis and 0.98 for advanced fibrosis [16].

3.1. ARCHITECTURE

Currently, Machine Learning Architecture holds significant attention in various industries as organizations strive to optimize resources and enhance output based on historical data. Furthermore, when integrated with data science technology, machine learning provides substantial benefits in terms of data forecasting and predictive analytics. The architecture of machine learning delineates the distinct layers in the machine learning cycle and encompasses the crucial steps involved in converting raw data into training datasets, facilitating effective decision-making for systems.

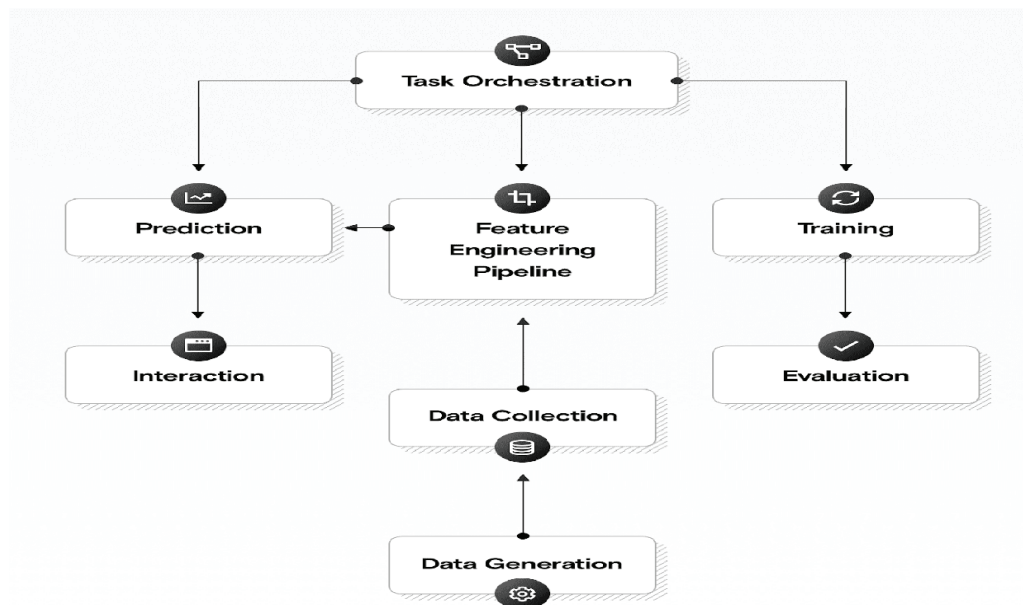


Fig. 1. Machine Learning Architecture

3.2. METHODOLOGY

The objective is to determine an appropriate machine learning algorithm capable of discerning whether an individual has liver disease or not. This constitutes a binary classification problem to be addressed through supervised learning. Each data point is characterized by ten features, accompanied by a label indicating whether the patient is afflicted by liver disease. To resolve this, the goal is to train diverse supervised learning models on the dataset. The aim is to develop a robust model that can effectively classify new data points as positive or negative with a satisfactory level of accuracy, outperforming established benchmarks.

FEATURES: -

- Age of the patient
- Gender of the patient
- Total Bilirubin
- Direct Bilirubin
- Alkaline Phosphatase
- Alamine Aminotransferase
- Aspartate Aminotransferase
- Total Proteins
- Albumin
- Albumin and Globulin Ratio
- Dataset: field used to split the data into two sets (patient with liver disease, or no disease)

3.2.1. ALGORITHM

In the realm of medical diagnosis, healthcare professionals employ various pathological methods to analyse medical reports and assess patients' conditions. In the contemporary era, the advent of computers and technology allows for data collection and the uncovering of concealed insights. Specific statistical machine learning algorithms tailored to particular issues can aid in decision-making processes. These data-driven algorithms not only validate existing methods but also assist researchers in proposing novel decisions. This paper employs multiple imputation by chained equations to address missing data and Principal Component

Analysis for dimensionality reduction. Data visualizations were employed to highlight significant findings. The study presents and compares several binary classifier machine learning algorithms, including Artificial Neural Network, Random Forest, and Support Vector Machine. These algorithms were utilized to classify individuals as blood donors or non-blood donors with hepatitis, fibrosis, and cirrhosis. Leveraging data from UCI-MLR [1], these techniques were applied to identify an optimal method for classifying blood donors and non-blood donors with the mentioned conditions, aiming to support healthcare professionals in making informed decisions in a laboratory setting.

