



ISSN: 0067-2904

Code-Talker Paradox: Machine Learning Algorithms with Differential Privacy to Classify Heart Failure Patients

Idrees A. Zahid¹, Samer Alaa Hussein^{1*}, Shakir Mahmood Mahdi²

¹Information Technology Center, University of Technology- Iraq, Baghdad, Iraq

²Department of Higher Education, University of Technology- Iraq, Baghdad, Iraq

Received: 1/4/2024

Accepted: 15/8/2024

Published: 30/8/2025

Abstract

The Code-Talker Paradox concept is applied to test the ability to achieve differential privacy while maintaining an acceptable level of machine learning accuracy. A real-world dataset for heart failure patients is used to test the accuracy. Four different machine learning algorithms, namely: decision tree, logistic regression, random forest, and Naïve Bayes, are employed. Laplace noise is added to the raw dataset to protect private and sensitive user data. This research aims to: first, find a balanced noise scale where differential privacy is achievable with an acceptable accuracy result. Second, evaluate the four machine learning classifiers and introduce the one that best fits the current heart failure dataset. Hyperparameter tunings have been applied to the employed algorithms. Different levels and scenarios are tested with the Laplace noise scale and added to the raw data. The accuracy results are recorded and compared. Laplace noise between 1 and 4 does not affect accuracy, while 5 to 7 results in regularization and increases the accuracy accordingly. A Laplace noise value of 28 and above significantly reduces the accuracy value. Finally, the decision tree shows the more stable algorithm regarding the added noise. While logistic regression is the more fluctuating algorithm, it still presents the highest accuracy. Potential future research and study limitations are discussed in order to contribute to a more comprehensive study.

Keywords: Differential Privacy, Accuracy, Laplace Noise, Heart Failure, Machine Learning, Security.

Code-Talker Paradox : خوارزميات التعلم الآلي مع الخصوصية التفاضلية لتصنيف مرضى قصور القلب

ادريس علاء زاهد¹, سامر علاء حسين^{1*}, شاكر محمود مهدي²

¹ مركز تكنولوجيا المعلومات, الجامعة التكنولوجية, بغداد, العراق

² قسم الدراسات العليا, الجامعة التكنولوجية, بغداد, العراق

الخلاصة

يتم تطبيق مفهوم Code-Talker Paradox لاختبار القدرة على تحقيق الخصوصية التفاضلية مع الحفاظ على مستوى مقبول من دقة التعلم الآلي في نفس الوقت مع الحفاظ على أفضل توازن بين الجانبين.

*Email: samer.a.hussein@uotechnology.edu.iq

يتم استعمال مجموعة بيانات حقيقية لمرضى قصور القلب لاختبار الدقة. وباستعمال أربع خوارزميات مختلفة للتعلم الآلي؛ وهي: Random Forest Naive Bayes، Logistic Regression، Decision Tree. تم إضافة Laplace Noise إلى مجموعة البيانات الأولية لحماية بيانات المستخدم الخاصة والحساسة. يهدف هذا البحث إلى: أولاً، إيجاد مقياس الضوضاء المتوازن حيث يمكن تحقيق الخصوصية التضاضلية بنتيجة دقة مقبولة. ثانياً يتم تقييم مصنفات التعلم الآلي الأربعة وتقديم أفضل ما يناسب الحالة الحالية لمجموعة بيانات قصور القلب. تم اختبار مختلف الاعدادات للمتغيرات الخاصة بخوارزميات التعلم الآلي المستعملة. تم اختبار مستويات وسيناريوهات مختلفة باستعمال مقياس ضوضاء لابلاس وإضافتها إلى البيانات الأولية. تم تسجيل نتائج الدقة ومقارنتها. تم إيجاد ما يلي: لا تؤثر Laplace Noise بين 1 و 4 على الدقة، في حين أن 5 إلى 7 تؤدي إلى التنظيم وتزيد الدقة وفقاً لذلك. تمثل قيمة Laplace Noise البالغة 28 وما فوق انخفاضاً ملحوظاً في قيمة الدقة. وأخيراً، تعرض Decision Tree الخوارزمية الأكثر استقراراً فيما يتعلق بالضوضاء المضافة. في حين أن Logistic Regression هو خوارزمية الأكثر تقلباً فإنها لا تزال تقدم أعلى دقة. تم مناقشة التوجهات المستقبلية الممكنة للدراسة الحالية والجوانب المحدودة الممكن تطويرها من أجل المشاركة في إنتاج دراسة شاملة.

1. Introduction

The code-talker paradox introduces an interesting concept. It was first implemented in the United States during World War II, when the indigenous language of the Navajo tribe was used for cipher communication [1]. The enemy can decipher the traditional coding of the messages. The indigenous language comes into rescue, as very few of the numbers belonging to this indigenous tribe use it. From here, the concept emerges as the language, which is a means of communication, begins to serve as a means of ambiguity and coding. The code-talker paradox concept is heavily utilized by linguists and historians in their research. The social production of races and its impact have been discussed in the context of the code-talker paradox [2]. Cultural-related issues are researched using this concept, as in [3], where film and visual impact are discussed. The historians, on the other hand, explored the extent of the code-talker paradox in different paradigms. Incorporating the historical context of the concept into military applications adds additional potential and interesting values, as it aids in military communications and enemy counter-interceptions [4]. Using the concept for military purposes ignites the idea of exploring all possible applications of the code-talker paradox. One of which is privacy preservation for various applications with guaranteed accuracy. In this regard, the concept is employed to hide sensitive information or private data, while at the same time exposing the information for different purposes.

