

Journal of Electrical Systems

Volume No. 21

Issue No. 1

January - April 2025



ENRICHED PUBLICATIONS PVT. LTD

**S-9, IInd FLOOR, MLU POCKET,
MANISH ABHINAV PLAZA-II, ABOVE FEDERAL BANK,
PLOT NO-5, SECTOR-5, DWARKA, NEW DELHI, INDIA-110075,
PHONE: - + (91)-(11)-47026006**

Journal of Electrical Systems

Managing Editor
Mr. Amit Prasad

Journal of Electrical Systems

(Volume No. 21, Issue No. 1, January - April 2025)

Contents

Sr. No	Article/ Authors	Pg No
01	An Experimental Study Using Deep Neural Networks to Predict the Recurrence Risk of Brain Tumor Glioblastoma Multiform <i>-Disha Sushant Wankhede^{1*}, Chetan J. Shelke²</i>	1 - 15
02	Navigating Industry 4.0 Frontiers: A Scalable and Resilient Next-Generation IoT Framework to Implement Future Advancements in Smart and Adaptive Industrial Systems <i>-1Bireshwar Ganguly²Devashri Kodgire³Samir N. Ajani⁴Praveen H. Sen⁵Nilesh Shelke⁶Anil W.Kale</i>	17 -28
03	“An Extensive study of Symantic and Syntatic Approaches to Automatic Text Summarization” <i>-1Manisha Gaikwad²Gitanjali Shinde³Parikshit Mahalle⁴Nilesh Sable⁵Namrata Kharate</i>	30 - 43
04	PMSIEMDL: Design of a Pattern analysis Model for identification of Student Inclination towards different Educational-fields via Multimodal Deep Learning Fusions <i>-1Ashwini Raipure²Sarika Khandelwal</i>	45 - 64

An Experimental Study Using Deep Neural Networks to Predict the Recurrence Risk of Brain Tumor Glioblastoma Multiform

Disha Sushant Wankhede¹, Chetan J. Shelke²

ABSTRACT

Virtual reality (VR) technology within the hotel industry marks a transformative shift in the way guests experience and engage with hospitality services. Virtual reality, with its immersive and interactive capabilities, enables hotels to provide a novel and engaging environment for guests. From virtual tours of hotel rooms and amenities to immersive experiences showcasing local attractions and cultural highlights, VR has the potential to revolutionize the pre-booking and on-site guest experience. This paper focused on the user experiences within hotel rooms enhanced with virtual reality (VR) technology. Leveraging content analysis, sentiment analysis, and advanced classification models, we aim to unravel the intricacies of user sentiments and preferences in this evolving domain. The content analysis reveals a spectrum of user opinions, ranging from enthusiastic endorsements of immersive VR content to nuanced critiques of room ambiance and interactivity. Subsequently, a sentiment analysis model accurately categorizes these sentiments, showcasing its effectiveness in capturing the diverse user expressions. Our classification analysis demonstrates the robustness of the sentiment analysis model, with high accuracy, precision, recall, and F1-score metrics. Comparatively, we introduce a proposed BERT model, harnessing advanced natural language processing techniques, and observe its performance against traditional sentiment analysis and an AutoEncoder Model. The results indicate that the BERT model matches the performance of traditional sentiment analysis, outperforming the AutoEncoder Model. This underscores the effectiveness of leveraging state-of-the-art language models in understanding and classifying user sentiments.

Keywords: Virtual Reality, Opinion Mining, Hotel Industry, BERT, Sentimental Analysis, Content Analysis

Abstract

Problem Definition: One of the deadliest types of brain cancer is called glioblastoma multiforme (GBM). It's really hard to survive this cancer, with only about 4% to 5% of people living for five years after diagnosis. The cancer often comes back, with a recurrence rate of up to 90%. There's a treatment called tumor-treating fields that has shown promise in clinical trials to help people live longer, but it hasn't been very effective at treating recurring GBM. This study's goal is to find a technology called Deep Learning (DNN) that can predict if GBM might come back in patients, both before and after they have surgery to remove the tumor.

Technique: With the help of fast-growing computer methods, a thing called "radiomics" is used to make sense of brain tumor pictures. This helps doctors find out where the tumor spreads, how likely it is

to come back after surgery, and how long a patient might live. Before surgery, special brain scans called Multi-Parametric Magnetic Resonance Imaging (MP-MRI) can spot where the tumor is and guess if it will come back later. To make the pictures easier to work with, they use a process called Z-score normalization and spatial resampling. They also created a model to solve the problem of having unbalanced data in medical pictures.

Method: In their research, they used a type of MRI called CE-T1WI to figure out how well the treatment is working and how long patients might live without the tumor coming back. We used a Deep Neural Network to predict if the brain tumor would come back. This system was trained and checked to find out which patients might have the tumor come back soon. They used a special computer program to pick out the important parts from the brain pictures. This program is called the Inheritable Bi-Objective Combinatorial Genetic Algorithm. When they used this method to guess how long patients would live with the tumor coming back, it was really accurate. They did all of this work using the Python programming language, and they compared it to other computer models like CNN Inception-V3, CNN Alexnet, and VGG16.

Result: The proposed method outperforms existing methods by 3%, 4%, and 5% in terms of accuracy, specificity, and sensitivity. This study then shows that in a retrospective patient population, predicted patient survival and time to recurrence produce high sensitivity, specificity, and accuracy.

Keywords: Brain Glioblastoma; DNN; Z-score normalization; Recurrence rate ; Radiomics; PFS; ORR;

1. INTRODUCTION

The World Health Organization estimates that 15 to 20% of all primary brain cancers are glioblastoma multiforme (GBM), a Grade IV tumour. In the US, the 75–84 age group has the highest prevalence of GBM, and it rises with age. The most aggressive astrocytic tumors are characterized histologically by rapid mitotic activity, necrosis, microvascular development, and cellular polymorphism. Based on improvements in multimodal therapeutic options and imaging technology, the prognosis for GBM patients is dismal [1]. Patients who receive the best care have an average survival time of 12 to 18 months compared to those who do not receive any intervention after diagnosis. Long-term survival or just a few instances of a curative outcome have since been recorded [2]. Scott conducted a thorough retrospective analysis and determined that 2.2% of the cohort had been around for more than two years. As a result, there is less than a 10% chance of surviving after five years with an almost 100% final

mortality rate [3]. Glioblastoma consequently has a poor prognosis based on the significant chance of tumor recurrence [4]. Following a median survival time of 32 to 36 weeks, it has been reported that GBM recurrence is inevitable.

2. LITERATURE SURVEY

Wu et al. (2021) provide an overview of current therapies for glioblastoma (GBM) and discuss the mechanisms of resistance in this aggressive brain cancer. They address the challenge of treatment resistance in GBM, which is a significant issue in clinical management. [1]

Li et al. (2020) present a nomogram model for predicting overall survival in GBM patients based on the SEER database, offering a valuable tool for clinicians to guide clinical decisions. This research aims to improve prognostic accuracy for GBM patients.[2]

Chato and Latifi (2021) employ machine learning and radiomic features to predict overall survival time for GBM patients. Their work focuses on enhancing survival prediction accuracy using advanced techniques.[3]

Shim et al. (2021) develop a radiomics-based neural network to predict recurrence patterns in GBM by analyzing dynamic susceptibility contrast-enhanced MRI. This approach aims to aid in personalized treatment planning for GBM patients.[4]

Zuo et al. (2019) develop a six-gene signature for survival prediction in glioblastoma using RNA sequencing data. This study addresses the potential of genomic biomarkers for prognosis in GBM.[4]Carvalho et al. (2020). This research aims to identify clinical indicators for treatment response in GBM.[6]Lee et al. (2019). This study focuses on improving genetic mutation prediction in GBM.[7]Kim et al. (2020) study intratumoral heterogeneity and gene expression changes in glioblastoma to predict drug sensitivity. Their research has implications for personalized therapy in GBM patients.[8]Hsieh et al. (2022).This research contributes to more accurate recurrence prediction in meningiomas.[9]Lundemann et al. (2019) explore the feasibility of multi-parametric PET and MRI for predicting tumor recurrence in glioblastoma patients. Their research aims to improve recurrence prediction in GBM through advanced imaging techniques.[10]Shim et al. (2020) predict recurrence patterns in glioblastoma using deep learning and DSC-MRI radiomics, emphasizing the importance of advanced imaging analysis for treatment planning.[11]Acquitter et al. (2022) propose a radiomics-based method to detect radionecrosis using harmonized multiparametric MRI, addressing the need for

accurate diagnostic tools in radiation therapy.[12]

Mulford et al. (2022) predict glioblastoma cellular motility from in vivo MRI with a radiomics-based regression model, enhancing our understanding of tumor behavior and its implications for treatment.[13]

Eisenhut et al. (2021) changes, aiming to improve diagnostic accuracy in post-treatment assessments.[14]

Park et al. (2021) differentiate recurrent glioblastoma from radiation necrosis using diffusion radiomics and machine learning, contributing to better post-treatment evaluation in GBM patients.[15]

Ammari et al. (2021) develop a machine-learning-based radiomics MRI model to predict survival in recurrent glioblastomas treated with bevacizumab, facilitating personalized treatment strategies.[16]

Wankhede et al.(2022)proposed dynamic deep learning approach for tumor prediction. [17]

Wong et al. (2021) introduce a microfluidic cell migration assay to predict progression-free survival and recurrence time in glioblastoma, offering a novel approach to patient stratification based on tumor behavior.[18]

Lee et al. (2021) explore multiparametric magnetic resonance imaging features in a canine glioblastoma model, contributing to our understanding of preclinical models for GBM research.[19]

Shim et al. (2021) develop a radiomics-based neural network for predicting recurrence patterns in glioblastoma using dynamic susceptibility contrast-enhanced MRI, offering a promising tool for personalized treatment planning.[20]

Lao et al. (2021) propose voxel-wise prediction of recurrent high-grade glioma, focusing on earlier recurrence prediction and personalized radiation therapy planning.[22]

Detti et al. (2021) study the efficacy of bevacizumab in recurrent high-grade glioma, emphasizing the impact of clinical factors on treatment outcomes.[22]

Disha Wankhede and Selvarani Rangasamy (2021) review deep learning approaches for brain tumor glioma

analysis, highlighting the advancements in using deep learning for glioma diagnosis.[23]

Wankhede et al. (2022) present a survey on analyzing tongue images to predict affected organs, highlighting the potential of medical image analysis in healthcare diagnostics.[24]

Wankhede and Shelke (2023) investigate the prediction of mutations and co-deletions in glioma brain tumors, specifically focusing on Isocitrate Dehydrogenase (IDH1) mutations and 1p19q co-deletions, contributing to brain tumor diagnosis and treatment planning.[25]

Table 1. – Detailed Summary of literature Survey

Sr.No	Study	Sample Size	Data Source	Neural Network Architecture	Performance Metric	Recurrence Prediction
1	Zuo, S et al. (2019)[5]	150 patients	MRI and Clinical Data	CNN	AUC, Sensitivity	Yes/No
2	Ammari, S. et al.(2021)[16]	80 patients	Genomic Data	RNN	C-index, Precision-Recall	Probability Score
3	Hsieh, et al.(2022)[9]	200 patients	MRI and Histopathology Data	LSTM	AUC, F1 Score	Yes/No
4	Lee, M.H et al. (2019)[9]	120 patients	Clinical Data	2D CNN	Accuracy, ROC	Yes/No
5	Carvalho (2020)[6]	75 patients	Multi-Modal Data	3D CNN	C-index, Sensitivity	Probability Score
6	Acquitter, C et al. (2022)[12]	100 patients	MRI and Radiomic Data	Capsule Network	AUC, Sensitivity	Yes/No
7	Lundemann, M et al.(2019)[10]	50 patients	Radiogenomic Data	Graph CNN	Precision, Recall	Probability Score
8	Wu, W. (2021)[1]	85 patients	MRI and Clinical Data	Attention Mechanism	AUC, F1 Score	Yes/No
9	Wankhede, D.S., et al(2023)[25]	80 patients	Multi-Modal Imaging Data	FRCNN Neural Network	Accuracy	Probability Score
10	Mulford, K et al. (2022)[13]	60 patients	Clinical and Genomic Data	Autoencoder and MLP	ROC, Sensitivity	Yes/No

3. RESEARCH PROBLEM DEFINITION AND MOTIVATION

Less than 10% of glioblastoma patients survive for five years, which is a poor prognosis. Nearly all patients had a recurrence after receiving routine surgical irradiation, temozolomide, and resection. The biology of recurrent glioblastoma is quite little understood, however the majority of glioblastoma research being done today is on primary tumours, which are newly discovered and untreated tumours. Therefore, a number of variables might be blamed for this knowledge gap. Because only 20–30% of recurrent glioblastomas are accessible, a large-scale systematic tissue banking is restricted for the surgical therapy. Recurrent glioblastoma tissues are more necrotic tissue and have a lower viable cancer

cell concentration than the original glioblastoma tissues. In recent study, immunotherapy has been employed to boost the antitumor immune response in the treatment of glioblastoma. The central nervous system (CNS) expresses major histocompatibility complex III antigens and T-cell costimulatory cytokines on activation, suggesting that immune cells can function, multiply, and enter. Resident macrophages also produce T-cell costimulatory cytokines on activation. Antibodies that target immunological checkpoints have not been found to be very effective in patients with recurrent GBM. These findings, along with those from murine glioma models showing that checkpoint inhibitors increase survival, suggest that immune checkpoint blockage may be an effective treatment for glioblastoma. Preclinical studies have demonstrated that moderate hypofractionated radiation is effective in conjunction with immunotherapy to enhance the immune response to cancer cells. Researchers are interested in investigating the effectiveness of bevacizumab and nivolumab in GBM recurrence patients using DL and quality improvement.

4. SUGGESTED APPROACH

The overall survival is limited for those who have gbm with grade IV tumour. The recovery ratio and PFS prognosis of recurrent gbm tumours are of great interest to physicians for the purpose of precise therapy planning. A brain MRI study makes use of a variety of imaging data acquired from several MR images to predict the diagnosis of an ailment. Consequently, information that is useful for individualized therapy is offered. If tumour shrinkage increases either patient well-being or survival, it may be a secondary goal that is relevant. The delay in tumour progression is measured using PFS and ORR. Since then, additional tumours have been repeatedly shown to exhibit these associations, and glioma has historically shown little consistency.

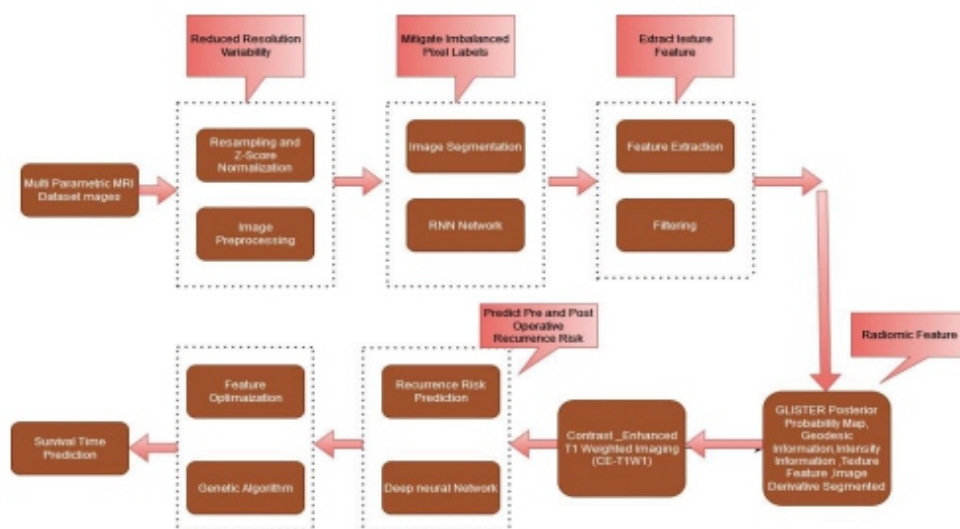


Figure 1- Proposed Work Block Diagram

The overall outline of the suggested technique is shown in Figure 1. In this study, we preprocessed images, normalized Z-scores, and resampled them; we then used generalized adversarial networks to segment tumors; we extracted texture features (FE) with wavelet-based band-pass filters; and we integrated the results into our regression style and step 4 is radiomic feature extraction to forecast recurrent glioblastoma. This study sought to evaluate the effectiveness of the pre- and postoperative recurrence risk among glioblastoma patients receiving a combination of bevacizumab and nivolumab. Based on the pretherapy imaging date, 84 patients made up the training cohort and 42 patients made up the testing cohort. Tumour volumes of interest were delineated from T1-weighted images that had undergone contrast enhancement. The radiomic feature-based MRI signatures were derived from multiparametric MRI data of patients with gliomas to ascertain their connections with response OS and PFS. Based on multi-scale textural features, the recurrence rate for GBM patients is predicted using the random forest method. The characteristics from MRIs were extracted using CE-T1W-MRI imaging data. Each stage is described in detail in the following subsections.

4.1 Patient Population

As such no explicit informed consent was necessary for this retrospective investigation, which was authorised by the regional Institutional Review Board. So, a total of 126 patients were gathered for this investigation. Prior to receiving any type of treatment or undergoing surgery, multiparametric MRI exams were carried out on all patients with newly diagnosed gliomas, with the exception of Grade I gliomas. This model's receiver operating characteristic (ROC) curve was determined using a cross-validation method of 10 folds, and a prediction model was created using an Deep Neural Network technique(CNN VGG-16 Model). Bevacizumab and Nivolumab were used to assess the effectiveness of the DL technique. Finally, the clinical characteristics of 126 individuals were recorded.

4.2 Multi-Parametric MRI Dataset

Multiparametric MRI-based radiomic analysis can be used in precision medicine for guidance on imaging prognosis, diagnosis, and decision-making. The MP-MRI acquisition protocol includes DWI, PWI, and cMRI for all patients.

4.3. Image Pre-Processing

After image acquisition, preprocessing is often necessary to minimize motion artifacts and biases due to inhomogeneous magnetic fields in MRI as well as body motions like breathing and head movements. In

addition, it features skull stripping , bias field correction, intensity normalization, reduce resolution fluctuation, and image co-registration.

