

Article

# Detection of Influential nodes using Hybrid Deep learning methods in IIOT environment

Lakshmi Patil<sup>1</sup>, Ahmed A. Elngar<sup>2</sup>

<sup>1</sup> Department of Electronics & Communication Engineering, Sharnbasva University, Kalaburagi, India

<sup>2</sup> Faculty of Computers and Artificial Intelligence, Beni-Suef University, Kurnool, Beni-Suef City, 62511, Egypt

\*Corresponding author: Lakshmi Patil (email: patillakshmi192@gmail.com)

Received 10 September 2024; Accepted 19 October 2024; Published 15 December 2024

**Abstract:** Influential nodes in an Industrial Internet of Things (IIoT) environment using Beluga Whale Optimization Algorithm (BWO) integrated with Residual Long Short-Term Memory (LSTM) networks. In IIoT networks, identifying influential nodes is crucial for optimizing data transformation, minimizing latency, and improving overall network effectiveness. The proposed method leverages the exploration and exploitation capabilities of the Beluga Whale Optimization (BWO) Algorithm to optimize the parameters of an Recurrent Long Short-Term Memory (RLSTM) model, which is used to predict the behavior of nodes and identify key influencers within the network. The integration of BWO with RLSTM helps improve the accuracy of node predictions by dynamically adjusting the RLSTM's hyperparameters based on the network's evolving data. Extensive experiments conducted in a simulated IIoT environment highlight the performance of the proposed model in enhancing prediction accuracy, reducing computational overhead, and improving network efficiency compared to traditional methods. The results highlight the potential of this hybrid optimization technique for real-time applications in smart manufacturing, predictive maintenance, and other IIoT-driven sectors.

**Keywords:** Beluga Whale Optimization Algorithm, Long Short-Term Memory, IIoT-driven sectors

Copyright © 2024 Journal of Smart Internet of Things (JSIoT) published by Future Science for Digital Publishing and Sciencedo . This is an open access article license CC BY (<https://creativecommons.org/licenses/by/4.0/>)

## 1. Introduction

The identification of influential nodes has gained considerable focus in recent times owing to its importance in understanding the functional characteristics and practical applications of various networks [1]. By identifying these influential nodes, critical aspects such as transmission, control, security, vulnerability, and network resilience can be analyzed. Research in this area has proven to be essential in the development of emergency logistics networks, social network communication, transportation systems, biological virus containment, and power grid protection. For example, in logistics networks, selecting influential communication nodes ensures the efficient and timely

transportation of emergency supplies. In social networks, identifying key communicators helps accelerate information dissemination and control the spread of rumors. In infrastructure networks like urban transportation, railway, aviation, and communication systems, determining the critical nodes provides valuable insights for network management and optimization. In the context of biological virus control, pinpointing key transmission nodes can prevent the early spread of diseases, while in power networks, protecting crucial transmission lines enhances overall network robustness and invulnerability. Overall, the identification of influential nodes serves a pivotal role in securing the safety and efficiency of real-world networks [4].

Recent studies on influential node identification have introduced voting-based methods inspired by real-world voting systems in social networks. Compared to non-voting algorithms, voting algorithms are simpler to implement, and their outcomes efficiently reflect the importance of nodes. Additionally, the voting power of neighboring nodes decreases after voting, helping to balance the distribution of influential nodes and minimize the overlap in their influence zones. This mechanism ensures efficient information spread by selecting a limited number of influential nodes. To address the low discrimination and accuracy found in recent algorithms, a new method called the Adaptive Adjustment of Voting Ability (AAVA) is introduced. The AAVA algorithm enhances node influence identification by evaluating both the node's self-influence based on its local attributes and the voting contributions of neighboring nodes. It dynamically adjusts the voting ability among nodes based on their similarity, without requiring predefined parameters, allowing nodes to exert varying levels of influence on their neighbors. A toy network is used as an example to illustrate the core idea and voting process of the AAVA method, which operates in multiple rounds.

The recognition of influential nodes in networks has long been a challenging problem due to the complexity and large scale of modern networks. Traditional methods often struggle with accurately distinguishing the most influential nodes, especially in dynamic environments where node interactions and network conditions change rapidly. Existing techniques often face issues such as low discrimination power, inadequate handling of large networks, and poor adaptability to dynamic changes. These problems are particularly prominent in areas like social networks, emergency logistics, and transportation systems, where the effective identification of influential nodes is critical for efficient operation and decision-making.

To address these challenges, a hybrid technique combining the BWO and RLSTM models has been proposed. The Beluga Whale Optimization Algorithm, motivated by the intelligent foraging behavior of beluga whales, is a global optimization method known for its ability to find optimal solutions in complex, high-dimensional search spaces. It is used here to intelligently select a subset of influential nodes by evaluating their potential contributions to the overall network. By leveraging the BWO's optimization capabilities, the technique ensures that the most critical nodes are identified based on their connectivity and influence in the network.

On the other hand, the RLSTM, which reinforcement learning with LSTM (Long Short-Term Memory) networks, adds an adaptive, dynamic element to the identification process. RLSTM allows the model to learn from past network states and predict the evolution of node influence over time, taking into account the temporal dependencies in the data. This hybrid approach not only optimizes node selection using the BWO but also allows the system to continuously adjust its decisions based on real-time data and changing network conditions. By combining these two powerful techniques, the hybrid model addresses the existing limitations of traditional methods, improving the accuracy and efficiency of influential node identification in complex, dynamic networks.

