

INTELLIGENT SEPSIS DETECTION IN ICU: A REAL-TIME DEEP LEARNING FRAMEWORK WITH HYBRID OPTIMIZATION

*Hema Ambiha. A*¹, Mukesh Kanna K², M. Inbavel³*

ABSTRACT

Without rapid medical intervention, sepsis the body's catastrophic reaction to infection can cause organ failure and increased mortality. There is an immediate need for real-time and automated detection methods because conventional diagnostic methods are often inaccurate or have delays. In this research, to present a deep learning-based sepsis detection system that uses a CNN for early prediction in conjunction with a Modified Bidirectional Gated Recurrent Unit (MBiGRU). Convolutional Neural Networks (CNNs) extract spatially rich and clinically significant features, while MBiGRUs efficiently model temporal patterns in physiological time-series data. The framework incorporates a modified Chaotic Zebra Optimization Algorithm to optimise feature selection and generalization, significantly enhancing model efficiency and eliminating redundant information. Using open-source sepsis datasets, to trained and validated the MBiGRU CNN-CZOA model. In comparison to more conventional models such simple CNN, LSTM, and GRU, the model outperformed them with a 93.82% accuracy, 94.24% precision, 93.82% recall, and 93.68% F1-score. Fast decision-making and efficient calculation are guaranteed by the system's optimization for real-time ICU deployment through edge computing. Based on these results, the suggested approach seems to be well-suited for clinical use in the real world, where it could facilitate prompt intervention and drastically cut down on mortality risks. In order to ensure the security of patient data when used in remote healthcare settings, future research will centre on improving interpretability and implementing federated learning strategies.

Keywords: Sepsis Detection, Deep Learning, Modified Bidirectional Gated Recurrent Unit, Convolutional Neural Network, Modified Chaotic Zebra Optimization, Real-time Monitoring, Edge Computing and ICU.

I. INTRODUCTION

An infection inside the body can lead to sepsis, or the syndrome of systemic inflammatory reaction, which can cause potentially deadly organ failure [1]. Sepsis was first defined in 1991 as a diagnostic criterion for Systemic Inflammatory Response Syndrome (SIRS). The terms sepsis and septic shock were revised in 2016 by the third worldwide consensus document, Sepsis-3. The report also suggested that patients be evaluated for organ dysfunction using the Sequential Organ Failure Assessment (SOFA) [2] score as well as the quick SOFA score (qSOFA). Because of how far medicine has come, something has to be done. The global epidemic of sepsis is currently a major public health concern. The World Health Organization estimates that around 30 million instances of infectious sepsis occur annually globally. There is a 22.5% mortality rate linked to these instances, which is equivalent to 20% of the global mortality rate [3]. The likelihood of the patient dying rises by about four to eight percent in the event that sepsis therapy is postponed by even one hour. Therefore, in order to lower the patient fatality rate, it is crucial to detect sepsis early and to act quickly [4].

Scoring systems such as SOFA and qSOFA are frequently utilized by medical professionals to assess the severity of sepsis and to forecast possible negative effects [5]. But sepsis sufferers aren't the target audience for these grading systems. Furthermore, these systems might struggle to conduct tailored evaluations, especially when patient health data is highly out of the ordinary. Moreover, there is a marked decline in their predictive accuracy [6].

Recent years have seen tremendous advancements in the use of machine learning for the early prediction of sepsis. Results obtained by training machine learning on data from

Department of Computer Science¹,
Karpagam Academy of Higher Education Coimbatore, India¹
hemaambiha.aravindakshan@kahedu.edu.in¹
Department of Computer Applications²,
Karpagam Academy of Higher Education Coimbatore, India²
mukesh982kanna@gmail.com²
Department of Computer Science³,
Sri Vasavi College, Erode- 638 316³
inbavel@gmail.com³

* Corresponding Author

EHR for sepsis prediction are far better than those obtained by currently used clinical scoring methods, demonstrating exceptional performance [7]. Traditional diagnostic approaches sometimes fall short of expectations when it comes to sepsis because of how complex and variable its presentations can be, even if early detection is crucial for better clinical outcomes [8]. Early and accurate identification of sepsis has been made possible in recent years by the application of deep learning (DL) and machine learning (ML) approaches. The advent of widespread Electronic Health Records (EHRs) and the fast expansion of data-driven healthcare have made this a reality. Creating algorithms that can learn from data automatically, without human intervention, is called machine learning [9]. With its help, researchers have analyzed organized clinical data, results, and demographic information. Algorithms including logistic regression, Support Vector Machines (SVM), Random Forests (RF), and gradient boosting machines (XGBoost and LightGBM) have shown promise in terms of early warning sign identification and sepsis risk stratification [10]. The feature engineering process, which requires clinical knowledge to manually extract and choose relevant data, is heavily relied upon.

In contrast, deep learning incorporates multilayered neural networks to represent complicated and non-linear data interactions. Automatic learning of high-level feature representations from raw data is another capability of deep learning [11]. Analyzing time-series data from Intensive Care Units (ICUs), like heart rate, blood pressure, besides respiratory patterns, is a specialty of digital learning models. Because of this, to can model patients' health over time and in real-time [12]. Architectures including Recurrent Neural Networks (RNNs), Long Short-Term Memory (LSTM) networks, and Gated Recurrent Units (GRUs) that can capture trends and sequential dependencies over time are critical for sepsis progression detection [13]. Detection problems involving the transformation of time-series signals into representations similar to images have also made use of Convolutional Neural Networks (CNNs). Furthermore, models that use attention mechanisms and transformers are being studied for their ability to zero focus on the most

insightful aspects of a patient's history. This could lead to better accuracy and outcomes that are easier to understand [14]. Anomaly detection and data augmentation are two applications of unsupervised deep learning techniques, such as autoencoders and Generative Adversarial Networks (GANs). When dealing with imbalanced datasets, which are often marked by a small number of septic cases, these methods are employed [15].

Although these technologies have several clear advantages, like the ability to inform users in real-time, increase diagnostic accuracy, and decrease physician workload, they also face many obstacles [16]. Major roadblocks to clinical adoption include worries about data quality, model interpretability, lack of healthcare system standardization, and regulatory difficulties. However, deep learning and machine learning integrated into sepsis detection systems is a huge leap forward for predictive medicine [17][18]. These approaches have the potential to revolutionize patient monitoring by providing early, personalized, and data-driven insights that can guide timely interventions, ultimately saving lives and reducing the burden on healthcare systems. Early patient observations could help achieve this goal.

A model for sepsis detection based on MBiGRU-CNN-CZOA is presented here:

- ❖ It is a framework for hybrid deep learning. In order to effectively capture both temporal and spatial features from clinical time-series data, proposed model incorporates a Modified Bidirectional Gated Recurrent Unit (MBiGRU) and a Convolutional Neural Network (CNN). This results in a significant improvement in accuracy and robustness of sepsis prediction.
- ❖ process of optimising features Utilising a modified version of Chaotic Zebra Optimisation Algorithm: incorporation of modified version of Chaotic Zebra Optimisation Algorithm improves performance of model by reducing features that are redundant and irrelevant. This allows model to retain only most important predictive attributes and ensures better generalisation.