CONFUSION MATRIX: -

Confusion Matrix plays a crucial role in assessing the performance of a model. Once we obtain the data and go through the steps of cleaning, pre-processing, and wrangling, the next critical step is to input it into a robust model, generating output probabilities. Evaluating the effectiveness of the model becomes essential for optimising performance. The Confusion Matrix becomes a focal point in this context, serving as a key tool for performance measurement.

		Actual Values	
		Positive (1)	Negative (0)
Predicted Values	Positive (1)	TP	FP
	Negative (0)	FN	TN

Fig.2. Confusion matrix

In the context of machine learning classification, a confusion matrix serves as a performance measurement tool, particularly when the output can belong to two or more classes. It consists of a table with four distinct combinations of predicted and actual values, proving highly valuable for evaluating Recall, Precision, Specificity, and Accuracy.

- True Positive: This occurs when a positive prediction is made, and it is accurate. For instance, predicting a woman is pregnant, and she indeed is.
- True Negative: This signifies a correct negative prediction. For instance, predicting a negative outcome, and it proves to be accurate.
- False Positive (Type 1 Error): This transpires when a positive prediction is made, but it is inaccurate. For instance, predicting a man is pregnant when he is not.
- False Negative (Type 2 Error): This arises when a negative prediction is made, but it is inaccurate. For

example, predicting a woman is not pregnant when she actually is.

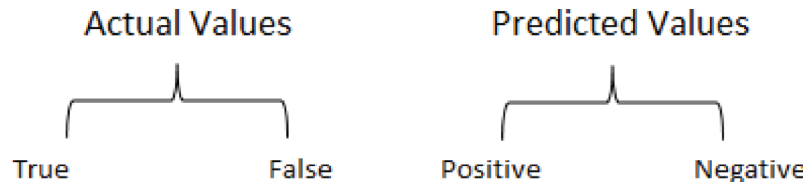


Fig.3. Actual vs Predicted values

DEEP LEARNING:

Deep learning, a subset of artificial intelligence and a branch of machine learning, is grounded in the concept of neural networks, which emulate the human brain. In contrast to explicit programming, deep learning relies on the utilization of numerous nonlinear processing units to perform feature extraction and transformation. Each subsequent layer in deep learning takes the output from the preceding layer as input, mirroring the interconnected nature of neural networks.

Deep learning models exhibit the capability to autonomously focus on pertinent features without extensive guidance from programmers, making them particularly effective in addressing the challenge of dimensionality. These models are especially valuable when dealing with a large number of inputs and outputs.

As an evolution of machine learning, itself a subset of artificial intelligence, deep learning aligns with the overarching goal of artificial intelligence—to replicate human behavior. The fundamental idea behind deep learning is to construct algorithms that mimic the human brain. This implementation occurs through neural networks, inspired by biological neurons, the building blocks of brain cells. Essentially, deep learning is realized through deep networks, which are neural networks characterized by multiple hidden layers.