Privacy concerns and data protection awareness start to spread to multiple disciplines, not just for military or war purposes. Especially in the health sector and patients' records obtained from diagnoses and scanning medical devices. Security awareness has evolved extensively, and it has been a research focus recently [5]. Security and privacy are pivotal for all domains. Particularly for patients and healthcare data, privacy is crucial because it pertains to sensitive information. Artificial intelligence (AI) is employed in the health sector to assist paramedical staff [6]. AI also articulates and affirms priorities and decisions produced by decision-makers, including various aspects and directions ranging from diagnosis to triaging patients to even assisting in the final categorization and recommendation [7] [8] [9].

Sun et al. [10] discussed medical data privacy, while research [11] also explored health data privacy. Given that heart failure ranks among the leading causes of death, approximately 6 million individuals in the United States suffer from this condition [12]. The necessity of studying heart failure and the impact of breaching such sensitive data is huge. With the breakthrough of science and especially artificial intelligence, employing its algorithms is spreading widely for better outcomes. Research collaboration produced more machine learning (ML) and power. AI algorithms and platforms that could be utilized for more

accurate, optimized, and enhanced results [13], [14]. Various sectors, medical and non-medical, extensively employ deep learning algorithms for the purpose of accurate and enhanced classification [15], [16], [17]. A heart failure dataset from the University of California, Irvine, is implemented in this study [18]. Private data was altered by adding noise to those specific items for each record, and four different machine learning classifiers were employed to test the evaluation metrics, especially the classification accuracy (CA). The argument behind this study is to find a balance between privacy and accuracy for the dataset studied in this manuscript. Evaluation metrics results were examined, and those models' hyperparameters have been tuned for better results.

In this research, we aim to:

1. Find a balanced noise scale in order to achieve the protection required for private data and maintain acceptable classification accuracy, amongst other metrics.
2. The second goal of this paper is to mark the best fit among the four machine learning classifiers that are more compatible with the heart failure dataset and achieve the requirements in point 1.

This paper is organized as follows: The next section discusses the literature review. Work procedures and tools are introduced in the third section, i.e., methodology. Results and discussion are presented in the fourth section. The fifth section discusses future research directions and potential limitations. The conclusion is drawn in the last section.

2. Related Work

Data privacy and information security have gained more concern nowadays. As the code talker paradox is not widely used within the medical field, differential privacy is still commonly used in multidisciplinary fields. Differential privacy is widely used in security research, cyber-physical systems, blockchain, and Internet of Thing encryption research [19]–[21]. The medical and health sectors have extensively researched and implemented the concept of privacy protection. Some of this research covers sharing patient-sensitive information [22]. Li et al. researched preserving privacy in brain tumor segmentation [23]. E-healthcare, medical big data, and health recommender systems from the user perspective have all embedded the differential privacy or privacy-preserving concept [24]–[26]. Other diseases and treatments heavily rely on modern hybrid models for classification and detection, such as skin cancer detection and autism triaging, among others [8], [27]. Diagnoses, triage, and prioritization in the healthcare sector are extensively explored with the employment of AI techniques, and results accuracy holds a major concern [28]. Classification accuracy is then considered a focal point, especially in the medical and healthcare sectors [7]. Increasing and even maintaining accuracy is essential to almost all systems and sectors. The healthcare system is one of the vastly researched areas in that manner. In [29], the study discussed an efficient deep-learning approach for the classification of pneumonia. The accuracy of medical imaging and electronic classification is widely used in the healthcare system [30]. Diabetes, disease, and health big data are all concerns with classification accuracy [31], [32], and [33]. Among all researched diseases, heart failure and cardiovascular diseases are intensively researched, especially from a classification accuracy perspective [34]. Implementing deep learning approaches or machine learning algorithms as well as neural networks to accurately classify heart disease [35]–[39]. Several studies have examined the scope of heart failure disease and the differential privacy concept. Additionally, Islam et al. [40] adopted a differential privacy approach to confuse the local model before transmission to the extra privacy layer to achieve a practical heart failure/cancer disease predictor while ensuring privacy. Furthermore, Grama et al. [41] proposed Federated Learning as differential privacy

and used two real-world datasets and observed that differential privacy did not have a significant impact on the learning convergence for the aggregation strategies adopted. As most ML and AI optimization algorithms search for optimal solutions, a balanced setting is required [42]. Generally, most of the studies related to health care and heart diseases specifically focus mainly on classification accuracy and maintaining the highest classification metrics. Other studies used differential privacy while testing for accuracy. Still, the research gap covered in this manuscript clearly illustrates that no research study tested the amount of noise added to maintain acceptable accuracy, and no comparison of machine learning algorithms' behavior with differential privacy is presented.

3. Methodology

The heart failure dataset from the University of California, Irvine, is used in our approach. Starting with sensitive and private data, the data attributes are identified. Laplace noise was then added to the extracted attributes from the retrieved dataset of heart failure. Four machine learning algorithms, namely Decision Tree, Logistic Regression, Random Forest, and Naïve Bayes, were used to test and evaluate the metrics. Classification accuracy, precision, and recall were the evaluation metrics implemented in the methodology. To demonstrate the overall working methodology, Figure 1 illustrates the stages of the methodology involved. As depicted in Figure. 1, the methodology starts with the patient's dataset. This dataset is fed into the four machine learning algorithms, and the metrics are evaluated afterward. Parallely, sensitive and private data is marked and identified. Laplace noise with different scales is then applied to the already-specified private data. Hyperparameters have been tuned for both sections of the data—the noisy data and the unaltered ones. Then the noisy data is fed into the four machine learning algorithms, and the resultant metrics are evaluated as well.

