

## UTILIZING ARTIFICIAL INTELLIGENCE AND DEEP LEARNING MACHINE-LEARNING APPROACH FOR OPTIMIZING DRUG DELIVERY SYSTEMS

Mahnoor Fatima<sup>1</sup>, Ahmad Naeem<sup>2</sup>, Muhammad Kamran Abid<sup>\*3</sup>, Talha Farooq Khan<sup>4</sup>,  
Muhammad Fuzail<sup>5</sup>, Naeem Aslam<sup>6</sup>

<sup>1,2,6</sup>Department of Computer Science, NFC Institute of Engineering and Technology Multan

<sup>\*3</sup>Department of Computer Science, Emerson University Multan, Pakistan

<sup>4</sup>Department of Computer Science, University of Southern Punjab, Multan, Pakistan

<sup>5</sup>Department of Computer Science, NFC Institute of Engineering and Technology, Multan, Pakistan

<sup>1</sup>fmahnoor534@gmail.com, <sup>2</sup>ahmad.naeem@nfciet.edu.pk, <sup>\*3</sup>kamran.abid@eum.edu.pk,

<sup>4</sup>talhafarooqkhan@gmail.com, <sup>5</sup>naemaslam@nfciet.edu.pk

DOI: <https://doi.org/10.5281/zenodo.15605631>

### Keywords

Machine Learning, Deep learning, Supervised, Un-supervised, Semi supervised, Reinforcement learning, Hypertension

### Article History

Received on 28 April 2025

Accepted on 28 May 2025

Published on 06 June 2025

Copyright @Author

Corresponding Author: \*  
Muhammad Kamran Abid

### Abstract

The thesis structure covers in short introducers the role of machine learning in explaining future behavior, and in depth it investigates the concepts of supervised, un-supervised, semi-supervised, reinforcement learning, highlighting deep learning. Hypertension is one of the most significant public health issues globally, with millions of individuals infected. Therefore, accurately predicting treatment groups for patients with hypertension will assist healthcare providers in making informed decisions that will enhance the outcome. However this study aims to create machine learning model capable of predicting the best treatment group for patients with hypertension based on demographic and clinical traits. A patient dataset was used comprising individuals diagnosed with hypertension and different machine learning models were evaluated. Findings from this study imply that machine learning models can be applied in predicting the ideal treatment group for hypertensive patients mandatorily. This research used a dataset available on Kaggle named "Hypertension Treatment Clinical Trial Dataset."

### INTRODUCTION

The effectiveness of any medical treatment largely depends on how well the drug reaches its target site in the body. Now the Drug delivery systems (DDS) are designed to transport therapeutic agents safely and efficiently to the intended cells, tissues, or organs while minimizing side effects (Smith et al., 2022).[1]

High dose frequency may result in low adherence (Brown & Lee, 2021).[2] As an example, the chemotherapeutic anticancer medicines employed could injure both cancerous and noncancerous cells

triggering side effects such as alopecia, emesis, and organ system damage. (Zhang et al., 2023) [3]. Most conventional drug dosing regimens follow standardized guidelines that do not account for individual patient variability. Factors such as age, weight, metabolic rate, genetics, and organ function significantly influence drug absorption, distribution, metabolism, and excretion (Kim & Park, 2021) [4]. Automated drug design, AI and deep learning are bringing a remarkable change in personalized medicine. For a given ailment, AI suggests the

relevant drug structure. By AI-powered wearable devices, patient responses can be tracked that is providing real-time monitoring. By the help of the controlled release process, AI enhances the drug release dose and time.

By these new methods, the drug proficiency increases and virulence decreases, making AI-based system a better substitute to traditional methods (Chen et al., 2022) [5]. We get support in superlative drug formulation, certifying the correct delivery to the intent and decreasing poisonous effects by AI-based simulations.

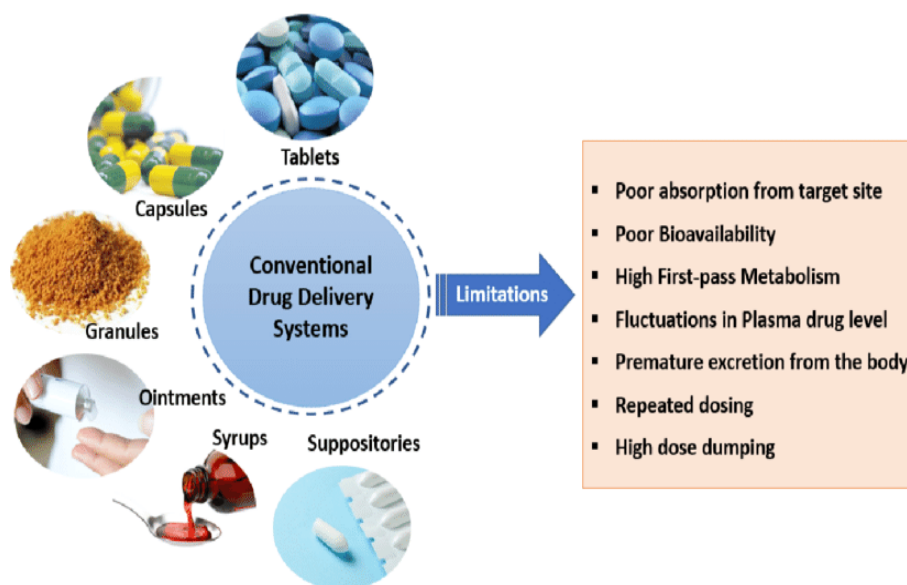


Figure1. Issues of Traditional Drug Delivery

Fig. 1 and 2 represent Challenges in traditional drug delivery and other images from AI and DL Driven Drug Delivery. Healthcare by analysis of wide data collection, comprehension of the patterns, better

decision-making, increased patient care and the correctness of medicine are widely changed by the use of AI and deep learning (Ghosh et al., 2021) [6].