### ***1.1 Contribution of the Research***

1. The research proposes the new methodology of integrating the Beluga whale optimization technique over the LSTM which can be utilised for better prediction.
2. The proposed algorithm is compared with the different existing artificial intelligence networks.

### ***1.2 Organisation of the Paper***

This paper is structured as follows: Section-2 explores the relevant studies by multiple authors. The primary views of LSTM, BWO are discussed in Section-3, And outlines the working principle of the suggested architecture. experimentations, along with an analytical discussion of the findings are discussed in Section-4. At last, the study concludes with the future development in Section -5.

## **2. Related works**

Camacho et al. (2020) [1] offer a thorough summary of social network analysis by introducing four key dimensions: research techniques, applications, and software tools. The paper explores the fundamental techniques and challenges associated with social network analysis, emphasizing how social media platforms, digital communications, and various online environments are increasingly studied through these lenses. They discuss the importance of understanding social connections in the context of influence, information dissemination, and network behavior. Their work also covers various computational methods and software tools that have emerged to facilitate the analysis of large-scale social networks, offering insights into the tools that are essential for researchers in the field of data science and computational social science.

Zareie et al. (2019) [2] introduce a new technique for recognizing influential users in social networks based on user interests. They address a common challenge in social network analysis: how to effectively determine influential nodes (users) within the network, not only based on user activity or interaction but by considering personal interests. Their approach highlights how user interests influence the spread of information and social dynamics, providing a more nuanced understanding of social influence. This methodology is especially relevant in the context of personalized content delivery and targeted marketing, where understanding user preferences can drastically improve the accuracy of influence prediction.

Al-Garadi et al. (2018) [3] conduct a survey on the recognizing of influential users within online social networks, identifying key research issues and open questions in the field. The paper reviews various techniques, including graph-based methods, centrality measures, and machine learning approaches, while focusing on challenges such as the complexity of real-world social networks and the dynamic nature of influence. They highlight the importance of recognizing influential users for a wide range of applications, like marketing, political influence, and network security. The paper also discusses the gaps in the current literature, pointing out the need for more sophisticated algorithms capable of adapting to changing network structures and real-time data.

Peng et al. (2016) [4] explore the opportunities and challenges of social influence analysis in the context of big data in social networks. With the expansion of social media and online interactions, they highlight the difficulty in analyzing vast amounts of data to discern meaningful patterns of influence. The paper discusses various methodologies for extracting and analyzing influence, focusing on the challenges posed by big data, including the scalability of models and the difficulty in handling noisy data. They argue for the application of sophisticated artificial intelligence methods to enhance the precision and effectiveness of influence detection, specifically in the context of large, complex networks.

Mohammadi and Saraee (2018) [6] investigate the recognizing of influential users in social

networks, particularly focusing on different time bounds using multi-objective optimization. Their study emphasizes the need for models that can account for the temporal aspect of social influence, where the impact of users changes over time. By employing a multi-objective optimization approach, they offer a more comprehensive model that simultaneously considers multiple factors, such as user activity and the evolution of influence over time. The research contributes to the development of more adaptive and dynamic methods for identifying influential users, which can be applied in areas like targeted advertising and real-time information propagation.

Saoud and Moussaoui (2016) [7] introduce a community identification algorithm relying on minimum spanning trees and modularity, which is critical for understanding the structural components of social networks. This paper focuses on the identification of subgroups or communities within large networks, which is an essential step in analyzing social networks. The proposed method enhances traditional community detection techniques by incorporating both the graph structure and the concept of modularity, offering improved performance in identifying tightly-knit communities. These insights are vital for applications such as information dissemination, marketing strategies, and understanding social dynamics within various online environments.

Zhao et al. (2016) [8] develop a technique for recognizing influential nodes in social networks that incorporates community structure through label propagation. This approach leverages the inherent structure of social networks, where nodes tend to form tightly-knit communities. The authors argue that by considering the community structure, their model improves the identification of influential nodes, as influence often spreads more effectively within communities. Their work highlights the significance of community-aware influence analysis, which can be applied in scenarios such as viral marketing or spreading information during emergencies.

Vathi et al. (2017) [9] explore the use of Twitter communities for mining and categorizing interesting topics, which is essential for understanding how influence operates within specific interest groups. This paper highlights how social media platforms like Twitter can be leveraged to identify influential topics and key users within distinct communities. By applying machine learning techniques and analyzing user interactions within these communities, the authors present a method for categorizing and understanding how information flows across different groups. This research is crucial for applications in trend analysis, targeted marketing, and understanding social dynamics in digital spaces.

Moosavi et al. (2017)[10] focus on community identification in social networks utilising user frequent pattern mining, which can be particularly useful for identifying influential users within specific communities. By analyzing the frequent patterns of interactions between users, the authors propose a method that identifies users who are likely to be influential within their communities. This approach enhances traditional community detection techniques by incorporating patterns of behavior, making it more effective for understanding how influence spreads within a network. The study contributes valuable insights into the identification of key communicators, which can be applied in fields like online marketing and social network analysis.