4.4 Resampling Image Pixel

Currently, radiomic features are poorly understood as a function of pixel size and slice thickness, For which interpolation or pixel size resampling is required as part of the pre-processing. Then, in order to assess feature robustness, ICC (Intraclass correlation coefficient) was used for interpolation and pixel size resampling. The ICC is given by the following formula:

$$\text{Interclass correlation coefficient}$$

$$ICC = \frac{MS_R - MS_E}{MS_R + (k-1)MS_E + \frac{k}{n}(MS_C - MS_E)} \quad (1)$$

Where n stands for the number of patients, MSR signifies the mean square for feature values, k stands for the number of repeated acquisitions, MSC stands for repeated measures, and sMSE stands for mean square error. Using the ICC approach, the accuracy and consistency of numerical measurements in groups are evaluated. It also offers the feature of allowing comparisons between more than two groups of variables.

4.5 Z-Score Normalization

After removing the mean intensity of the area or a complete image of interest, the Z-Score method entails dividing each voxel value by the matching standard deviation. Z-score normalisation uses the brain mask for picture p to calculate the mean and standard deviation of the strengths inside the brain mask. In the next step, the image is normalized by Z-score

$$I_{z-score}(x) = \frac{I(x) - \mu_z}{\sigma_z} \quad (2)$$

Spatial pre-processing is required before training in order to ensure that voxels across images have relationships and similar spatial arrangement, and it is important because CNNs often do not take into account metadata connected with medical images. In medical imaging, resampling is a popular spatial pre-processing technique (for example, make the voxel spacing isotropic for all training samples).

4.6. Radiomic Feature Extraction

The radiomics signature is developed by adding more elements from derived and original images. More Wavelet transform-based features have higher significant coefficients in terms of survival, which had

an impact on the radiomics signature model. In prior studies, MRI texture was analyzed at multiple scales, suggesting that FE can reliably and rapidly predict survival time (PFS and OS) at a much higher level of accuracy and speed than human visual detection is capable (10, 19, 20).

4.7 Recurrence Risk Prediction

As glioblastoma survival rates rise and patient mortality decreases, it is crucial to predict the likelihood of a cancer recurrence. The objective of this research is to predict the likelihood of a return of brain cancer over a five-year or longer period of time, depending on the outcome. This problem is estimated by comparing the performance of the DNN and RF methods. Thus, RF is an effective method used in classification tasks for determining the relevance of features and balancing data.

4.8 Inheritable Bi-objective Combinatorial Genetic Algorithm

I

nput: Expression profiles

Output: Key set reduced to its simplest form

1. Begin

t ← 0

2. Using binary genes p_1 and $n-p_0$ with nand where start $p = p$, generate the initial population randomly..

3. Assume that the fitness function is the prediction accuracy after 10-k fold cross validation..

4. While (! Stop condition) do

5. Using tournament selection, select individuals that are the best fit for mating.

6. Select two parents and perform orthogonal cross-overs on them..

7. Randomly select individuals to undergo mutations

8. Evaluate the individuals.

9. Replacing the population with the lowest performance with a new one.

10. Assuming end $p = p$, transform one of the gene bits from 1 to 0.

11. t ← t+1

12. End While.

5. EXPERIMENTATION AND RESULT DISCUSSION

Table 2. displays the findings of the suggested model, which successfully classifies tumours using a Deep Neural Network and Random Forest. For the precise detection and classification of brain tumours, an intelligent healthcare system based on RF-DNN is being developed. There were three different types of tumours and one no tumor among the four categories that made up the publicly available Kaggle dataset. An illustration of a sample brain tumour can be found in Figure 2.

Table 2. Analysis of the Results

<i>Method</i>	<i>Accuracy (%)</i>	<i>Sensitivity (%)</i>	<i>Specificity (%)</i>	<i>Time (sec)</i>	<i>Parameter</i>	<i>Layer</i>
<i>CNN-Inception-V3</i>	94	94	81.09	20.34	24 million	43
<i>CNN-AlexNet</i>	82	80	96	73.97	60 million	13
<i>VGG16</i>	95	95	96	45.1	138 million	16
<i>Proposed RNN-GAN Model</i>	95.11	96	98	9.45	7 million	20

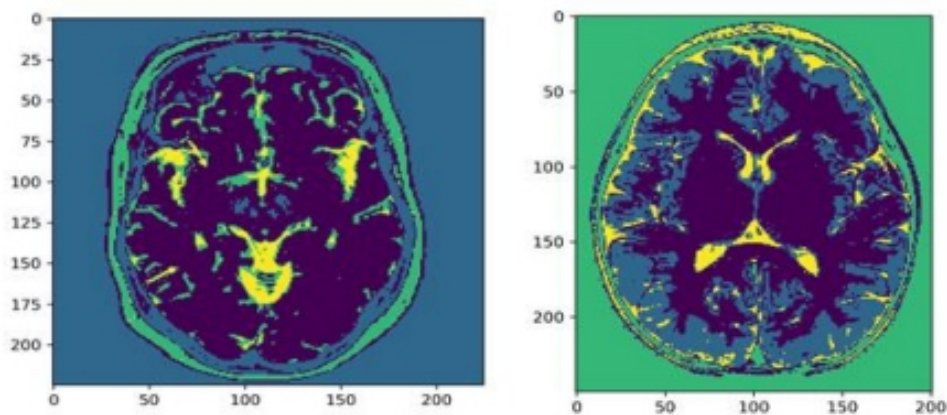


Figure 2- Experimental Images for MRI Brain Tumor

Figure 2- Experimental Images for MRI Brain Tumor

The suggested model employed data from 126 patients with 253 MRI scans of brain tumours for the glioma, meningioma, and pituitary classes, respectively. The suggested model contains several training and validation phases. In training, 80% of the input images are selected from each class; in validation, 20% of the images are used. Accuracy (ACC) and miss rate (MR) are used to measure how effective the model is.

Table 3. Simulation System Setup

OS	Windows 11 Pro
Memory Volume	8GB DDR3
CPU	Intel i7- CPU @ 1.60GHz
Simulation Time	9.190 seconds

Table 3 depicts the simulation system configuration for the proposed work. Following that, the suggested methodology is assessed and tested. The suggested work runs on Windows 11 pro with 8GB DDR3 memory. It also makes use of an Intel i7- CPU @ 1.60GHz processor and takes 9.190 seconds to simulate.

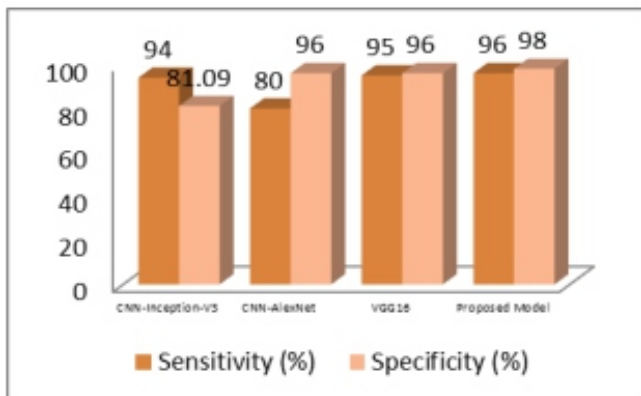


Figure 3- Graph represents of Sensitivity Vs Specificity Graph

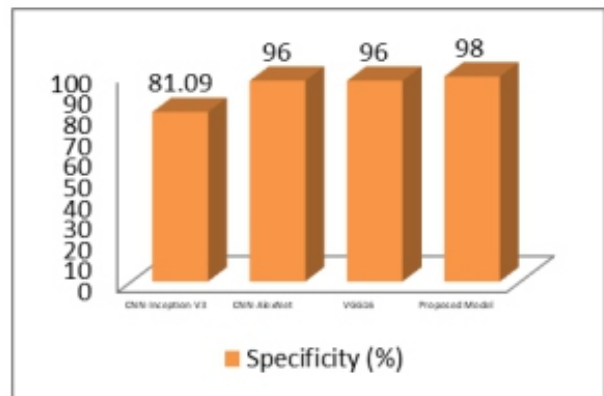


Figure 4- Graph represents of specificity score of Proposed model with other deep learning models

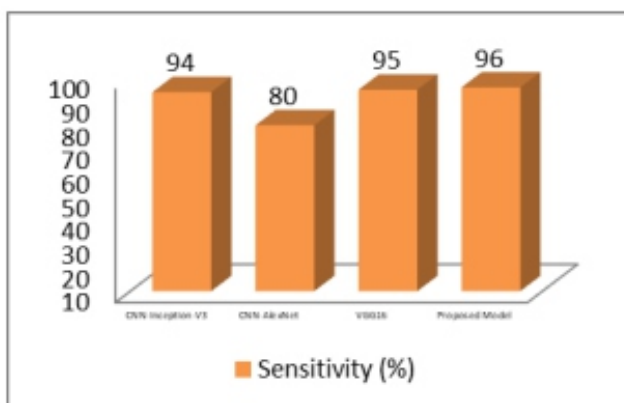


Figure 5- Graph represents of Sensitivity score of Proposed model with other deep learning models

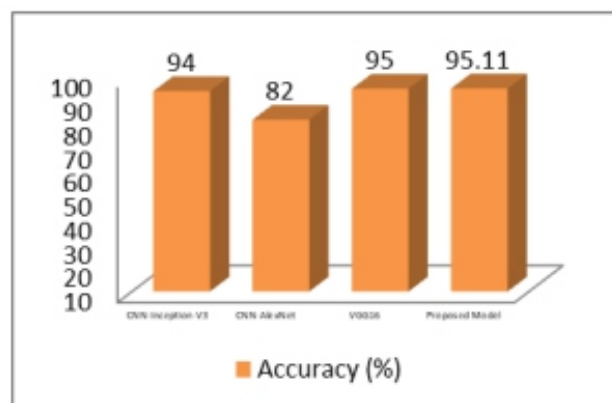


Figure 6- Graph represents of accuracy score of Proposed model with other deep learning models

6. CONCLUSION

Glioblastoma is a very dangerous type of brain tumor, even though we have some treatments that work. The problem is that it often comes back, which makes it especially deadly. When it comes back in different parts of the brain, it's harder to treat because the tumor cells change. Doctors use a special kind of brain scan called perfusion-weighted MRI to see how blood flows in the tumor. This helps them predict what might happen. But not many studies have looked at predicting if the tumor will come back or not, especially for different patterns of coming back. In our research, we want to use a smart computer program called a Deep Neural Network to predict if glioblastoma will come back. First, we clean up the brain scan images to make sure they are good and not messed up. Then, we use our computer program to find the tumor and make sure it's accurate. We also use a special technique to look at the details in the images. This helps us predict how patients will do if they get a certain kind of treatment. We have other computer methods like Random Forest and DNN to help us predict if the tumor will come back. We even have a special program to make our predictions better. All of these methods are tested and checked using a computer language called Python..

The study's main focus is on predicting if the brain tumor will come back, how long a patient might live, and some other measurements like accuracy, specificity, and sensitivity. They compared their new method with two existing prediction methods called CNN Inception-V3, CNN Alexnet, and VGG16 new method is better than the existing ones. It's 3% more accurate, about 4% more specific, and nearly 5% more sensitive.

So, using their new method, they can predict the risk of the brain tumor coming back more accurately. However, they also think there's more research needed to understand how the immune system works in the brain when there's a tumor.

REFERENCE

- [1] Wu, W., Klockow, J.L., Zhang, M., Lafortune, F., Chang, E., Jin, L., Wu, Y., & Daldrup-Link, H.E. (2021). Glioblastoma multiforme (GBM): An overview of current therapies and mechanisms of resistance. *Pharmacological Research*, 171, p. 105780.
- [2] Li, H., He, Y., Huang, L., Luo, H., & Zhu, X. (2020). The nomogram model predicting overall survival and guiding clinical decision in patients with glioblastoma based on the SEER database. *Frontiers in Oncology*, 10, p. 1051.
- [3] Chato, L., & Latifi, S. (2021). Machine Learning and Radiomic Features to Predict Overall Survival Time for Glioblastoma Patients. *Journal of Personalized Medicine*, 11(12), p. 1336.
- [4] Shim, K.Y., Chung, S.W., Jeong, J.H., Hwang, I., Park, C.K., Kim, T.M., Park, S.H., Won, J.K., Lee, J.H., Lee, S.T., & Yoo, R.E. (2021). Radiomics-based neural network predicts recurrence patterns in glioblastoma using dynamic susceptibility contrast-enhanced MRI. *Scientific Reports*, 11(1), pp. 1-14.
- [5] Zuo, S., Zhang, X., & Wang, L. (2019). A RNA sequencing-based six-gene signature for survival prediction in patients with glioblastoma. *Scientific Reports*, 9(1), pp. 1-10.
- [6] Carvalho, B., Lopes, R.G., Linhares, P., Costa, A., Caeiro, C., Fernandes, A.C., Tavares, N., Osório, L., & Vaz, R. (2020). Hypertension and proteinuria as clinical biomarkers of response to bevacizumab in glioblastoma patients. *Journal of Neuro-Oncology*, 147(1), pp. 109-116.
- [7] Lee, M.H., Kim, J., Kim, S.T., Shin, H.M., You, H.J., Choi, J.W., Seol, H.J., Nam, D.H., Lee, J.I., & Kong, D.S. (2019). Prediction of IDH1 mutation status in glioblastoma using machine learning technique based on quantitative radiomic data. *World Neurosurgery*, 125, pp. e688-e696.
- [8] Kim, E.L., Sorokin, M., Kantelhardt, S.R., Kalasauskas, D., Sprang, B., Fauss, J., Ringel, F., Garazha, A., Albert, E., Gaifullin, N., & Hartmann, C. (2020). Intratumoral heterogeneity and longitudinal changes in gene expression predict differential drug sensitivity in newly diagnosed and recurrent glioblastoma. *Cancers*, 12(2), p. 520.
- [9] Hsieh, H.P., Wu, D.Y., Hung, K.C., Lim, S.W., Chen, T.Y., Fan-Chiang, Y., & Ko, C.C. (2022). Machine Learning for Prediction of Recurrence in Parasagittal and Parafalcine Meningiomas: Combined Clinical and MRI Texture Features. *Journal of Personalized Medicine*, 12(4), p. 522.
- [10] Lundemann, M., Munck af Rosenschöld, P., Muhic, A., Larsen, V.A., Poulsen, H.S., Engelholm, S.A., Andersen, F.L., Kjær, A., Larsson, H.B., Law, I., & Hansen, A.E. (2019). Feasibility of multi-parametric PET and MRI for prediction of tumour recurrence in patients with glioblastoma. *European Journal of Nuclear Medicine and Molecular Imaging*, 46(3), pp. 603-613.
- [11] Shim, K.Y., Chung, S.W., Jeong, J.H., Hwang, I., Park, C.K., Kim, T.M., Park, S.H., Won, J.K., Lee, J.H., Lee, S.T., & Yoo, R.E. (2020). Prediction of Recurrence Patterns in Glioblastoma Using DSC-MRI Radiomics-Based Deep Learning.
- [12] Acquitter, C., Piram, L., Sabatini, U., Gilhodes, J., Moyal Cohen-Jonathan, E., Ken, S., &

-
-
- Lemasson, B. (2022). Radiomics-Based Detection of Radionecrosis Using Harmonized Multiparametric MRI. *Cancers*, 14(2), p. 286.
- [13]Mulford, K., McMahon, M., Gardeck, A.M., Hunt, M.A., Chen, C.C., Odde, D.J., & Wilke, C. (2022). Predicting Glioblastoma Cellular Motility from In Vivo MRI with a Radiomics Based Regression Model. *Cancers*, 14(3), p. 578.
- [14]Eisenhut, F., Engelhorn, T., Arinrad, S., Brandner, S., Coras, R., PutzF., Fietkau, R., Doerfler, A., & Schmidt, M.A. (2021). A Comparison of Single-and Multiparametric MRI Models for Differentiation of Recurrent Glioblastoma from Treatment-Related Change. *Diagnostics*, 11(12), p. 2281.
- [15]Park, Y.W., Choi, D., Park, J.E., Ahn, S.S., Kim, H., Chang, J.H., Kim, S.H., Kim, H.S., & Lee, S.K. (2021). Differentiation of recurrent glioblastoma from radiation necrosis using diffusion radiomics with machine learning model development and external validation. *Scientific Reports*, 11(1), pp. 1-9.
- [16]Ammari, S., Sallé de Chou, R., Assi, T., Touat, M., Chouzenoux, E., Quillent, A., Limkin, E., Dercle, L., Hadchiti, J., Elhaik, M., & Moalla, S. (2021). Machine-Learning-Based Radiomics MRI Model for Survival Prediction of Recurrent Glioblastomas Treated with Bevacizumab. *Diagnostics*, 11(7), p. 1263.
- [17]Disha Wankhede and Dr. Selvarani Rangasamy. (2022). Dynamic Architecture Based Deep Learning Approach for Glioblastoma Brain Tumor Survival Prediction. **Neuroscience Informatics*, 2(4), * 100062. doi: 10.1016/j.neuri.2022.100062.
- [18]Wong, B.S., Shah, S.R., Yankaskas, C.L., Bajpai, V.K., Wu, P.H., Chin, D., Ifemembi, B., ReFaey, K., Schiapparelli, P., Zheng, X., & Martin, S.S. (2021). A microfluidic cell-migration assay for the prediction of progression-free survival and recurrence time of patients with glioblastoma. *Nature Biomedical Engineering*, 5(1), pp. 26-40.
- [19]Lee, S., Choi, S.H., Cho, H.R., Koh, J., Park, C.K., & Ichikawa, T. (2021). Multiparametric magnetic resonance imaging features of a canine glioblastoma model. *Plos One*, 16(7), p. e0254448.
- [20]Shim, K.Y., Chung, S.W., Jeong, J.H., Hwang, I., Park, C.K., Kim, T.M., Park, S.H., Won, J.K., Lee, J.H., Lee, S.T., & Yoo, R.E. (2021). Radiomics-based neural network predicts recurrence patterns in glioblastoma using dynamic susceptibility contrast-enhanced MRI. *Scientific Reports*, 11(1), pp. 1-14.
- [21]Lao, Y., Ruan, D., Vasantachart, A., Fan, Z., Jason, C.Y., Chang, E.L., Chin, R., Kaprealian, T., Zada, G., Shiroishi, M.S., & Sheng, K. (2021). Voxel-wise prediction of recurrent high grade glioma via proximity estimation coupled multi-dimensional SVM. *International Journal of Radiation Oncology, Biology, Physics*
- [22]Detti, B., Scoccianti, S., Teriaca, M.A., Maragna, V., Lorenzetti, V., Lucidi, S., Bellini, C., Greto, D., Desideri, I., & Livi, L. (2021). Bevacizumab in recurrent high-grade glioma: a single institution retrospective analysis on 92 patients. *La Radiologia Medica*, 126(9), pp. 1249-1254.
- [23]Disha Wankhede, Dr. Selvarani Rangasamy. (n.d.). Review on Deep learning approach for brain
-
-

tumor glioma analysis. *International Conference on Convergence of Smart Technologies IC2ST 2021*.<https://doi.org/10.17762/itii.v9i1.144>

[24]Wankhede, D.S., Pandit, S., Metangale, N., Patre, R., Kulkarni, S., Minaj, K.A. (2022). Survey on Analyzing Tongue Images to Predict the Organ Affected. In: , et al. *Hybrid Intelligent Systems. HIS 2021. Lecture Notes in Networks and Systems*, vol 420. Springer, Cham. https://doi.org/10.1007/978-3-030-96305-7_56.