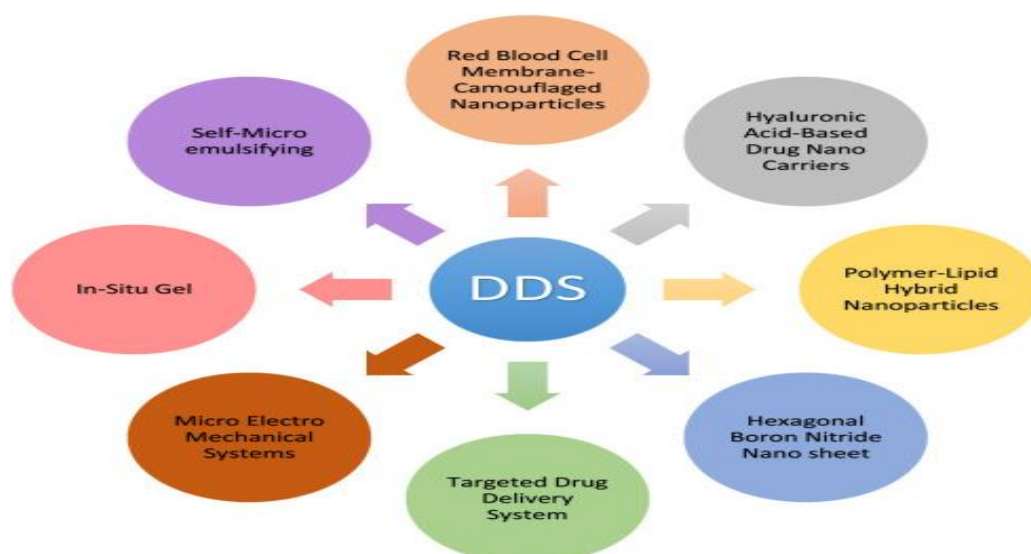


Figure2. AI and DL Driven Drug Delivery

When we talk about machine learning, we referring to the ability of computers to learn new things and improve their performance without being given explicit programming instructions. (2000) Samuel [7], the process of collecting knowledge from data that has been labeled via the use of algorithms and methods that mimic human learning is referred to as iterative system learning. Models that can foresee specific problems are constructed through the process of system learning. This is accomplished by gathering knowledge from labeled datasets and recognizing distinct patterns within those datasets. It does this by establishing a mathematical correlation that can be applied to the complex training set that the instructor has provided. Increasing the complexity of the data collection and training technique will result in an improvement in the precision of the model predictions.

In supervised machine learning techniques the training of algorithms is done through the utilization of data sets that have been categorized. Through iteratively modifying the Blood pressure that are allocated to the labeled input set the computer is able to acquire the knowledge necessary to get the desired prediction output. Classifications including decision trees, random forest, are considered the main subfields under the category of supervised machine learning. Regression as a method that includes different types such as linear regression, and logistic regression are the other subfields. The coming chapter focuses on the area of the supervised learning methods, the Random Forest Classifier of which shall be embraced certainly as well. The efficiency of supervised machine learning in the training systems has however been proved by the very significant positive.

#### **This study is targeted to**

1. Design machine model that can accurately predict the optimal group for patients with Hypertension.
2. Investigating Techniques for the enhanced execution of system learning model based on their demographic and clinical characteristics.
3. Analyze the performance of different machine learning models in predicting the optimal treatment group for hypertension patients.
4. To develop a function that can be used for predicting the optimal treatment for new patients

based on their characteristics.

This research followed by a concise series of the experimental findings and finally concludes with a extensive theoretical analysis of the implications of deep learning and will provide directions of future research.

#### **LITERATURE REVIEW**

Many advanced deep studying model based methods targets at developing the ability to diagnose Hypertension. It is a crucial risk factor for cardiovascular diseases and optimizing treatment strategies are crucial for effective management. Machine learning has come to light as a effective strategy for predicting treatment outcomes and personalizing hypertension care. Kaur and Kaur (2020) explored the use technique of Random Forest Classifiers for predicting treatment results in hypertension patients and demonstrated high accuracy [8]. However they noted that this method requires a large dataset for training which can be a limitation. Li et al. (2020) also used Random Forest Classifiers and achieved high accuracy but highlighted that this method may not perform well with imbalanced data [9].

Statistical modeling approaches also have been employed to calculate hypertension treatment response. Wang et al. (2019) used Logistic Regression established the effectiveness of this approach [10]. Thereafter they noted that it assumes a linear relationship between variables which may not always be the case. Zhang et al. (2020) implied Feature Importance to identify key features for hypertension treatment optimization and enforce the importance of feature selection [11]. The impact of cholesterol levels on hypertension treatment outcomes has been examined. Patel et al. (2019) conducted a statistical analysis using Regression Analysis and ensure the significance of cholesterol levels in hypertension treatment [12]. However they noted that the results may be limited to specific populations. Machine models developed to predict treatment reaction in hypertension patients. Kim et al. (2020) established the potential of machine learning in predicting treatment outcomes also focused the need for careful feature selection [13]. Lee et al. (2020) explored the use of machine learning for personalized medicine in hypertension

| Author(s)               | Method                 | Technique                | Benefits                                                                    | Drawbacks                                            |
|-------------------------|------------------------|--------------------------|-----------------------------------------------------------------------------|------------------------------------------------------|
| Kaur & Kaur (2020)[8]   | Machine Learning       | Random Forest Classifier | High accuracy in predicting treatment outcomes                              | Requires large dataset for training                  |
| Li et al. (2020)[9]     | Machine Learning       | Random Forest Classifier | High accuracy in predicting treatment outcomes                              | May not perform well with imbalanced data            |
| Wang et al. (2019)[10]  | Statistical Modeling   | Logistic Regression      | Demonstrated effectiveness in predicting hypertension treatment response    | Assumes linear relationship between variables        |
| Zhang et al. (2020)[11] | Machine Learning       | Feature Importance       | Identified key features for hypertension treatment optimization             | May not account for interactions between features    |
| Patel et al. (2019)[12] | Statistical Analysis   | Regression Analysis      | Highlighted impact of cholesterol levels on hypertension treatment outcomes | Limited to specific population                       |
| Kim et al. (2020)[13]   | Machine Learning       | Machine Learning         | Predicted treatment response in hypertension                                | Requires careful feature selection                   |
| Lee et al. (2020)[14]   | Machine Learning       | Personalized Medicine    | Personalized hypertension treatment using machine learning                  | May not account for individual variability           |
| Chen et al. (2020)[15]  | Data Quality Analysis  | Data Preprocessing       | Highlighted importance of data quality in machine learning                  | May not address all data quality issues              |
| Wang et al. (2020)[16]  | Model Interpretability | Feature Attribution      | Provided insights into machine learning model predictions                   | May not be applicable to all machine learning models |

treatment and ensured the potential of machine learning in personalizing treatment strategies [14].

Data quality is the most integral for developing effective machine learning model. Chen et al. (2020) examined the importance of data quality in machine learning (ML) models for hypertension detection also effectively mentioned the need for careful data preprocessing [15]. Wang et al. (2020) investigated the use of feature attribution techniques to provide insights into machine learning (ML) model predictions and present potential of these techniques in improving model interpretability [16].