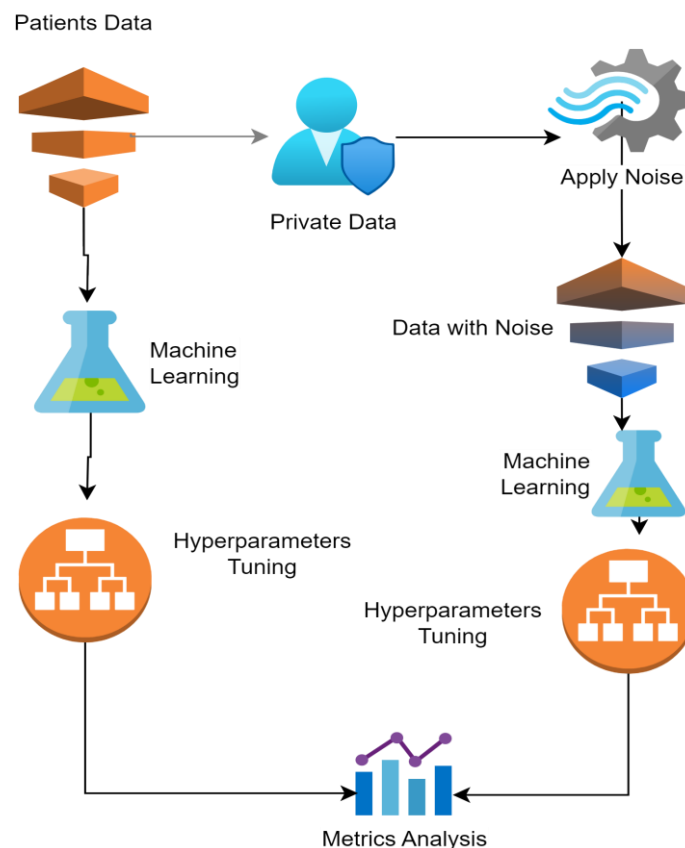


Figure 1: Methodology Procedure and Stages

3.1 Dataset

The dataset used in this research was obtained from the University of California Irvine Machine Learning Repository [18]. This dataset represents 299 heart failure patients, each with 13 features (attributes), including a target value that determines the presence or absence of a death event during the follow-up procedure. The dataset has been examined, and it has no missing values. Two versions of the dataset have been implemented according to the required methodology. The dataset, in its raw form, is the heart failure dataset. The noisy dataset (the Laplace noise) is added to the markedly sensitive data.

3.2 Laplace Noise Addition

Laplace noise is one of the statistical methods used to be applied when differential privacy needs to be implemented. Hence, the Laplace scale parameter was used to test the differential privacy concept for the used dataset. The scale will determine the amount of noise we add to the dataset. Accordingly, two datasets were used in our approach. The first dataset contains raw data. Then, using the developed Python script, the second dataset was created by adding Laplace noise. The noise was added to the features that represent private data points for each patient. A new dataset was generated accordingly and applied to the tested algorithms.

Laplace noise formula is presented in the equation below:

$$Scale = \frac{Sensitivity}{Epsilon} \quad (1)$$

Where:

Sensitivity: represents the data sensitivity, i.e., the maximum change as one unit of noise is added to the data.

Epsilon: Also called privacy budget, represents the amount of added noise to the data.

Considering some of the dataset columns as private information about the patients, representing related information about the patients, or standing for some history-related information. Laplace noise was added to those columns to make the dataset more private.

Developed code is freely available at the GitHub repository at: <https://github.com/iajzahid/laplace-noise.git>.

3.3 Machine Learning Algorithm

A set of machine learning algorithms is used to test the accuracy, precision, and recall evaluation metrics. Those machine learning algorithms were deployed in two stages:

- Raw dataset of the heart failure disease classification of patients was applied to train and test the four machine learning classifiers.
- Noisy dataset of the patients is tested as input to the trained and tuned machine learning classifiers.

Both pipelines were tested via the following machine learning algorithms: Decision Tree, Logistic Regression (L. Regression), Naïve Bayes, and Random Forest.

Figure 2 shows the machine learning testbed, where the data was first trained and tested using the mentioned machine learning algorithms. Equation (1) is then utilized to modify the sensitivity and Epstein values, resulting in the creation of multiple Laplace noise scales. Before testing the data using the four trained machine learning classifiers, forward those different scale values.

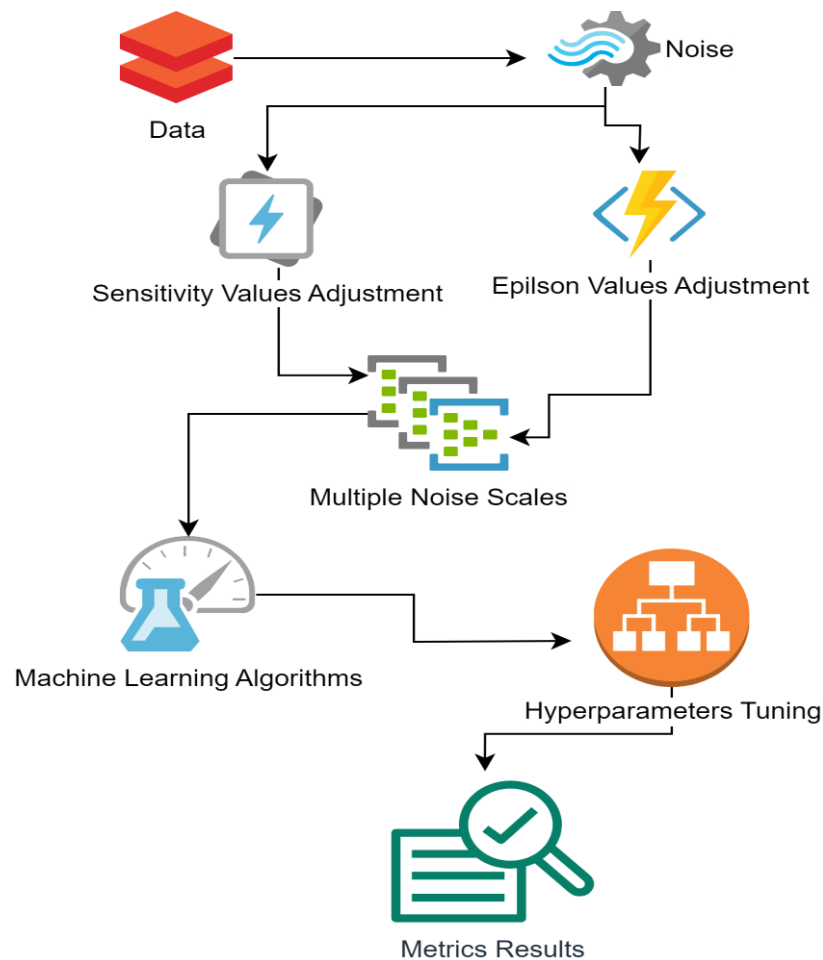


Figure 2: Laplace Noise Multiple Scales

4. Results and Discussion

As the framework flows, results have been recorded for each case, i.e., the noise scale. Evaluation metrics according to the machine learning algorithms used are recorded alongside the Sensitivity and Epsilon values, as well as the metrics as follows:

4.1 Evaluations on Raw Dataset

In this case, raw data is fed into the machine learning algorithms to test the classification accuracy, as no noise is added. Table 1 presents the results of this scenario:

Table 1: Evaluation Results on Raw Dataset

Model	Accuracy	Precision	Recall	Epsilon	Sensitivity
Decision Tree	0.6333	0.5789	0.44	N/A	N/A
L. Regression	0.8	0.8824	0.6	N/A	N/A
R. Forest	0.75	0.8571	0.48	N/A	N/A
Naive Bayes	0.7333	0.9091	0.4	N/A	N/A

Table 1 presents evaluation metrics: classification accuracy, precision, and recall for the four studied machine learning classifiers. In this case, raw data with no noise has been used to train and test the evaluation metrics for those machine learning algorithms. Logistic regression provides the highest accuracy, at 80%. While the decision tree gives us the lowest

accuracy with 63.33%. Random Forest and Naïve Bayes come in the middle with 75% and 73.33%, respectively.

4.2 Evaluations on Noisy Dataset

26 variation scenarios and adjusted values for the column's sensitivity and privacy budget were used, and multiple cases were produced as follows:

4.2.1 Infinitesimal Noise:

The infinitesimal values set for sensitivity and privacy budget, represented by the Epsilon value, do not alter the evaluation metrics used to compare and evaluate the results of the machine learning classifiers used in this approach. Table 2 displays the results for classification accuracy, precision, and recall for each of the four machine learning algorithm variations.

Table 2: Infinitesimal Noise

Model	Accuracy	Precision	Recall	Epsilon	Sensitivity
Decision Tree	0.6333	0.5789	0.44	0.1	0.2
L. Regression	0.8	0.9333	0.56	0.1	0.2
R. Forest	0.75	0.8571	0.48	0.1	0.2
Naive Bayes	0.7333	0.9091	0.4	0.1	0.2

As we begin to add noise using equation 1, we set the sensitivity value in Table 2 to 0.2 and the privacy budget represented by Epsilon to 0.1. No change has occurred to the registered values of the evaluation metrics as the Laplace noise added is infinitesimal.

4.2.2 Small noise

A small amount of noise scale is added to the fed dataset, and the resulting metrics are presented in Table 3. As we compare the accuracy metric as well as others for this case scenario with raw data, an increase in those metrics is noticed. Accuracy, precision, and recall values are improved despite the added noise. In this case, a regularization occurs when a system is too complex, and adding a specific amount of Laplace noise could result in an improvement in machine learning classification. Table 3 presents the improved evaluation metrics for the machine learning classifiers used.

Table 3: Small Noise Values

Model	Accuracy	Precision	Recall	Epsilon	Sensitivity
Decision Tree	0.6667	0.619	0.52	0.1	0.6
L. Regression	0.8167	0.8889	0.64	0.1	0.6
R. Forest	0.7667	0.9231	0.48	0.1	0.6
Naive Bayes	0.75	0.9167	0.44	0.1	0.6

In Table 3, a slight increase is added to the Laplace noise scale. The sensitivity value is set to 0.6, while the privacy budget remains at 0.1. As the noisy dataset is fed into the trained model, accuracy metrics show an increase in their values for all the classifiers. The accuracy of the decision tree increased from 63% to 66%. Logistic regression accuracy also increased, going from 80% to 81.6%. As well as Random Forest and Naïve Bayes, both increased from 75% to 76.6% and from 73% to 75%, respectively. This phenomenon indicates that the trained system was too complex, and overfitting is reduced with the regularization approach as we added a slight amount of noise.

4.2.3 Considerable Noise

As we proceed through the testing scenarios and start to increase both the sensitivity value and Epsilon values, there is a reduction in accuracy, precision, and recall value. As the amount of noise added to the private information increases, evaluation metrics become more reeducational. Table 4 presents a sample result, utilizing Sensitivity and Epsilon to represent the added amount of noise according to Equation 1.

Table 4: Considerable Noise Value

Model	Accuracy	Precision	Recall	Epsilon	Sensitivity
Decision Tree	0.6333	0.5652	0.52	0.2	10.0
L. Regression	0.7333	0.68	0.68	0.2	10.0
R. Forest	0.6667	0.6	0.6	0.2	10.0
Naive Bayes	0.65	0.5667	0.68	0.2	10.0

Table 4 shows that the trained model has been subjected to a significant amount of noise. Sensitivity is set to 10.0, and the privacy budget is set to 0.2. In Table 1, the results of the metrics with no noise dataset showed a slight decrease in accuracy. Logistic regression went down from 80% to 73%. Random Forest went from 75% to 66.7%. And Naïve Bayes went from 73% to 65%. The decision tree maintains its accuracy at 63%.

To provide a more comprehensive analysis and to illustrate the trade-off between privacy and the tested accuracy, Table 5 summarizes the privacy-accuracy trade-off with different noise levels.