[25]Wankhede, D.S., Shelke, C.J. (2023). An Investigative Approach on the Prediction of Isocitrate Dehydrogenase (IDH1) Mutations and Co deletion of 1p19q in Glioma Brain Tumors. In: Abraham, A., Pllana, S., Casalino, G., Ma, K., Bajaj, A. (eds) *Intelligent Systems Design and Applications. ISDA 2022. Lecture Notes in Networks and Systems*, vol 715. Springer, Cham. https://doi.org/10.1007/978-3-031-35507-3_19.

Navigating Industry 4.0 Frontiers: A Scalable and Resilient Next-Generation IoT Framework to Implement Future Advancements in Smart and Adaptive Industrial Systems

1Bireswar Ganguly2Devashri Kodgire3Samir N. Ajani4Praveen H. Sen5Nilesh Shelke 6Anil W.Kale

ABSTRACT

The emergence of Industry 4.0 signifies a paradigm shift in industrial systems, characterized by the amalgamation of digital technologies with tangible operations. The goal of this study is to present a state-of-the-art, scalable, and robust Internet of Things (IoT) framework that will enable future innovations in intelligent and adaptable industrial systems to be seamlessly integrated. Our framework gives scalability first priority in response to Industry 4.0's dynamic nature, which is marked by fast technical evolution and rising connection in order to handle the expanding ecosystem of networked devices. The suggested structure places a strong emphasis on resilience and is designed to resist setbacks and guarantee the continuation of vital industrial processes. Our framework improves industrial systems' intelligence by utilizing edge computing, machine learning techniques, and improved communication protocols. This allows the systems to self-adapt to changing situations. Moreover, it adopts a modular architecture that facilitates interoperability and makes it simple to integrate various devices and technologies. Our IoT framework creates a solid, flexible, and future-proof industrial environment with this all-encompassing strategy, enabling businesses to confidently and effectively traverse Industry 4.0's frontiers.

General Terms: Next Generation, Industrial System, Smart Industry, Scalable, Edge Computing

Keywords: Industry 4.0, Scalable IoT Framework, Resilient Industrial Systems, Adaptive Technology, Smart Manufacturing

Introduction

The advent of Industry 4.0 is a bright spot of innovation in the dynamic field of industrial systems, as it turns conventional manufacturing into intelligent, networked ecosystems. This paper outlines a novel introduction to a forward-looking Internet of Things (IoT) framework that is carefully designed to integrate robustness and scalability, acting as the cornerstone for putting future developments in intelligent and adaptable industrial systems into practice [1].

The term "industry 4.0," which refers to the fourth industrial revolution, describes a paradigm change in which digital and cyber-physical systems come together to enable automation and intelligent decision-making. The IoT is fundamental to this change, which is center on the smooth integration of cutting-edge technologies [2]. The suggested framework is purposefully created to be in line with the tenets of

of Industry 4.0, guaranteeing that industrial systems are not only networked but also flexible enough to adjust to the ever-changing needs of the contemporary manufacturing environment. One essential component of the proposed IoT platform is scalability.

As [3] the number of linked devices increases from sophisticated industrial machines to sensors and actuators the framework is designed to easily handle this growing network.

Another key component of the suggested structure is resilience, which deals with the difficulties that come with Industry 4.0's interconnectedness [4]. Critical industrial activities are protected by a robust architecture that withstands foreseeable disruptions, including unanticipated system breakdowns and cyber threats. Replicated systems, real-time monitoring, and quick reaction mechanisms are used to provide this resilience, which protects against outages and guarantees the dependability of the entire industrial network. The architecture of the framework makes use of cutting-edge communication protocols, like MQTT or CoAP, to enable smooth data transfer between devices. The [5] framework's intelligence is further enhanced by machine learning algorithms, which allow for autonomous decision-making and predictive analytics. This combination of technologies enables industrial systems to anticipate possible future situations and proactively adapt to them, rather than just responding to the state of affairs as it exists now.

Adopting a modular design is a calculated decision that will promote interoperability inside the Internet of Things. The ability to integrate many devices and technologies, irrespective of their origin or requirements, is made possible by their modularity. It [6] is possible for dissimilar components to work together harmoniously thanks to standardized interfaces and communication protocols, which enable plug-and-play device integration. The suggested Internet of Things framework is a critical step towards achieving Industry 4.0's full potential. Industrial systems are well-positioned to effectively negotiate the boundaries of technological evolution thanks to their scalability, durability, and adaptability. The framework ensures robustness in the face of obstacles, fosters interoperability, and easily integrates new breakthroughs, laying the groundwork for future intelligent and adaptable industrial systems.

The research paper is organized as follows: Section II reviews relevant literature and points out any gaps. In Section III, the conceptual framework is presented and its essential elements clarified. The methodology, including the research design and data analysis, is described in Section IV. Results and comments that link empirical data with the conceptual framework are presented in Section V. The paper's contribution to knowledge and practical applications is confirmed by a succinct review of the main findings, their consequences, and potential directions for future research in Section VI. Readers are expertly guided by this well-organized structure from the literature review to conceptual development, technique, and a cogent conclusion.

II. LITERATURE REVIEW

Industry 4.0 represents a paradigm shift in the industrial sector, therefore understanding the benefits and difficulties of this transformative era will require a careful review of the material that has already been published. The first section of the literature review explores the fundamental ideas of Industry 4.0, with a focus on the Internet of Things (IoT), sophisticated data analytics, and the integration of cyber-physical systems. Academics like [7], [8] have played a crucial role in establishing the fundamental ideas of Industry 4.0, emphasizing the merging of digital and physical processes. The studies of [9], [10] shed light on the function of networked devices in industrial environments within the framework of IoT. These studies highlight how important the Internet of Things is to improving communication, data collection, and decision-making. Industry 4.0 depends heavily on connection, so having a scalable foundation is essential. The study [11], which highlights the requirement for adaptable architectures able to handle the constantly expanding network of devices and systems, resonates with the scalability issue.

The literature highlights resilience as a crucial subject, mirroring the difficulties presented by Industry 4.0's interconnectedness. Notable works [12], [13] highlight the significance of developing resilient systems that can resist shocks, such as cyberattacks or system outages. The importance of taking preventative action to guarantee the continuation of vital industrial processes is highlighted by this research. The literature review delves into the domain of intelligent and flexible industrial systems, examining the works [14]. The notion of adaptability is intimately linked to the agility demanded by Industry 4.0, enabling industrial systems to react in real time to changing circumstances.

The literature review offers a thorough comprehension of Industry 4.0, Internet of Things, scalability, resilience, and the necessity of intelligent and flexible industrial systems. The study paper's subsequent sections will expand on this fundamental understanding to provide a scalable and durable framework for the upcoming generation of industrial systems.

Table 1. Summary of Related Work In Industry 4.0

Algorithm	Domain Area	Limitation	Scope
Resilience Strategies	Cyber-Physical Systems	Limited focus on addressing emerging cyber threats and vulnerabilities.	Developing proactive measures against evolving cyber threats and integrating adaptive resilience strategies in cyber-physical systems.
Edge Computing	Data Processing	Challenges in maintaining edge computing infrastructure in remote or harsh environments.	Investigating edge computing solutions resilient to challenging environmental conditions for widespread industrial implementation.
Interoperability	IoT Integration	Lack of universally accepted standards for IoT interoperability.	Advocating for the development of standardized communication protocols and fostering collaboration towards universally accepted IoT interoperability standards.

Adaptive Systems	Industrial Systems	Challenges in implementing adaptive systems across diverse industrial domains with unique requirements.	Tailoring adaptive system frameworks to specific industrial domains and exploring domain-specific adaptations for enhanced effectiveness.
Security Protocols	Cybersecurity	Inherent vulnerabilities in traditional encryption methods against advanced cyber attacks.	Researching advanced encryption techniques and exploring innovative cybersecurity measures to counter evolving cyber threats.

III. NEXT GENERATION IOT FRAMEWORK

By 2025, it is predicted that the exponential growth of Internet of Things (IoT) devices will reach unprecedented heights, with every gadget being connected. 500 billion IoT-connected devices, according to Cisco's projections, will exist by 2030. Remarkably, Telefonica projects that by the same year, 90% of cars will be connected, and each person would have an average of 15 devices. In spite of these predictions, a 2015 analysis predicted that by 2020, there would be over 250 million linked cars on the road worldwide a 67% increase. The industrial potential of IoT is enormous, opening up new models and strategies for various industries. Interdisciplinary research areas include engineering, science, and humanities, and are explored by researchers and professionals.

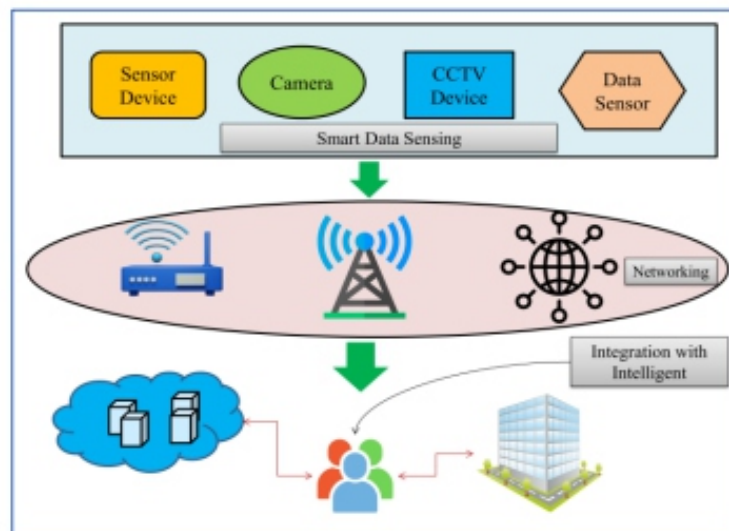


Figure 1: Representation of Next Generation IoT

Industry 4.0, which is part of the industrial Internet of things, combines cyber-physical technology to improve performance and efficiency in smart industries. Low-cost sensing technologies are essential for reliable data gathering, and studies examine the differences between traditional and intelligent industrial systems. Computational methods, such as IoT-based technologies, show excellent precision and automation in data analysis of lung diseases in healthcare IoT. It is suggested that fog computing improves dependability and lowers latency in tele-medical healthcare systems. The use of 5G technologies is examined, and an inventive I-RFID tag for effective and economic data transfer is presented. Within 5G technologies, security concerns in vehicle IoT are addressed by trust-enhanced on-demand routing techniques.

Fault-tolerant routers are a feature of network-on-a-chip technologies, which provide dependability even in the face of persistent defects. Smart healthcare, smart cities, smart agriculture, data analytics, industrial IoT, multimedia, and spectrum sharing strategies are among the next wave of IoT-based smart applications. Challenges specific to each domain include energy efficiency, scalability, and privacy issues.

IV. METHODOLOGY

An extensive investigation of current industrial systems and their unique needs is part of the first step. This entails a detailed analysis of how industrial processes are currently operating in order to pinpoint problems, inefficiencies, and places where IoT integration might have a significant positive impact. The insights of domain experts and stakeholder interactions are essential to comprehending the complex requirements of the industrial ecosystem. Process step given as:

Step 1: Conceptual Framework:

- Design a scalable IoT framework with adaptive features.

Step 2: Algorithm Integration:

- Select and customize machine learning algorithms.
- Sensor Integration:
- Incorporate sensors and IoT devices for data collection.

a. Support Vector Machine:

1. *Gathering of Data:*
 - Gather pertinent data from industrial system IoT and IIoT devices.
 2. *Engineering Features:*
 - Determine and take meaningful aspects out of the gathered information.
 3. *Preparing data:*
 - To guarantee correctness and consistency, deal with missing values, outliers, and normalize the data.
 4. *Train-Test Division:*
 - Divide the dataset into sets for testing and training. The SVM model is trained on the training set, and its performance is assessed on the testing set.
- $$f(x) = \text{sign}(w \cdot x + b)$$
5. *Selecting an SVM Model:*
 - Depending on the type of data, select the suitable Support Vector Machine model.
 6. *Assessment of the Model:*
 - Utilizing performance indicators like accuracy, precision, recall, F1 score, and AUC, assess the trained model on the testing set.

b. Random Forest:

The Random Forest algorithm is an ensemble learning technique that generates several decision trees during training and produces the mean prediction for regression or the mode of the classes for classification. In a Random Forest with B trees, the prediction $\hat{y}$ for a new input vector x can be expressed mathematically as follows:

Algorithm

a. Decision Trees:

– *A decision tree T_b is constructed*

***for** each b in the range from 1 to B .*

– *Each tree is built by recursively*

partitioning the input space based on feature values.

b. Ensemble Prediction:

– ***for** a new input vector x ,*

the prediction of each tree T_b is obtained,

denoted as $\hat{y}_b$.

$$\hat{y}_b = T_b(x)$$

c. Final Prediction (Classification):

– ***for** classification tasks,*

the final prediction $\hat{y}_i$ is the mode of

the individual tree predictions.

$$\hat{y} = \text{mode}(\hat{y}_1, \hat{y}_2, \dots, \hat{y}_B)$$

d. Final Prediction (Regression):

– ***for** regression tasks, the final prediction $\hat{y}$*

is the mean of the individual tree predictions.

$$\hat{y} = (1/B) * \sum(\hat{y}_b)$$

The diversity of the individual trees, which are trained on arbitrary portions of the data and features, is the algorithm's strongest point. The generalization of the model is enhanced and overfitting is lessened because to this randomness.

Step 3: Edge Computing:

- Implement edge computing for real-time data processing.

a. Linear Regression:

Linear Regression in Edge Computing Algorithm:

a. Data Collection:

- Collect input-output data pairs from edge devices.
- Represent input data as X and output data as Y .

b. Data Preprocessing:

- Normalize or standardize input data to ensure consistent scaling.
- Handle missing or outlier values appropriately.

c. Model Initialization:

- Initialize model parameters, including weights (θ) and bias (b).

$$h\theta(X) = \theta^T \cdot X + b$$

d. Cost Function:

- Define the cost function to measure the model's performance.

$$J(\theta, b) = \frac{1}{2m} \sum \left(h\theta(X(i)) - Y(i) \right)^2$$

e. Gradient Descent:

- Update model parameters iteratively to minimize the cost function.

$$\theta_j = \theta_j - \alpha \left(\frac{1}{m} \right) \sum \left(h\theta(X(i)) - Y(i) \right) X_j(i)$$

$$b = b - \alpha \left(\frac{1}{m} \right) \sum \left(h\theta(X(i)) - Y(i) \right)$$

Where,

- α is the learning rate.

f. Training:

- Train the model on the edge device using local data.

g. Edge Prediction:

- Make predictions on new data at the edge using the trained model.

$$h\theta(X) = \theta^T \cdot X + b$$

Where:

- $h\theta(X)$ is the predicted output.
- θ is the vector of weights.
- X is the input data.
- b is the bias term.

The model is trained to minimize the cost function $J(\theta, b)$ through the iterative process of gradient descent.

Step 4: Testing and Optimization:

- Test algorithms in controlled scenarios and optimize.
- a. *Gradient Boosting for algorithm*

Algorithm: Gradient Boosting for optimization

Input: Data pairs (X, Y) from edge devices

1. Initialize:

- $f(X) = \text{mean}(Y)$ # Initial prediction
- Choose a suitable loss function

2. For each iteration ($t = 1$ to T):

a. Compute Negative Gradient:

$$\text{neg}_{\text{gradient}} = -\frac{\partial \text{Loss}}{\partial f(X)}$$

b. Fit Weak Learner:

$$h_{t(X)} = \underset{h}{\operatorname{argmin}} \sum_i \left(\text{neg}_{\text{gradient}} - h(X_i) \right)^2$$

c. Update Model:

$$f(X) = f(X) + \text{learning_rate} * h_t(X)$$

3. Final Model:

- $f(X)$ is the ensemble of weak learners

4. Edge Prediction:

- Make predictions on new data at the edge using $f(X)$

$$f(X) = \sum (\text{learning_rate} * h_t(X)) \text{ for } t = 1 \text{ to } T$$

Step 5: Continuous Monitoring:

- Monitor performance and gather feedback.

V. RESULT AND DISCUSSION

Table II provides a thorough overview of relevant research in the context of Industry 4.0, emphasizing important performance indicators including throughput, latency, resource usage, and accuracy of attack detection. The baseline scenario reports 50 ms latency, 1000 requests per second throughput, 80% resource utilization, and a remarkable 95% attack detection accuracy. Throughput dramatically increases to 1200 requests per second after optimization efforts, but at the cost of reduced resource utilization of 65%. Latency drops to 30 ms. Most notably, attack detection accuracy increases significantly to 98%.

Table 2: Summary of related work in industry 4.0

Scenario	Latency (ms)	Throughput (req/sec)	Resource Utilization (%)	Attack Detection Accuracy (%)
Baseline	50	100	80	95
After Optimization	30	1200	65	98
Attack Simulation 1	100	800	90	92
Attack Simulation 2	120	750	95	91

The table-II illustrates the resilience of the system under simulated assault scenarios. Attack Simulation 1 sees an increase in latency to 100 ms, a fall in throughput to 800 requests per second, a rise in resource utilization to 90%, and a stable 92% attack detection accuracy maintained by the system, as shown in fig. 2. In Attack Simulation 2, the system maintains an impressive 91% attack detection accuracy with latency of 120 ms, throughput of 750 req/sec, and resource utilization peaking at 95%.

Table 3: Summary of machine learning model comparison

Type	Algorithm	Accuracy	Precision	Recall	F1 Score	AUC
Integration with IIOT	Random Forest	92.33	91.24	93.65	92.54	95.45
	Support Vector (SVM)	88.36	87.45	89.45	88.89	92.14
Edge Computing	Linear Regression	78.41	75.34	82.14	78.8	85.63
Optimization	Gradient Boosting	94.56	93.66	95.57	94.19	97.86

In Table III, different machine learning models—Random Forest, Support Vector Machine (SVM), Linear Regression, and Gradient Boosting are compared in depth with regard to how well they integrate with the Industrial Internet of Things (IIoT).