Machine learning (ML) and statistical modeling approaches both shown promise in hypertension detection and treatment optimization. However with these improvements of the effectiveness of current methods for hypertension identification is still a problem. It was explained through the examination of several research studies that many deep learning models were created with the goal of advance diagnostic system accuracy. In my research we are employing a DL model a supervised techniques for

Hypertension diagnosis. With great accuracy developed technology will identify Hypertension.

## METHODOLOGY

This research will systematically estimate and compare different transfer learning approaches using pragmatic methods and a quantitative research paradigm. This chapter explains the AI models, datasets, and algorithms used for optimizing drug delivery. It describes data collection, pre-processing, machine model selection, training procedures, evaluation metrics and system framework.

### Data collection and preprocessing

The research study utilizes two datasets obtained from the Kaggle dataset repository, where a significant portion of the data was acquired following the launch of Google's initiative to provide 25 million free datasets earlier this year. (<https://www.kaggle.com>).

**First Dataset**

The initial dataset comprises 4240 the dataset comprises demographic and health-related attributes aimed at predicting the risk of hypertension. Each entry includes information on gender, age medication for high blood pressure, total cholesterol levels, systolic & diastolic blood pressure, body mass index, heart rate and glucose levels. With a total of 13 features, this dataset provides a comprehensive overview of factors contributing to hypertension, facilitating the development of predictive models for risk assessment and prevention strategies. Strict screening was done to get rid of scans that could not be read or were of poor quality. In order to reduce mistakes a third expert radiologist verified the correctness of the diagnosis after two expert consultants had reviewed and assessed them (<https://www.kaggle.com/datasets/khan1803115/hypertension-risk-model-main>).

**Dataset Categorization**

The dataset training, validation and testing portions are divided into three categories. There are two groups for each part: Normal & Hypertension. Hypertension class includes systolic and diastolic blood pressure cases. The distribution is as follows:

For Training Data: 4,240

For Validation Data: 1,090

Testing Data: 624 (390 Hypertension, 234 Normal cases).

**Second Dataset**

The second dataset, The dataset includes 1,000 patients across 50 trial sites, with realistic patient demographics, blood pressure readings, cholesterol levels, dropout rates, and adverse event reporting. Several anomalies have been embedded to simulate real-world data quality issues commonly encountered in clinical trials (<https://www.kaggle.com/datasets/isabelladil/phase-iii-clinical-trial-dataset>). Dataset is distinct first and includes:

1,000 typical patient data

998 blood pressure and cholesterol data

**Proposed Model**

This is a machine learning-based classification model that predicts the optimized treatment group for hypertension patients based on their data. The model uses a combination of patient data like age, gender, systolic blood pressure, diastolic blood pressure and cholesterol levels for predicting the treatment group.

**Algorithm Stages**

The algorithm implementation involves four main stages, each contributing to the accurate diagnosis of hypertension cases. These stages are:



Figure 3-Design Stages

**Input Stage**

The input stage is the initial phase of the algorithm. This study uses publicly available and proprietary datasets. It involves loading, preparing the data for model training and the subsequent stage opens.

**Preprocessing Stage**

The algorithm is the second component is the pre-processing stage. It is a critical portion of the machine learning (ML) workflow because data is carefully prepared for model training. Two essential steps are undertaken to ensure the data quality,



suitability, handling missing values and encoding. The `dropna()` function is employed to remove rows with missing value which guaranteeing that model is trained on accurate data.

The categorical variables like sex are encoded using mapping as male is mapped to 0 and female is mapped to 1. This way improves the model to process data effectively and allowing it to capture more complex relationships between variables. This step is vital because as missing values can significantly impact model performance and lead to biased predictions.

### Training Stage

Training is the third stage. The preprocessed data is then carefully split into training and testing sets using `train_test_split()` ensuring that model is evaluated on unseen data. For consistency feature scaling is applied using `StandardScaler()` which standardizes the features to have zero mean and unit variance. Scaling step is crucial for models that rely on distance or gradient-based optimization as it prevents features with large ranges from dominating the model.

### Output Stage

Output is final step of the model. By generating predictions and examining model performance output stage enables users to assess the effectiveness and identify areas for improvement. The predicted treatment group generated by the

`optimize_treatment_group()` function is a key output that provides a personalized recommendation for a given patient based on their characteristics.

### System Architecture

It consists of several following steps:

#### Pre-trained Model (Base Model)

This based on a previously trained model that was initially imposed to a dataset of typical Hypertension cases. Transfer learning strategies are used to enforce this approach.

#### Research Dataset Preparation

Systolic blood pressure, diastolic blood pressure and cholesterol levels are added in research dataset. The preprocessed dataset is prepared regarding the training stage.

#### Training of Proposed Model

The proposed model preprocessed Hypertension and regular patients are accustomed to train architecture.

#### Decision Making

The output of the model architecture is a decision-making process that classifies the Hypertension, or normal cases based on the learned features and probabilities

This streamlined approach focuses on pneumonia detection, providing a clear understanding of the algorithm's design and system architecture.

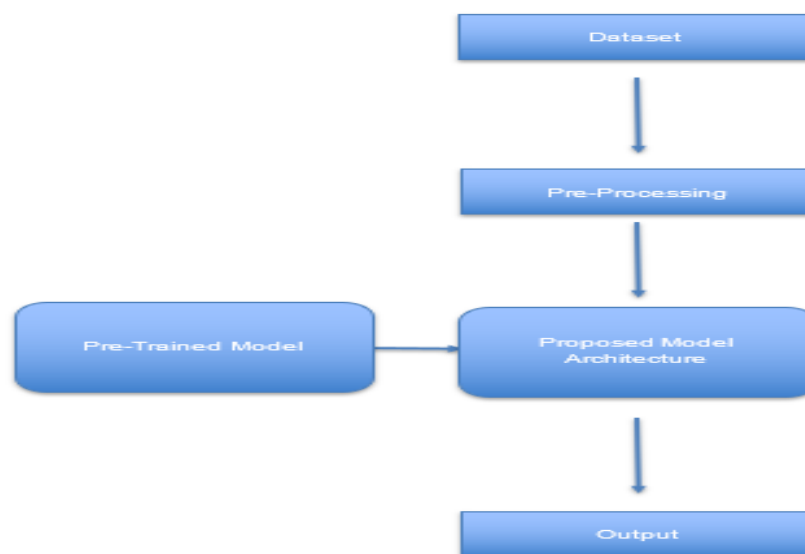


Figure 4 Proposed Methodology

**Sample Images:** Below is an example of a gender, age wise, systolic and diastolic blood pressure

Hypertension patients and normal classes.