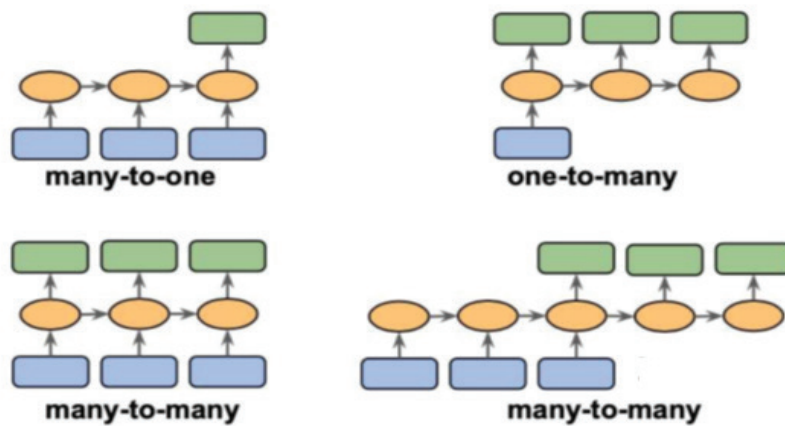


Fig. 4. Sequential neural network

Sequence models refer to machine learning models that either take in or produce data sequences. Examples of sequential data include text streams, audio clips, video clips, and time-series data. Among the various methods used in sequence models, Recurrent Neural Networks (RNNs) stand out. The need to analyze sequential data, such as text sentences and time-series data, led to the development of Sequence Models. These models demonstrate a superior capacity to handle sequential data compared to Convolutional Neural Networks, which are more adept at processing spatial data.

4. RESULTS AND DISCUSSION:

The implementation utilizes the following system specifications: Operating System - Windows 11, RAM - 16 GB, Processor - Intel Core i5, and Memory - 500 GB. Google Colab, short for Colaboratory, is employed as the coding platform. Developed by Google Research, Colab enables users to write and execute Python code seamlessly through a web browser. It is particularly well-suited for tasks such as machine learning, data analysis, and education. Technically, Colab functions as a hosted Jupyter notebook service, requiring no setup while offering free access to computing resources, including GPUs. The Python libraries employed in this implementation include NumPy, Pandas, Matplotlib, Scikit-learn, and Seaborn. The steps involved here in implementing a) Importing the libraries b) Read and Load the Data c) Commence Data Processing d) Checking if there are any NULL values in our Dataset and Removing the NULL values d) Transform the data.

```
#selecting and dividing dataset to data and labels

x=dataset.iloc[:,0:-1].values
y=dataset.iloc[:, -1].values

#splitting dataset in to train and test

from sklearn.model_selection import train_test_split
x_train,x_test,y_train,y_test = train_test_split (x,y,test_size = 0.2, random_state=0)
```

Fig. 5. Screenshot of code

```
[ ] #Exploratory Data Analysis

#Determining the healthy-affected split
print("Positive records:", data['Dataset'].value_counts().iloc[0])
print("Negative records:", data['Dataset'].value_counts().iloc[1])

Positive records: 404
Negative records: 162

[ ] #The above output confirms that we have 414 positive and 165 negative records. This indicates that this is a highly unbalanced dataset.

#Determine statistics based on age
plt.figure(figsize=(12, 10))
plt.hist(data[data['Dataset'] == 1]['Age'], bins = 16, align = 'mid', rwidth = 0.5, color = 'black', alpha = 0.8)
plt.xlabel('Age')
plt.ylabel('Number of Patients')
plt.title('Frequency-Age Distribution')
plt.grid(True)
plt.savefig('fig1')
plt.show()
```

Fig. 6. Screenshot showing details of exploratory data analysis

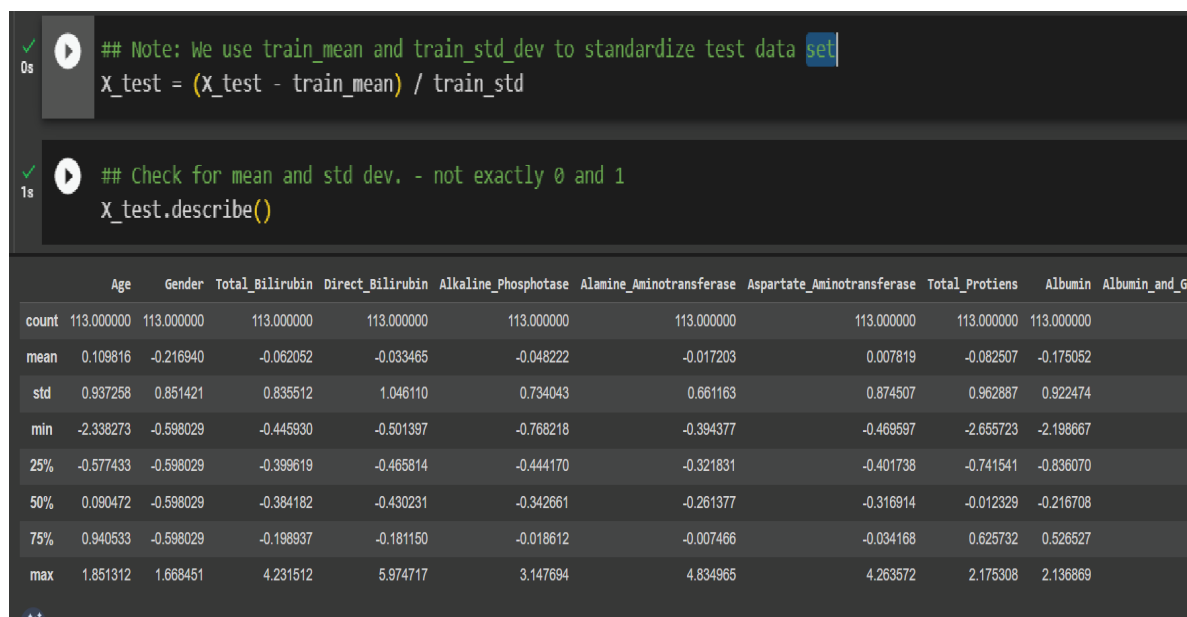
```
[47] train_mean = X_train.mean()
     train_std = X_train.std()

[48] ## Standardize the train data set
     X_train = (X_train - train_mean) / train_std

[49] ## Check for mean and std dev.
     X_train.describe()
```