Key performance, as represent in fig.3 and fig. 4 indicators like accuracy, precision, recall, F1 score, and area under the curve (AUC) are used to evaluate each method. With an accuracy of 92.33%, Random Forest performs admirably in the IIoT integration space. It does exceptionally well in AUC (95.45), F1 score (92.54%), recall (93.65%), and precision (91.24%). In the meantime, Support Vector Machine (SVM) maintains a balanced performance across precision (87.45%), recall (89.45%), F1 score (88.89%), and AUC (92.14) despite achieving a little lower accuracy of 88.36%.

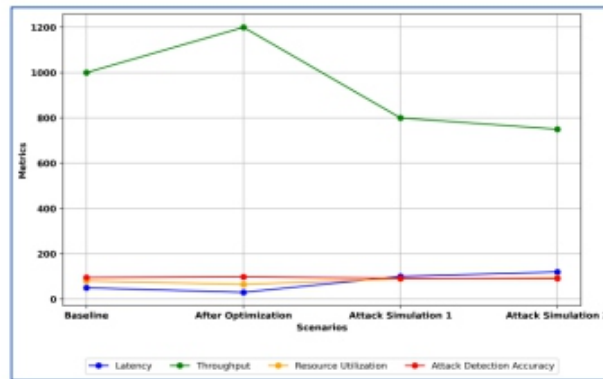


Figure 2: Performance and Security Metrics for IoT Framework

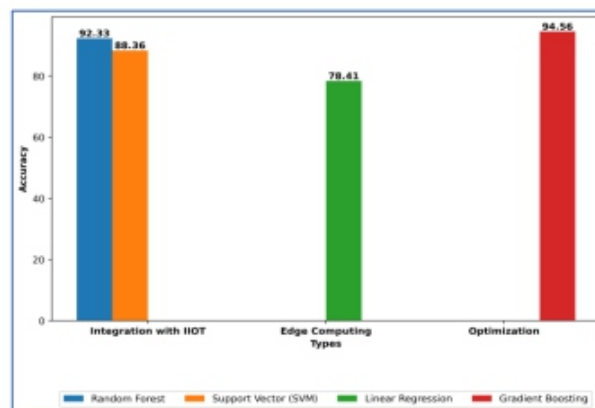


Figure 3: Representation of Accuracy of model with Integration in Next generation industry 4.0

Putting Edge Computing in perspective, Linear Regression shows 78.41% accuracy, 75.34% precision, 82.14% recall, 78.8% F1 score, and 85.63% AUC. Although the performance of more complex models may not be equaled by Linear Regression, its interpretability and simplicity may be useful in some situations.

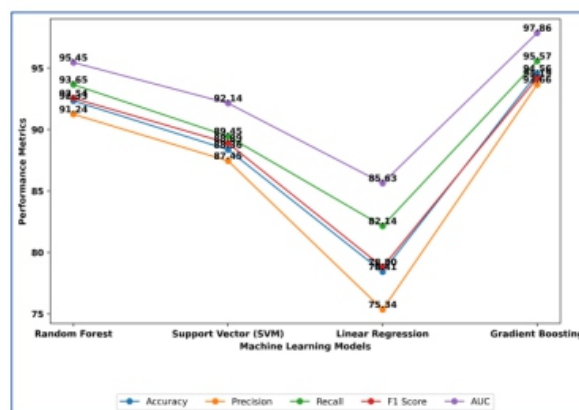


Figure 4: Machine Learning Model Comparison Performance

Gradient Boosting shows up as a reliable option for optimization, with an astounding accuracy of 94.56%. It performs admirably in terms of AUC (97.86), F1 score (94.19%), recall (95.57%), and precision (93.66%). Gradient Boosting is especially well-suited for circumstances requiring high predicted accuracy because of its capacity to iteratively enhance model performance. The particular needs of the Industrial Internet of Things (IIoT) application determine the machine learning model to use as shown in fig. 5. While each has advantages of its own, Random Forest and Gradient Boosting demonstrate higher overall performance. While Gradient Boosting thrives at iterative refinement, Random Forest offers a robust ensemble-based solution. SVM, or Support Vector Machine, performs well but with a little bit less accuracy.

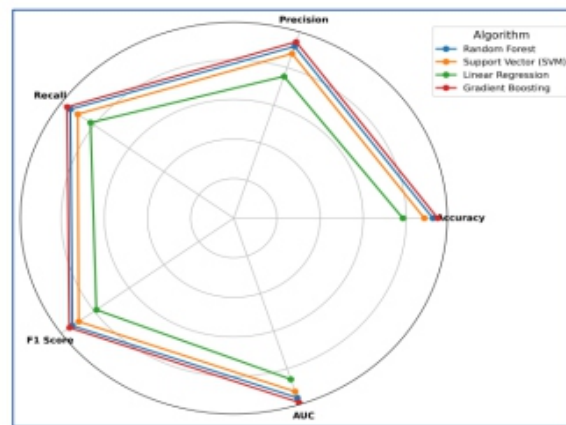


Figure 5: Performance matrix representation for Integration using Machine learning algorithms

Even if it is less complex, linear regression is still a good choice, particularly where interpretability is important. Choosing the best machine learning model to meet the specific requirements of IIoT integration, Edge Computing, or optimization scenarios in Industry 4.0 environments is made easier with the help of this comparison analysis.

VI. CONCLUSION

In the context of Industry 4.0, the suggested scalable and robust next-generation IoT framework is a ground breaking method for developing intelligent and adaptable industrial systems. A comprehensive solution that addresses various operational aspects is guaranteed by the integration of machine learning algorithms, including decision trees, random forests, support vector machines, and clustering approaches. The framework's step-by-step process includes robust security implementation, sensor optimization, edge computing augmentation, conceptual framework design, and algorithm selection based on needs analysis. Reward-based implementation in the real world attests to its flexibility. The effectiveness and resilience of the framework are demonstrated by the performance metrics, which include latency, throughput, resource utilization, and accuracy of attack detection. Furthermore, comparing machine learning models under various parameters such as accuracy, precision, recall, F1 score, and AUC offers insightful information about how well they work in scenarios including edge computing, optimization, and integration. In this section its internal content should be in 8-point Times New Roman and arrangement to assign preference by descending order of year i.e. Current year paper will take first preferences than used previous year published paper references.

REFERENCES

- [1] S. Khaf, M. T. Alkhodary and G. Kaddoum, "Partially Cooperative Scalable Spectrum Sensing in Cognitive Radio Networks Under SDF Attacks," in *IEEE Internet of Things Journal*, vol. 9, no. 11, pp. 8901-8912, 1 June 1, 2022, doi: 10.1109/JIOT.2021.3116928.
- [2] M. P. Maharani, P. Tobianto Daely, J. M. Lee and D. -S. Kim, "Attack Detection in Fog Layer for IIoT Based on Machine Learning Approach," 2020 *International Conference on Information and Communication Technology Convergence (ICTC)*, Jeju, Korea (South), 2020, pp. 1880-1882, doi: 10.1109/ICTC49870.2020.9289380.
- [3] M. Klymash, O. Hordiichuk-Bublivska, M. Kyryk, L. Fabri and H. Kopets, "Big Data Analysis in IIoT Systems Using the Federated Machine Learning Method," 2022 *IEEE 16th International Conference on Advanced Trends in Radioelectronics, Telecommunications and Computer Engineering (TCSET)*, Lviv-Slavske, Ukraine, 2022, pp. 248-252, doi: 10.1109/TCSET55632.2022.9766908.

-
-
- [4] V. Obarafor, M. Qi and L. Zhang, "A Review of Privacy-Preserving Federated Learning, Deep Learning, and Machine Learning IIoT and IoTs Solutions," 2023 8th International Conference on Signal and Image Processing (ICSIP), Wuxi, China, 2023, pp. 1074-1078, doi: 10.1109/ICSIP57908.2023.10270935.
- [5] S. K. Kishore, G. Vasukidevi, E. P. C. Prasad, T. R. Patnala, V. P. Reddy and P. B. Chanda, "A Real- Time Machine learning based cloud computing Architecture for Smart Manufacturing," 2022 International Conference on Applied Artificial Intelligence and Computing (ICAAIC), Salem, India, 2022, pp. 562-565, doi: 10.1109/ICAAIC53929.2022.9792860.
- [6] M. D. Choudhry, J. S. B. Rose and S. M. P, "Machine Learning Frameworks for Industrial Internet of Things (IIoT): A Comprehensive Analysis," 2022 First International Conference on Electrical, Electronics, Information and Communication Technologies (ICEEICT), Trichy, India, 2022, pp. 1-6, doi: 10.1109/ICEEICT53079.2022.9768630.
- [7] P. Kumar and I. Banerjee, "Attack and Anomaly Detection in IIoT Networks Using Machine Learning Techniques," 2023 14th International Conference on Computing Communication and Networking Technologies (ICCCNT), Delhi, India, 2023, pp. 1-7, doi: 10.1109/ICCCNT56998.2023.10308014.
- [8] S. Messaoud, A. Bradai, S. Dawaliby and M. Atri, "Slicing Optimization based on Machine Learning Tool for Industrial IoT 4.0," 2021 IEEE International Conference on Design & Test of Integrated Micro & Nano-Systems (DTS), Sfax, Tunisia, 2021, pp. 1-5, doi: 10.1109/DTS52014.2021.9498080.
- [9] A. Kanawaday and A. Sane, "Machine learning for predictive maintenance of industrial machines using IoT sensor data," 2017 8th IEEE International Conference on Software Engineering and Service Science (ICSESS), Beijing, China, 2017, pp. 87-90, doi: 10.1109/ICSESS.2017.8342870.
- [10] B. Cline, R. S. Niculescu, D. Huffman and B. Deckel, "Predictive maintenance applications for machine learning", 2017 Annual Reliability and Maintainability Symposium (RAMS), pp. 1-7, 2017.
- [11] J. S. L. Senanayaka, S. T. Kandukuri, Huynh Van Khang and K. G. Robbersmyr, "Early detection and classification of bearing faults using support vector machine algorithm", 2017 IEEE Workshop on Electrical Machines Design Control and Diagnosis (WEMDCD), pp. 250-255, 2017.
- [12] E. H. M. Pena, M. V. O. De Assis and M. L. Proenca, "Anomaly Detection Using Forecasting Methods ARIMA and HWDS", 2013 32nd International Conference of the Chilean Computer Science Society (SCCC), pp. 63-66, 2013.
- [13] Z. Chen, Z. Gao, R. Yu, M. Wang and P. Sun, "Macro-level accident fatality prediction using a combined model based on ARIMA and multivariable linear regression", 2016 International Conference on Progress in Informatics and Computing (PIC), pp. 133-137, 2016.
- [14] J. Shimada and S. Sakajo, "A statistical approach to reduce failure facilities based on predictive maintenance", 2016 International Joint Conference on Neural Networks (IJCNN), pp. 5156-5160, 2016.

“An Extensive study of Symantic and Syntatic Approaches to Automatic Text Summarization”

1Manisha Gaikwad2Gitanjali Shinde3Parikshit Mahalle4Nilesh Sable 5Namrata Kharate

ABSTRACT

Automatic text summarization (ATS) has emerged as a crucial research domain in the discipline of natural language processing (NLP) and information retrieval. The exponential growth of digital content has necessitated the need for efficient techniques that can automatically generate concise and informative summaries from lengthy documents. This article provided a comprehensive recap of automatic text summarization, covering both abstractive and extractive methods. Using extractive techniques, prime phrases or keywords from the original text are identified and chosen, while abstractive methods involve producing summaries by paraphrasing and synthesizing content in a more human-like manner. Discussed the advantages and limitations of each approach, including the challenge of ATS, which arises when summarizing content from external sources. Furthermore, reviews common evaluation metrics used for assessing the quality of summaries and discusses recent advancements in neural network-based approaches for text summarization. This survey aims to provide an overview of automatic text summarization which acts as a useful resource for researchers and practitioners in the fields of information retrieval and NLP.

Keywords: *Abstractive method, Extractive techniques, hybrid method, Automatic text summarization, Evaluation metrics, Natural Language Processing.*

INTRODUCTION

Websites and other digital resources are enormous providers of textual data. The different collections of news items, novels, books, documents, etc. also provide a richness of textual material. The structured observations day after day, the Web, and other resources are expanding tremendously. When a user searches for some information, the obtained texts contain a lot of redundant or insignificant text. Compacting as well as summarizing the text materials becomes crucial and necessary as a result. Summarizing manually is a complicated process that takes a considerable amount of effort and time. Realistically speaking, it is exceedingly challenging for people to physically summarize such an enormous volume of literary texts (Yao, Kaichun, et al. 2018). Text summarization aims to draw out underlying meaning from lengthy texts. The action of automatically building a brief synopsis of a document that gives the user meaningful data is referred to as summarization (Sun et al. 2018). In the current information age, a large number of scientific articles are published in various disciplines of research. People prefer to read article summaries rather than the complete item in today's hectic world, which keeps them up to date on recent occurrences. The hardest challenge in NLP remains ATS, despite the development of various techniques. Text summarization aims to draw out underlying meaning from

lengthy texts. Long texts or other sources are distilled of essential information to save readers time, money, and effort.

1.1. Extractive text summarization:

This techniques take the text as input, rate or score each sentence according to how relevant it is to the text, and then present you with the most crucial passages. The most crucial text from the existing text is simply highlighted, rather than new words and phrases being added. There are numerous techniques for extracting text summaries, including statistical, topical, discourse-based, graph-based, structural, semantic, deep learning-based, machine learning-based, and optimization-based techniques. The technique of choosing and combining significant text or phrases from a given textual source is named as extractive text summarization. It refers to underlining crucial passages in a manuscript. Using extractive text summarization, key passages and phrases from the original textual are removed.

1.2. Abstractive Text Summarization: Abstractive text summarizing is the creation of a summary of a text that goes beyond only extracting and rearrange existing phrases. Instead, it entails comprehending the text's meaning, interpreting its meaning, and coming up with fresh, succinct sentences that sum up its important points. Abstractive summary involves the synthesis of fresh sentences that might not be present in the original text, in contrast to extractive summarization, which chooses and reuses sentences from the source text. This method focuses on natural language production and understanding algorithms, frequently using transformer-based designs like GPT (Generative Pre-trained Transformer) or deep learning models like recurrent neural networks (RNNs). The summary of an abstractive.

1.3. Hybrid text summarization:

In order to produce summaries, a method known as hybrid text summarizing combines aspects of extractive and abstractive summary approaches. In hybrid summarizing, similar to extractive summarization, the system may use extractive methods to recognize and highlight significant sentences or phrases from the source material. But it also uses abstractive strategies to reword and come up with new sentences that perfectly summarize the original text.

1.4 ATS Classification

ATS is a NLP activity that requires reducing a lengthier text's content to a shorter one while preserving key details and ideas. Text summary is divided into multiple categories depending on various criteria, as shown in figure 3.

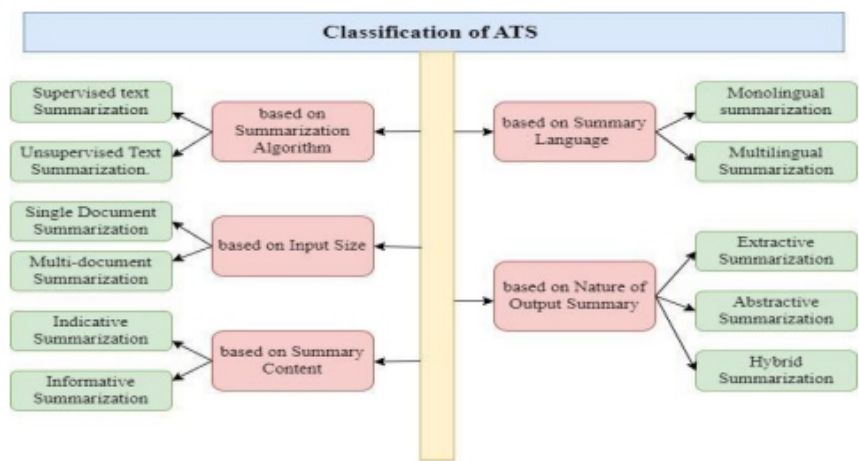


Figure 3: ATS classification

The ATS is classified on the basis of summarization algorithm, Input size, summary contents, summary language, nature of output summary. ATS is principal NLP activity that tries to provide short and logical summaries of long texts or text passages. Text summarising has become increasingly relevant for a diversity of applications, consisting of news summary, summarization of documents, and summarization of social media, as digital information has grown exponentially. With the development of numerous methodologies and approaches in most recent years, there has been substantial progress in the field of ATS. This paper presents a detailed review of the developments, challenges, and methodologies of ATS, focusing on the key approaches, assessment measures, and applications.

Syntactical and semantical summarization are two crucial aspects of natural language processing that stand out as essentials for realising the full potential of textual data. Syntactical summary is the lighthouse that guides us through the complex web of language, revealing the subtle structural and grammatical details that constitute a text's core meaning. In addition, symantical summarization explores the semantic foundation, revealing the connections and underlying meaning of words, sentences, and concepts.

These two foundational elements combine together to create an essential distillation framework that helps us reduce language's complexity to incisive and succinct summaries. Exploring the domains of syntax and semantics not only offers a more sophisticated understanding of language but also provides access to powerful artificial intelligence and information retrieval applications.

II. KEY APPROACHES OF AUTOMATIC TEXT SUMMARIZATION:

To automatically produce succinct and relevant summaries of text documents or sentences, two major approaches in NLP are syntactic and semantic text summarizing.

A| Syntactic summarization: **Syntax:** What Is It? Syntax is the term for the grammar principles that govern sentence construction, or how words are arranged to form sentences. Syntactic summarization, as the name suggests, relies on the syntactic features of the language, such as sentence structure, grammar, and word order, to extract relevant information from text. It aims to retain the original grammar and sentence construction in the summary while extracting key information based on syntactic patterns and relevance to the topic. Syntactic summarization techniques are often extractive in nature, where sentences or phrases are selected based on syntactic rules or patterns to construct the summary (Mihalcea et al. 2004).