Table 5: Privacy-Accuracy Trade-off

Noise Level	Accuracy (Metric)	(Privacy Loss)	Privacy Guarantee
Low	High	High	Weak
Medium	Medium	Medium	Moderate
High	Low	Low	Strong

4.3 Hyperparameter Tuning Analysis

To provide a more comprehensive evaluation of the accuracy and privacy trade-off, the impact of different hyperparameter settings on the machine learning algorithms should be explored. Hyperparameter tuning can have a significant impact on models' performance and sensitivity to noise, affecting privacy guarantees. Specifying the hyperparameters used for each machine learning model and their optimization process is essential. Below is a list of the key hyperparameters for each model used in our study. Table 6 presents the tuning process and optimal settings for each hyperparameter per model.

Table 6: Hyperparameters Tuning Process

Model	Hyperparameters	Tuning Process	Optimal Settings
Decision Tree	Maximum depth, minimum samples split, minimum samples leaf	Grid search over a range of values for each hyperparameter	max_depth=10 min_samples_split=5 min_samples_leaf=2.
Logistic Regression	Regularization strength (C), solver	Grid search over different values of C and various solvers	C=0.5, solver='liblinear
Random Forest	Number of estimators, maximum depth, minimum samples split, minimum samples leaf	Random search over a wide range of values for each hyperparameter	n_estimators=100, max_depth=15, min_samples_split=4 min_samples_leaf=2.
Naïve Bayes	None (parameters are inherently defined by the algorithm and dataset characteristics)	Naive Bayes does not require extensive tuning, but model selection between Gaussian, Bernoulli, or Multinomial Naive Bayes is performed	Multinomial Naive Bayes performed best for the given dataset.

The analysis reveals the impact of hyperparameter tuning on the performance of machine learning models in terms of accuracy and privacy trade-offs. Below is a comparative summary of the results before and after hyperparameter tuning under different noise levels, based on the settings and tuning applied in Table 6. Optimal settings for each model presented in Table 6 are applied, and a comparative and detailed analysis is performed. Accuracy improvement is analyzed as a privacy-accuracy trade-off as well as the impact of the noise level.

4.3.1 Accuracy Improvements:

- **Raw Dataset:** Hyperparameter tuning resulted in improved accuracy for all models. For example, the decision tree's accuracy increased from 63.33% to 68.33%, and the logistic regression's from 80% to 81.33%. Other machine-learning algorithms followed suit.
- **Infinitesimal Noise:** Similar improvements were observed when infinitesimal noise was added. The decision tree's accuracy increased from 63.33% to 68.33%, as did logistic regression, from 80% to 81.33%.
- **Small Noise:** The benefits of hyperparameter tuning were more explicit with small noise. Decision tree accuracy rose from 66.67% to 70%, and logistic regression rose from 81.67% to 83%.
- **Considerable Noise:** The accuracy of the models remained similar with and without tuning when considerable noise was added. This indicates that, beyond a certain noise threshold, hyperparameter tuning does not significantly impact your performance.

4.3.2 Privacy-Accuracy Trade-off:

- **With Hyperparameter Tuning:** The models achieved higher accuracy with small and infinitesimal noise levels while maintaining the same privacy guarantees (and sensitivity values). This demonstrates that hyperparameter tuning can help achieve better performance without compromising privacy.
- **Without Hyperparameter Tuning:** The models generally performed worse, indicating a higher trade-off between accuracy and privacy.

4.3.3 Impact of Noise Levels:

- **Low Noise Levels:** Hyperparameter tuning was effective in improving accuracy without a significant drop in performance.
- **High Noise Levels:** The performance gains from hyperparameter tuning were less significant. This indicates that beyond a certain point, the added noise cancels the benefits of tuning for the hyperparameter.

Collectively, hyperparameter tuning enhances the performance of machine learning models. This is especially true when handling unprocessed or low-noise datasets. This optimization helps achieve higher accuracy while maintaining privacy guarantees. However, as the noise level increases significantly, the impact of tuning decreases. Incorporating hyperparameter tuning into the analysis provides a more comprehensive understanding of the trade-offs between accuracy and privacy, thereby offering valuable insights for practical applications where both are critical and important for applicable studies.

4.4 Accuracy Trending Line

Multiple scenarios with different noise scales are tested. A trending line representing the accuracy metric is drawn using those 26 scenarios along with each of the Laplace noise scale values.

As shown in Figure 3, the x-axis is represented by the Laplace noise scale values. The classification accuracy for the four classifiers is presented. The Y-axis represents the accuracy metric. According to the trending line represented by the accuracy metric, a noticeable decrease occurs when the value of Laplace noise is between 28 and above. When the Laplace noise scale value is between 5 and 7, regularization occurs. Amongst all the four classifiers, the Decision Tree algorithm shows the most stable accuracy with the changeable Laplace noise. The logistic regression algorithm exhibits significant fluctuations in accuracy metrics when the Laplace noise scale shifts.