Fig. 7. Screenshot of data standardization

During the data standardization process, we execute zero mean centering and unit scaling, meaning that we set the standard deviation to one and the mean of all features to zero. Therefore, we utilize each feature's mean and standard deviation. Because we utilize the same mean and standard deviation in the test set, it is crucial to record the mean and standard deviation for every feature from the training set.



The screenshot shows a Jupyter Notebook interface with two code cells. The first cell contains a comment and a line of code to standardize test data. The second cell contains a comment and a line of code to check the mean and standard deviation of the test data. Below the code cells is a summary statistics table for the test data.

```

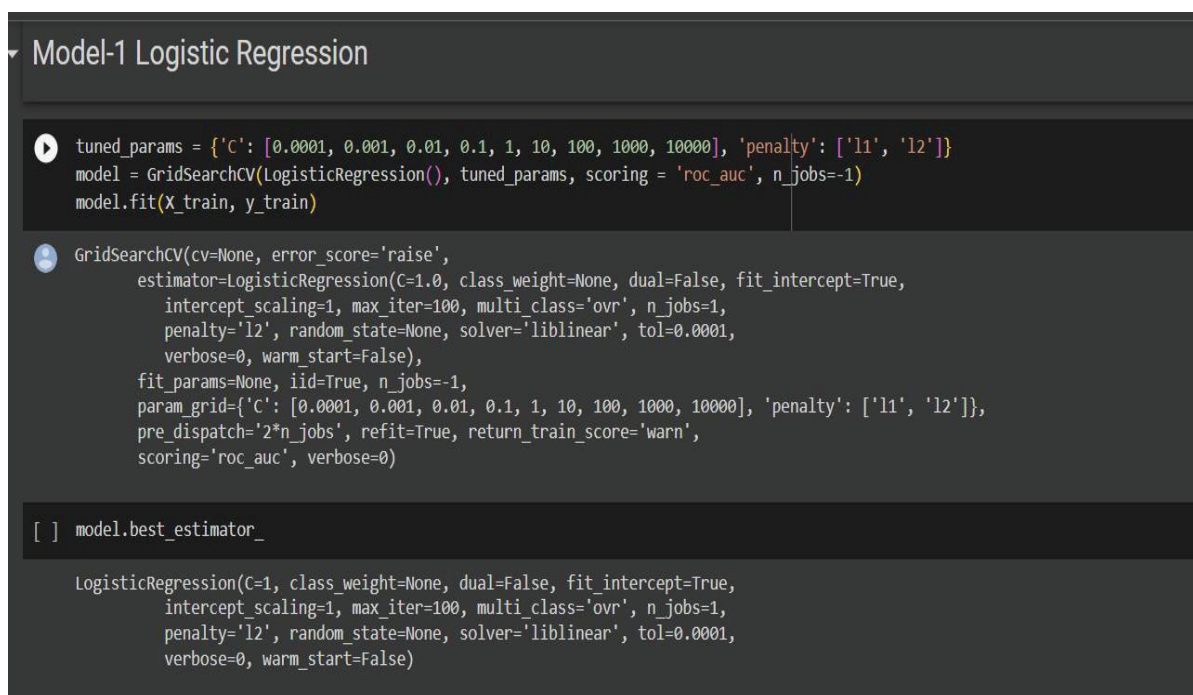
0s ✓  ## Note: We use train_mean and train_std_dev to standardize test data set
X_test = (X_test - train_mean) / train_std

1s ✓  ## Check for mean and std dev. - not exactly 0 and 1
X_test.describe()