### 3. Proposed Methodology

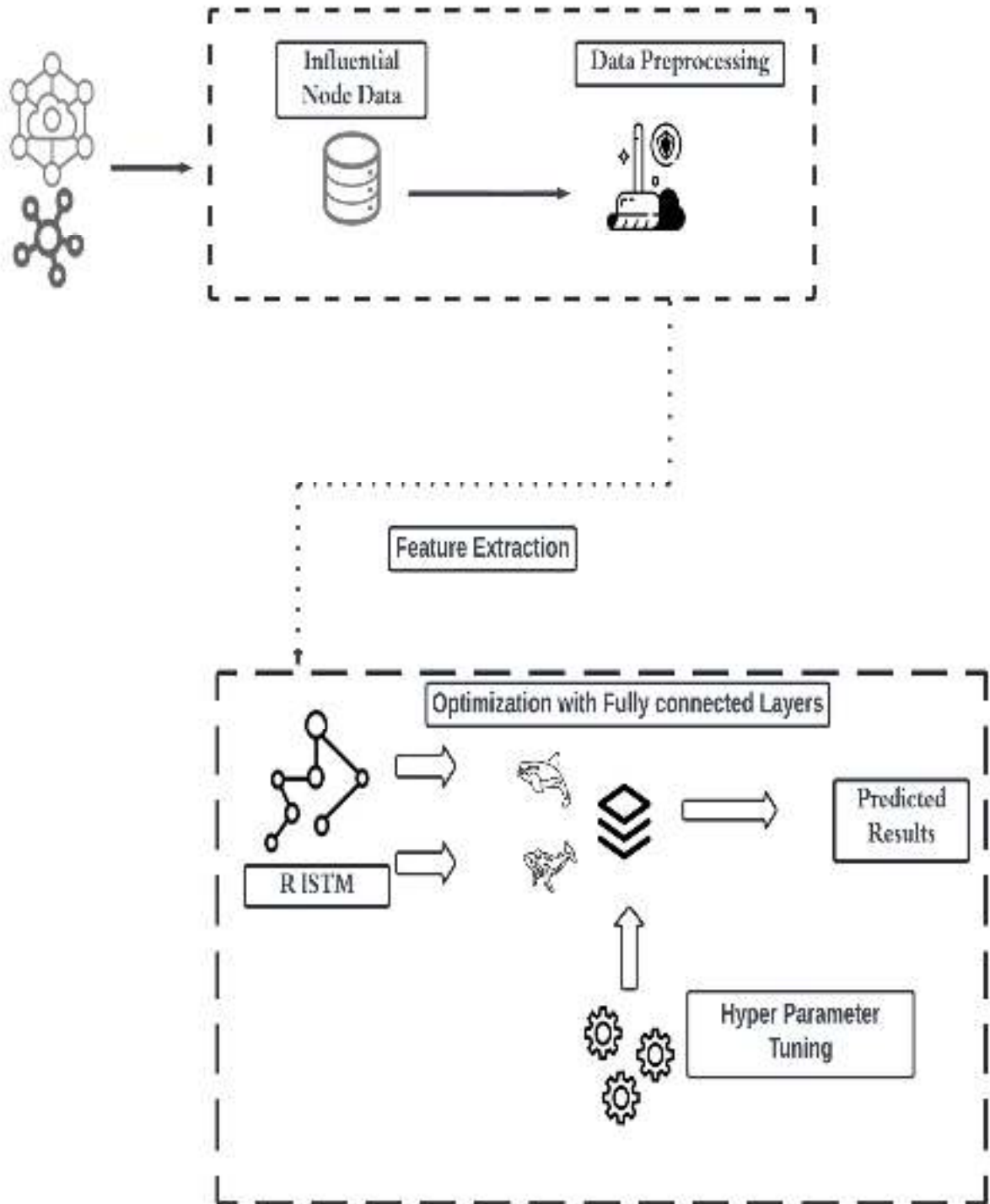


Figure 1: Proposed architecture

### 3.1 Materials and Method

In the context of identifying influential nodes within a network, the data typically consists of various network attributes and metrics that help evaluate the importance of each node. This data may include connectivity measures such as the count of direct neighbors (degree), the strength of connections (edge weights), and the centrality of each node, which reflects its position in the network in terms of influence or control. Other factors like the node's betweenness, closeness, or eigenvector centrality might also be analyzed to determine how well a node connects distinct parts of the network. In some cases, temporal data, such as the frequency or timing of interactions between nodes, can also be considered. This data is crucial for understanding the role of each node within the network and recognizing which nodes are most critical for maintaining the network's functionality, ensuring security, or facilitating rapid communication. The identified influential nodes are often used to optimize network operations, improve security measures, or manage resource allocation in various domains like logistics, communication, transportation, and biological or power systems.

### 3.2 Data preprocessing

It is an essential phase in preparing unprocessed data for examination or machine learning operations. It includes various techniques to clean and transform the data into a format which is appropriate for modeling. First, missing values can be handled by either imputing them utilising mean, median, or mode values, or by removing rows with missing data. Next, outliers can be detected and either removed or adjusted to prevent them from skewing the results. Data normalization or standardization is often applied to assure that the features are on a similar scale, which helps improve the performance of machine learning models. Categorical variables may be encoded utilising methods such as one-hot encoding or label encoding to convert them into numerical values. Additionally, features that are irrelevant or redundant can be dropped through feature selection techniques. Finally, data may be split into training and testing sets to assess the algorithm's effectiveness. By performing these preprocessing steps, the quality and accuracy of the analysis or predictive model are greatly improved.