- ❖ framework exhibits superior predictive performance, surpassing conventional models such as CNN, LSTM, and GRU. It achieves high performance with an accuracy of 93.82%, precision of 94.24%, recall of 93.82%, and F1-score of 93.68%, ultimately establishing itself as a dependable solution for early sepsis detection.
- ❖ A Real-Time Intensive Care Unit Deployment Utilising Edge Computing: Utilising edge computing, model was developed for real-time implementation in intensive care units (ICUs). Its purpose is to guarantee low-latency predictions, computational efficiency, and rapid clinical decision-making at point of care.

Following is structure of remaining parts of paper: In second section, related works are discussed; in third section, proposed methodology is presented in detail; in fourth section, results are analyzed; and in fifth section, conclusion is presented.

II. RELATED WORKS

To overcome these obstacles, Abualigah et al. [19] laid out a new AI-based architecture that combines GBMs and DNNs. MIMIC-IV database, which includes clinical data of critically ill patients, and UK Biobank, a collection of genetic, clinical, and lifestyle data from 500,000 participants, were used to evaluate framework. Networks are some of more established models. suggested framework outperformed them noticeably. So, compared to Neural Networks' performance of 0.92 on UK Biobank dataset, model's AUROC of 0.96 is noticeably better. Also, training on MIMIC-IV took just 32.4 seconds, hence framework was clearly effective. Because of its short prediction latency, it was also well-suited for use in real-time processing applications. With its suggested AI-based framework, important problems in translational medicine are effectively addressed while also providing higher predicted accuracy and efficiency. Its reliability across many datasets suggests it could be useful in real-time clinical decision support systems, which could facilitate adoption of personalized medicine and improve health outcomes for patients. main

goal of future research will be to make results more scalable and easier to comprehend so they may be used in a wider range of therapeutic contexts.

Musanga et al. [20] offer a modern hybrid AI model that combines best features of both methods. Some of these benefits include Symbolic AI's interpretability and logical reasoning and Deep Learning's automated feature extraction and categorization capabilities. approach relies on attention-based encoder, which improves feature saliency by zeroing in on important areas in CT images. An additional crucial part of model is adaptive deformable module, which enhances spatial feature extraction by taking into account differences in lung anatomy. Through experimental validation utilizing performance metrics like F1-score, accuracy, precision, and recall, model achieves near-perfect accuracy (99.16%) and an F1-score of 0.9916. This shows that compared to baseline setups, model is far better. In addition to achieving state-of-the-art diagnostic performance, this hybrid AI system ensures interpretability through its symbolic reasoning layer, making it easier to use in healthcare settings. These results show how important it could be to combine symbolic approaches with advanced machine learning methods to create AI systems that are transparent and resilient enough to handle important medical tasks.

A prediction model utilizing clinical data from first twenty-four hours after an individual is admitted to intensive care unit was built by Shi et al. [21] using machine learning. This model aims to enable rapid screening and early management for sepsis patients. electronic medical records of patients with sepsis were examined using machine learning techniques in this retrospective cohort study that was carried out across various sites. To evaluated models in American and Chinese healthcare environments for their ability to forecast sepsis outcomes within first twenty-four hours of ICU admission. Using a battery of 31 clinical parameters, machine learning models outperformed more conventional methods in predicting sepsis outcomes. Contrasted with linear regression's poor test scores of 0.25, machine learning algorithms produced AUCs greater than 0.8 and scores of 0.78. These models' consistent high performance in external validation (scores between 0.63 and

0.77) should not be overlooked. Clinical decision-making could be greatly improved with use of machine learning-based sepsis prediction models, which could lead to better patient outcomes through earlier diagnosis and faster treatment in critical first twenty-four hours after ICU admission.

In order to anticipate sepsis using just one time-point and non-invasive vital indicators, Yang et al. [22] created an early warning system. Also, researchers created this method to see how it relates to other biomarkers, like C-Reactive Protein (CRP) and Procalcitonin (PCT). The goal of developing and validating a machine learning algorithm for sepsis prediction is to build on top of XGBoost. Physio net and four Taiwanese medical centers provided retrospective data used to construct this algorithm. data included 46,184 patients treated in ICU. Furthermore, non-invasive vital indicators gathered at a particular time point were intended to be used in model's development. It was found out whether CRP and PCT levels were correlated with sepsis AI prediction model. constructed model showed consistent performance across several datasets, with an average recall of 0.908 and precision of 0.577. Another dataset from Cheng-Hsin General Hospital confirmed model's performance (recall: 0.986, precision: 0.585). Temperature, systolic blood pressure, and respiratory rate were most influential parameters in model's forecast. model's sepsis predictions were significantly correlated with high C-reactive protein levels, whereas trend in PCT was less stable. In order to improve care in hospitals and other critical care settings, our solution integrates artificial intelligence algorithms with vital sign data and its clinical relevance to CRP level. This results in a more accurate and timely detection of sepsis.

Thibou et al. [23] created a machine learning algorithm that can identify beginning of sepsis in any hospital department. A total of 45,127 patients from all departments at France's Valenciennes Hospital were included in a retrospective collection of sepsis predictors for training purposes. To build binary classifier SEPSI Score for sepsis prediction, a gradient boosted tree technique was used. Next, this score was assessed using study dataset, which included 5270 patient stays. Out of the total, 121 cases of

sepsis, or 2.3%, were identified. To concluded by comparing the model's performance to that of previous sepsis scoring systems and assessing its ability to detect the early stages of sepsis. On one hand, to have the SEPSI Score with an average positive predictive value of 0.610; on the other hand, to have the SOFA score with an average positive predictive value of 0.174. The SEPSI Score had an average area under the precision-recall curve of 0.738 while the most efficient score (SOFA) had an average area under the curve of 0.174. Both the sensitivity (0.845) and specificity (0.987) of the SEPSI Score were found to be quite high. There was no score that could match the model's accuracy up to three hours prior to sepsis start. The program correctly predicted the start of sepsis in 50% of instances at least 48 hours before medical experts verified it. The SEPSI Score model outperformed competing scoring systems and demonstrated accurate prediction of sepsis onset in the early stages. As a predictive tool, it could help in the early diagnosis and treatment of sepsis in all hospital departments. There has to be further research into its effects on the linked morbidity and death.

A. Problem Statement

Despite the need of early detection for efficient treatment, Prompt action is key in preventing sepsis. There is a tendency for traditional diagnostic tools, such as severity scoring systems, to rely on static thresholds and limited clinical features, which can result in delayed or inaccurate diagnoses. Despite the fact that machine learning and deep learning models have emerged as potentially useful alternatives, there are still a number of significant obstacles to overcome. The temporal dynamics of physiological signals, which are essential for recognizing the progression towards sepsis, are difficult to capture in many of the models that are currently in use.

There are some methods that make use of intricate architectures that, despite their precision, are not ideally suited for environments with limited resources because of the high computational demands the methods require. These individuals rely on data from a single time point or only a subset of vital signs, which can result in a decrease in the predictive precision and an increase in the number of false

positives. In addition, certain systems exhibit satisfactory performance on particular datasets, but they are unable to generalise their findings across a wide range of clinical settings, which restricts their applicability in situations that occur in the real world. When it comes to many models, feature selection is either done manually or based on conventional optimization techniques. This can result in the retention of redundant information and a reduction in the efficiency of the model.

Due to these limitations, there is an urgent requirement for a sepsis prediction model that is not only lightweight but also accurate and interpretable. This model should be able to effectively process time-series data, function in real time, and be deployed across a variety of intensive care unit environments in order to facilitate timely medical intervention.