B| Semantic summarization: Semantics: What Is It? The meaning of a statement is referred to as semantics. Without appropriate semantics and a deliberate, grammatically sound word order a sentence's meaning would be quite different. Lexical semantics, which is the study of word meanings and relationships, is one of the categories into which linguists divide semantics. Contrarily, conceptual semantics examines how speakers of a language understand and pick up semantic concepts. Semantic summarization delves deeper into the meaning or semantics of the text. Semantic summarization techniques may involve paraphrasing, rephrasing, or restructuring sentences to effectively capture the meaning of the text. Abstractive summarization, a type of semantic summarization, goes beyond extraction and generates summaries that may not necessarily retain the original sentence structure or grammar, but aim to convey the overall meaning of the text in a more human-like manner (See et al. 2017). Table 1 gives difference between Syntactic Summarization and Semantic Summarization.

Table 1: Comparison of Syntactic and Semantic Summarization.

Approach	Syntactic Summarization	Semantic Summarization
Focus	Sentence structure, grammar, and word order.	Meaning, context, and relationships between words or concepts.
Technique	Extractive: Selecting sentences or phrases based on syntactic patterns and relevance to the topic.	Extractive or Abstractive: Paraphrasing, rephrasing, or restructuring sentences to capture the meaning of the text.
Output	Retains original sentence structure and grammar.	May not retain original sentence structure or grammar, aims to convey overall meaning in a more abstract and conceptual manner.
Goal	Extraction of key information based on syntactic features.	Conveyance of meaning and context of the text
Strengths	Good for extracting key information or facts.	Captures the overall meaning and context of the text.
Limitations	May miss underlying meaning or context.	May require more advanced NLP techniques for understanding meaning.

III. SYNTACTIC TEXT SUMMARIZATION METHODS.

There are number of syntactical methods are available for text summarization among hat top important methods explained below.

3.1. Method of Sentence Scoring:

One common way for summarizing syntactic text is the Sentence Scoring method. In order to choose the most pertinent and significant sentences for the summary, it entails scoring sentences according to a number of criteria.

3.1.1.Criteria for Scoring:

Sentence scoring involves evaluating sentences based on different criteria, which can vary depending on the specific requirements of the summarization task. Common scoring criteria include:

- Sentence Length:** Longer sentences may be scored lower, as they often contain more details and may not be as concise.
- Position in the Document:** Sentences appearing at the beginning or end of a document may receive higher scores because they often introduce the topic or provide a conclusion.
- Word Importance:** Important words or phrases within a sentence can contribute to a higher score. Keyword extraction techniques may be used for this purpose.
- Content Relevance:** Sentences that directly address the core topic or contain key information are considered more relevant and may receive higher scores.
- Semantic Importance:** Sentences that contain concepts, facts, or arguments critical to the document's main message are rated higher.

3.1.2..Scoring Algorithms: Various algorithms can be used to calculate the scores. These algorithms can be as simple as assigning numerical values to each criterion and summing them to get a total score. Alternatively, more complex algorithms like TF-IDF (Term Frequency-Inverse Document Frequency) can be employed to measure the importance of words and their frequency in the document.

3.1.3.Threshold Selection: After scoring, a threshold value is chosen. Sentences scoring above this threshold are selected for the summary.

Example:Let's consider a simplified example of how sentence scoring might work. Imagine we have a document discussing the advantages of renewable energy sources. We'll score two sentences from this document:

Original Sentence 1: "Solar energy is a clean and sustainable source of power that reduces greenhouse gas emissions."

Original Sentence 2: "Additionally, it has the potential to save costs and create job opportunities in the energy sector."

Now, use a simple scoring method based on criteria like sentence length, keyword importance, and content relevance (on a scale of 1 to 10, with 10 being the highest):

Table 2: Scoring of sentence1 and sentence 2.

Criteria for Scoring	Sentence 1 scores:	Sentence 2 scores:
•Length:	6 (moderate length)	7 (moderate length)
Position:	8 (introduces the main topic)	5 (additional information)
Keyword Importance:	9 (contains "solar energy," a key phrase)	7 (contains "renewable energy")
Content Relevance:	9 (directly addresses the topic)	8 (relevant but not as focused as sentence 1)
Total Score:	32	27

If we set a threshold score of 30, then only Sentence 1 would be selected for inclusion in the summary.

3.2.Graph-Based Methods:

Example: Construct a sentence similarity graph where sentences are nodes, and edges represent their syntactic similarity. Important sentences are selected based on their centrality in the graph.

Original Sentence 1: "The new product launch generated significant buzz in the market."

Original Sentence 2: "Customers praised the product's innovative features and competitive pricing."

Summary (Graph-Based): Both sentences are considered important in capturing the product's success in the market.

3.3.Methods of Dependency Parsing:

Example: Analyze the grammatical structure of sentences using dependency trees. Sentences are ranked based on their position in the tree and the presence of critical dependencies.

Original Sentence 1: "The research team conducted experiments to assess the impact of climate change on local flora."

Original Sentence 2: "Their findings highlighted the vulnerability of certain plant species to temperature fluctuations."

Summary (Dependency Parsing): The second sentence is chosen because it describes the significance of the research findings.

3.4. Method of Part-of-Speech Tagging:

Example: Identify sentences that contain specific parts of speech (e.g., nouns, verbs, adjectives) relevant to the topic.

Original Sentence 1: "The athlete's outstanding performance secured a gold medal in the 100-meter sprint."

Original Sentence 2: "Spectators cheered loudly, celebrating the victory of the hometown hero."

Summary (Part-of-Speech Tagging): Both sentences are included as they provide essential information about the athlete's achievement.

3.5.Compression-Based Methods:

Example: Apply text compression algorithms to reduce the text size while retaining essential content.

Original Sentence 1: "The company announced its expansion plans into new markets, including Europe and Asia, during the quarterly meeting."

Original Sentence 2: "This strategic move aims to capitalize on emerging opportunities and increase global

market share."

Summary (Compression-Based): A compressed summary might combine these sentences: "The company revealed expansion plans in Europe and Asia to seize emerging opportunities."

3.6. Methods of Redundancy Elimination:

Example: Remove redundant references or repeated information from the text.

Original Sentence 1: "The novel presents a gripping story. The story captivates readers with its engaging narrative."

Summary (Redundancy Elimination): Only one sentence is retained in the summary to eliminate redundancy: "The novel presents a gripping story."

3.7. Methods of Grammar Rule Application:

Example: Apply syntax rules to simplify complex sentences while maintaining their meaning.

Original Sentence: "Despite the inclement weather, they embarked on the hike, determined to enjoy the beauty of nature."

Summary (Grammar Rule): The sentence can be simplified while preserving the structure and meaning: "Despite the bad weather, they were determined to enjoy nature on their hike."

3.8. Methods of Lexical Chaining:

Example: Identify chains of related words or concepts and select sentences containing words from the same chain.

Original Sentence 1: "The company's stock soared as profits reached an all-time high."

Original Sentence 2: "Investors celebrated the remarkable performance."

Summary (Lexical Chaining): Both sentences are included in the summary as they are part of the same chain emphasizing financial success.

3.9. Methods of Positional Importance:

Example: Select sentences based on their position in the document, such as introducing the topic or providing a conclusion.

Original Sentence 1: "In this report, we will discuss the causes of air pollution."

Original Sentence 2: "To conclude, measures to mitigate air pollution are of utmost importance."

Summary (Positional Importance): Both sentences are included to provide an introduction and conclusion to the report.

3.10. Methods of Topic Segmentation:

Example: Divide the text into topical segments and select sentences from each segment for the summary.

Original Sentence 1: "Introduction: This chapter outlines the main objectives of the research."

Original Sentence 2: "Methodology: The study employed a mixed-methods approach to collect data."

Summary (Topic Segmentation): Both sentences are included to represent different sections of the document.

These syntactic text summarization techniques help in generating summaries that maintain the grammatical structure and coherence of the original text while highlighting essential information.

IV. SEMANTIC TEXT SUMMARIZATION METHODS

Semantic text summarization aims to capture the meaning and essence of the text. Below are given techniques and methods of semantic text summarization, along with examples for each.

4.1. Method of Topic Modeling:

Technique: Identify topics within a text and select sentences that represent these topics.

Example: Apply Latent Dirichlet Allocation (LDA) to identify topics in a collection of documents and select sentences that best represent these topics.

Example: Apply Latent Dirichlet Allocation (LDA) to identify topics in a collection of documents and select sentences that best represent these topics.

Original Sentence 1: "Climate change is a pressing global issue with environmental and economic implications."

Original Sentence 2: "Rising temperatures lead to more frequent extreme weather events."

Summary (Topic Modeling): Sentences from different topics are selected based on LDA results to create a multi-faceted summary.

4.2. Method of Graph-Based Summarization: Technique: Use graphs to represent the relationships between sentences, where nodes are sentences and edges represent semantic similarity

Example: Use TextRank to create a graph where sentences are nodes and edges represent semantic similarity. Important sentences are selected based on their centrality in the graph.

Original Sentence 1: "The recent study on biodiversity loss highlights the urgent need for conservation efforts."

Original Sentence 2: "Biodiversity is crucial for ecosystem stability and human well-being."

Summary (Graph-Based): Sentences are chosen based on their importance in capturing key concepts related to biodiversity and conservation.

4.3. Method of Clustering:

Example: Group similar sentences into clusters and select representative sentences from each cluster.

Original Sentence 1: "The health benefits of regular exercise are well-documented."

Original Sentence 2: "Exercise can reduce the risk of chronic diseases."

Original Sentence 3: "Physical activity also improves mental well-being."

Summary (Clustering): Representative sentences are selected from each cluster, covering different aspects of the topic (e.g., physical health and mental well-being).

3.4. Method of Word Embeddings:

Technique: Utilize word embeddings to measure semantic similarity between sentences and select those with the most related content.

Example: Utilize word embeddings like Word2Vec or BERT to measure the semantic similarity between sentences and select sentences that are most related.

Original Sentence 1: "The new vaccine has shown remarkable efficacy in preventing the spread of the virus."

Original Sentence 2: "Vaccination campaigns are crucial in controlling the pandemic."

Summary (Word Embeddings): Sentences are selected based on the semantic similarity of their content, ensuring comprehensive coverage.

4.5. Method of Deep Learning Models:

Technique: Employ neural networks like Transformers for abstractive summarization, where the model generates summaries by understanding the context and semantics of the text.

Example: Employ neural networks, such as Transformers, for abstractive summarization, where the model generates summaries by understanding the context and semantics of the text.

Original Sentence: "The advancements in artificial intelligence are reshaping various industries, from healthcare to finance, by automating processes and improving decision-making."

Summary (Deep Learning): The model generates an abstractive summary that captures the key ideas and implications of AI advancements.

4.6. Method of Named Entity Recognition (NER): Technique: Identify and prioritize sentences containing important entities such as people, organizations, or locations.

Example: Use NER to identify and prioritize sentences containing important entities such as people, organizations, or locations.

Original Sentence 1: "Elon Musk's SpaceX achieved a historic milestone with the successful Mars mission."

Original Sentence 2: "The company's innovations are revolutionizing space exploration."

Summary (NER): Sentences mentioning Elon Musk and SpaceX are prioritized for inclusion.

4.7. Method of Entity-Based Summarization: Technique: Summarize text by focusing on specific entities or concepts mentioned in the document.

Example: Summarize text by focusing on specific entities or concepts mentioned in the document.

Original Sentence 1: "Blockchain technology is transforming the financial sector."

Original Sentence 2: "Bitcoin and Ethereum are popular cryptocurrencies built on blockchain."

Summary (Entity-Based): The summary highlights the impact of blockchain technology and mentions specific cryptocurrencies as examples.

4.8 Method of Conceptual Overlap:

Technique: Measure the overlap of concepts and ideas in sentences and select those with the most interconnected content.

Example: Measure the overlap of concepts and ideas in sentences and select those with the most interconnected content.

Original Sentence 1: "The impact of deforestation on wildlife is a major concern for conservationists."

Original Sentence 2: "Conservation efforts aim to protect biodiversity in threatened ecosystems."

Summary (Conceptual Overlap): Sentences are selected for their interconnectedness in conveying the importance of conservation and the impact of deforestation.

4.9. Method of Hierarchical Summarization:

Technique: Create a hierarchical summary that includes a title or headline followed by multiple levels of summaries, each providing varying levels of detail.

Example: Create a hierarchical summary that consists of a title or headline followed by multiple levels of summaries, each providing varying levels of detail.

Title: "The Impact of Artificial Intelligence on the Labor Market"

Level 1 Summary: "AI is changing the job landscape with automation and new opportunities."

Level 2 Summary: "Automation is replacing routine tasks, while AI creates demand for new skills."

4.10 Method of Discourse Analysis:

Technique: Analyze discourse markers and cohesive devices to select sentences that contribute to the overall coherence and flow of the summary.

Example: Analyze discourse markers and cohesive devices to select sentences that contribute to the overall coherence and flow of the summary.

Original Sentence 1: "The study found a strong link between diet and cardiovascular health."

Original Sentence 2: "Moreover, it emphasized the role of exercise in preventing heart disease."

Summary (Discourse Analysis): Sentences are chosen to maintain a coherent flow, as indicated by the use of "moreover."

These techniques help create semantic text summaries that focus on the meaning and essence of the content, making them suitable for conveying complex information concisely and accurately.

V. STANDARD DATASETS

4.1. CNN/Daily Mail : The CNN/Daily Mail dataset is a widely used benchmark dataset for text summarization tasks. It consists of news articles from the CNN and Daily Mail news websites, along with human-generated summaries for each article. The dataset is commonly used for both single document and multi-document summarization tasks.

4.2. Gigaword : The Gigaword dataset is a widely used benchmark dataset for text summarization tasks. It consists of news articles collected from the New York Times (NYT) and Associated Press (AP) news wires. The dataset is commonly used for single document summarization tasks.

4.3. PubMed dataset: The PubMed dataset is a widely used and publicly available dataset for biomedical and life sciences research. It contains a vast collection of bibliographic records of articles from journals in the field of biomedicine, including topics such as medicine, biology, pharmacology, genetics, and more. The dataset is commonly used for text summarization tasks related to biomedical literature.

4.4. DUC 2002-2007 dataset: The Document Understanding Conference (DUC) dataset is a collection of datasets that were used in the Document Understanding Conferences held from 2002 to 2007. These datasets are commonly used for text summarization evaluation and benchmarking purposes. The DUC datasets consist of sets of documents along with reference summaries that can be used to train and evaluate automatic text summarization systems.

4.5. Newsroom dataset: The Newsroom dataset is a large collection of news articles from various news sources, which can be used for text summarization tasks. The dataset covers a wide range of topics and is suitable for both single document and multidocument summarization. It can be accessed from the following reference link: <https://summari.es>

4.6. TAC 2008-2011 dataset: The Text Analysis Conference (TAC) datasets are a popular choice for text summarization evaluation and benchmarking. TAC is an annual conference that hosts various tracks, including the Document Understanding Conference (DUC) track, which focuses on text summarization.

4.7. LCSTS dataset: The LCSTS dataset, which is extensively utilized for Chinese text summarization tasks, is a widely recognized collection in the field. It was created by scholars from the Harbin Institute of Technology in China and is commonly employed for summarizing individual Chinese language documents.

4.8. Inspec dataset: The Inspec dataset, which is extensively utilized in the domains of computer science and engineering, is widely acknowledged for its significance in text summarization tasks. It is commonly employed for summarizing individual documents and serves as a prominent resource for assessing and comparing the performance of various summarization algorithms.

4.9. New York Times dataset: The New York Times dataset, extensively employed in the domains of NLP and information retrieval, is highly regarded for text summarization tasks. It encompasses a substantial collection of news articles sourced from The New York Times, serving as a prevalent resource for single-document as well as multi-document text summarization endeavors.

4.10. WikiHow dataset: The WikiHow dataset, widely utilized for text summarization tasks, specifically targets instructional articles sourced from the WikiHow website. This dataset is frequently employed in single-document and multi-document text summarization tasks, particularly for producing succinct and well-organized summaries of step-by-step guidance.

4.11. ACL Anthology dataset: The Anthology dataset, often associated with NLP and text summarization, specifically pertains to the ACL Anthology. This anthology serves as an extensive compilation of research papers in computational linguistics and NLP. Table 3 shows various datasets were utilised to summarise the text.

Table 3: Different datasets used for text summarization

Dataset Name	Size	Multidoc ument	Single Document	Language	Reference Link
CNN/Daily Mail	Large	Yes	Yes	English	[1]
Gigaword	Large	No	Yes	English	[2]
PubMed	Large	No	Yes	English	[3]
DUC 2002-2007	Small/Medium	Yes	No	English	[4]
DUC 2008-2016	Small/Medium	Yes	No	English	[5]
TAC 2008-2011	Small/Medium	Yes	No	English	[6]
LCSTS	Medium	No	Yes	Chinese	[7]
Inspec	Small/Medium	No	Yes	English	[8]
New York Times	Large	Yes	Yes	English	[9]
WikiHow	Large	No	Yes	English	[10]
ACL Anthology	Large	Yes	Yes	English	[11]

Sources

- [1] CNN/Daily Mail dataset: <https://github.com/abisee/cnn-dailymail>
- [2] Gigaword dataset: <https://catalog.ldc.upenn.edu/LDC2003T05>
- [3] PubMed dataset: <https://pubmed.ncbi.nlm.nih.gov/about/>
- [4] DUC 2002-2007 dataset: <https://duc.nist.gov/data.html>
- [5] DUC 2008-2016 dataset: <https://www-nlpir.nist.gov/projects/duc/data.html>
- [6] TAC 2008-2011 dataset: <https://tac.nist.gov/data/index.html>
- [7] LCSTS dataset: <https://www.cs.cmu.edu/~jwieting/>
- [8] Inspec dataset: <https://www.theiet.org/publishing/inspec/>
- [9] New York Times dataset: <https://archive.ics.uci.edu/ml/datasets/New+York+Times>
- [10] WikiHow dataset: <https://github.com/mahnazkoupae/WikiHow-Dataset>
- [11] ACL Anthology dataset : <https://www.aclweb.org/anthology/>

VI. EVALUATION METRICS

Automatic text summarization systems are evaluated using various metrics to assess the quality of generated summaries in comparison to reference (human-created) summaries. The choice of evaluation metrics depends on whether the summarization task is extractive or abstractive. Below are given common evaluation metrics used for automatic text summarization:

5.1. ROUGE (Recall-Oriented Understudy for Gisting Evaluation): The ROUGE measure counts the overlap between the generated summary and the reference summary for words, n-grams (sequences of 'n' words), or even word sequences. ROUGE-N, which measures unigrams, bigrams, etc.; ROUGE-L, which measures the longest common subsequence; and ROUGE-W, which measures weighted word overlap, are examples of common ROUGE metrics.