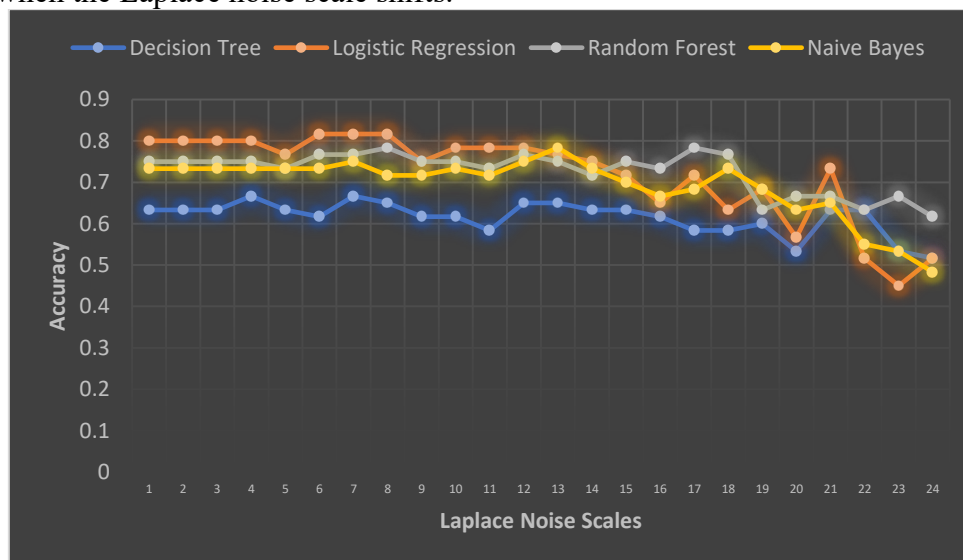


Figure 3: Classifiers Accuracy with Laplace Noise

5. Future Direction and Study Limitations

This section discusses possible future directions and limitations of the study to gain a thorough understanding. Potential improvements and possible enhancements in our differential privacy study employing the code-talker paradox could be listed in the following directions:

1. Various Real-world applications: The current technique could be applied to different real-world applications in future research and help validate its effectiveness in multiple domains. It is applicable in other healthcare sectors and other diseases, and it also has a lot of potential in finance.

2. Hybrid Models: Limitations and drawbacks could possibly overcome by combining multiple machine learning models. This would open possibilities for more accurate results within the privacy limitations of the specific application.

3. Alternative Noise Addition Techniques: future directions could utilize other noise addition techniques. An alternative to Laplace noise addition was used in this study. Gaussian noise and adaptive noise techniques could be employed and might produce enhanced and balanced results between privacy and accuracy.

For the limitations of the study, the following points have been addressed: Examining these points could potentially lead to improvements in the study and enhance the research. The following points outline the observed limitations:

1. Limitation of noise levels: The current study explored limited levels of noise. Delving into a wider range of noise levels added to the dataset could produce different results and remarks.

2. Algorithm Generalizations: This study explores the balance of privacy and accuracy using only four ML algorithms. Evaluating the current method for a different and larger set of machine learning algorithms would result in more generalizations. This generalization would have an impact on the approach's applicability.

6. Conclusion

Preserving user data privacy is applicable while maintaining an acceptable classification accuracy level. Using the four machine learning algorithms, namely: decision tree, logistic regression, random forest, and Naïve Bayes, with a real-world heart failure dataset and evaluated with accuracy, precision, and recall metrics. Tuning hyperparameters for the employed ML algorithms produced enhanced results. For the noisy dataset and the unaltered ones, each ML's hyperparameters have been addressed. Applying different scale values for Laplace noise to achieve differential privacy for sensitive and private data, classification accuracy fluctuates. The accuracy of classifiers remains unaffected by the Laplace noise scale, which ranges from 1 to 4. Interestingly, when adding a Laplace noise of values between 5 and 7, classification accuracy increased as regularization occurred. A noticeable decrease in accuracy was recorded when the Laplace noise level was 28 and above. As tested with the heart failure dataset, the decision tree algorithm is more resistant to Laplace noise and shows almost stable output. Logistic regression, on the other hand, presents the highest fluctuations in classification accuracy value amongst other algorithms. Potential limitations and future directions are addressed to identify a more comprehensive study.

References

- [1] A. Pubill Ambros and C. N. Buzinde, "Indigenous self-representations in the touristic sphere," *Ann. Tour. Res.*, vol. 86, p. 103099, Jan. 2021, doi: 10.1016/J.ANNALS.2020.103099.
- [2] J. Romanow, "Mediating Whiteness: Triangular Racialization in the Anglo-Indian Picaresque," *Vic. Lit. Cult.*, pp. 1–23, 2023, doi: 10.1017/S1060150323000372.
- [3] A. Griffiths, "Amateur Film, Cultural Memory and the Visual Legacy of the 1920s Inter-Tribal Indian Ceremonial," *Vis. Anthropol.*, vol. 36, no. 3, pp. 201–228, 2023, doi: 10.1080/08949468.2023.2203294.
- [4] T. P. Galvin and C. D. Allen, "Diversity Management and the Postdiversity Vision: An Applied Pragmatist Approach," *Armed Forces Soc.*, vol. 47, no. 1, pp. 48–76, Apr. 2020, doi: 10.1177/0095327X20920311.
- [5] I. Zahid, S. Hussein, and S. Mahdi, "Measuring Individuals Cybersecurity Awareness Based on