III. PROPOSED METHODOLOGY

To address the existing challenges in early and accurate sepsis detection, a novel deep learning-based framework is projected that combines the strengths of temporal sequence modeling, spatial feature extraction, and intelligent feature optimization. The architecture is designed to function efficiently in real-time ICU environments, providing timely and precise predictions to support clinical decision-making and it is visually shown in Figure 1.

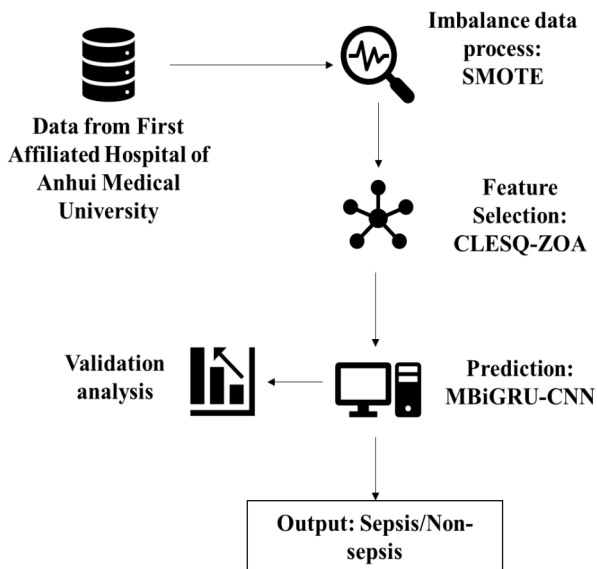


Figure 1: Workflow of the Research Model

A Modified Bidirectional Gated Recurrent Unit (MBiGRU) is incorporated into the core of the model. This unit is capable of effectively capturing both long-term and short-term temporal dependencies in physiological time-series data. This ensures that a comprehensive understanding of the progression of the patient's state is achieved. A CNN layer is incorporated into the architecture in order to improve the learning of spatial features from clinical data sequences. This combination makes it possible for the model to acquire knowledge of both temporal transitions and intricate data patterns, both of which are essential for accurate classification. An altered version of the Chaotic Zebra Optimization Algorithm is utilized for the purpose of feature selection in order to further optimise the performance of the model and reduce the amount of computational overhead. In order to improve the accuracy and generalization of the model, this metaheuristic algorithm intelligently prunes features that are redundant and irrelevant while maintaining characteristics that are critical predictors. The framework is trained and evaluated with the help of sepsis datasets that are freely accessible to the public. It demonstrates superior performance in key metrics such as accuracy, precision, recall, and F1-score when compared to conventional models such as CNN, LSTM, and GRU. Additionally, the system is designed to be deployed at the edge of the network, which enables real-time execution and rapid inference within intensive care unit configurations. This integrated approach not only improves the accuracy of predictive outcomes, but it also enables deployment that is both scalable and interpretable in a variety of clinical settings.

A. Data source and study population

A total of 2,385 patients from Anhui Medical University's First Affiliated Hospital and affiliated hospitals participated in this study [24]:

- The basic info;
- The life support employed;
- The outcome of blood test;
- The infection besides the use of antibiotics;
- The immunomodulatory nutritional support;
- The use of the analgesia/sedation.

To go over all the specifics, including demographics, lab results, illness scores, and basic vital signs. Columns in the system were labelled "sepsis information collection" and "non-sepsis information collection" according to the type of selection made. The two primary criteria for the official diagnosis of sepsis were laid out by the Third International Consensus Definition of Sepsis-3 (Sepsis): First, there has to be a suspicion of infection, which can be confirmed by prescribing antibiotics and collecting bodily fluids for microbiological cultures. Second, after the infection has been confirmed, there must be organ dysfunction, as shown by a SOFA score increase of 2 points or higher. The model analysis is now complete with data from 1968 patients, including 310 patients with sepsis. Johnson et al. (2016) and Pollard et al. (2018) used the MIMIC-III and eICU databases for external validation. Information regarding the characteristics of the intensive care unit in the twenty-four hours preceding hospitalization was gathered for both internal and external validation. The MIMIC-III dataset, which includes 7,230 cases of data, was used as an external validation source for a single centre. The dataset encompasses almost 60,000 hospital admissions from 2001 to 2012 (inclusive). eICU, on the other hand, offered external validation across many centers;

- 1) The test set comprised data from 11,900 patients and covered almost 139,000 hospital admissions (inclusive) from 2014 to 2015. The following were the inclusion criteria:
- 2) First, patients had to be at least 18 years old. Second, they had to have been in the intensive care unit (ICU) for at least 24 hours and have enough data.
- 3) Third, patients had to have a SOFA score of 2 contamination diagnosis of sepsis in either the MIMIC-III or eICU datasets, according to the Third International Consensus Definition of Sepsis (Sepsis-3). Database access and data extraction were managed by a single author (ZLY, ID: 11706576).

B. Process of Imbalance Data using SMOTE

Oversampling a dataset that is used in a typical classification problem (using a accomplished through a variety of different methods. One of the most widely used

methods is referred to as SMOTE, which stands for Synthetic Minority Over-sampling Technique [25]. In order to demonstrate how this method operates, let's take a look at some training data that contains s space of the data. Take note that, for the sake of simplicity, these characteristics are continuous. Taking a sample from the dataset and taking into account is the next step in the process of oversampling. In order to generate a synthetic data point, you must first obtain the vector that connects the current data point to one of the k existing neighbours. Using a random number x that falls somewhere between 0 and 1, multiply this vector by that number. In order to generate the new, synthetic data point, you must first add this to the existing data point.

C. Feature selection CLESQ-ZOA

As a crucial pre-development step for Machine Learning (ML) models, feature selection helps boost ML speed, accuracy, and performance. Excluded from the analysis are anomalous data such as mean arterial pressure and procalcitonin, as well as irrelevant features such as patient ID, admission that are not easily codable. Based on the optimizer described in this work, the remaining features are handled.

This section gives a comprehensive explanation of the proposed CLESQ-ZOA [26] by combining multiple strategies to overcome the drawbacks of the original ZOA algorithm. Like many other metaheuristics, ZOA suffers from drawbacks such as falling into local minima, loss of population diversity, lack of exploration capability, and imbalance between exploration and exploitation phases. Three fundamental improvements have been integrated into the algorithm to improve ZOA's performance and overcome these drawbacks. The improvement of the diversity of the initial population, the logarithmic spiral strategy, and the enhanced solution quality (ESQ) strategy are all clarified comprehensively in the ensuing subsections. In CLESQ-ZOA, chaotic mapping is used to generate the initial population. With chaotic mapping, population diversity is increased, which improves the algorithm's discovery ability. In the next stage, LSS is integrated into the algorithm to improve the search ability of ZOA in the early stages. ESQ is incorporated into ZOA to boost the quality of solutions and

overcome the drawback of imbalance between exploration and exploitation.