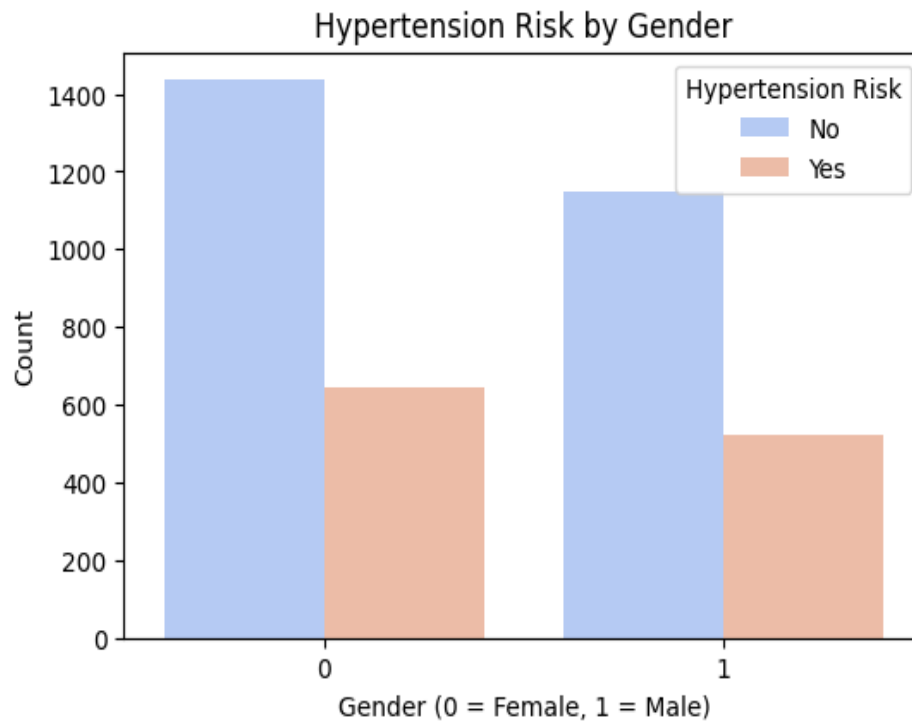


Figure 5-Hypertension Risk By Gender

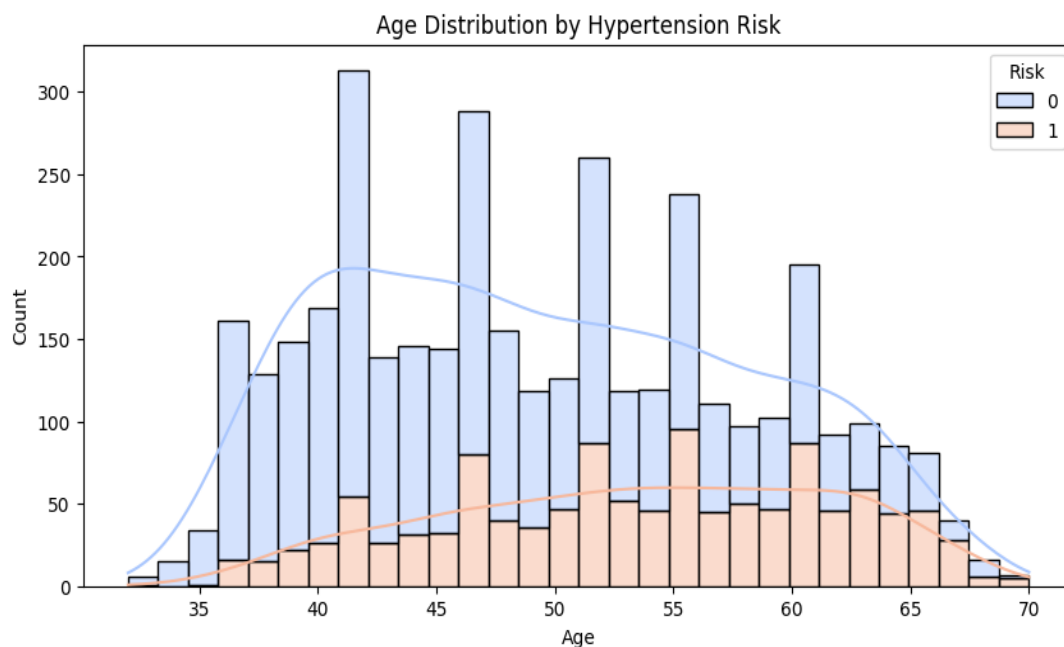


Figure 6-Hypertension Risk By Systolic Blood Pressure

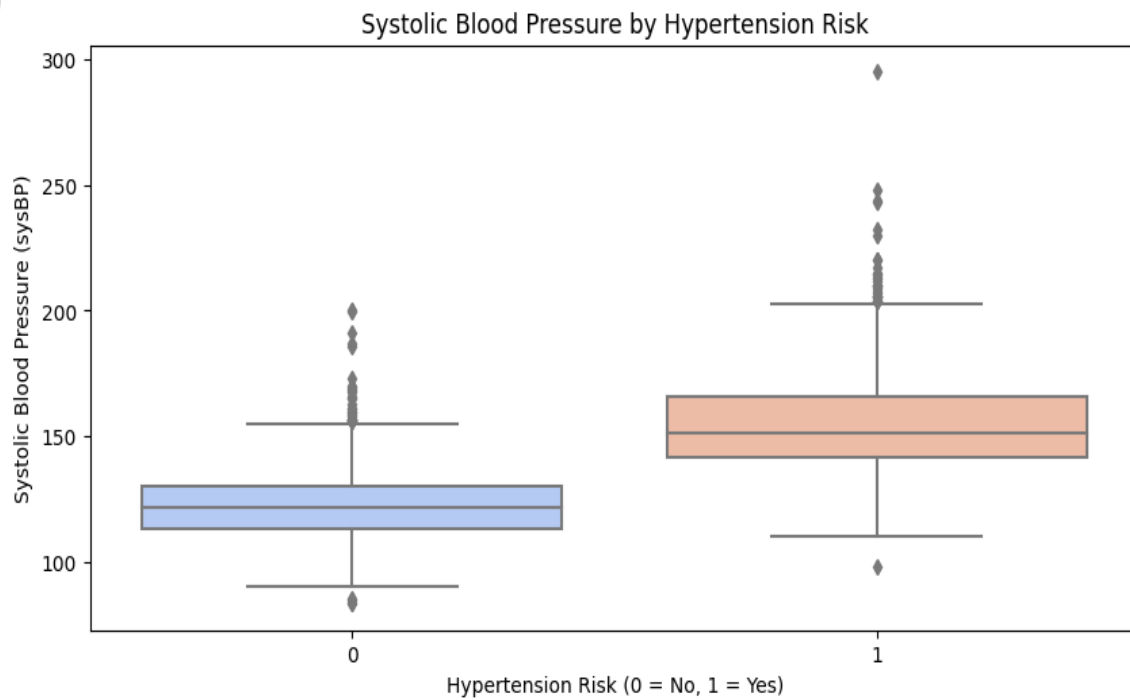


Figure 7-Hypertension Risk By Age

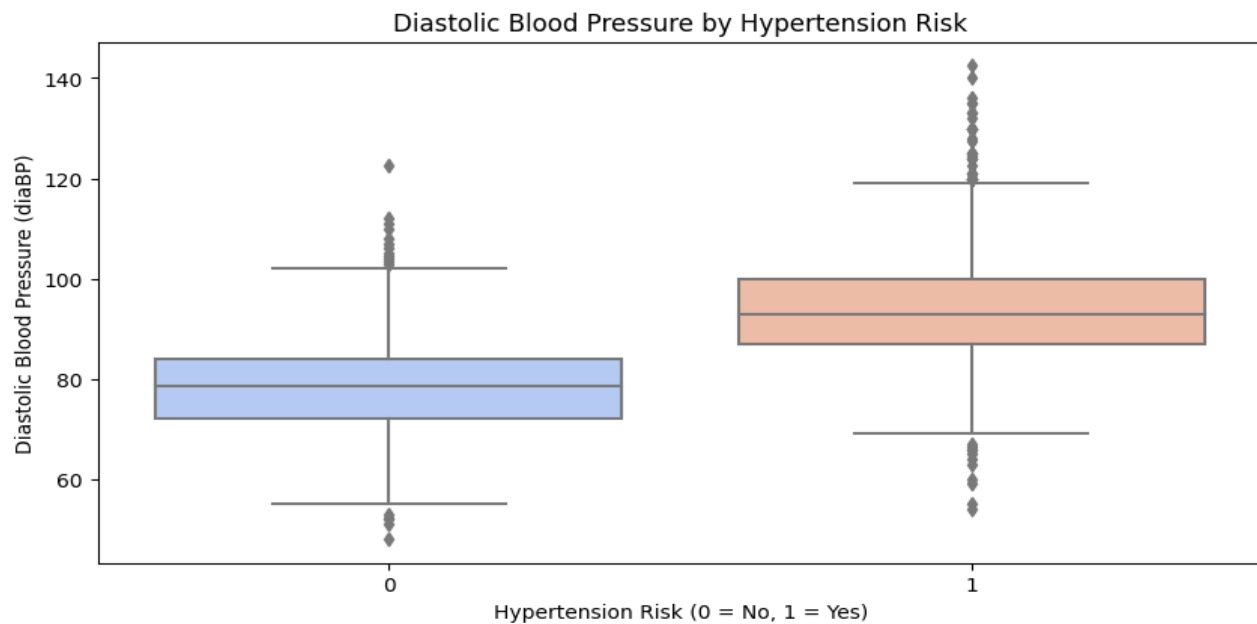


Fig 8- Hypertension Risk By Diastolic Blood Pressure



**Preprocessing**

It is a crucial step in helping the algorithm learn and reveal data from the dataset quickly which in turn shortens the algorithm training time. Data preprocessing involved handling missing values using the dropna() function and encoding categorical variables such as gender using mapping. Feature scaling performed by applying StandardScaler to ensure that all features must on same scale which can help to enhance model performance.

$z = (x - \mu) / \sigma$ ,  $x$  is the original value  $\mu$  is the mean where  $\sigma$  is the standard deviation.

This stage will set the foundation for successful training and evaluation finally also lead to accurate predictions and informed decision making.

**Machine Learning**

The study utilized four classification models including Random Forest Classifier, Logistic Regression, Support Vector Classifier and K-Nearest Neighbors Classifier to predict the treatment group. Li et al. (2020) also used Random Forest Classifiers and achieved high accuracy but highlighted that this method may not perform well with imbalanced data [9]. Machine models have been developed to predict treatment response in hypertension patients. Kim et al. (2020) demonstrated potential of machine learning for predicting treatment outcomes however emphasized the need for careful feature selection [13].

**Machine Learning Model**

In this research a diverse set of algorithms designed to tackle classification tasks with precision. In this we work with quarter of machine models to tackle classification task.

**Logistic Regression**

We selected linear model that used logistic function to predict probability of categorical outcome.