- Demographic Features,” *Iraqi J. Electr. Electron. Eng.*, vol. 20, no. 1, pp. 58–67, Jun. 2024, doi: 10.37917/IJEEE.20.1.6.
- [6] S. S. Joudar , A. S. Albahri, R. A. Hamid, I. A. Zahid, M. E. Alqaysi, O. S. Albahri, and A. H. Alamoody., “Artificial intelligence-based approaches for improving the diagnosis, triage, and prioritization of autism spectrum disorder: a systematic review of current trends and open issues,” *Artificial Intelligence Review*, vol. 56, pp. 53 - 117, 2023, doi: 10.1007/s10462-023-10536-x.
- [7] S. S. Joudar, A. S. Albahri, and R. A. Hamid, “Intelligent triage method for early diagnosis autism spectrum disorder (ASD) based on integrated fuzzy multi-criteria decision-making methods,” *Informatics Med. Unlocked*, vol. 36, p. 101131, 2023, doi: 10.1016/j.imu.2022.101131.
- [8] A. S. Albahri, S. S. Joudar, R. A. Hamid, I. A. Zahid, M. E. Alqaysi, O. S. Albahri, A. H. Alamoody, G. Kou, and I.M. Sharaf., “Explainable Artificial Intelligence Multimodal of Autism Triage Levels Using Fuzzy Approach-Based Multi-criteria Decision-Making and LIME,” *Int. J. Fuzzy Syst.*, vol. 26, no. 1, pp. 274–303, Feb. 2024, doi: 10.1007/S40815-023-01597-9/METRICS.
- [9] A. Y. Yousif, S. M. Younis, S. A. Hussein, N. Mohammed, and G. Al-Saidi, “AN INTELLIGENT COMPUTING FOR DIAGNOSING COVID-19USING AVAILABLE BLOOD TESTS,” *Int. J. Innov. Comput.*, vol. 18, no. 1, pp. 57–72, 2022, doi: 10.24507/ijic.18.01.57.
- [10] Z. Sun, Y. Wang, M. Shu, R. Liu, and H. Zhao, “Differential Privacy for Data and Model Publishing of Medical Data,” *IEEE Access*, vol. 7, pp. 152103–152114, 2019, doi: 10.1109/ACCESS.2019.2947295.
- [11] O. Choudhury *et al.*, “Differential Privacy-enabled Federated Learning for Sensitive Health Data,” Oct. 2019, Accessed: Apr. 13, 2023. [Online]. Available: <https://arxiv.org/abs/1910.02578v3>.
- [12] “Heart Failure - | NHLBI, NIH.” <https://www.nhlbi.nih.gov/health/heart-failure> (accessed Apr. 13, 2023).
- [13] I. A. Zahid and S. S. Joudar, “Does Lack of Knowledge and Hardship of Information Access Signify Powerful AI? A Large Language Model Perspective,” *Appl. Data Sci. Anal.*, vol. 2023, pp. 150–154, Dec. 2023, doi: 10.58496/ADSA/2023/014.
- [14] S. A. Hussein and A. Y. Yousif, “An Improved Meerkat Clan Algorithm for Solving 0-1 Knapsack Problem”, *Iraqi J. Sci.*, vol. 63, no. 2, pp. 773–784, Feb. 2022, doi.org/10.24996/ij.2022.63.2.32.
- [15] M. M. Mijwil, R. Doshi, K. K. Hiran, O. J. Unogwu, I. Bala, and A. History, “MobileNetV1-Based Deep Learning Model for Accurate Brain Tumor Classification,” *Mesopotamian J. Comput. Sci.*, vol. 2023, pp. 29–38, Mar. 2023, doi: 10.58496/MJCSC/2023/005.
- [16] R. T. Hameed and O. A. Mohamad, “Federated Learning in IoT: A Survey on Distributed Decision Making,” *Babylonian J. Internet Things*, vol. 2023, pp. 1–7, Jan. 2023, doi: 10.58496/BJIOT/2023/001.
- [17] G. G. Shayeia *et al.*, “Fuzzy Evaluation and Benchmarking Framework for Robust Machine Learning Model in Real-Time Autism Triage Applications,” *Int. J. Comput. Intell. Syst.*, vol. 17, no. 1, pp. 1–27, Dec. 2024, doi: 10.1007/S44196-024-00543-3/TABLES/9.
- [18] D. Chicco and G. Jurman, “Machine learning can predict survival of patients with heart failure from serum creatinine and ejection fraction alone,” *BMC Med. Inform. Decis. Mak.*, vol. 20, no. 1, pp. 1–16, Feb. 2020, doi: 10.1186/S12911-020-1023-5/TABLES/11.
- [19] M. U. Hassan, M. H. Rehmani, and J. Chen, “Differential Privacy Techniques for Cyber Physical Systems: A Survey,” *IEEE Commun. Surv. Tutorials*, vol. 22, no. 1, pp. 746–789, Jan. 2020, doi: 10.1109/COMST.2019.2944748.
- [20] Y. Liu, J. J. Q. Yu, J. Kang, D. Niyato, and S. Zhang, “Privacy-Preserving Traffic Flow Prediction: A Federated Learning Approach,” *IEEE Internet Things J.*, vol. 7, no. 8, pp. 7751–7763, Aug. 2020, doi: 10.1109/JIOT.2020.2991401.
- [21] M. Shen, X. Tang, L. Zhu, X. Du, and M. Guizani, “Privacy-Preserving Support Vector Machine Training over Blockchain-Based Encrypted IoT Data in Smart Cities,” *IEEE Internet Things J.*, vol. 6, no. 5, pp. 7702–7712, Oct. 2019, doi: 10.1109/JIOT.2019.2901840.
- [22] H. Jin, Y. Luo, P. Li, and J. Mathew, “A Review of Secure and Privacy-Preserving Medical Data