### 3.3 Beluga whale optimization

In this section, we present the Beluga Whale Optimisation Mechanism (BWOM), a suggested method for efficiently choosing predictive features. BWOM is a meta-heuristic algorithm inspired by nature that replicates beluga whale social behaviours, cooperative hunting, and echolocation. The main Finding the most pertinent information to improve prediction accuracy while maintaining computational efficiency is the aim of this approach. BWOM, like other meta-heuristic algorithms, seeks to get an ideal result by striking a balance between exploitation (fine-tuning local solutions) and exploration (broadly seeking for global answers). This balance is frequently difficult for traditional optimisation techniques to achieve, which can result in early convergence or high computing needs.

BWOM tackles this problem by classifying search agents into roles according to the behaviour of beluga whales. Follower agents adaptively modify their routes in response to the leaders' signals, while leading agents employ echolocation signals to guide group movement. The adaptive exploration of the search space made possible by this echolocation-driven navigation and collaboration raises the possibility of arriving at global optima while maintaining stability

during the exploitation phase. By mimicking these By minimising computing overhead and avoiding frequent traps like becoming stuck in local optima, BWOM shows a strong capacity to efficiently converge towards optimal solutions in natural dynamics. Consequently, in feature selection and other optimisation tasks, BWOM delivers high computational efficiency and robustness.

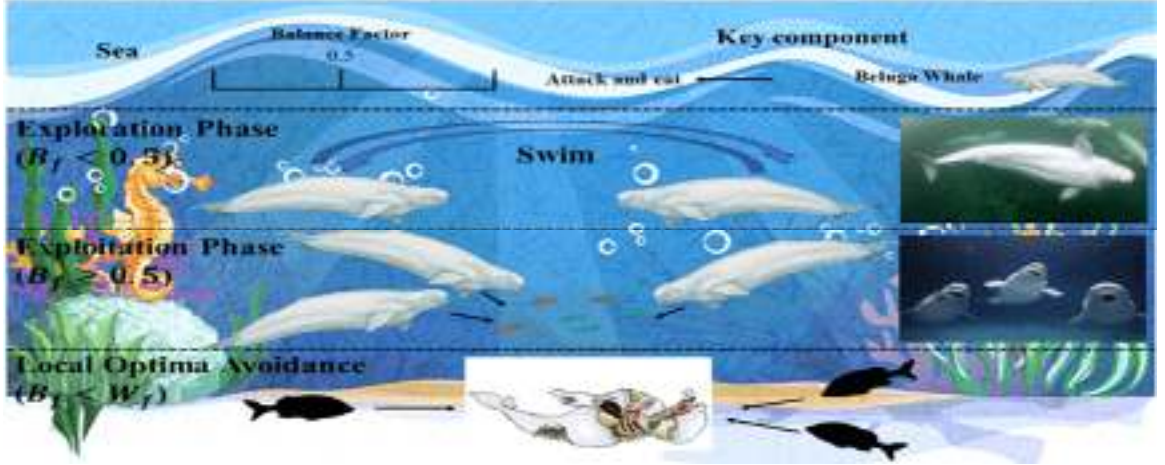


Figure 2: Beluga Whale optimization

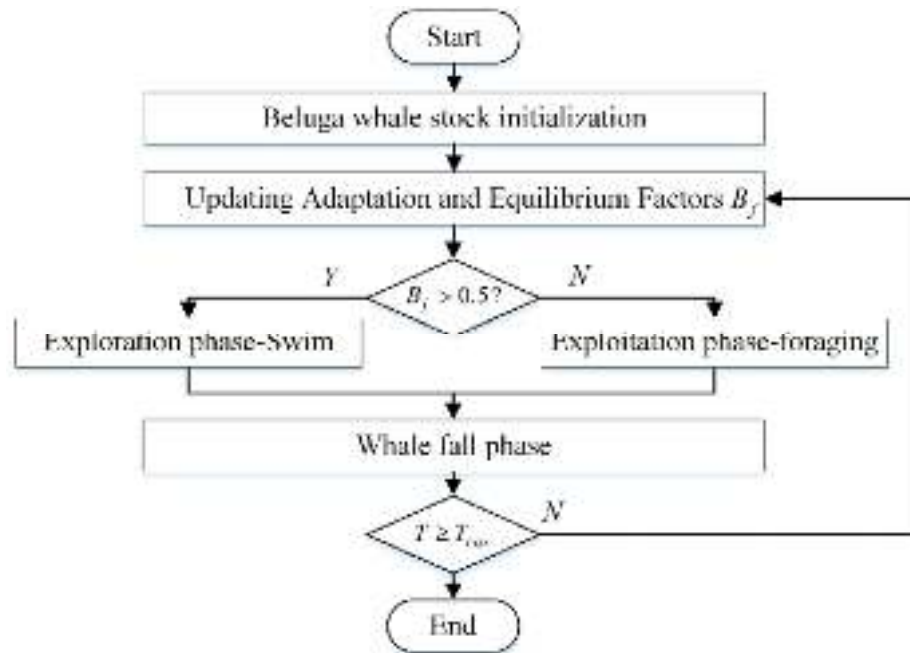


Figure 3: Flowchart of the optimization

### 3.3.1 : PROCESS MECHANISM

1) Initialization Phase: Equation (1) is utilized to initialize the various agents within the starting population of the BWO method.