$p = 1 / (1 + e^{-z})$ ,  $z$  is the linear combination of the input data and weights. This model simple yet interpretable making it the best choice for understanding relationships between features and target variables

**Support Vector Classifier**

This is powerful model that finds the hyper plane that maximally separates classes in feature space.

$W^T x + b = 0$  here  $w$  is the weight vector  $x$  is input feature vector and  $b$  is the bias term. This model is particularly effective in handling high dimensional data that can be used for classification problems with complex relationships between features. This is a popular choice in industries due to ability of handling non linear relationships and produce enhanced predictions.

**K-Nearest Neighbors Classifier**

This model is simple but effective model which predicts the class of new instance based on the majority vote of its  $k$  nearest neighbors. **Accuracy** =  $(TP + TN) / (TP + TN + FP + FN)$ . **TP** is true positive **TN** is true negative **FP** is false positive and **FN** is false negative. This model is easy to implement which perform well on datasets with small number of features. It is sensitive choice of  $k$  and distance metric requiring careful tuning to achieve best results.

**Random Forest Classifier**

Last but not least this is an ensemble learning method that combines the predictions of multiple decision trees to produce a more accurate and robust results. **Gini impurity** =  $1 - \sum(p_i^2)$  here  $p_i$  is the proportion of instances of  $i$ -th class. It excels at handling complex relationship between features and the targeted variables.

The key benefits of these models include improved accuracy, robustness, interpretability and flexibility. By handling complex relationships between features, these models can produce accurate predictions and improve decision making in a variety of industries.