1) Population Initialization with Chaotic Mapping

As in most optimization algorithms, ZOA generates the initial population in the search space at random. Random initialization can result in a reduction of population diversity, and it does not ensure a uniform distribution of the population across the search space. A population with a uniform distribution helps broaden the search spaces, which enhances the algorithm's convergence speed. It is crucial to improve techniques for population initialization since the optimization algorithm's effectiveness is affected by the quality of the initial individual. Recently, the use of chaotic mapping, which has stochastic and ergodic properties, in various optimization algorithms has gathered much attention. By using chaotic mapping, which produces chaotic sequences as the initial population, loss of population diversity is prevented, and the distribution is made more uniform, thus increasing the search efficiency of the algorithm. Owing to the advantages of chaotic mapping, in this paper, chaotic sequences are utilized to generate the initial population of CLESQ-ZOA. Chaotic sequences are commonly mapped using a variety of chaotic models, such as the Chebyshev map, the logistic map, the circle map, the sine map, and the cubic map. It has been claimed that the population distribution generated by the cubic map indicates better uniformity compared to other chaotic mapping methods. Therefore, CLESQ-ZOA uses the cubic map to enhance the initial population's quality. The mathematical formula is expressed in (1) and (2).

$$X_{i,j} = lb_j + (ub_j - lb_j) \times \frac{(c_{i,j} + 1)}{2} \quad (1)$$

$$c_{i+1,j} = 4c_{i,j}^3 - 3c_{i,j}, -1 < c_{i,j} < 1, c_i \neq 0, i = 0, 1, \dots, N \quad (2)$$

For an optimization problem with a number of solutions N , the search space's lower and upper bounds are indicated by lb_j and ub_j respectively. In the population of dimension j , $X_{i,j}$ represents the i th solution. The initial step in solving a d -dimensional optimization issue is to vector, where each dimension has a value from the range $[-1, 1]$. The remaining $N-1$ individuals are then obtained by applying (8) to repeat

over each dimension of the initial individual. Finally, (2) is applied to create a mapping of the values of the operators generated through the cubic mapping.

Initializing the initial population in ZOA with a cubic chaotic map improves the CLESQ-ZOA's capacity to explore the solution space effectively during the early stages of optimization. The cubic chaos map generates initial populations that cover a diverse range of solutions, helping to avoid premature convergence to local optima.

2) Logarithmic Spiral Strategy (LSS)

The metaheuristic algorithm's strong search capability in the early stages plays a critical role in its superior performance. In the original ZOA, the search agents follow a straight line to the best solution with the best exploitation potential, updating their positions at each iteration. The straight-line search approach reduces the ability to discover solutions of better quality in the search space, reduces population diversity, and falls more easily into a local optimum. In this paper, LSS is integrated into ZOA to overcome these problems. The mathematical model of the LSS is described by (3).

$$X_{i,j}^{t+1} = PZ_j + e^{a\theta} (\cos(2\pi\theta)) |PZ_j - X_{i,j}^t| \quad (3)$$

where PZ_j indicates the best member called the pioneer zebra in the j th dimension. a represents a fixed value that identifies the form of the logarithmic spiral. θ is a linearly decreasing parameter from 1 to -1 calculated using (4).

$$\theta = 2(1 - t/T_{max}) - 1 \quad (4)$$

where t and T_{max} represent the current iteration and maximum number of iterations, respectively.

By integrating the logarithmic spiral strategy into CLESQ-ZOA, each individual is enabled to explore the search space an effective and controlled manner. Better positioning of solutions with LSS provides higher search efficiency and therefore improved performance for finding optimum and near-optimum solutions.

3) Enhanced Solution Quality (ESQ)

One of the best strategies for increasing search efficiency in optimization algorithms is the ESQ strategy. ESQ is

utilized to improve solution quality in the RUN optimization algorithm, which is grounded in the Runge-Kutta mathematical model. In ZOA, relying on the exiting optimal solution for population updates at each iteration can easily lead to the algorithm getting stuck in a local optimum. This problem is solved by adding the ESQ strategy to ZOA, which improves the solution quality and guarantees that each one advances to a better position before the start of the subsequent iteration. Additionally, ESQ strikes the balance between exploration and exploitation in ZOA by enabling the algorithm to both search for potentially better solutions and improve existing solutions. (5) outlines the mathematical formulation of the EQS strategy.

$$\begin{aligned}
 & \text{if } rand < 0.5 \\
 & \quad \text{if } w < 1 \\
 & \quad \quad X_{new2} = X_{new1} + r * \omega * |(X_{new1} - X_{avg}) + randn| \\
 & \quad \text{else} \\
 & \quad \quad X_{new2} = (X_{new1} - X_{avg}) + r * \omega * |(X_{new1} - X_{avg}) + randn| \\
 & \quad \text{end if} \\
 & \text{end if}
 \end{aligned} \tag{5}$$

where X_{new2} is the new solution produced by ESQ. The variable r is an integer whose value is equal to 1, 0, or -1. As expressed in (6), w represents a random number that declines as the number of iterations rises. The average of three random solutions in the population is represented by X_{avg} and its value is calculated by (7). X_{new2} calculated with (8) shows the new solution produced by using X_{avg} and X_{best} which is the current best value.

$$\omega = \exp(-(5 * rand)(t/T_{max}) * rand(0,2)) \tag{6}$$

where the current and the maximum number of iterations are denoted by t and T_{max} respectively.

$$X_{avg} = \frac{X_{r1} + X_{r2} + X_{r3}}{3} \tag{7}$$

where X_{r1} , X_{r2} and X_{r3} represent three random solutions.

$$X_{new1} = (1 - \beta) * X_{best} + \beta * X_{avg} \tag{8}$$

where β is a randomly produced value within the range of [0,1]. It's possible that X_{new2} fitness value isn't higher than the existing solution's. In this case, the algorithm is provided with another opportunity to generate a new better solution

X_{new3} as defined in (9).

$$\begin{aligned}
 & \text{if } rand < \omega \\
 & \quad X_{new3} = (X_{new2} - rand * X_{new2}) + SF * (rand * X_{RK} + ((2 * rand) * X_b - X_{new2})) \\
 & \quad \quad SF = (1 - 2 * rand) * v_1 * \exp(-v_2 * rand * (t/T_{max}))
 \end{aligned} \tag{9}$$

where SF is an adaptive factor which is a random value uniformly distributed within the interval [0,1]. v_1 and v_2 represent constant numbers whose values are 12 and 20, respectively. X_b indicates the initial best position while X_{RK} represents a solution derived from the Runge Kutta method.

The exploration phase is improved by employing the ESQ approach in CLESQ-ZOA, which allows the algorithm to explore new search regions. Additionally, the balance between exploration and exploitation in CLESQ-ZOA reduces the likelihood of getting trapped in local optima while searching for global solutions.

4) Architecture of CLESQ-ZOA

Firstly, the initial population in CLESQ-ZOA is initialized with cubic chaotic mapping to reduce the randomness in the initial population, which contributes to the search efficiency of the algorithm. Then, LSS is added to the algorithm to update the search agents' positions. This strategy enables CLESQ-ZOA to discover higher-quality solutions in the search space, thus reducing the probability of falling into local optima. The ZOA phases continue as normal after this approach is used. Next, the Enhanced Solution Quality is applied to ZOA, aiming to achieve significant improvements in the exploration stage and increase solution quality.

5) Computational Complexity of CLESQ-ZOA

In this section, the computational complexity of the projected CLESQ-ZOA is examined. Assume that the maximum number of iterations is T , the number of populations is n , and the dimensionality size of the optimization problems is d . The computational complexity of CLESQ-ZOA to initialize the population with the chaotic map and evaluate the fitness function is $O(n*d)$. Each member of the population is updated in two stages according to the fitness function value. The computational complexity of this update is $O(2 * n * T * d)$. Therefore, the overall

complexity of CLESQ-ZOA is. As a result, the complexity of the proposed CLESQ-ZAO in this work is the same $O(n * d * (1+2*T))$ as the original ZOA.

Finally, the proposed feature selection technique is selecting the 20 features. In that maximum number of times selecting the same features such as, HR, O2Sat, Temp, MAP, Resp, SBP, DBP, Resp_sys, Creatinine, Platelets, Bilirubin total, Age and Sepsis Label max.