$$X_{i,j} = \alpha(UB-LB) + LB, i \in \{1, \dots, N\}, j \in \{1, \dots, D\} \quad (1)$$

2) Exploration Phase: The positions of search agents are defined by the coupled swimming behavior, where two beluga whales swim collectively in a coordinated or reflective manner. This method enables search agents to investigate the search domain more effectively and productively,

resulting in the identification of novel and potentially superior solutions. The locations of the various agents are modified utilising

$$X_{i,j,t+1} = X_{i,p,t} + (X_{r,1,t} - X_{i,p,t})(1+r1)\sin(2\pi r2), j=2k \quad (2)$$

$$X_{i,j,t+1} = X_{i,p,t} + (X_{r,1,t} - X_{i,p,t})(1+r1)\cos(2\pi r2), j=2k+1$$

3) Exploitation Phase: The search participants exchange data about their present locations and evaluate the optimal solution, as well as other adjacent solutions, when adjusting their positions. This approach assists the participants in effectively navigating towards favorable areas of the search domain. This tactic aids in improving participants' convergence towards the universal solution of the search space. The locations of participants are adjusted utilising

$$X_{it+1} = r3X_{bestt} - r4X_{it} + 2r4(1-tT_{max}) \quad (3)$$

$$LF(X_{rt} - X_{it})LF = 0.05(v \times \sigma |v|1\beta)$$

$$\sigma = (\Gamma(1+\beta) \times \sin(\pi \times \beta/2) \Gamma(1+\beta/2) \times \beta \times 2^{\beta-12})1\beta$$

4) Balance between Exploration and Exploitation: Throughout the search procedure, the equilibrium between the exploration and exploitation stages is governed by a switching coefficient, referred to as  $B_f$ , which is computed using Eq. (4). Consequently, if the value of  $B_f$  is larger than or equal to a predefined threshold (e.g., 0.5), the search domain is investigated using Eq. (2); otherwise, the search domain is utilized using Eq. (3).

$$B_f = B_0(1-t2T_{max}) \quad (4)$$

5) Whale Fall: The whale fall phase denotes the idea of beluga whales' demise and transforming into a nourishment resource for other organisms. In the BWO technique, the whale fall stage functions as an arbitrary operator to infuse variety and avert the exploration procedure from becoming confined to local optima. The mathematical representation of this phase is conveyed through

$$X_{it+1} = r5X_{it} - r6X_{rt} + r7 \quad (5)$$

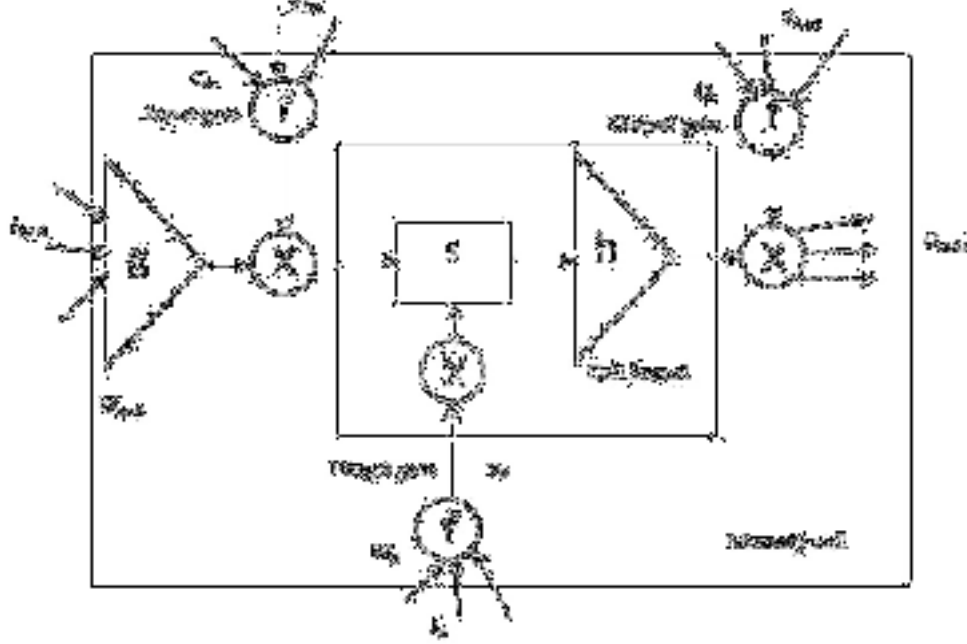
$$X_{step}X_{step} = (UB-LB)\exp(-C2tT_{max})$$

$$C2 = 2 \times W_f \times N(6)$$

$$W_f = 0.1 - 0.05tT_{max} \quad (6)$$

### 3.4 LSTM

It can handle memory in a flexible way and is appropriate for large datasets. The LSTM network is depicted in the image below.