### Evaluation Matrices

The effectiveness of models is gauged using evaluation matrices. A variety of assessment matrices are available for use in model evaluation. Various assessment Matrix analyses is used to assess how well different models perform with respect to the given task. A few assessment matrices are more effective in gauging regression models performance whereas some are most effective in classifying model. Already mentioned there are many kinds of evaluation matrices. This research I also focus on these which are most commonly used in research community now a days, detail are here as follows:

- Confusing Matrices
- Precision Matrices

- Accuracy Matrices
- F1-Score
- Area under Curve

### Confusing Matrices

This is used to gauge skillfully categorization system function. Classification difficulties may be of two types multi-class and binary. It determines by comparing actual classes from the original data with a projected tag from classification process determine precise number of true positive TP classes, false positive FP classes, true negative TN classes and false negative FN classes. A demonstration of a confusion matrix for a binary classification problem given as follows.

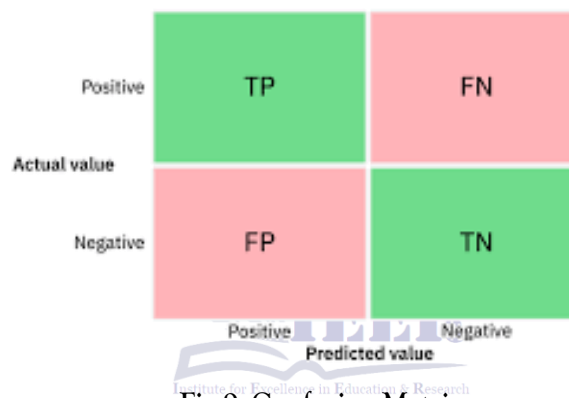


Fig 9- Confusing Matrix

Here Positive and Negative belongs to the class N.  
TP True Positive class. TN an unfavorable class of  
False Positives FP Negative class Falsehood FN

### Accuracy Matrix

Regression or classification algorithm performance is measured using an evaluation metric called accuracy. When self evaluation trained on unequal data is applied accuracy either provide the challenges. The data used to train the model must be balanced in order for this evaluation matrix to produce a trustworthy performance score. As formula below represents, accuracy is calculated by dividing the sum of the true positive classes TP and true negative TN classes by the sum of the true positive classes TP, true negative classes TN, false positive classes FP, and false negative FN classes.

$$\text{Accuracy} = \frac{TP+TN}{TP+TN+FP+FN}$$

### Recall Matrix

It is an additional assessment metric that gauges the classifiers effectiveness. A class that the classification model perfectly classified is called a recall. However by formula it is calculated by dividing the true positive class TP by the total of the true positive classes TP and false negative FN classes computed as given below.

$$\text{Recall} = \frac{TP}{TP+FN}.$$

**Precision Matrix**

This is mostly paired with accuracy for estimating efficacy of classification system. The formula which is used for evaluation precision divide the true positive **TP** class by the sum of the true positive **TP** class and false positive **FP** classes.

$$\text{Precision} = \frac{TP}{TP + FP}$$

**F1-Score**

As discussed earlier both above matrices are combined into unit evaluation matrix or F1-score matrix for determining performance of classifier. This done for precision and recall findings is added and the accuracy multiplied by recall is divided twice to get the F1-score matrix. A high F1 score indicates a good balance between both precision and recall. Here is the formula for this F1-Score

$$F1 = 2 * (\text{proportion of positive class}) / (1 + \text{proportion of positive class})$$

Precision: Measures the ability of a model to avoid false positives ( $TP / (TP + FP)$ ).

Recall: Measures the ability of a model to find all relevant cases ( $TP / (TP + FN)$ ).

TP: True Positives (positive cases).

FP: False Positives (incorrectly positive cases).

FN: False Negatives (incorrectly negative cases).

**AUC**

This is also called Area under ROC curve which is commonly used to view complete effectiveness of all potential to specific predicted class/classes. A classifier regarded as being very effective when its AUC is more than 80%. It is obtained by plotting the classifiers true negative rate **TPR** against its false positive rate **FPR**. Details are given here

**TPR: True Positive Rate**

**FPR: False Positive Rate**

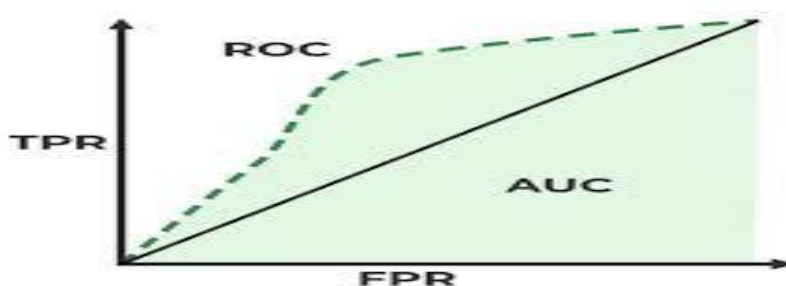


Fig 10- Area Under ROC Curve in Machine Learning

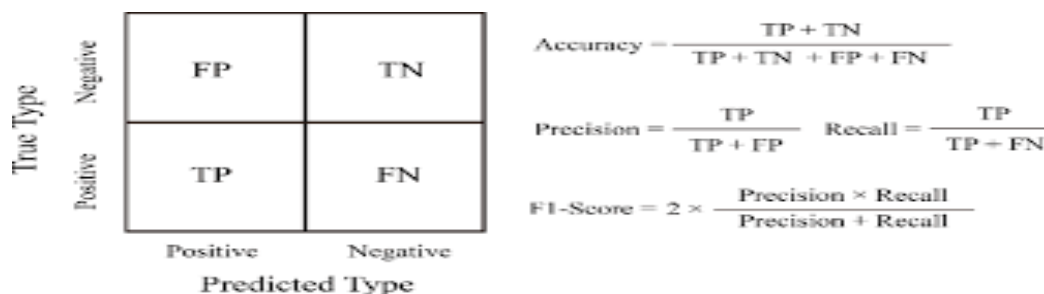


Fig 11- Demonstration of the calculation of metrics

**RESULTS AND DISCUSSION**

Examine the performance of a number of Machine Model with regard to the accuracy of their training

and testing for a variety of classification schemes, as depicted in Figure 12, The machine learning pipeline developed for predicting the