D. Proposed classification

1) Convolutional neural network

The neural network (CNN) was initially utilized in the field of computer vision. Subsequently, it was introduced into the text field by modifying the width besides height of the filter [27]. This particular implementation was commonly referred to as text CNN. The five layers that are listed below are the primary components that make up the convolutional neural networks. One-hot encoding vectors or word vectors, such as Glove, were utilized at the input layer in order to map each vector with dimensions ranging from fifty to three hundred, and continuous tokens were used to organize the text matrix. This particular convolutional layer had a filter that was not square in shape, and its width was set to be the same as the current word vector. Additionally, the height of the filter could be customized. The word was used as the smallest granularity, and multiple different filters were slid on the text matrix at the same time in order to extract the information that was local to the text. At the pooling layers, max-pooling was utilized to extract and keep the most utilized to average all of the features; consequently, it represented the overall level of text features. Following the fully connected layer, the features that were generated by the layer below it was combined into a single dimension for the purpose of subsequent classification. The category that corresponds to the highest probability is the one that represents the discrimination result at the output layer.

2) Dilated convolution

Computer vision and speech processing are two areas that routinely make use of dilated convolution methods. It is common practice to employ the three approaches that are detailed below in order to obtain the long-distance info of

the text. The first possibility is to increase the size of the filter, the second possibility is to increase the sum of layers of the network, and the third possibility is to increase the number of neural units; however, all of these approaches will result in an increase in the number of parameters. There are reports that it is possible to increase the size of disguised form. This is something that has been reported. The receptive field would be expanded in the dilated convolution filter if "holes" were inserted between the pillars of the filter. When a hole of size three was inserted into a filter, for instance, the filter's field of vision increased from three by three to five by five. Therefore, the incorporation of holes made it possible for dilated convolution to acquire information about the text that was located at a great distance.

The following is the formula that represented the calculation for the receptive field:

$$F = k + (d - 1) * (k - 1) \quad (10)$$

The size of the filter is denoted by the letter k, and the dilation rate, denoted by the letter d, is the number of intervals that exist between the blocks. When d equals one, the standard for convolution is dilated convolution. The gridding effects will take place when multiple filters use the same dilation rate which will result in the loss of vital local correlations and words, which will ultimately have an effect on the semantic sentences. By using multi-layer varying dilation rates, it is possible to effectively eliminate the gridding effect.

3) Depth separable convolution (DSC)

DSC splits the conventional convolution calculation procedure in half. Step one is the depth wise process, and it entails applying a single filter to each channel separately. In the second stage, known as the pointwise process, a filter of size 1*1 is used to combine the outputs of the first step before they are finally produced. Consequently, DSC shortens the convolution operation's running time and decreases the number of parameters without substantially affecting the model's performance.

The sizes of assumed to be $D_F * D_F * M$, and those of map to be $D_G * D_G * N$, and the size of the filter is $D_K * D_K$, where D_R and D_G represent width and maps, correspondingly; M

and N denote the number of input and output channels, and the following is the calculation for ordinary convolution:

$$C = D_k * D_k * M * N * D_F * D_F \quad (11)$$

The DSC is the total of the depth wise and pointwise processes' computations. Here is the computation of the convolution in each input channel when just one filter is used:

$$C_1 = D_k * D_k * M * D_F * D_F \quad (12)$$

Here is how to compute the computational convolution utilizing $1*1$ filters to generate new features by linearly combining the outputs:

$$C_2 = M * N * D_F * D_F \quad (13)$$

Then cost of DSC is sum of C_1 and C_2 .

$$C_D = C_1 + C_2 = D_k * D_k * M * D_F * D_F + M * N * D_F * D_F \quad (14)$$

To discovered a specific proportional discrepancy in the computations.

$$r = \frac{C_D}{C} = \frac{D_k * D_k * M * D_F * D_F + M * N * D_F * D_F}{D_k * D_k * M * N * D_F * D_F} + \frac{1}{N} + \frac{1}{D_k^2} \quad (15)$$

The DSC was 7/96 of filter size in this study was $4*4$. In theory, this finding proved that DSC may successfully decrease the number of limits and increase computing efficiency.

4) Self-attention

First and foremost, the self-attention apparatus is primarily concerned with the internal dependence of input. It has been reported that the output of the present time is likely to be affected by the words that are very close to it as well as words that are very far away. A variety of weight parameters are assigned to words in accordance with the degree of influence. This is done in order to model is able to pay attention to the relevant information regarding the text's most important words and sentences. In this particular investigation, the scaled dot-product attention model was utilized.

$$Attention(Q, K, V) = softmax\left(\frac{QK^T}{\sqrt{d_k}}\right)V \quad (16)$$

The input vector's dimension is denoted by d_k , while the corresponding query, key, besides value vectors make up the matrices Q , K , and V , respectively. As far as self-attention is concerned, the inputs to Q , K , and V are identical. To determined how similar each word was in the given statement. The dependence inside the sentence can be captured by observing the degree of similarity among the terms. The stronger the connection, the more apparent it will be.

5) Focal loss

Modelling classification usually involves changing the model's architecture and making use of function; nonetheless, this loss function applies the same logic to every sample. Since the model does not account for the sample size per class or the complexity of sample classification, it cannot learn features that can affect. Because of this, the model has a hard time learning feature. To address these problems, researchers in the area of dense object detection first developed the idea of focus loss. Here is the formula for focus loss:

$$focal\ loss = -(1 - p_t)^{\gamma} \log(p_t) \quad (17)$$

Where $\gamma = [0, 5]$ helps to differentiate between samples that are easy to identify and those that are not; it is a focusing parameter. The binary-class model explains the focal loss function.

$$p_t = \begin{cases} p & \text{if } y = 1 \\ 1 - p & \text{otherwise,} \end{cases} \quad (18)$$

Where $y \in \{0, 1\}$ is the sample's class label, and p is the estimated probability of belonging to that class. The focusing limit is a function that separates the function's easy-to-classify samples from its hard ones. The modulation factor is near 1 when the categorization is incorrect and p_t is modest.

$(1 - p_t)^{\gamma}$ is close to 1. When the p_t is quite big, and this regulates it is almost zero. The modulating factor's smooth and derivable procedure can alter the loss function's easy/hard sample proportion and increase the range of

samples with low loss. In an imbalanced sample, a larger number of samples always makes it easier for the model to learn their characteristics; these samples also make up most of the loss and control the gradient; and a smaller number of samples makes classification more difficult. In order to accomplish the impact of balancing sample categories, focal loss modifies the weight of classify tasters.

6) Our MBiGRU-CNN model

The model's main components are the pooling layer, focused loss, and dilated convolution layer.

- 1) The layer that receives input. As an input for the subsequent layer, the model makes use of particular features that were selected from the optimiser.
- 2) Layer 2 of the BiGRU. In addition to being utilised for the purpose of receiving contextual semantics, BiGRU is also utilised for the purpose of processing the long-distance emotional material comprised of the input data.
- 3) It is the attention layer. Through the utilisation of self-attention, it is possible to determine the degree of similarity between data over any distance without relying on any external information. The text places an emphasis on the key words that convey powerful feelings, and the similarity calculation results in an increase in the amount of global information contained within the text.
- 4) There is a convolution layer. For the purpose of extracting local features, a DSC algorithm with a filter size of four and a single layer is utilised. This algorithm's computation amount is lower than that of a standard convolution.
- 5) The fifth layer is a dilated convolution. The convolution method used is called DSC, and it is a superposition of three layers of dilated dilation rates. There is a step of one, the size of the filter is five, and the dilation rate of each layer is [1, 2, 3]. Not only does this parameter setting allow the coverage of the layer to reach 20, which not only satisfies the length prevents the influence of irrelevant ultra-long-distance text. From the perspective of the dilation rate of [1, 2, 3], the info of the phrase is taken by a relatively low dilation

rate, whereas the information of the sentence is captured by a relatively high dilation rate.