LSTM networks and an BWO Optimiser are combined in the proposed hybrid learning model. The four distinct components that comprise LSTM networks are the input gate (I.G. ), forget gate (F.G. ), cell input (C.I. ), and output gate (O.G.). LSTM networks are essentially memory-based neural architectures that retain information during iterations. Let "ht" be the current hidden state, "ht-1" be the prior hidden state, and xt be the input layer output in this instance. "Ct" represents the current state of the cell, "Gt" the output state of the cell, and "Gt-1" the state of the cell prior to that. The gate states are represented by  $j_t$  [ , T ] \_f. The LSTM unit updates its memory by combining the previous unit's output with the current input state, utilizing the output and forget gates. The calculation of Gt and ht is performed using the following equations.

$$I.G: j_t = \theta(G_l^i.O_t + G_h^i.e_{t-1} + s_i) \quad (7)$$

$$F.G: T_f = \theta(G_l^f.O_t + G_h^f.e_{t-1} + s_f) \quad (8)$$

$$O.G: T_o = \theta(G_l^o.O_t + G_h^o.e_{t-1} + s_o) \quad (9)$$

$$C.I: \widetilde{T}_c = \tanh(G_l^c.O_t + G_h^c.e_{t-1} + s_c) \quad (10)$$

The vanishing gradient problem is resolved in the proposed model thanks in large part to residual connections, which are then used to find redundant data and preserve the important features in LSTM networks. Figure 4 illustrates the proposed RLSTM architecture. A residual connection network is built into the standard LSTM. The activation function Relu layers (ReLU), residual block, and batch normalisation (BN) come after the two convolutional layers that comprise the residual connection in this model. The blocks employed in the architecture of the proposed R-

LSTM networks determine the ultimate outcome, and the data must be entered into the proposed R-LSTM connection.

Mathematically the output from R-LSTM networks are expressed as

$$\mathbf{O}(t) = \text{ReLU}(\text{Res}(\mathbf{IG.FG} * \mathbf{OG})) \quad (11)$$

$$O_k = \text{Softmax}(\mathbf{O}(t)) \quad (12)$$

The primary goal of this hybrid approach is computationally efficient algorithm by combining the Beluga whale Optimization Model with R-LSTM networks.

### 3.5 IIOT Environment

It refers to the integration of connected devices, sensors, and machines within industrial environments to enhance operational efficiency, reduce downtime, and improve safety. IIoT enables the collection and transaction of real-time data across machines, devices, and systems in sectors like manufacturing, energy, transportation, and agriculture. This interconnected infrastructure allows businesses to monitor critical processes, track assets, and manage complex systems in a more automated and intelligent manner. With IIoT, industrial operations can be transformed into smart systems capable of making decisions based on data, ultimately leading to more efficient production, predictive maintenance, and improved overall system performance.

The IIoT ecosystem is built on a variety of technologies, like sensors, cloud computing, edge computing, and machine learning algorithms. Sensors and devices installed on industrial equipment generate large volumes of information, like temperature, vibration, pressure, and humidity, which is transmitted to cloud platforms or edge devices for analysis. Cloud computing provides centralized storage and processing capabilities, while edge computing brings data processing closer to the source, reducing latency and enabling real-time decision-making. This integration of smart devices and computing resources creates a robust framework that supports automation, remote monitoring, and optimization across different industries.

However, the deployment of IIoT systems also introduces potential limitations, especially in terms of security, data management, and interoperability. Given the vast number of connected devices and the critical nature of the data involved, ensuring secure communications and protecting sensitive information are vital concerns. Additionally, data collected by IIoT devices must be processed and assessed effectively to extract actionable insights, requiring the incorporation of advanced analytics and deep learning methods. Furthermore, as IIoT systems often involve diverse devices and platforms, ensuring interoperability and seamless communication among distinct systems is essential for the smooth operation of the entire network. Despite these challenges, IIoT continues to be a transformative force in industries around the world, driving innovations in automation, efficiency, and sustainability.

## 4. Results and Discussion

These criteria are evaluated alongside computational overhead to demonstrate that the suggested model is more efficient and has lower overhead expenses. Table 1 displays the formulas used to calculate the performance measures. The issues of poor generalisation and overfitting are likewise addressed by the early halting technique. This technique halts the training process when the validation performance of the method no longer improves over time.

### 4.1 Implementation Details

Python 3.19 was used to create the entire model, and matplotlib, numpy, pandas, Scikit-Learn, and seaborn are among the tools used to assess the suggested model. The PC workstation used for the experiment has an i7 CPU, an operating frequency of 3.2 GHz, and 16GB of RAM.

#### 4.2 Evaluation Metrics

| SL.NO | Performance Measures | Expression                                              |
|-------|----------------------|---------------------------------------------------------|
| 1     | Accuracy             | $\frac{TP + TN}{TP + TN + FP + FN}$                     |
| 2     | Recall               | $\frac{TP}{TP + FN} \times 100$                         |
| 3     | Specificity          | $\frac{TN}{TN + FP}$                                    |
| 4     | Precision            | $\frac{TP}{TP + FP}$                                    |
| 5     | F1-Score             | $2 \cdot \frac{Precision * Recall}{Precision + Recall}$ |

Predictions in categorisation tasks are often evaluated using four categories. The True Positive (TP) category is used when both the expected and actual values are positive. The second category is False Positive (FP), where a positive forecast is made but a negative number is obtained. The third form is called False Negative (FN), where the prediction is negative but the actual value is positive. True Negative (TN) is the fourth and last category, where the prediction and the actual value are negative.