```

	Age	Gender	Total_Bilirubin	Direct_Bilirubin	Alkaline_Phosphotase	Alamine_Aminotransferase	Aspartate_Aminotransferase	Total_Protiens	Albumin	Albumin_and_Gl
count	113.000000	113.000000	113.000000	113.000000	113.000000	113.000000	113.000000	113.000000	113.000000	113.000000
mean	0.109816	-0.216940	-0.062052	-0.033465	-0.048222	-0.017203	0.007819	-0.082507	-0.175052	-0.175052
std	0.937258	0.851421	0.835512	1.046110	0.734043	0.661163	0.874507	0.962887	0.922474	0.922474
min	-2.338273	-0.598029	-0.445930	-0.501397	-0.768218	-0.394377	-0.469597	-2.655723	-2.198667	-2.198667
25%	-0.577433	-0.598029	-0.399619	-0.465814	-0.444170	-0.321831	-0.401738	-0.741541	-0.836070	-0.836070
50%	0.090472	-0.598029	-0.384182	-0.430231	-0.342661	-0.261377	-0.316914	-0.012329	-0.216708	-0.216708
75%	0.940533	-0.598029	-0.198937	-0.181150	-0.018612	-0.007466	-0.034168	0.625732	0.526527	0.526527
max	1.851312	1.668451	4.231512	5.974717	3.147694	4.834965	4.263572	2.175308	2.136869	2.136869

Fig. 8. Screenshot showing details of code



The screenshot shows a Jupyter Notebook interface with a code cell titled "Model-1 Logistic Regression". The code defines the tuned parameters for GridSearchCV, fits the model, and prints the best estimator.

```

Model-1 Logistic Regression

 tuned_params = {'C': [0.0001, 0.001, 0.01, 0.1, 1, 10, 100, 1000, 10000], 'penalty': ['l1', 'l2']}
model = GridSearchCV(LogisticRegression(), tuned_params, scoring = 'roc_auc', n_jobs=-1)
model.fit(X_train, y_train)

 GridSearchCV(cv=None, error_score='raise',
  estimator=LogisticRegression(C=1.0, class_weight=None, dual=False, fit_intercept=True,
    intercept_scaling=1, max_iter=100, multi_class='ovr', n_jobs=1,
    penalty='l2', random_state=None, solver='liblinear', tol=0.0001,
    verbose=0, warm_start=False),
  fit_params=None, iid=True, n_jobs=-1,
  param_grid={'C': [0.0001, 0.001, 0.01, 0.1, 1, 10, 100, 1000, 10000], 'penalty': ['l1', 'l2']},
  pre_dispatch='2*n_jobs', refit=True, return_train_score='warn',
  scoring='roc_auc', verbose=0)

[ ] model.best_estimator_

LogisticRegression(C=1, class_weight=None, dual=False, fit_intercept=True,
  intercept_scaling=1, max_iter=100, multi_class='ovr', n_jobs=1,
  penalty='l2', random_state=None, solver='liblinear', tol=0.0001,
  verbose=0, warm_start=False)