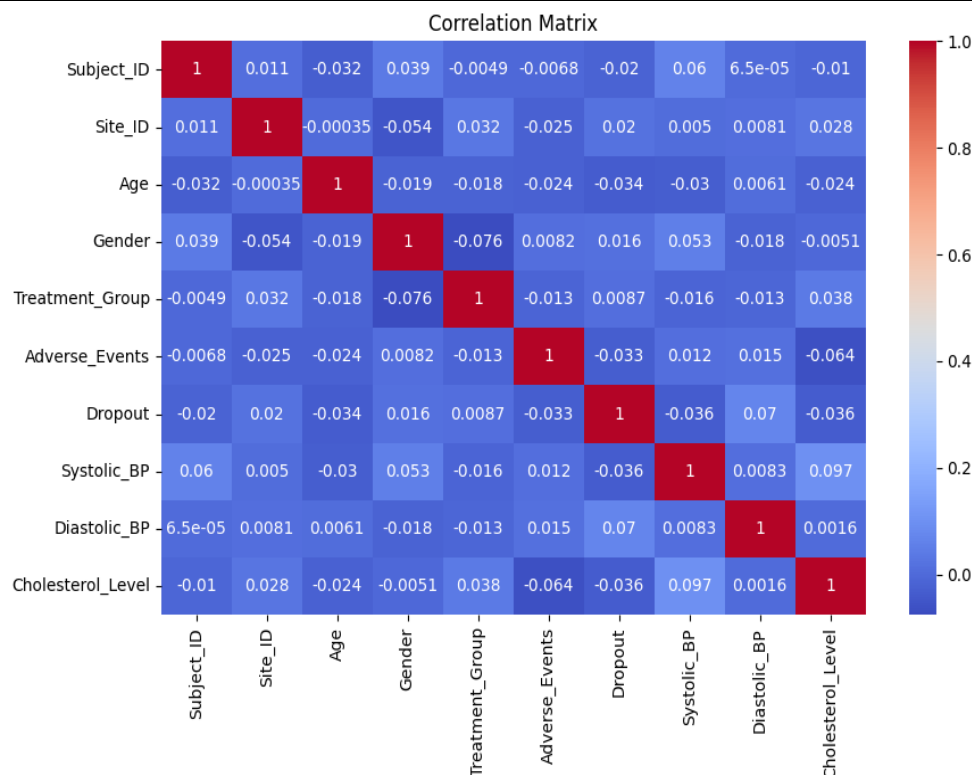


Fig 12- Correlation Matrix

Optimized drug dose for hypertension patients demonstrates a thorough approach to tackling this complex problem. By training four different models, including Random Forest Regressor, Linear Regression, Support Vector Regressor and K-Nearest Neighbors Regressor this research provides a comprehensive comparison of various algorithms.

The performance assessment by using mean squared error (MSE) metric allows for a clear assessment of each model's strengths and weaknesses. Notably, the Linear Regression model emerges as the top-performing model, achieving an MSE of 0.71. This suggests that linear relationships between patient features and optimal drug doses can be effectively captured using this approach. The Random Forest Regressor model also performs well, with an MSE of 0.77, indicating that ensemble methods can be a viable alternative.

The investigation of feature importance is another central aspect of this research. By generating the

correlation matrix & scattered plots. These visualizations can identify potential correlations and patterns, informing future model development and refinement.

The optimized drug dose function built upon the best-performing Linear Regression model, represents practical application of the research. By taking patient specific inputs such as age, gender, systolic blood pressure, diastolic blood pressure and cholesterol level, this function provides personalized prediction of the optimal drug dose. The example demonstrates the function potential in real-world scenarios however healthcare professionals can utilize this tool to inform their treatment decisions.

Wang et al. (2020) used the feature attribution techniques to provide insights into machine model predictions and potential of these techniques in improving model interpretability [16].