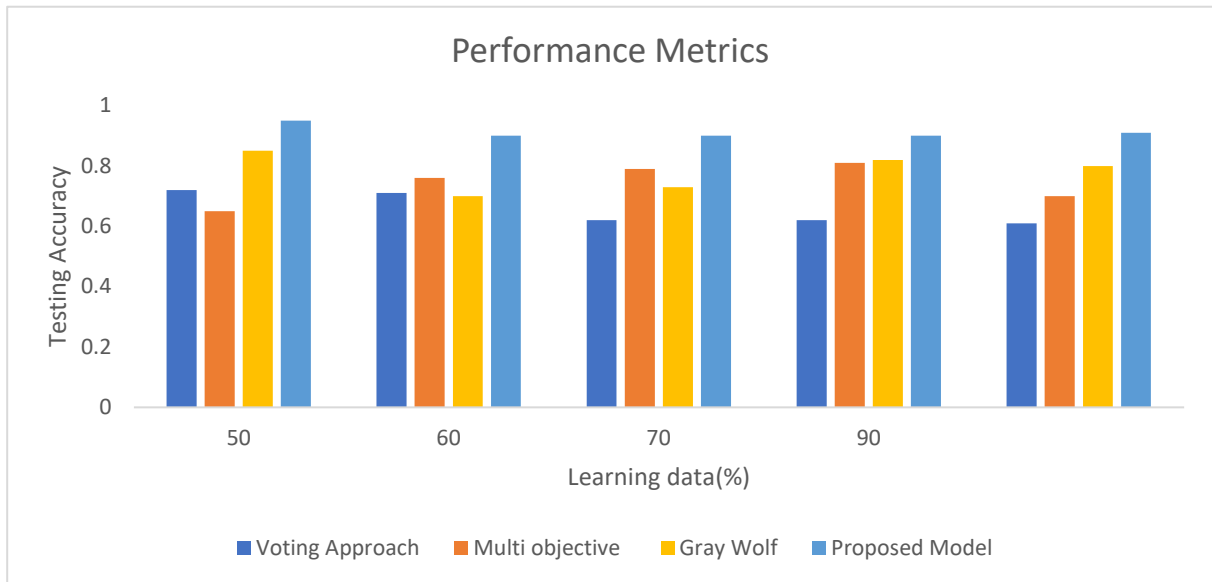


Figure 5: Performance metrics comparison of different approach

Performance metrics analysis for comparing different models involves evaluating key metrics

like accuracy, precision, recall, F1-score, and AUC-ROC to assess each method's performance. These metrics help in understanding the model's ability to correctly classify data, minimize false positives/negatives, and handle imbalanced datasets. By comparing these metrics, we can identify which model performs best in terms of predictive power, robustness, and generalization. A higher accuracy may not always be the best metric if the model performs poorly in other areas, such as precision or recall. Ultimately, the goal is to select the model that optimally balances all metrics for the given task.

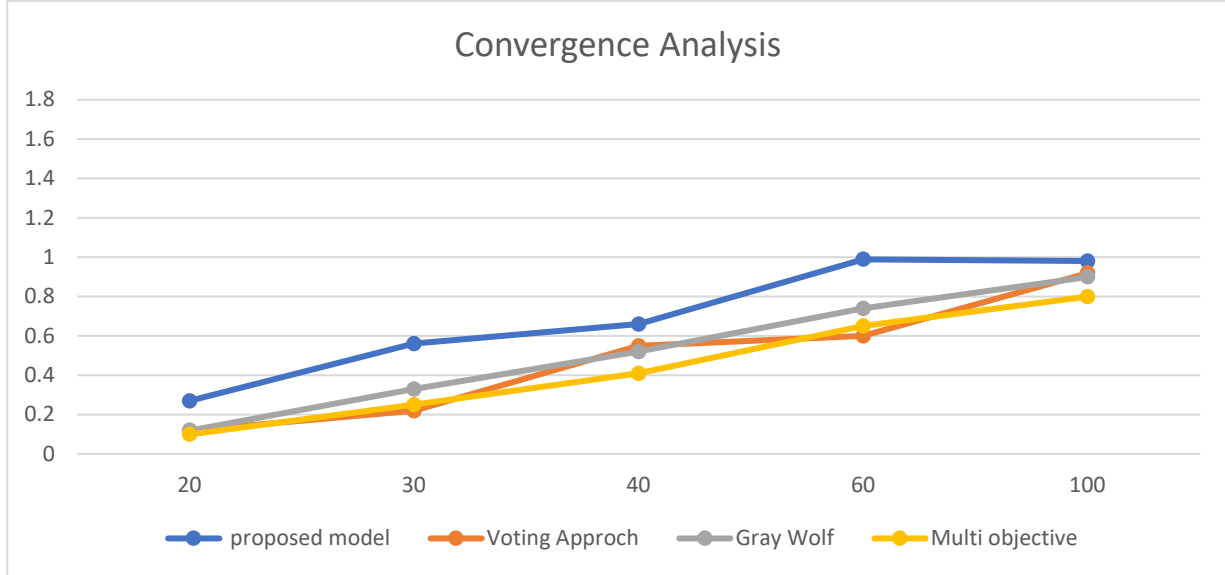


Figure 6: Convergence Analysis comparison of different approach

Convergence analysis in the comparison of different models involves evaluating how quickly each model reaches a stable solution during training or optimization. It examines whether the model's performance improves consistently over time or if it stalls, indicating a lack of convergence. The analysis includes comparing the convergence rates, accuracy, and stability of different models, often using metrics like loss reduction or performance benchmarks. Faster convergence typically indicates a more efficient model, while slower convergence might suggest a need for parameter tuning. By analyzing convergence, we can determine which model offers better performance in terms of efficiency and accuracy for a specific task.