- 6) The following three benefits emerge from the utilisation of this model. To begin, there is an increase in the total number of filters. It is important to note that the dilated convolution, which has a filter convolution. This indicates that the model contains filters with sizes of 4 and 5, respectively. The second point is that the convolution layers, which have filter dilation significant amount of ground. Therefore, sentence-level information that is typically obtained only by complex networks in the traditional method can now be obtained through the simple supplementation to the BiGRU layer. This is a significant advancement in the field. A third point to consider is that it is possible to acquire multi-scale feature information by employing a single extraction, which requires fewer parameters and maximises efficiency.
- 7) The GAP layer: the pooling process is carried out with the help of the GAP layer; the layer and the dilated convolution layer are GAP layer. The GAP layer is responsible for concentrating on map. The feature map is responsible for activating a particular value for each sample class. After that, the GAP layer sends the vector to the SoftMax layer. It is not necessary to configure the parameters in the fully connected layer in order to use GAP. This reduces the number of parameters, which helps to prevent over-fitting. Additionally, GAP summarises spatial info, which helps to make model more precise.
- 8) The layer of type. Classifying the data that is being input is the responsibility of this layer. Due to the fact that the model makes use of focal function, modifies the quantitative difference among samples belonging to various categories, proportion of diverse texts function while the training process is taking place, the model is able to achieve a more effective classification effect.

7) 3.4.7. Parameter setting

Using the control variable method, to debugged the model parameters multiple times. Table 1 displays the final parameters.:

Table 1: Model limit surroundings.

Parameters	Values	Parameters	Values
Self-attention	64/64/192	Maxlen	500/500/140
Standard filters size	4	Batch size	128
Dilation filters size	5	Activation function	LeakyReLU
BiGRU components	32/32/96	Optimizer	Adam
Dropout	0.3/0.3/0.5	Learning amount	0.001
Dilated rates	[1, 2, 3]	Epochs	15
Filters	96	L2 Regularization	0.01

IV. RESULTS AND DISCUSSION

The proposed model performs binary classification, categorizing patients into two classes such as, Non-Sepsis (Healthy/At-Risk Patients) and Sepsis (Patients diagnosed with sepsis). The confusion matrix presented visualizes the performance of projected binary classification model, which categorizes patients into two classes: Non-Sepsis (Healthy/At-Risk Patients) and Sepsis (Patients diagnosed with sepsis) and it is shown in Figure 2.

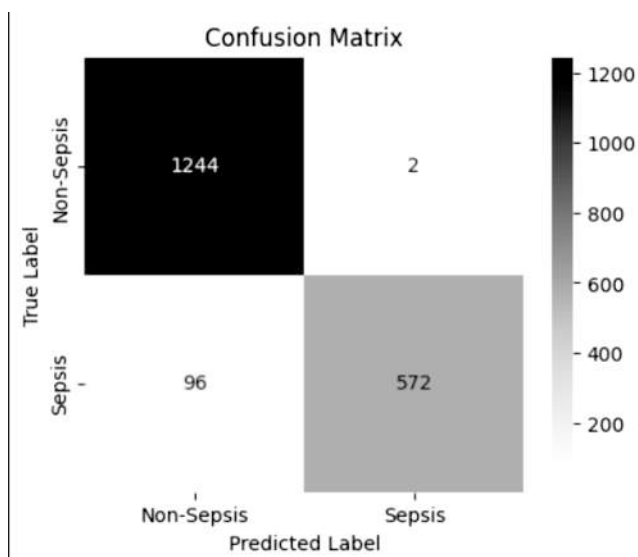


Figure 2. Confusion Matrix

Here's a thorough description of the matrix: TP: The perfect correctly predicted 572 patients as having sepsis, TN: The model correctly identified 1244 patients as non-septic, FP: Only 2 non-septic patients were incorrectly classified as septic and FN: 96 patients who actually had sepsis were wrongly classified as non-septic. This confusion matrix

indicates that the model performs exceptionally well in identifying non-septic patients (very high true negative rate), and also maintains a strong performance in correctly classifying septic cases. However, there are some false negatives, meaning that a few septic cases are missed by the model, which could be critical in medical diagnosis. Nevertheless, the overall balance of sensitivity (recall) and specificity appears to be well-optimized for this binary classification task.

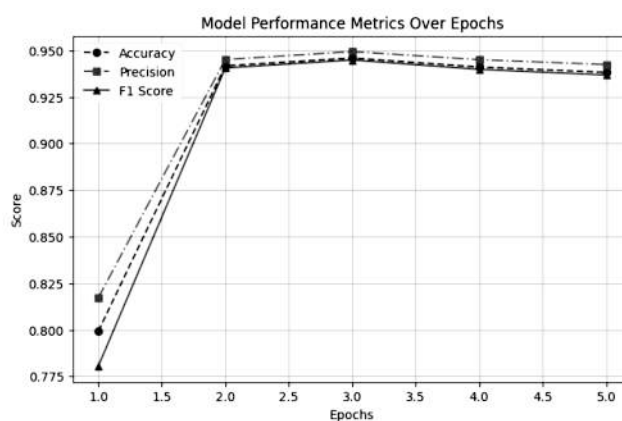


Figure 3: Proposed model over Epochs

The graph of Figure 3 illustrates the performance progression of the projected model across five training epochs using three key metrics: score. Initially, in the first epoch, the model shows moderate performance, with all metrics below 0.83. However, there is a significant improvement in the second epoch, where accuracy, precision, and F1 score sharply rise to approximately 0.94, indicating rapid learning and effective training. The metrics continue to improve slightly in the third epoch, peaking close to 0.95. From the fourth to fifth epoch, a slight decline is observed, suggesting early signs of overfitting. Overall, the graph demonstrates that the model achieves high and stable performance early in training, validating its efficiency and reliability. The two graphs illustrate the training progress of the proposed model over five epochs in terms of classification performance and model optimization and it is exposed in Figure 4.