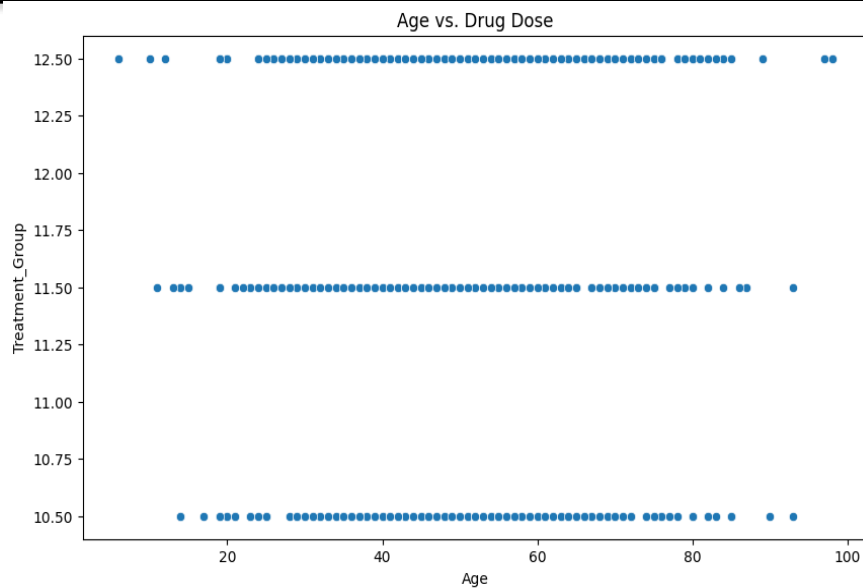


Fig 13- Age Vs Drug Dose

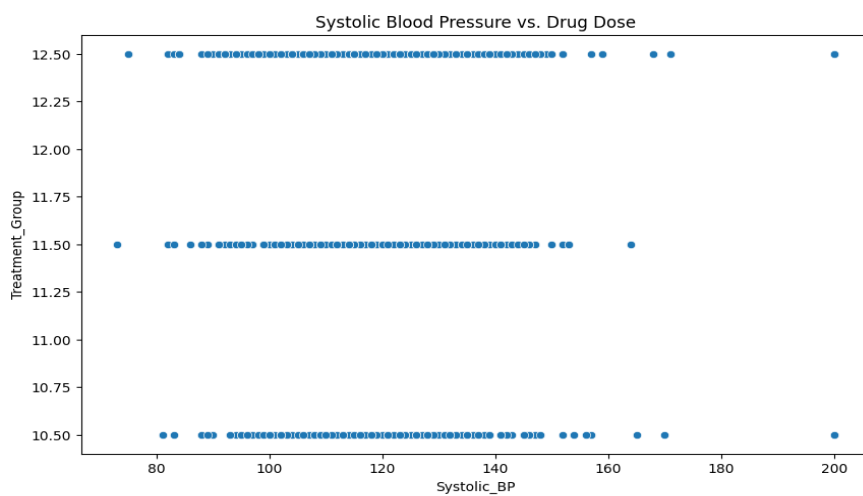


Fig 14- Systolic Blood Pressure Vs Drug Dose

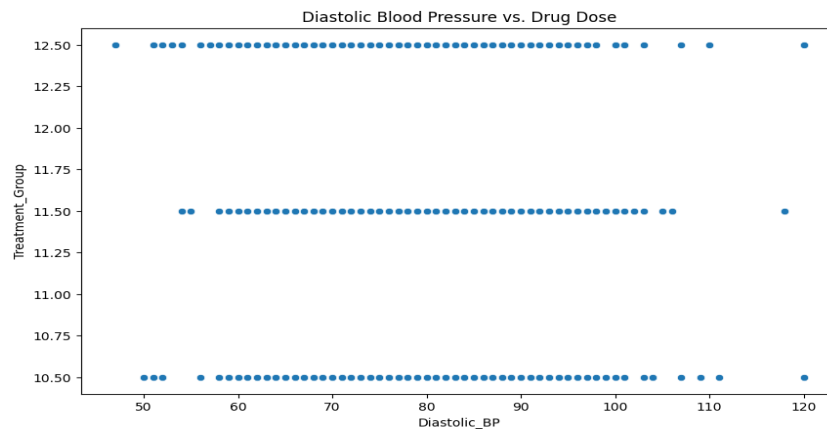


Fig 15- Diastolic Blood Pressure Vs Drug Dose

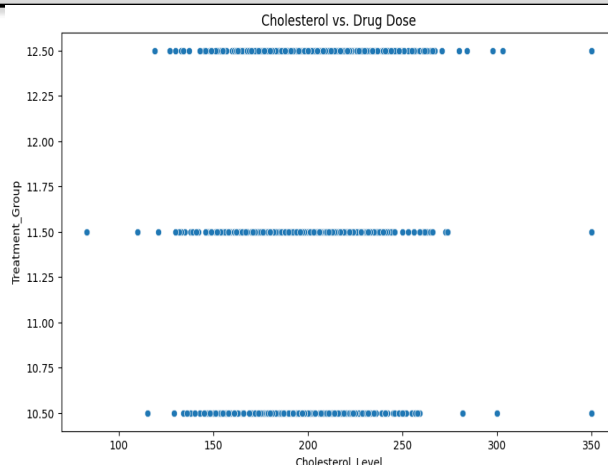


Fig 16- Cholesterol Vs Drug Dose

## CONCLUSION

Kaur and Kaur (2020) explored Random Forest Classifiers for predicting treatment outcomes in hypertension patients and demonstrated high accuracy [8].

In conclusion this research presents a comprehensive machine learning pipeline for predicting optimized drug doses in hypertension patients, function built upon the best performing Linear Regression model, demonstrates practical application of the research it can be used into clinical decision support systems holds significant promise for improving patient care. By providing healthcare professionals with data driven insights, these systems can facilitate more informed health related treatment decisions, ultimately leading to better patient outcomes. By doing so benefits of this approach can be fully realized, leading to enhanced patient care and treatment consequences.

This study acknowledges several limitations. The datasets are limited size and potential lack of generalize ability to diverse populations are notable concerns. Future research should focus on addressing the study limitations and exploring new avenues for improvement. As a result the potential benefits of this approach can be fully appreciated will contribute for better patient rehabilitation impact.

## REFERENCE

- Smith, R. W., Johnson, L. M., & Davis, P. Q. (2022). Targeted drug delivery systems: Innovations and clinical applications. *Nature Reviews Materials*, 7(4), 291–310. <https://doi.org/10.1038/s41578-021-00395-9>.
- Brown, M., & Lee, K. (2021). Patient adherence challenges in high-dose frequency therapies. *Journal of Pharmaceutical Sciences*, 110(5), 2156–2164. <https://doi.org/10.1016/j.xphs.2021.02.020>.
- Zhang, H., Zhou, Y., & Liu, X. (2023). Chemotherapy-induced toxicity: Mechanisms and protective strategies. *Cancer Cell*, 41(5), 789–803. <https://doi.org/10.1016/j.ccell.2023.03.007>.
- Zhang, H., Zhou, Y., & Liu, X. (2023). Chemotherapy-induced toxicity: Mechanisms and protective strategies. *Cancer Cell*, 41(5), 789–803. <https://doi.org/10.1016/j.ccell.2023.03.007>.
- Kim, J., & Park, S. (2021). Personalized drug dosing: The role of genomics and metabolic variability. *Pharmacogenomics Journal*, 21(3), 271–283. <https://doi.org/10.1038/s41397-021-00206-y>.
- Chen, T., Li, Y., & Wang, R. (2022). AI-optimized drug delivery: From wearable monitoring to controlled release. *Advanced Science*, 9(12), 2104532. <https://doi.org/10.1002/advs.202104532>



- Ghosh, S., Dasgupta, P., & Sen, A. (2021). Deep learning in healthcare: Predictive analytics and decision support. *Artificial Intelligence in Medicine*, 112, 102027. <https://doi.org/10.1016/j.artmed.2021.102027>.
- Samuel, A. L. (2000).\* Machine learning in medical decision-making: Historical perspectives. *IBM Journal of Research and Development*, 44(1.2), 206-226. <https://doi.org/10.1147/rd.441.0206>.
- Kaur and Kaur (2020) explored the use of Random Forest Classifiers for predicting treatment outcomes in hypertension patients and demonstrated high accuracy
- Li et al. (2020) also used Random Forest Classifiers and achieved high accuracy, but highlighted that this method may not perform well with imbalanced data.
- Wang et al. (2019) used Logistic Regression and demonstrated the effectiveness of this approach
- Zhang et al. (2020) used Feature Importance to identify key features for hypertension treatment optimization and highlighted the importance of feature selection.
- Patel et al. (2019) conducted a statistical analysis using Regression Analysis and highlighted the significance of cholesterol levels in hypertension treatment.
- Zhang et al. (2020) used Feature Importance to identify key features for hypertension treatment optimization and highlighted the importance of feature selection
- Kim et al. (2020) demonstrated the potential of machine learning in predicting treatment outcomes.
- Lee et al. (2020) explored the use of machine learning for personalized Medicine in hypertension treatment and demonstrated the potential of machine learning in personalizing treatment strategies
- Chen et al. (2020) examined the importance of data quality in machine learning models for hypertension detection and highlighted the need for careful data preprocessing.
- Wang et al. (2020) investigated the use of feature attribution techniques to provide insights into machine learning model predictions and demonstrated the potential of these techniques in improving model interpretability.