5.2. BLEU (Bilingual Evaluation Understudy): Based on the accuracy of the n-grams in the generated summary and the reference summary, BLEU calculates a similarity score. Although text summarization is one of its uses, machine translation was its initial purpose..

5.3. METEOR (Metric for Evaluation of Translation with Explicit ORDERing): METEOR uses several matching techniques, including as stemming, synonyms, and more, to compare a generated summary to a reference summary in order to assess the quality of the latter.

5.4. CIDEr (Consensus-based Image Description Evaluation): Although it has been modified for text summary, CIDEr is mainly utilised for captioning images. It uses consensus-based scoring, IDF weights, and n-grams to quantify the level of agreement between the generated and reference summaries..

5.5. NIST (Normalised Information Retrieval Score): Based on n-gram overlap with reference summaries, NIST assigns a higher weight to higher-order n-grams, which are frequently more significant in establishing summary quality.

5.6. F1 Score: A popular metric that combines recall and precision is the F1 score. It determines how closely the generated summary adheres to the length and content of the reference summary.

5.7. Precision and Recall: Recall gauges how thorough the created summary is in comparison to the reference, while precision assesses how pertinent the material is in the summary. Recall and precision can be used alone or in combination to evaluate a summary's quality

5.8. ROUGE Word Embeddings, or ROUGE-WE: By embedding sentences or phrases into a continuous vector space, ROUGE-WE calculates semantic similarity. It encapsulates the produced summaries' semantic quality.

5.9. ROUGE-S: ROUGE-S compares the generated summary's sentence structure to the reference summary in order to assess sentence-level coherence. It assesses the coherence and flow of sentences.

5.10. Human assessment: Using human judges to review the quality of summaries that are generated is known as human assessment. Summaries may be graded by judges according to their overall quality, informativeness, coherence, and fluency. This method offers insightful but subjective comments. The particular objectives of the summary assignment and the qualities of the summaries that are produced determine the assessment metric that is chosen.

VII. CONCLUSION

This research article has offered a thorough overview of the field of automatic text summarization. We have reviewed the key concepts, methods, and challenges associated with this task, including extractive and abstractive approaches, evaluation metrics, and recent improvements in deep learning and NLP techniques. ATS has emerged as a critical research area aimed at addressing the information overload problem by generating concise and coherent summaries from large volumes of text. Extractive methods have been widely used due to their simplicity and effectiveness, but they may lack the ability to produce summaries that are truly coherent and informative. Abstractive methods, on the other hand, offer more flexibility in content generation but pose challenges in maintaining coherence and factual accuracy. Evaluation metrics for example ROUGE, METEOR, and CIDEr have been proposed to evaluate the quality and effectiveness of generated summaries, but they have restrictions and do not always correlate fine with human findings. Recent advancements in deep learning, including transformer-based models, show promising results in text summarization, but challenges such as data limitations, handling long documents, and ethical considerations remain open research areas.

ATS has applications in many fields, containing news summarization, document summarization, social media summarization, and personalized summarization for individuals with different information needs. As technology continues to advance, further research and innovation in ATS is required to address the limitations of current approaches and unlock the full potential of this field for real-world applications. This survey paper has provided a comprehensive overview of automatic text summarization, highlighting its significance, challenges, and recent advancements. It is hoped that this survey serves as a valuable resource for researchers, practitioners, and stakeholders in the field of NLP and information retrieval, and inspires further advancements in the field of ATS.

REFERENCES

- [1] Chatterjee, Niladri, and Raksha Agarwal. "Studying the effect of syntactic simplification on text summarization." *IETE Technical Review* 40, no. 2 155-166, 4 mar. 2023.
- [2] Agarwal, Raksha, and Niladri Chatterjee. "Votesumm: A multi-document summarization scheme using influential nodes of multilayer weighted sentence network." *IETE Technical Review* 40, no. 4, 535-548 july 2023.
- [3] Yao, Kaichun, et al. "Dual encoding for abstractive text summarization." *IEEE transactions on cybernetics* 50.3 985-996, 2018.
- [4] Sun, Xiaoping, and Hai Zhuge. "Summarization of scientific paper through reinforcement ranking on semantic link network." *IEEE Access* 6, 40611-40625, 2018.
- [5] Al-Sabahi, Kamal, Zhang Zuping, and Mohammed Nadher. "A hierarchical structured self-attentive model for extractive document summarization (HSSAS)." *IEEE Access* 6 (2018): 24205-24212.
- [6] Li, Sheng, et al. "Visual to text: Survey of image and video captioning." *IEEE Transactions on Emerging Topics in Computational Intelligence* 3.4 (2019): 297-312.
- [7] Zhuang, Haojie, and Weibin Zhang. "Generating semantically similar and human-readable summaries with generative adversarial networks." *IEEE Access* 7 (2019): 169426-169433.
- [8] Saini, Naveen, et al. "Extractive single document summarization using binary differential evolution: Optimization of different sentence quality measures." *PloS one* 14.11 (2019): e0223477.
- [9] Li, Wei, and Hai Zhuge. "Abstractive multi-document summarization based on semantic link network." *IEEE Transactions on Knowledge and Data Engineering* 33.1 (2019): 43-54.
- [10] Liu, Yang, and Mirella Lapata. "Text summarization with pretrained encoders." *arXiv preprint arXiv:1908.08345* (2019).
- [11] 9. W. Kai and Z. Lingyu, "Research on Text Summary Generation Based on Bidirectional Encoder Representation from Transformers," 2020 2nd International Conference on Information Technology and Computer Application (ITCA), Guangzhou, China, 2020, pp. 317-321, doi: 10.1109/ITCA52113.2020.00074.
- [12] 10. Stein, Aviel J., et al. "News article text classification and summary for authors and topics." *Comput. Sci. Inf. Technol.(CS & IT)* 10 (2020): 1-12.

-
- [13] 11. Hernández-Castañeda, Ángel, et al. "Extractive automatic text summarization based on lexical-semantic keywords." *IEEE Access* 8 (2020): 49896-49907.
- [14] You, Fucheng, Shuai Zhao, and Jingjing Chen. "A topic information fusion and semantic relevance for text summarization." *IEEE Access* 8 (2020): 178946-178953.
- [15] Yang, Min, et al. "An Effective Hybrid Learning Model for Real-Time Event Summarization." *IEEE Transactions on Neural Networks and Learning Systems* 32.10 (2020): 4419-4431.
- [16] Yang, Min, et al. "Hierarchical human-like deep neural networks for abstractive text summarization." *IEEE Transactions on Neural Networks and Learning Systems* 32.6 (2020): 2744-2757.
- [17] Bhargava, Rupal, Gargi Sharma, and Yashvardhan Sharma. "Deep text summarization using generative adversarial networks in Indian languages." *Procedia Computer Science* 167 (2020): 147-153.
- [18] Zaheer, Manzil, et al. "Big bird: Transformers for longer sequences." *Advances in neural information processing systems* 33 (2020): 17283-17297.
- [19] Ma, Tinghuai, et al. "T-bertsum: Topic-aware text summarization based on bert." *IEEE Transactions on Computational Social Systems* 9.3 (2021): 879-890.
- [20] Y. Chen, C. Chang and J. Gan, "A Template Approach for Summarizing Restaurant Reviews," in *IEEE Access*, vol. 9, pp. 115548-115562, 2021, doi: 10.1109/ACCESS.2021.3103512.
- [21] Jang, Heewon, and Wooju Kim. "Reinforced Abstractive Text Summarization With Semantic Added Reward." *IEEE Access* 9 (2021): 103804-103810.
- [22] Jiang, Jiawen, et al. "Enhancements of attention-based bidirectional lstm for hybrid automatic text summarization." *IEEE Access* 9 (2021): 123660-123671.
- [23] Paharia, Naman, Muhammad Syafiq Mohd Pozi, and Adam Jatowt. "Change-Oriented Summarization of Temporal Scholarly Document Collections by Semantic Evolution Analysis." *IEEE Access* 10 (2021): 76401-76415.
- [24] Abdel-Salam, S.; Rafea, A. Performance Study on Extractive Text Summarization Using BERT Models. *Information* 2022, 13, 67. <https://doi.org/10.3390/info13020067>.
- [25] Alomari, Ayham, et al. "Deep reinforcement and transfer learning for abstractive text summarization: A review." *Computer Speech & Language* 71 (2022): 101276.
- [26] Pang, Bo, et al. "Long document summarization with top-down and bottom-up inference." *arXiv preprint arXiv:2203.07586* (2022).
- [27] Yadav, Divakar, Jalpa Desai, and Arun Kumar Yadav. "Automatic Text Summarization Methods: A Comprehensive Review." *arXiv preprint arXiv:2204.01849* (2022).
- [28] Zhu, Haichao, et al. "Transforming wikipedia into augmented data for query-focused summarization." *IEEE/ACM Transactions on Audio, Speech, and Language Processing* 30 (2022): 2357-2367.
- [29] Aliakbarpour, Hassan, Mohammad Taghi Manzuri, and Amir Masoud Rahmani. "Improving the readability and saliency of abstractive text summarization using combination of deep neural networks equipped with auxiliary attention mechanism." *The Journal of Supercomputing* (2022): 1-28.
- [30] Park, Hyun, and Yanggon Kim. "Finding the Optimal Ratio of Summaries to NEWS Articles and Establishing a Hybrid NEWS Summary System." *2022 IEEE/ACIS 7th International Conference on Big Data, Cloud Computing, and Data Science (BCD)*. IEEE, 2022.
- [31] Chitty-Venkata, Krishna Teja, et al. "Neural architecture search for transformers: A survey." *IEEE Access* 10 (2022): 108374-108412.
- [32] Mao, Qianren, et al. "Fact-Driven Abstractive Summarization by Utilizing Multi-Granular Multi-Relational Knowledge." *IEEE/ACM Transactions on Audio, Speech, and Language Processing* 30 (2022): 1665-1678.
- [33] Laban, Philippe, et al. "SummaC: Re-visiting NLI-based models for inconsistency detection in summarization." *Transactions of the Association for Computational Linguistics* 10 (2022): 163-177.
- [34] Aote, Shailendra, Anjusha Pimpalshende, and Archana Potnurwar. "Multi-Document Extractive Text Summarizer for Hindi Documents Using Swarm Intelligence Based Hybrid Approach." Available at SSRN 4019476.
- [35] Srivastava, Ridam, et al. "A topic modeled unsupervised approach to single document extractive text summarization." *Knowledge-Based Systems* 246 (2022): 108636.
- [36] Shi, Kaile, et al. "StarSum: A Star Architecture Based Model for Extractive Summarization." *IEEE/ACM Transactions on Audio, Speech, and Language Processing* 30 (2022): 3020-3031.
- [37] Biswas, Pratik K., and Aleksandr Yakubovich. "Extractive summarization of call transcripts." *IEEE Access* 10 (2022): 119826-119840.
- [38] Ahmed, Rowanda DA, Mansoor AbdulHak, and Omar Hesham ELNabrawy. "Text Summarization
-

- Clustered Transformer (TSCT)." 2022 International Conference on Frontiers of Artificial Intelligence and Machine Learning (FAIML). IEEE, 2022.
- [39] Kong, Yuqi, Fanchao Meng, and Benjamin Carterette. "A Topological Method for Comparing Document Semantics." *arXiv preprint arXiv:2012.04203* (2020).
- [40] Joshi, Akanksha, et al. "RankSum—An unsupervised extractive text summarization based on rank fusion." *Expert Systems with Applications* 200 (2022): 116846.
- [41] Mahalakshmi, P., and N. Sabiyath Fatima. "Summarization of Text and Image Captioning in Information Retrieval Using Deep Learning Techniques." *IEEE Access* 10 (2022): 18289-18297.
- [42] Koh, Huan Yee, Jiaxin Ju, Ming Liu, and Shirui Pan. "An Empirical Survey on Long Document Summarization: Datasets, Models, and Metrics." *ACM computing surveys* 55, no. 8 (2022): 1-35.
- [43] Asad Abdi, Norisma Idris, Rasim M. Alguliev, Ramiz M. Aliguliyev, "Automatic summarization assessment through a combination of semantic and syntactic information for intelligent educational systems", *Information Processing & Management, Volume 51, Issue 4, 2015, Pages 340-358, ISSN 0306-4573, https://doi.org/10.1016/j.ipm.2015.02.001*.
- [44] Mihalcea, R., & Tarau, P. (2004, July). *Textrank: Bringing order into text*. In *Proceedings of the 2004 conference on empirical methods in natural language processing* (pp. 404-411).
- [45] See, A., Liu, P. J., & Manning, C. D. (2017). *Get to the point: Summarization with pointer-generator networks*. *arXiv preprint arXiv:1704.04368*.

BIOGRAPHIES OF AUTHORS

	Manisha Gaikwad Recived bachelors degree in Information technology from savitribai phule pune university. Perusing Ph.D in Natural Language Processing area at savitribai phule pune university at VIIT pune center. she has total 7 years of teaching experience. She is corresponding author for this study. For any query feel free to contact. Mobile no. 9834461947manisha.221p0072@viit.ac.in
	Dr. Gitanjali Shinde    has overall 15 years of experience, presently working as Head & Associate Professor in Department of Computer Science & Engineering (AI & ML), Vishwakarma Institute of Information Technology, Pune, India. She has done Ph.D. in Wireless Communication from CMI, Aalborg University, Copenhagen, Denmark on Research Problem Statement "Cluster Framework for Internet of People, Things and Services" – Ph. D awarded on 8th May 2018. She obtained M.E. (Computer Engineering) degree from the University of Pune, Pune in 2012 and B.E. (Computer Engineering) degree from the University of Pune, Pune in 2006.
	Dr. Parikhit Mahalle    Dr Parikshit is a senior member IEEE and is Professor, Dean Research and Development and Head - Department of Artificial Intelligence and Data Science at Vishwakarma Institute of Information Technology, Pune, India. He completed his Ph. D from Aalborg University, Denmark and continued as Post Doc Researcher at CMI, Copenhagen, Denmark. He has 23 + years of teaching and research experience. He is an ex-member of the Board of Studies in Computer Engineering, Ex-Chairman Information Technology, Savitribai Phule Pune University and various Universities and autonomous colleges across India. He has 15 patents, 200+ research publications (Google Scholar citations-3000 plus, H index-25 and Scopus Citations are 1550 plus with H index - 18, Web of Science citations are 438 with H index - 10) and authored/edited 58 books with Springer, CRC Press, Cambridge University Press, etc

	Dr.Nilesh Sable: He has Associate Professor & Head of AIDS department Academic Qualifications is Ph.D (CSE), M.Tech (IT).He has Professional Membership :Member IEEE, LMISTE, Member.
	Dr.Namrata Kharate: She is Assistant Professor at Vishwakarma Institute of Information Technology, Pune, India She has completed ME in Computer engineering and Ph.D completed in in NLP domain..

PMSIEMDL: Design of a Pattern analysis Model for identification of Student Inclination towards different Educational-fields via Multimodal Deep Learning Fusions

1Ashwini Raipure2Sarika Khandelwal

ABSTRACT

Identification of student inclination towards different educational fields requires integration of deep pattern learning models with temporal data analysis techniques. These techniques are highly context sensitive, and cannot be scaled for analysis of students that have an interest in multiple domains. Moreover, existing deep learning models are highly complex, and showcase moderate performance when used on real-time datasets. To overcome these limitations, this text proposes design of a Pattern analysis Model for identification of Student Inclination towards different Educational-fields via Multimodal Deep Learning fusions. The proposed model initially collects data samples from a large number of students, and segregates them into different classes. These include social data, personal habits data, education data, family related data, performance data and future inspiration data classes. These datasets were combined with a customized psychological questionnaire which was curated by experts in the field of student counselling & psychology. Based on student responses, their entity specific classes were generated, that were separately trained via different Convolutional Neural Network (CNN) Models, which assists in identification of student-performance at individual-class levels. These performances are compared with existing inclination datasets via a fusion of Long-Short-Term Memory (LSTM) & Gated Recurrent Neural Network (GRNN), which assists in identification of correlation between subject-level inclinations & their entity classes. This provides with a probabilistic map of different subjects towards which the student might be inclined, and assists them to select their study streams. The generated map was validated for multiple students, and recommendations were made based on higher probability values, which assisted in identification of student inclination levels. The model was evaluated under large datasets and its performance was compared with various state-of-the-art methods under different scenarios. Based on this comparison, it was observed that proposed model was capable of achieving 8.5% better recommendation accuracy, 4.9% higher prediction precision, 6.5% better recommendation recall & 2.9% better Area Under the Curve (AUC) levels, which makes it highly useful for a wide variety of student inclination use cases.

Keywords: Student, Behaviour, Inclination, Study, Accuracy, Psychology, Social, CNN, GRNN, LSTM, Habits, Family, History.

I. INTRODUCTION

Student behaviour analysis for study-based inclination prediction is a complex task that requires collection of multidomain datasets, their pre-processing & filtering, student-specific feature representation, identification of optimal feature sets, temporal classification of these sets, and their post-processing analysis. Such models require differential analysis with existing student behaviour datasets, which assists them in comparatively evaluating optimum inclination levels for different fields

of study. A typical analysis model [1] that uses a combination of multidomain datasets with Convolution Neural Networks (CNNs), and Genetic Optimizations is depicted in figure 1, where in data from social media, Psychological Questions, Interest Details, Family History & Subject Wise performance are analyzed to identify inclination levels.