- Sharing,” *IEEE Access*, vol. 7, pp. 61656–61669, 2019, doi: 10.1109/ACCESS.2019.2916503.
- [23] W. Li, Fausto Milletari, Daguang Xu, Nicola Rieke, Jonny Hancox, Wentao Zhu, Maximilian Baust, Yan Cheng, Sébastien Ourselin, M. Jorge Cardoso, and Andrew Feng., “Privacy-Preserving Federated Brain Tumour Segmentation,” *Lect. Notes Comput. Sci. (including Subser. Lect. Notes Artif. Intell. Lect. Notes Bioinformatics)*, vol. 11861 LNCS, pp. 133–141, 2019, doi: 10.1007/978-3-030-32692-0_16/COVER.
- [24] W. N. Price and I. G. Cohen, “Privacy in the age of medical big data,” *Nat. Med.* 2019 251, vol. 25, no. 1, pp. 37–43, Jan. 2019, doi: 10.1038/s41591-018-0272-7.
- [25] W. Tang, J. Ren, K. Deng, and Y. Zhang, “Secure Data Aggregation of Lightweight E-Healthcare IoT Devices with Fair Incentives,” *IEEE Internet Things J.*, vol. 6, no. 5, pp. 8714–8726, Oct. 2019, doi: 10.1109/JIOT.2019.2923261.
- [26] A. Calero Valdez and M. Ziefle, “The users’ perspective on the privacy-utility trade-offs in health recommender systems,” *Int. J. Hum. Comput. Stud.*, vol. 121, pp. 108–121, Jan. 2019, doi: 10.1016/J.IJHCS.2018.04.003.
- [27] H. J. Mohammed, A. A. Nafea, H. K. Almulla, S. A. S. Aliesawi, and M. M. Al-Ani, “An Effective Hybrid Model for Skin Cancer Detection Using Transfer Learning,” *Proc. - Int. Conf. Dev. eSystems Eng. DeSE*, pp. 840–845, 2023, doi: 10.1109/DESE60595.2023.10468994.
- [28] S. S. Joudar, A. S. Albahri, and R. A. Hamid, “Triage and priority-based healthcare diagnosis using artificial intelligence for autism spectrum disorder and gene contribution: A systematic review,” *Comput. Biol. Med.*, vol. 146, p. 105553, Jul. 2022, doi: 10.1016/J.COMPBIOMED.2022.105553.
- [29] O. Stephen, M. Sain, U. J. Maduh, and D.-U. Jeong, “An Efficient Deep Learning Approach to Pneumonia Classification in Healthcare,” *J. Healthc. Eng.*, vol. 2019, pp. 1–7, Mar. 2019, doi: 10.1155/2019/4180949.
- [30] S. C. Huang, A. Pareek, S. Seyyedi, I. Banerjee, and M. P. Lungren, “Fusion of medical imaging and electronic health records using deep learning: a systematic review and implementation guidelines,” *npj Digit. Med.* 2020 31, vol. 3, no. 1, pp. 1–9, Oct. 2020, doi: 10.1038/s41746-020-00341-z.
- [31] A. S. Albahri, A. A. Zaidan, O. S. Albahri, B. B. Zaidan, A. H. Alamoodi, Ali H. Shareef, Jwan K. Alwan, Rula A. Hamid, M. T. Aljbory, Ali Najm Jasim, M. J. Baqer, and K. I. Mohammed., “Development of IoT-based mhealth framework for various cases of heart disease patients,” *Health Technol. (Berl.)*, vol. 11, no. 5, pp. 1013–1033, Sep. 2021, doi: 10.1007/S12553-021-00579-X/METRICS.
- [32] M. Maniruzzaman, M. J. Rahman, B. Ahammed, and M. M. Abedin, “Classification and prediction of diabetes disease using machine learning paradigm,” *Heal. Inf. Sci. Syst.*, vol. 8, no. 1, pp. 1–14, Dec. 2020, doi: 10.1007/S13755-019-0095-Z/TABLES/13.
- [33] W. Xing and Y. Bei, “Medical Health Big Data Classification Based on KNN Classification Algorithm,” *IEEE Access*, vol. 8, pp. 28808–28819, 2020, doi: 10.1109/ACCESS.2019.2955754.
- [34] J. P. Li, A. U. Haq, S. U. Din, J. Khan, A. Khan, and A. Saboor, “Heart Disease Identification Method Using Machine Learning Classification in E-Healthcare,” *IEEE Access*, vol. 8, pp. 107562–107582, 2020, doi: 10.1109/ACCESS.2020.3001149.
- [35] R. K. Sevakula, W. T. M. Au-Yeung, J. P. Singh, E. K. Heist, E. M. Isselbacher, and A. A. Armoundas, “State-of-the-Art Machine Learning Techniques Aiming to Improve Patient Outcomes Pertaining to the Cardiovascular System,” *J. Am. Heart Assoc.*, vol. 9, no. 4, Feb. 2020, doi: 10.1161/JAHA.119.013924.
- [36] N. I. Hasan and A. Bhattacharjee, “Deep Learning Approach to Cardiovascular Disease Classification Employing Modified ECG Signal from Empirical Mode Decomposition,” *Biomed. Signal Process. Control*, vol. 52, pp. 128–140, Jul. 2019, doi: 10.1016/J.BSPC.2019.04.005.
- [37] C. R. Olsen, R. J. Mentz, K. J. Anstrom, D. Page, and P. A. Patel, “Clinical applications of machine learning in the diagnosis, classification, and prediction of heart failure,” *Am. Heart J.*, vol. 229, pp. 1–17, Nov. 2020, doi: 10.1016/J.AHJ.2020.07.009.
- [38] Y. Isler, A. Narin, M. Ozer, and M. Perc, “Multi-stage classification of congestive heart failure based on short-term heart rate variability,” *Chaos, Solitons & Fractals*, vol. 118, pp. 145–151, Jan. 2019, doi: 10.1016/J.CHAOS.2018.11.020.
- [39] A. Çınar and S. A. Tuncer, “Classification of normal sinus rhythm, abnormal arrhythmia and

- congestive heart failure ECG signals using LSTM and hybrid CNN-SVM deep neural networks,” <https://doi.org/10.1080/10255842.2020.1821192>, vol. 24, no. 2, pp. 203–214, 2020, doi: 10.1080/10255842.2020.1821192.
- [40] T. U. Islam, R. Ghasemi, and N. Mohammed, “Privacy-Preserving Federated Learning Model for Healthcare Data,” *2022 IEEE 12th Annu. Comput. Commun. Work. Conf. CCWC 2022*, pp. 281–287, 2022, doi: 10.1109/CCWC54503.2022.9720752.
- [41] M. Grama, M. Musat, L. Muñoz-González, J. Passerat-Palmbach, D. Rueckert, and A. Alansary, “Robust Aggregation for Adaptive Privacy Preserving Federated Learning in Healthcare,” Sep. 2020, Accessed: Apr. 13, 2023. [Online]. Available: <https://arxiv.org/abs/2009.08294v1>.
- [42] S. A. Hussein and I. A. Zahid, “Improved Naked Mole-Rat Algorithm Based on Variable Neighborhood Search for the N-Queens Problem,” *Iraqi J. Sci.*, vol. 65, no. 1, pp. 528–545, Jan. 2024, doi: 10.24996/IJS.2024.65.1.41.