```

Fig.9. Model 1: Logistic Regression

```

Model-2 Random Forest

▶ tuned_params = {'n_estimators': [100, 200, 300, 400, 500], 'min_samples_split': [2, 5, 10], 'min_samples_leaf': [1, 2, 4]}
  model = RandomizedSearchCV(RandomForestClassifier(), tuned_params, n_iter=15, scoring = 'roc_auc', n_jobs=-1)
  model.fit(X_train, y_train)

▶ RandomizedSearchCV(cv=None, error_score='raise',
  estimator=RandomForestClassifier(bootstrap=True, class_weight=None, criterion='gini',
    max_depth=None, max_features='auto', max_leaf_nodes=None,
    min_impurity_decrease=0.0, min_impurity_split=None,
    min_samples_leaf=1, min_samples_split=2,
    min_weight_fraction_leaf=0.0, n_estimators=10, n_jobs=1,
    oob_score=False, random_state=None, verbose=0,
    warm_start=False),
  fit_params=None, iid=True, n_iter=15, n_jobs=-1,
  param_distributions={'n_estimators': [100, 200, 300, 400, 500], 'min_samples_split': [2, 5, 10], 'min_samples_leaf': [1, 2, 4]},
  pre_dispatch='2*n_jobs', random_state=None, refit=True,
  return_train_score='warn', scoring='roc_auc', verbose=0)

[ ] model.best_estimator_

RandomForestClassifier(bootstrap=True, class_weight=None, criterion='gini',
  max_depth=None, max_features='auto', max_leaf_nodes=None,
  min_impurity_decrease=0.0, min_impurity_split=None,
  min_samples_leaf=4, min_samples_split=2,
  min_weight_fraction_leaf=0.0, n_estimators=500, n_jobs=1,
  oob_score=False, random_state=None, verbose=0,
  warm_start=False)

[ ] y_train_pred = model.predict(X_train)

[ ] y_pred = model.predict(X_test)

[ ] # Get just the prediction for the positive class (1)
  y_pred_proba = model.predict_proba(X_test)[: ,1]

[ ] # Display first 10 predictions
  y_pred_proba[:10]

array([ 0.61562769,  0.60077365,  0.62661236,  0.58193463,  0.34953728,
        0.99093564,  0.51492834,  0.50708841,  0.95323413,  0.42825976])

[ ] # Calculate ROC curve from y_test and pred
  fpr, tpr, thresholds = roc_curve(y_test, y_pred_proba)

▶ # Plot the ROC curve
  fig = plt.figure(figsize=(8,8))
  plt.title('Receiver Operating Characteristic')

  # Plot ROC curve
  plt.plot(fpr, tpr, label='l1')
  plt.legend(loc='lower right')

  # Diagonal 45 degree line
  plt.plot([0,1],[0,1], 'k--')

  # Axes limits and labels
  plt.xlim([-0.1,1.1])
  plt.ylim([-0.1,1.1])
  plt.ylabel('True Positive Rate')
  plt.xlabel('False Positive Rate')
  plt.show()

```

Fig. 10. Screenshot showing Model 2: Random Forest



Fig. 11. Screenshot showing details of Model 3: KNN

Table 1: Accuracy table comparing different of models

MODEL	ACCURACY
Logistic regression	0.710262345679
Random forest	0.726466049383
KNN	0.628279320988
Decision tree	0.667438271605
Neural network	0.692515432099

5. CONCLUSION

Diseases associated with the liver and heart are on the rise, and continual technological advancements suggest a further increase in their prevalence. Despite growing health awareness, evidenced by increased participation in activities like yoga and dance classes, the persisting sedentary lifestyle and ongoing enhancements in luxuries indicate that these health issues will endure. The proposed application offers the capability to predict the presence of liver disease, which is particularly beneficial in low-income nations where medical infrastructure and specialized experts are scarce. In our study, there are several avenues for future exploration in this field. While we have examined some popular supervised machine learning algorithms, there is potential to select additional algorithms to construct a more precise model for predicting liver disease, thereby progressively enhancing performance. Furthermore, this work can play a significant role in healthcare research and medical facilities, aiding in the proactive prediction of liver diseases. In such a scenario, our paper stands to be exceptionally valuable to society. Based on the dataset used, it is evident that we can accurately predict and identify instances of liver diseases.