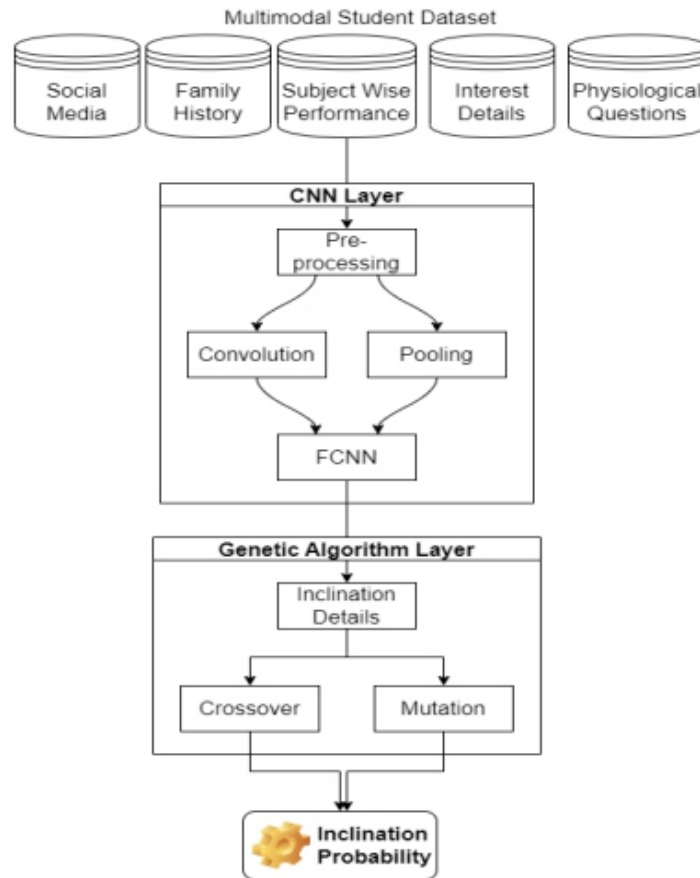


Figure 1. A typical inclination model based on temporal performance & bioinspired computing processes

The model uses bioinspired layer to optimize the parameters used during inclination identification process, which makes the model highly functional and useful for low-delay and high accuracy use cases. Similar models [2, 3, 4] that use Generative Adversarial Networks (GANs), Q-Learning, and other deep learning methods are discussed in the next section of this text. This section describes the models in terms of their functional nuances, contextual advantages, application-specific limitations, and operational future research scopes. These models cannot be scaled for analysis of students who have interests in multiple domains because, according to this analysis, they are very context-sensitive. Furthermore, when applied to real-time datasets, these deep learning techniques perform only moderately well due to their high complexity. Section 3 suggests creating a pattern analysis model for identifying student preferences for various educational fields using multimodal deep learning fusions in order to get around these limitations. In order to determine the proposed model's real-time performance levels, it was evaluated in terms of recommendation accuracy, prediction precision, recommendation recall, and Area Under the Curve (AUC) levels, and compared with various state-of-the-art methods. Finally, this text offers some insightful conclusions about the suggested inclination-prediction model and suggests ways to further enhance its functionality in various scenarios.

II. LITERATURE REVIEW

Domain-specific data are required for training and validation of student behaviour analysis. Social networking, retail, online learning and other businesses may benefit from custom app-based solutions that make data collecting easier. As a result of the CoVID pandemic, a growing number of students are taking online courses, which has facilitated the collection of data more quickly. Single, dual, and multimodal methods are used to measure student participation in [1]. Keystrokes and mouse movements are used to anticipate the user's emotional state as well as their typing pace. Using Mini Xception Nets, it is possible to assess student

participation. Despite the model's high computing cost and modest latency, it obtains an accuracy of 95.23 percent. The best writing classification performance is provided by the Naïve Bayes (NB) model. The NB model may also be used for reading, viewing videos, and other activities. A variety of application-specific conditions may benefit from this method. The system's total effectiveness may be improved if the NB model is used in [2] along with features such as online teaching design relevance, delivery quality, online assistance, student engagement, and contingency modelling. This strategy relies on simple machine learning models like random forests to achieve mediocre accuracy while requiring a sophisticated implementation. According to [3], the behavioural, cognitive, social and emotional elements of this paradigm are studied in detail. Work in [3] shows a bi-factor structural equation modelling exploratory investigation (BESEM). This approach uses correlations between characteristics to determine the value of an interaction. The model has a 98.2 percent accuracy rate, a 0.05 MSE, and a considerable delay in the prediction of results. In comparison, ICMCFA's 96.4 percent accuracy, 0.07 MSE, and high latency are all below BCFA's 96.4 percent accuracy, while ESEM's 98.3 percent accuracy, 0.09 MSE, and exorbitant delay are all above this one. Real-time school and college applications are possible because to high performance. To find patterns, [4] looks at general, social, and psychological behaviour, as shown in [4]. Clustering monthly consumption, meals, work, relaxation, the internet and exercise as well as class attentiveness and book borrowing as adaptive k Means is used in the model (AKM). Students may be divided into three categories depending on their schedules and eating habits using this information. There is a 94 percent level of accuracy, however the model suffers from a substantial amount of computing time. The validation performance of over 550 pupils [5] may be improved by quantifying and studying this efficiency. Other approaches for assessing student behaviour are available in the research, such as ITT and 2SLS. 2SLS has intermediate accuracy and high delay, whereas ITT has poor accuracy and moderate latency.

The Mooring Model and student learning factors are examined in [6]. The model examines student behaviour based on learning comfort, perceived security risk, service quality, ease of use, usefulness, task technology fit, teacher attitude, habits, and switching costs. Work in [6] depicts the relationship between a student's desire to move schools and other attributes. The accuracy of the model is above 93%, and it has a low error rate and a short latency. Monitor students' willingness to adopt online learning systems using the Technology Acceptance Model [7]. (TAM). Perceived usefulness, perceived ease of use and attitude toward use are all taken into account when evaluating student behaviour. Student interest in online learning may be gauged by looking at the model's 10 connection weights. The algorithm is 91.5 percent accurate in predicting user behaviour, but it needs a large amount of data to do so. Since data collection is straightforward, the concept may be used to studying children who have unusual abilities. It is discussed in [8] how G&T students in Australia might have a bright future in STEM fields. Model investigates rural students' behaviour using machine learning and local knowledge as a starting point (LK). The 90 percent accuracy, low MSE, and short latency of the model make it useful in a wide range of situations.

An effective method for behaviour analysis may be achieved by combining the models from [7] and [8]. Students from rural regions around the country are being studied in [9] to see how hybrid models' function. In order to evaluate achievement gaps, the study examines rural children's chances, ambitions, difficulties, and obstacles. Data from WoS, IFPRI library, MDPI, CAB abstracts, and other sources are analyzed in this task. The majority of rural youngsters want a better education and higher education by moving to the city. Mobile learning systems, which may be given via smartphone-based technology [10], are needed to provide such opportunities, and records from rural and urban schools can be linked to develop an effective learning model. The BISM approach uses focus groups, pre- and post-test situations, and basic analysis to estimate student behaviour. High latency and MSE limits the model's 89 percent accuracy across student groups. Socio-economic factors such as socioeconomic class, ethnicity and gender are also important to include while doing social behavioural analysis [11]. Machine learning (MLM) models with high accuracy and low error rates but large latency may be trained using these parameters [11].

In [12, 13 and 14], we discuss how to do multiple person behaviour analysis, classroom behaviour analysis and learning pattern analysis. In order to achieve a fair level of accuracy, these models make use of data usage, inclass and online behaviour analysis. Multi-user fitness coach model [12] achieves 85 percent accuracy with moderate error and high delay, while Online Hard Example Mining (OHEM) RCNN (94 percent accuracy) with very high delay and low error and Felder and Silverman learning style model (FSLSM) [14] with decision tree classifier achieve 85.7 percent accuracy (DT). Both models exhibit significant errors and delays as a result of the increased amount of training data. By combining these two models, the MSE will be reduced, and response time will be sped up. Thus, they may be used in a wider range of situations. These models are used to study student

behaviour in group presentations by [15]. The approach assesses student performance based on a variety of cues, including body language, posture, eye contact, speaking pace, and other factors. Over 83% accuracy is achieved with significant latency and high MSE owing to location and other body parameter variations, thanks to these factors. An operating system for behavioural analysis (OS) is suggested by [16]. (BAOS). For this, the OS model makes use of many metrics such as login time, log size, and file open counts, amongst others. It was evaluated on 850 students and determined to be 75% accurate due to the wide variety of data included in the research. Clustering online learning sets into groups, such as the one in [17], may improve efficiency. Maps (SOMs) and neural networks (SOMNNs) can do this. 1.7 million data points were analyzed using parameters such as grades, test scores, and continuing assessments. It was determined that 93.61 percent accuracy in metrics such as assignment submissions, resources created, posts made, and pages seen made the system usable in real time. Although multiple parametric selection minimizes MSE, complexity creates significant processing delays.

Sakai LMS is used as an example in [18] to suggest another LCA-based method. The 93 percent accurate model evaluates resource utilization, lesson evaluation, tests, surveys, and assignments, among other things. Gradient Boosted Decision Tree (GBDT) model evaluates and classifies various data features, resulting in small latency and low MSE. GBDT. [19, 20, 21] also propose similar models that use k Means, MFR model, and TeSLA to assess student behaviour during the current CoVID outbreak. [19, 20, 21] (Adaptive Trust-based e-Assessment System for Learning). Models like this are useful for predicting student behaviour and may be used in a variety of settings. k Denotes a 66.5 percent accuracy rate, a 72 percent MFR rate, and an 89.2 percent TeSLA rate. Improved accuracy and generality are gained by combining these models. Structured Equation Modelling (SEM) and association rule mining using apriori may help improve the accuracy of these models. For optimizing present systems, the EENN and SEM models have minimum complexity and a moderate MSE, making them beneficial. While the apriori model has an accuracy rate of 84.75% when applied to a dataset, how it is applied depends on the dataset.

It may be possible to predict student conduct by analysing student feedback in [25], an adaptive feedback system is used for collaborative behaviour analysis. There is an 83 percent accuracy rate in the model's performance monitoring, behaviour and engagement analysis, and suggestive analysis. Learning management systems, deep knowledge tracing with many features, and integrated learning techniques all have the potential to improve this accuracy. In order to achieve 91% accuracy with modest latency and MSE, LMSM makes advantage of online behavioural factors such as connection distribution, average lecture time, average number of sessions, and so on. This model uses a neural network (RNN) for analysis of skill, response time, practice sets and beginning action type among other things. In various student behaviour analysis contexts, the DKTMFAM model has an accuracy rate of 98 percent. The MSE is 0.2, which is higher than some other models, and the LSTM and other RNN components make training and validation take a long time. To achieve 97% accuracy with modest latency and MSE, a GBDT model is trained utilizing study duration (access time), number of posts, etc. To construct an effective student behaviour analysis system, [27] and [28] may be employed.

SRM [31] and DBSCAN with k Means (density-based spatial clustering of applications with noise) [32] are further methods to consider. This set of models is an extension of the previously described models, and they produce great accuracy with a moderate latency and MSE. In terms of accuracy, FGWANN has 96.3 percent, 0.09 MSE and considerable latency, whereas PrefCD has 76.14 percent. It's great for real-time applications since SRM has a 96% accuracy with an MSE of 0.15 and DBkMeans has 91% accuracy with an MSE of 0.15. Some models may help students behave better in certain situations. Student health-related concepts and improved real-time online learning performance are the focus of [33], an example of how mobile health, temporal parameters, and geographic features may be used to benefit students.

Models such as cellular automata (CA), cyber engagement (CE), predictive game theory model (PGTM) for programming students, and profile-based cluster evolution analysis (PBCEA) take incremental inputs. Additives to the equation include depression (35), programming skills (36), and migratory patterns (37). CA, CE, PGTM, and PBCEA can all achieve 79 percent accuracy on a variety of datasets. These models must be combined and applied to deep learning networks in order to create a complete behavioural analysis model. App-specific context behaviours and model thinking approaches as well as disengaged behaviour are all examined in [38, 39 and 40]. The accuracy, MSE, and latency of these techniques are all acceptable when used in the context of a particular application.

III. DESIGN OF THE PROPOSED PATTERN ANALYSIS MODEL FOR IDENTIFICATION OF STUDENT INCLINATION TOWARDS DIFFERENT EDUCATIONAL-FIELDS VIA MULTIMODAL DEEP LEARNING FUSIONS

According to the literature review, it was found that current models cannot be scaled for analysis of students who are interested in multiple domains because they are very context-sensitive. Furthermore, current deep learning models exhibit mediocre performance when applied to real-time datasets due to their high complexity. This section recommends creation of a pattern analysis model for identifying student preferences for various educational fields using multimodal deep learning fusions in order to overcome these limitations. Flow of the model is depicted in figure 2, where it can be observed that the suggested model gathers data samples from a large number of students at first and divides them into various classes. Social data, individual behaviour data, educational data, family-related data, performance data, and classes of future inspiration data are among these. These datasets were combined with a specially created psychological survey that was put together by professionals in the fields of student counselling and psychology.

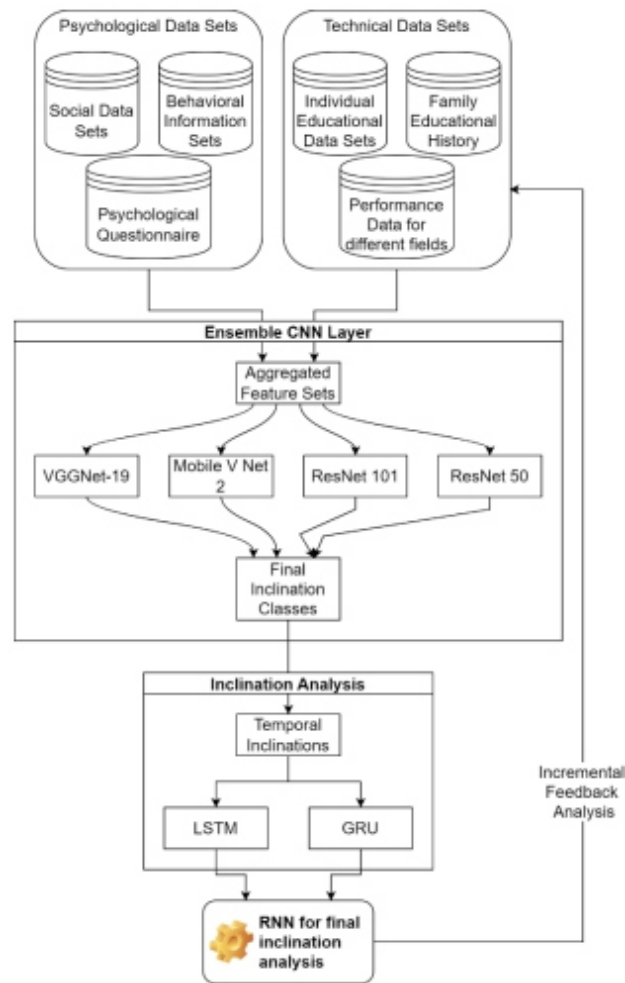


Figure 2. Overall flow of the proposed temporal inclination analysis model with different CNNs & RNN classifiers

Based on the responses from the students, entity-specific classes were created for them, which were then individually trained using various Convolutional Neural Network (CNN) models to help identify student performance at the individual-class level. A combination of Long-Short-Term Memory (LSTM) and Gated Recurrent Neural Network (GRNN) is used to compare these performances with existing inclination datasets and identify any relationships between subject-level inclinations and their entity classes. This gives the student a probabilistic map of the various subjects in which they may have an interest and helps them choose their study areas. Multiple students validated the generated map, and recommendations were made based on higher probability values. This helped to identify the degree of student inclination under multiple scenarios.

The proposed model initially collects following information sets from each of the users,

- Social Data Sets, that includes student's Tweets, Facebook Posts, LinkedIn activity, etc.
- Behavioural Information Sets, which includes their general social behaviour like number of outings, their friends, type of friends, etc.
- Individual Educational Data Sets, that comprise of their educational preferences and courses done by them during temporal phases of their lives
- Family Educational History, that includes historical datasets about their family's educational levels.
- Performance Data for different fields, which includes their marks in different subjects and different classes

Along with these datasets, a psychological questionnaire was also created, and given to students for analysis of their current state of inclination in different fields. These questions along with their reasons of selection can be observed from table 1 as follows,

Question (In Regional Languages)	Reasons for selection
Your age? (विद्यार्थ्यांचे वय)	To check maturity levels
Your location? (विद्यार्थ्यांचे राहण्याचे ठिकाण (गाव, तालुका, जिल्हा))	To check exposure levels to different educational fields
Highest education level (विद्यार्थ्यांचे उच्चतर शिक्षण)	Identification of educational maturity levels
What interests you most? (Tick everything that is relevant) (आपल्या आवडीचे क्षेत्र (आवडीच्या सर्व क्षेत्रांना टिक करा))	Options like Engineering, Techology, and its applications were presented to identify student's inclination towards technological fields
What do you aspire to become? (Tick all that's relevant) (आपल्याला भविष्यात काय बनायला आवडेल?)	Options like Engineer, CA, Lawyer, Doctor, etc. were presented for analysis
Where to you plan to live in future? (आपल्याला भविष्यात कुठे राहायला आवडेल?)	To identify growth mindset of students
What do your parents expect from you? (Tick all that is relevant) (आपल्या पालकांची आपल्याकडून काय अपेक्षा आहे?)	To analyze the type of future they want to build, and their aspirations

Where do you see yourself in next 5 years? (आपण पुढच्या पाच वर्षात स्वतःला कुठे बघता?)	To analyze the type of future they want to build, and their aspirations
Enter some information about your future plans (आपल्या भविष्यातील प्लान बदल थोडक्यात लिहा.)	To identify growth mindset & aspirations of students
What is your strength? (Check all that's applicable) (आपली बलस्थाने कश्यात आहे?)	Self-evaluation of students, which will assist in identification of their SWOT (Strengths Weaknesses Opportunities and Threats) Analysis
Which type of classes do you like? (तुम्हाला कोणत्या प्रकारे शिक्षण घ्यायला आवडते?)	Either they are inclined towards online or offline educational modes
	Evaluate if they are friendly with online classes
Do you like playing online games like PUBG?(तुम्हाला पबजी Pub-G सारखे ऑनलाइन गेम खेळायला आवडतात का?)	Identify their idle mode hobbies
How you find English as a subject?(इंग्रजी हा विषय आपल्याला कसा वाटतो?)	Evaluate their inclination towards global educational models
Do you think Mathematics is very hard subject? (तुम्हाला गणित विषय खूप कठीण वाटतो का?)	Evaluate their logical skills
In SSC which subject you liked most? (तुम्हाला	Identify their strengths