## 5. Conclusion and Future Enhancement

In conclusion, the integration of the Beluga Whale Optimization Algorithm (BWO) with Residual Long Short-Term Memory (RLSTM) networks for identifying influential nodes in IIoT environments proves to be an effective approach. The proposed method significantly improves prediction accuracy, optimizes parameter tuning, and enhances network performance by dynamically adjusting the RLSTM's hyperparameters. The results from extensive simulations highlight the potential of this hybrid optimization technique in real-time applications, such as smart manufacturing and predictive maintenance. Future enhancements could focus on further refining the model's scalability to handle larger, more complex networks, incorporating adaptive learning techniques for better generalization across diverse IIoT scenarios, and exploring the integration of additional optimization algorithms for multi-objective optimization tasks in real-world environments.

---

**Funding Statement:** The authors received no specific funding for this study.

**Conflicts of Interest:** The authors declare that they have no conflicts of interest to report regarding the present study.

---

## References

- [1] A. A Camacho, D., Panizo-LLedot, Á., Bello-Orgaz, G., Gonzalez-Pardo, A. & Cambria, E. The four dimensions of social network analysis: An overview of research methods, applications, and software tools. *Inf. Fusion* 63, 88–120 (2020).
- [2] Zareie, A., Sheikahmadi, A. & Jalili, M. Identification of influential users in social networks based on users' interest. *Inf. Sci.* 493, 217–231 (2019).
- [3] Al-Garadi, M. A. et al. Analysis of online social network connections for identification of influential users: Survey and open research issues. *ACM Comput. Surv. (CSUR)* 51(1), 1–37 (2018).
- [4] Peng, S., Wang, G. & Xie, D. Social influence analysis in social networking big data: Opportunities and challenges. *IEEE Netw.* 31(1), 11–17 (2016).
- [5] Zareie, A., Sheikahmadi, A. & Jalili, M. Identification of influential users in social network using gray wolf optimization algorithm. *Expert Syst. Appl.* 142, 112971 (2020).
- [6] Mohammadi, A. & Saraee, M. Finding influential users for different time bounds in social networks using multi-objective optimization. *Swarm Evol. Comput.* 40, 158–165 (2018).
- [7] Saoud, B. & Moussaoui, A. Community detection in networks based on minimum spanning tree and modularity. *Physica A* 460, 230–234 (2016).
- [8] Zhao, Y., Li, S. & Jin, F. Identification of influential nodes in social networks with community structure based on label propagation. *Neurocomputing* 210, 34–44 (2016).
- [9] Vathi, E., Siolas, G. & Stafylopatis, A. Mining and categorizing interesting topics in twitter communities. *J. Intell. Fuzzy Syst.* 32(2), 1265–1275 (2017).
- [10] Moosavi, S. A., Jalali, M., Misaghian, N., Shamshirband, S. & Anisi, M. H. Community detection in social networks using user frequent pattern mining. *Knowl. Inf. Syst.* 51(1), 159–186 (2017).
- [11] Wang, M., Zuo, W. & Wang, Y. An improved density peaks-based clustering method for social circle discovery in social networks. *Neurocomputing* 179, 219–227 (2016).
- [12] Hu, Y. & Yang, B. Enhanced link clustering with observations on ground truth to discover social circles. *Knowl. Based Syst.* 73, 227–235 (2015).
- [13] Chen, M., Kuzmin, K. & Szymanski, B. K. Community detection via maximization of modularity and its variants. *IEEE Trans. Comput. Soc. Syst.* 1(1), 46–65 (2014).
- [14] Lancichinetti, A., Fortunato, S. & Kertész, J. Detecting the overlapping and hierarchical community structure in complex networks. *New J. Phys.* 11(3), 033015 (2009).
- [15] J. Zhang, D. Chen, Q. Dong, Z. Zhao Identifying a set of influential spreaders in complex networks *Sci. Rep.*, 6 (1) (2016), Article 27823
- [16] S. Kumar, B. Panda Identifying influential nodes in social networks: neighborhood coreness based voting approach *Physica A*, 553 (2020), Article 124215
- [17] A. Ullah, B. Wang, J. Sheng, J. Long, N. Khan, Z. Sun Identification of nodes influence based on global structure model in complex networks *Sci. Rep.*, 11 (1) (2021), pp. 1-11
- [18] H. Zhang, S. Zhong, Y. Deng, H.C. Kang LFIC: identifying influential nodes in complex networks by local fuzzy information centrality *IEEE Trans. Fuzzy Syst.*, 30 (8) (2021), pp. 3284-3296
- [19] F. Yang, R. Zhang, Z. Yang, R. Hu, M. Li, Y. Yuan, K. Li Identifying the most influential spreaders in complex networks by an extended local K-shell sum *Int. J. Mod. Phys. C*, 28 (1) (2017), Article 1750014
- [20] H. Sun, D. Chen, J. He, E. Ch'ng A voting approach to uncover multiple influential spreaders on weighted networks *Physica A*, 519 (2019), pp. 303-312
- [21] L. Li, S. Fang, Y. Yao Identifying influential nodes in social networks: a voting approach *Chaos, Solit. Fractals*, 152 (2021), Article 111309