REFERENCES

1. Dua, Dheeru and Graff C 2017 {UCI} Machine Learning Repository Univ. California, Irvine, Sch. Inf. Comput. Sci.
2. Abdar, Moloud. "A survey and compare the performance of IBM SPSS modeler and rapid miner software for predicting liver disease by using various data mining algorithms." *Cumhuriyet Üniversitesi Fen Edebiyat Fakültesi Fen Bilimleri Dergisi* ۲۲۴۱-۲۲۳۰ : (۲۰۱۵) ۳۶,۳.
3. Mala, K., V. Sadasivam, and S. Alagappan. "Neural network based texture analysis of CT images for fatty and cirrhosis liver classification." *Applied Soft Computing* 32 (2015): 80-86.
4. Samir, Anthony E., et al. "Shear-wave elastography for the estimation of liver fibrosis in chronic liver disease: determining accuracy and ideal site for measurement." *Radiology* 274.3 (2015): 888-896.
5. Ayeldeen, Heba, et al. "Prediction of liver fibrosis stages by machine learning model: A decision tree approach." 2015 Third World Conference on Complex Systems (WCCS). IEEE, 2015.
6. Rau, Hsiao-Hsien, et al. "Development of a web-based liver cancer prediction model for type II diabetes patients by using an artificial neural network." *Computer methods and programs in biomedicine* 125 (2016): 58-65.
7. Pasha, Maruf, and Meherwar Fatima. "Comparative Analysis of Meta Learning Algorithms for Liver Disease Detection." *J. Softw.* 12.12 (2017): 923-933.
8. Loomba, Rohit, et al. "Gut microbiome-based metagenomic signature for non-invasive detection of advanced fibrosis in human nonalcoholic fatty liver disease." *Cell metabolism* 25.5 (2017): 1054-1062.
9. Chen, Yang, et al. "Machine-learning-based classification of real-time tissue elastography for hepatic fibrosis in patients with chronic hepatitis B." *Computers in biology and medicine* 89 (2017): 18-23.
10. Yip, TC-F., et al. "Laboratory parameter-based machine learning model for excluding non-alcoholic fatty liver disease

(NAFLD) in the general population.” *Alimentary pharmacology & therapeutics* 46.4 (2017): 447-456.

11. Abdar, Moloud, Neil Yuwen Yen, and Jason Chi-Shun Hung. “Improving the diagnosis of liver disease using multilayer perceptron neural network and boosted decision trees.” *Journal of Medical and Biological Engineering* 38 (2018): 953-965.
12. Nahar, Nazmun, and Ferdous Ara. “Liver disease prediction by using different decision tree techniques.” *International Journal of Data Mining & Knowledge Management Process* 8.2 (2018): 01-09.
13. Hashem, Somaya, et al. “Comparison of machine learning approaches for prediction of advanced liver fibrosis in chronic hepatitis C patients.” *IEEE/ACM transactions on computational biology and bioinformatics* 15.3 (2017): 861-868.
14. Książek, Wojciech, et al. “A novel machine learning approach for early detection of hepatocellular carcinoma patients.” *Cognitive Systems Research* 54 (2019): 116-127.
15. Singh, Amardip Kumar. “A comparative study on disease classification using machine learning algorithms.” *Proceedings of 2nd International Conference on Advanced Computing and Software Engineering (ICACSE)*. 2019.
16. Wang, Kun, et al. “Deep learning Radiomics of shear wave elastography significantly improved diagnostic performance for assessing liver fibrosis in chronic hepatitis B: a prospective multicentre study.” *Gut* 68.4 (2019): 729-741.
17. Samir, Anthony E., et al. “Shear-wave elastography for the estimation of liver fibrosis in chronic liver disease: determining accuracy and ideal site for measurement.” *Radiology* 274.3 (2015): 888-896.
18. Ayeldeen, Heba, et al. “Prediction of liver fibrosis stages by machine learning model: A decision tree approach.” *2015 Third World Conference on Complex Systems (WCCS)*. IEEE, 2015.
19. Rau, Hsiao-Hsien, et al. “Development of a web-based liver cancer prediction model for type II diabetes patients by using an artificial neural network.” *Computer methods and programs in biomedicine* 125 (2016): 58-65.
20. Pasha, Maruf, and Meherwar Fatima. “Comparative Analysis of Meta Learning Algorithms for Liver Disease Detection.” *J. Softw.* 12.12 (2017): 923-933.