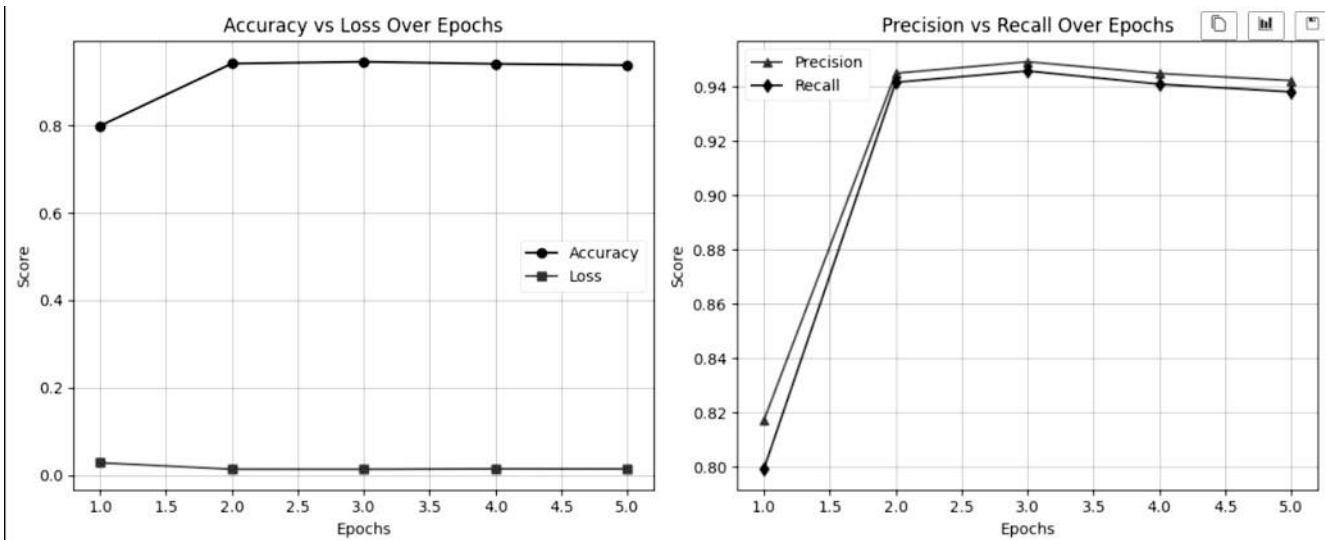


Figure 4: Investigation of proposed model a) Loss and Accuracy; b) Precision-Recall Curve

In the first graph (left), the Accuracy vs Loss Over Epochs plot shows a steady increase in accuracy from 0.80 in the first epoch to around 0.92 in the second, maintaining this level in the subsequent epochs. Concurrently, the loss decreases significantly from the first to the second epoch and remains consistently low, indicating that the model quickly converges and avoids overfitting. In the second graph (right), the Precision vs Recall Over Epochs plot reveals a sharp improvement between the first and second epochs, where both precision and recall exceed 0.94. From epochs three to five, the values remain high with slight fluctuations, suggesting a well-balanced performance between detecting true positives and minimizing false positives. Together, these graphs confirm that the model achieves stable and optimized learning early in training, ensuring high reliability and generalization for real-time sepsis prediction, where the achieved results of the proposed classical is shown in Figure 5.

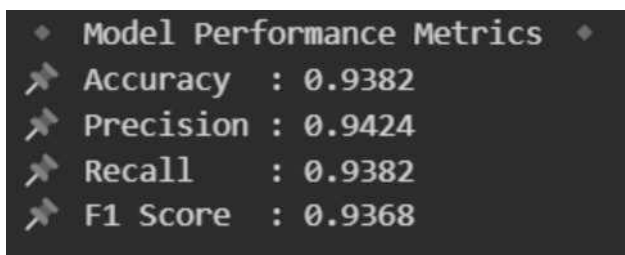


Figure 5: Proposed archived results.

4.1. Comparative analysis of proposed classical

Table 2 presents the comparative results of the projected model with existing baseline models in terms of different metrics besides it is visually publicized in Figure 6.

Table.2. Comparative analysis proposed with various existing baseline representations.

Model	Accuracy	Precision	Recall	F1 Score
Basic CNN	0.8810	0.8840	0.8760	0.8800
LSTM	0.9055	0.9125	0.9000	0.9062
GRU	0.9163	0.9208	0.9141	0.9174
Proposed MBiGRU-CNN (with CZOA)	0.9382	0.9424	0.9382	0.9368

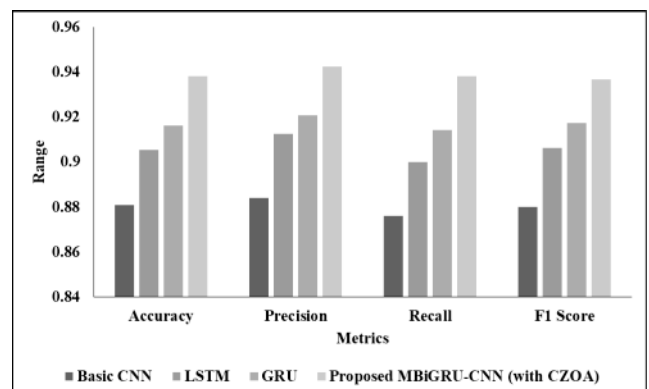


Figure 6: Visual Analysis of proposed model with existing models

Comparative study of the proposed MBiGRU-CNN model enhanced with CZOA (Chaotic Zebra Optimization Algorithm) against various existing baseline models: Basic CNN, LSTM, and GRU, using four performance Score. The proposed MBiGRU-CNN (with CZOA) achieves superior performance across all metrics, attaining the highest accuracy of 0.9382, precision of 0.9424, and a well-balanced recall and F1 score of 0.9382 and 0.9368, respectively. This indicates that the integration of bidirectional GRUs and CNN, combined with the optimization power of CZOA, effectively enhances the model's learning ability and generalization. In comparison, the GRU model performs slightly lower, with an accuracy of 0.9163 besides F1 score of 0.9174, indicating strong sequence learning but lacking the optimized hybrid features of the proposed system. The LSTM model follows with moderate results (accuracy: 0.9055, F1 score: 0.9062), while the Basic CNN model trails behind with the lowest scores across all metrics, particularly accuracy (0.8810) and F1 score (0.8800), reflecting its limited temporal modeling capability. This comparative evaluation clearly illustrates the effectiveness of the MBiGRU-CNN with CZOA, showcasing its robustness in capturing both spatial and sequential patterns while ensuring optimal parameter tuning through the use of metaheuristic optimization.

V. CONCLUSION AND FUTURE DIRECTION

This study introduces a highly effective and computationally efficient deep learning-based framework for the early discovery of sepsis, a condition that leftovers a critical test in intensive care units due to its rapid onset and high mortality rates. By integrating a Modified MBiGRU with a CNN, the proposed model successfully captures both temporal dependencies and spatial features from complex physiological time-series data. The addition of the CLESQ-ZOA further enhances model performance by optimizing feature selection, reducing redundancy, and maintaining crucial predictive attributes. The proposed MBiGRU-CNN-CLESQ-ZOA framework demonstrates superior classification accuracy of 93.82%, precision of 94.24%, recall of 93.82%, besides F1-score of 93.68%, clearly

outperforming existing baseline models such as CNN, LSTM, and GRU. Moreover, the system is designed for real-time application through edge computing, ensuring rapid prediction, low latency, and high adaptability in ICU settings. These outcomes underscore the model's significant potential to support timely clinical interventions, ultimately improving patient survival rates and reducing the burden on critical care resources. Looking forward, future research will explore the enhancement of model interpretability to foster greater trust and transparency in clinical environments. Additionally, the incorporation of federated learning mechanisms is planned to enable decentralized training across multiple healthcare institutions while preserving data privacy and patient confidentiality. This privacy-preserving approach will promote collaboration between institutions without the need for direct data sharing. Moreover, expanding the dataset to include a wider demographic and clinical variety will further validate the model's robustness and generalizability. Overall, the proposed system marks a promising step toward intelligent, data-driven sepsis management, with broad implications for real-time predictive healthcare applications.

REFERENCES

- [1] Verma, D. K., Singh, S., Dubey, S., & Raghuvanshi, K. (2024, May). Revolutionize Infectious Prevention Using Artificial Intelligence and Deep Learning. In International Conference on Advances in Computing and Data Sciences (pp. 334-345). Cham: Springer Nature Switzerland.
- [2] Zhang, G., Shao, F., Yuan, W., Wu, J., Qi, X., Gao, J., & Wang, T. (2024). Predicting sepsis in-hospital mortality with machine learning: a multi-center study using clinical and inflammatory biomarkers. *European Journal of Medical Research*, 29(1), 156.
- [3] Rahman, M. S., Islam, K. R., Prithula, J., Kumar, J., Mahmud, M., Alam, M. F., . & Chowdhury, M. E. (2024). Machine learning-based prognostic model for 30-day mortality prediction in Sepsis-3. *BMC medical informatics and decision making*, 24(1), 249.

- [4] O'Reilly, D., McGrath, J., & Martin-Loeches, I. (2024). Optimizing artificial intelligence in sepsis management: Opportunities in the present and looking closely to the future. *Journal of Intensive Medicine*, 4(1), 34-45.
- [5] Yeo, H. J., Noh, D., Kim, T. H., Jang, J. H., Lee, Y. S., Park, S., ... & Kang, H. K. (2024). Development and validation of a machine learning-based model for post-sepsis frailty. *ERJ Open Research*, 10(5).
- [6] Yong, L., & Zhenzhou, L. (2024). Deep learning-based prediction of in-hospital mortality for sepsis. *Scientific Reports*, 14(1), 372.
- [7] Bomrah, S., Uddin, M., Upadhyay, U., Komorowski, M., Priya, J., Dhar, E., ... & Syed-Abdul, S. (2024). A scoping review of machine learning for sepsis prediction-feature engineering strategies and model performance: a step towards explainability. *Critical Care*, 28(1), 180.
- [8] Patil, P., & Narawade, V. (2024). RESP dataset construction with multiclass classification in respiratory disease infection detection using machine learning approach. *International Journal of Information Technology*, 1-18.
- [9] Kallonen, A., Juutinen, M., Värri, A., Carrault, G., Pladys, P., & Beuchée, A. (2024). Early detection of late-onset neonatal sepsis from noninvasive biosignals using deep learning: A multicenter prospective development and validation study. *International Journal of Medical Informatics*, 184, 105366.
- [10] Gao, Y., Wang, C., Shen, J., Wang, Z., Liu, Y., & Chai, Y. (2024). Systematic review and network meta-analysis of machine learning algorithms in sepsis prediction. *Expert Systems with Applications*, 245, 122982.
- [11] Gao, J., Lu, Y., Ashrafi, N., Domingo, I., Alaei, K., & Pishgar, M. (2024). Prediction of Sepsis Mortality in ICU patients using machine learning methods. *BMC Medical Informatics and Decision Making*, 24(1), 228.
- [12] Pérez-Tome, J. C., Parrón-Carreño, T., Castaño-Fernández, A. B., Nievas-Soriano, B. J., & Castro-Luna, G. (2024). Sepsis mortality prediction with machine learning techniques. *Medicina Intensiva (English Edition)*, 48(10), 584-593.
- [13] Parmar, S., Shan, T., Lee, S., Kim, Y., & Kim, J. Y. (2024, February). Extending machine learning-based early sepsis detection to different demographics. In *2024 IEEE First International Conference on Artificial Intelligence for Medicine, Health and Care (AIMHC)* (pp. 70-71). IEEE.
- [14] Giordano, M., Dheman, K., & Magno, M. (2024). SepAI: Sepsis Alerts on Low Power Wearables with Digital Biomarkers and On-Device Tiny Machine Learning. *IEEE Sensors Journal*.
- [15] Boussina, A., Shashikumar, S. P., Malhotra, A., Owens, R. L., El-Kareh, R., Longhurst, C. A., ... & Wardi, G. (2024). Impact of a deep learning sepsis prediction model on quality of care and survival. *NPJ digital medicine*, 7(1), 14.
- [16] Yadgarov, M. Y., Landoni, G., Berikashvili, L. B., Polyakov, P. A., Kadantseva, K. K., Smirnova, A. V., ... & Likhvantsev, V. V. (2024). Early detection of sepsis using machine learning algorithms: a systematic review and network meta-analysis. *Frontiers in Medicine*, 11, 1491358.
- [17] Agnello, L., Vidali, M., Padoan, A., Lucis, R., Mancini, A., Guerranti, R., ... & Carobene, A. (2024). Machine learning algorithms in sepsis. *Clinica Chimica Acta*, 553, 117738.
- [18] Ambiha, A. H., S. S. K., Kokilamani, M. (2024). A novel development of medical technology and AI for intelligent healthcare. In *Smart healthcare and machine learning* (pp. 249-267). Springer.
- [19] Abualigah, L., Alomari, S. A., Almomani, M. H., Zitar, R. A., Saleem, K., Migdady, H., ... & Ezugwu, A. E. (2025). Artificial intelligence-driven translational medicine: a machine learning framework for predicting disease outcomes and optimizing patient-centric care. *Journal of Translational Medicine*, 23(1), 302.

- [20] Musanga, V., Viriri, S., & Chibaya, C. (2025). A Framework for Integrating Deep Learning and Symbolic AI Towards an Explainable Hybrid Model for the Detection of COVID-19 Using Computerized Tomography Scans. *Information*, 16(3), 208.
- [21] Shi, S., Zhang, L., Zhang, S., Shi, J., Hong, D., Wu, S., ... & Lin, W. (2025). Developing a rapid screening tool for high-risk ICU patients of sepsis: integrating electronic medical records with machine learning methods for mortality prediction in hospitalized patients—model establishment, internal and external validation, and visualization. *Journal of Translational Medicine*, 23(1), 97.
- [22] Yang, A. C., Ma, W. M., Chiang, D. H., Liao, Y. Z., Lai, H. Y., Lin, S. C., ... & Wang, C. Y. (2025). Early Prediction of Sepsis Using an XGBoost model with Single Time-Point Non-Invasive Vital Signs and Its Correlation with C-Reactive Protein and Procalcitonin: A Multi-Center Study. *Intelligence-Based Medicine*, 100242.
- [23] Thiboud, P. E., François, Q., Faure, C., Chaufferin, G., Arribe, B., & Ettahar, N. (2025). Development and Validation of a Machine Learning Model for Early Prediction of Sepsis Onset in Hospital Inpatients from All Departments. *Diagnostics*, 15(3), 302.
- [24] Zhou, L., Shao, M., Wang, C., & Wang, Y. (2024). An early sepsis prediction model utilizing machine learning and unbalanced data processing in a clinical context. *Preventive Medicine Reports*, 45, 102841.
- [25] Selvamuthukumar, N., Aravinda, K., Manjunatha, B., & Thirumalraj, A. (2025). Breast Cancer Detection Using Mother Optimisation Algorithm Based Chaotic Map with Private AI Model. In *Sustainable Development Using Private AI* (pp. 278-294). CRC Press.
- [26] Özbay, F. A. (2025). An Enhanced Zebra Optimization Algorithm with Multiple Strategies for Global Optimization and Feature Selection Problems: A Hepatocellular Carcinoma Case Study. *IEEE Access*.
- [27] Aruna, T. M., Kumar, P., Naresh, E., Divyaraj, G. N., Asha, K., Thirumalraj, A., ... & Yadav, A. (2024). Geospatial data for peer-to-peer communication among autonomous vehicles using optimized machine learning algorithm. *Scientific Reports*, 14(1), 20245.