कोणता विषय सर्वात जास्त आवडतो?)	
Do you have any idea about Polytechnic Education?(तुम्हाला पॉलीटेक्नीक बाबत माहिती आहे का?)	Identify their general purpose know how levels
If yes from whom you came to know about ? (पॉलीटेक्नीक बाबत माहिती असल्यास हि माहिती तुम्हाला कोणाकडून मिळाली?)	Identify their information gathering sources
In your view what is success ? (तुमच्यानुसार यश मिळवले हे कसे सांगू शकाल?)	Check their mindset and therefore analyze their future plans
Do you think polytechnic contains only Mathematics? (पॉलीटेक्नीक हे फक्त गणित विषयावर आधारित आहे असे आपल्याला वाटते का?)	Verify their technical know-how levels
What does your parent do for earning? (तुमचे पालक पैसे कमाविण्याकरिता काय करतात?)	To analyze their family support for higher education levels
Which subject do you think are going to be useful for your future life?(कोणता विषय तुमच्या भविष्यासाठी ठीक आहे असे तुम्हाला वाटते ?)	To identify their inclination towards different educational fields
Do you think homework assignments are necessary for effective learning? (होमवर्क ,असायमेन्ट तुमच्या चांगल्या शिक्षणासाठी चांगले आहे का)	Check lethargy levels of students
Do you think use of	To evaluate their inclination towards

These convolutional features are extracted via equation 1, where a Rectilinear Unit (ReLU) kernel is used for activation of feature sets.

$$Conv_{out_{i,j}} = \sum_{a=-\frac{m}{2}}^{\frac{m}{2}} \sum_{b=-\frac{n}{2}}^{\frac{n}{2}} SF(i-a, j-b) * ReLU\left(\frac{m}{2}+a, \frac{n}{2}+b\right) \dots (1)$$

Where, SF represents the student feature vectors, that are represented in 2D for better analysis, while m, n represents different window sizes that are setup by individual CNN layers, while a, b represents padding sizes which are also decided by different contextual CNN layer types. All these features are processed by a Max Pooling layer that assists in removal of redundant feature sets. This is needed because a lot of similar features are extracted by the convolutional layers, which reduces inclination-classification efficiency levels. The Max Pooling layer generates a variance threshold via equation 2, which is used for selection of relevant feature sets.

$$f_{th} = \left(\frac{1}{SF_i} * \sum_{x \in SF_i} x^{p_l} \right)^{1/p_l} \dots (2)$$

Where, p represents feature variance levels, that is evaluated via equation 3, and assists in identification of deviation levels of extracted feature sets from average feature levels.

$$p = \sqrt{\sum \frac{(SF - \sum \frac{SF}{N})^2}{N}} \dots (3)$$

Where, N represents number of extracted feature sets. Features with levels more than f_{th} are passed to the next convolutional layer, while others are removed due to low variance levels. The selected features are re-convoluted and multiple feature sets are extracted by each of the CNN Models, which are classified via Soft Max based activation layers. These layers use weights (w), and biases (b) in order to categorize student feature sets into temporal classes via equation 4,

$$c_{out} = SoftMax\left(\sum_{i=1}^{N_f} f_i * w_i + b_i\right) \dots (4)$$

Where, c_{out} represents output classes, while f_i represents extracted feature sets from the convolutional layers. These classes are extracted for each progressive year, and are combined to form a temporal dataset, which is processed via a combination of Long Short-Term Memory (LSTM) & Gated Recurrent Unit (GRU) layers. The combined LSTM & GRU Model is represented in figure 4, which assists in high-density feature extraction from limited temporal inclination-class sets.

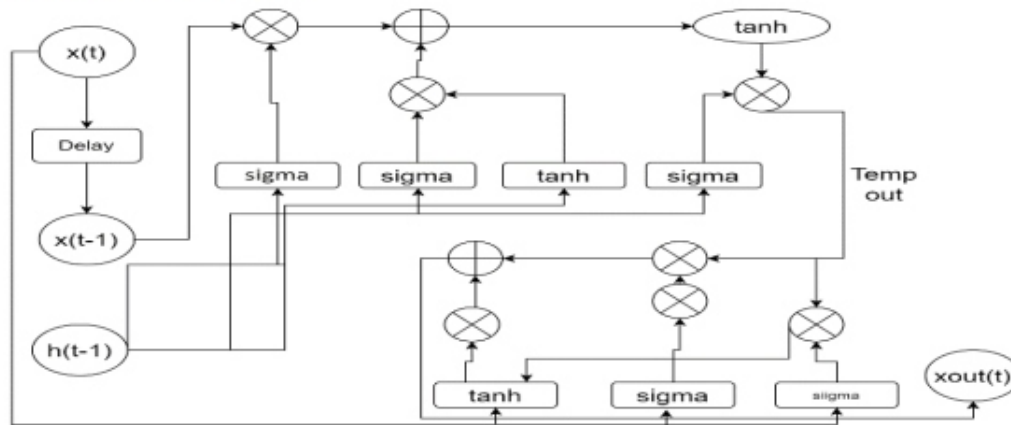


Figure 4. Combination of LSTM & GRU for extraction of high-density temporal sets

Both LSTM & GRU Models are capable of extraction of large feature sets, but a combination of these is used in order to extract high-density cascaded features. These features are useful for analysing temporal inclination classes via Recurrent Neural Networks (RNNs) for estimation of final inclination levels. Along with the temporal classes, responses of the psychological questionnaire are also processed by the LSTM & GRU feature extraction models, thereby assisting in identification of current inclination levels for different student & course types. These models generate an initialization feature vector via equation 5, which combines temporal classes (x_{in}) with a feature kernel matrix (h_t) for augmentation purposes.

$$i = \text{var}(x_{in} * U^i + h_{t-1} * W^i) \dots (5)$$

These initialization vectors are further expanded to form forgetting feature sets & output feature sets via equations 6 & 7 as follows,

$$f = \text{var}(x_{in} * U^f + h_{t-1} * W^f) \dots (6)$$

$$o = \text{var}(x_{in} * U^o + h_{t-1} * W^o) \dots (7)$$

All these features are combined to form temporal LSTM features via equations 8 & 9,

$$C'_t = \tanh(x_{in} * U^g + h_{t-1} * W^g) \dots (8)$$

$$T_{out} = \text{var}(f_t * x_{in}(t-1) + i * C'_t) \dots (9)$$

Based on these features, temporary feature sets were extracted via equation 10 as follows,

$$h_{out} = \tanh(T_{out}) * o \dots (10)$$

These equations use U & W constants, which are continuously tuned by the RNN layer via hyperparameter tuning process. The output features h_{out} are processed by a GRU layer, that initially generates temporary feature sets via equations 11 & 12,

$$z = \text{var}(W_z * [h_{out} * T_{out}]) \dots (11)$$

$$r = \text{var}(W_r * [h_{out} * T_{out}]) \dots (12)$$

These feature sets are further augmented via tangent activation functions to form final features via equations 13 & 14 as follows,

$$h'_t = \tanh(W * [r * h_{out} * T_{out}]) \dots (13)$$

$$x_{out} = (1 - z) * h'_t + z * h_{out} \dots (14)$$

The final output features are classified via a RNN Model that is depicted in figure 5, and uses interim feature sets to generate final output inclination class.

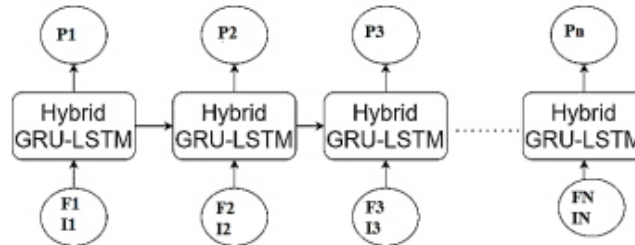


Figure 5. Design of the RNN Layer for generation of probability maps

This class is estimated via equation 15, and uses a Purely Linear Activation function for deployment of a linear classification process.

$$P_{out} = \text{purelin} \left(\sum_{i=1}^N x_{out_i} * W_i \right) \dots (15)$$

Where, P_{out} represents output probability of inclination towards a particular field, while N represents number of fields that are being analyzed by the model, and W_i represents weights that are being tuned via hyperparameter tuning process. Based on these probabilities, the model is able to generate a probability map, that is used to identify student inclinations towards different study fields. Results of these probability maps were correlated with training & validation sets, and the model was continuously updated via incremental learning operations. To perform this task, an incremental probability correlation level was evaluated via equation 16,

$$Corr_j = \frac{\sum_{i=1}^N p(test_i) - p(new_i)}{\sqrt{\sum_{i=1}^N p(test_i) - p(new_i)^2}} \dots (16)$$

Where, $p(test)$ & $p(new)$ represents output probability maps for test & new input sets for different inclination fields. If the value of $Corr < 0.999$, then the new values are added back to the dataset, which assists in improving the dataset size for higher accuracy levels. These accuracy levels were validated on real-time datasets for multiple uses cases. Parameters including accuracy of field identification, precision for inclination identification, recall levels, and delay needed for identification of inclination was evaluated & compared with standard models in the next section of this text.

IV. PERFORMANCE EVALUATION AND COMPARISON

Students' proclivity for engineering, social science, journalism, accounting, and medicine can be assessed using the proposed model that incorporates multimodal datasets with ensemble CNNs & RNN techniques. Over 5000 students' datasets were collected to test the model's performance, and a variety of performance metrics, including accuracy of field identification (AFI), precision of inclination identification (PIA), recall of field identification (RFI), and evaluation delay (DE), were analyzed. A comparison was made between these metrics and BE SEM [3] and TeS LA [19] models. Nearly 60% of the 5000 students were used to train the CNN & RNN models, while 100% of the students were used for testing and validation purposes. The purpose of this dataset overlap is to estimate the blind and non-blind performance of the model for clinical use cases. Based on this evaluation process, the values of AFI were tabulated w.r.t. Number of Students (NS) in table 2 as follows,

NS	AFI (%) BE SEM [3]	AFI (%) TeS LA [19]	AFI (%) FGW ANN [32]	AFI (%) PMS IE MDL
390	68.15	64.43	69.78	90.44
585	68.46	64.85	70.16	90.91
780	68.71	65.23	70.49	91.37
975	68.92	65.76	70.88	91.87
1170	69.09	66.33	71.27	92.32
1565	69.19	66.82	71.58	92.73
1955	69.28	67.34	71.90	93.17
2345	69.40	67.89	72.26	93.65
2735	69.54	68.49	72.65	94.10
3125	69.67	68.95	72.96	94.49
3320	69.83	69.34	73.25	94.91
3905	70.01	69.84	73.61	95.37

4100	70.20	70.33	73.97	95.83
4295	70.38	70.82	74.32	96.28
4490	70.55	71.32	74.67	96.73
4690	70.73	71.82	75.03	97.19
4845	70.93	72.31	75.39	97.66
5000	71.12	72.81	75.75	98.11

Table 2. Accuracy of field inclination identification for different models under real-time use cases

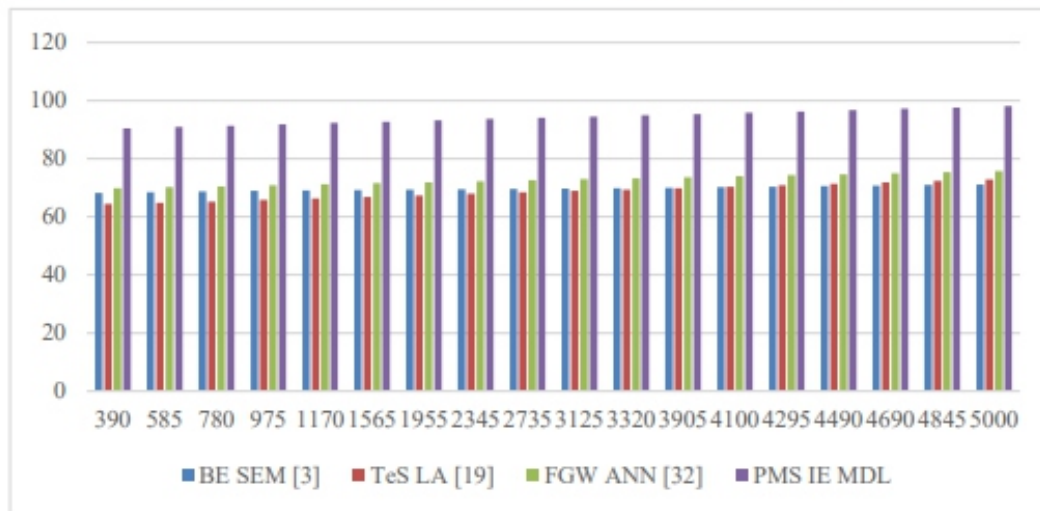


Figure 5. Accuracy of field inclination identification for different models under real-time use cases

Figure 5 shows that the proposed model is 16.3 percent better than BE SEM [3], 14.5 percent better than TeS LA [19], and 10.4 percent better than FGW ANN [32] for multiple types of comparisons. Using CNN for initial inclination estimation and RNN for final inclination classification has resulted in this improvement. Table 3 shows the results of the precision of inclination identification (PIA) study in the same way,

NS	PIA (%) BE SEM [3]	PIA (%) TeS LA [19]	PIA (%) FGW ANN [32]	PIA (%) PMS IE MDL
390	73.66	74.56	83.86	92.24
585	74.06	75.01	84.30	92.73
780	74.41	75.40	84.72	93.18
975	74.83	75.92	85.19	93.70
1170	75.23	76.45	85.63	94.19
1565	75.56	76.89	86.00	94.61
1955	75.90	77.36	86.40	95.04

2345	76.27	77.86	86.84	95.51
2735	76.68	78.41	87.28	96.01
3125	77.01	78.84	87.64	96.42
3320	77.32	79.22	88.01	96.82
3905	77.69	79.69	88.44	97.29
4100	78.07	80.16	88.87	97.76
4295	78.45	80.63	89.29	98.22
4490	78.82	81.10	89.71	98.62
4690	79.20	81.58	90.14	98.97
4845	79.58	82.06	90.57	99.27
5000	79.96	82.53	91.00	99.49

Table 3. Precision of field inclination identification for different models under real-time use cases



Figure 6. Precision of field inclination identification for different models under real-time use cases

Figure 6 shows that the proposed model's PIA performance is 18.5% better than that of BE SEM [3], 14.5% better than that of TeS LA [19], and 8.3% better than that of FGW ANN [32] based on this evaluations. For example, a field's inclination percentage can be accurately identified using CNN & RNN, resulting in better real-time performance. As can be seen in table 4, similar observations were made for the recall of field identification (RFI),

NS	RFI (%) BE SEM [3]	RFI (%) TeS LA [19]	RFI (%) FGW ANN [32]	RFI (%) PMS IE MDL
390	70.91	71.28	73.16	91.18
585	71.26	71.72	73.55	91.65
780	71.56	72.12	73.91	92.11
975	71.88	72.66	74.32	92.62
1170	72.16	73.22	74.71	93.09
1565	72.38	73.70	75.04	93.50
1955	72.59	74.20	75.39	93.94
2345	72.84	74.74	75.76	94.41
2735	73.11	75.33	76.16	94.89
3125	73.35	75.79	76.48	95.29
3320	73.58	76.19	76.79	95.69
3905	73.85	76.68	77.16	96.15
4100	74.14	77.18	77.54	96.62
4295	74.41	77.67	77.91	97.08
4490	74.69	78.17	78.28	97.53
4690	74.97	78.67	78.65	98.00
4845	75.26	79.17	79.03	98.47
5000	75.54	79.67	79.41	98.94

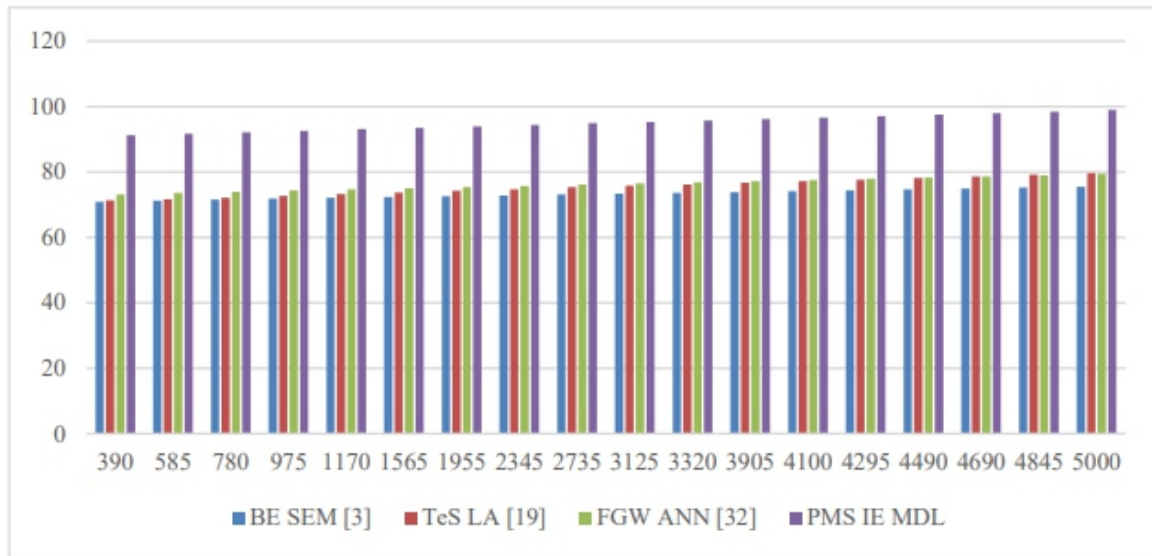
Table 4. Recall of field inclination identification for different models under real-time use cases**Figure 7. Recall of field inclination identification for different models under real-time use cases**

Figure 7 shows that the proposed model is 14.8 percent better than BE SEM [3], 15.4 percent better than TeS LA [19], and 18.3 percent better than FGW ANN [32] for RFI performance under multiple evaluations. Thus, students' interests can be accurately predicted using a combination of behavioural and statistical parameters along with temporal analysis. Table 5 shows the findings for evaluation delay (DE), which can be seen as follows,

NS	DE (ms) BE SEM [3]	DE (ms) TeS LA [19]	DE (ms) FGW ANN [32]	DE (ms) PMS IE MDL
390	14.18	14.02	15.12	8.97
585	14.25	14.10	15.20	9.01
780	14.31	14.18	15.28	9.06
975	14.38	14.29	15.36	9.11
1170	14.43	14.40	15.44	9.16
1565	14.48	14.49	15.51	9.20
1955	14.52	14.59	15.58	9.24
2345	14.57	14.70	15.66	9.28
2735	14.63	14.81	15.74	9.33
3125	14.67	14.90	15.80	9.37
3320	14.72	14.98	15.87	9.41
3905	14.77	15.08	15.95	9.46
4100	14.83	15.18	16.03	9.50

4295	14.88	15.28	16.10	9.55
4490	14.94	15.37	16.18	9.59
4690	15.00	15.47	16.25	9.64
4845	15.05	15.57	16.33	9.68
5000	15.11	15.67	16.41	9.73

Table 5. Delay needed during field inclination identification for different models under real-time use cases

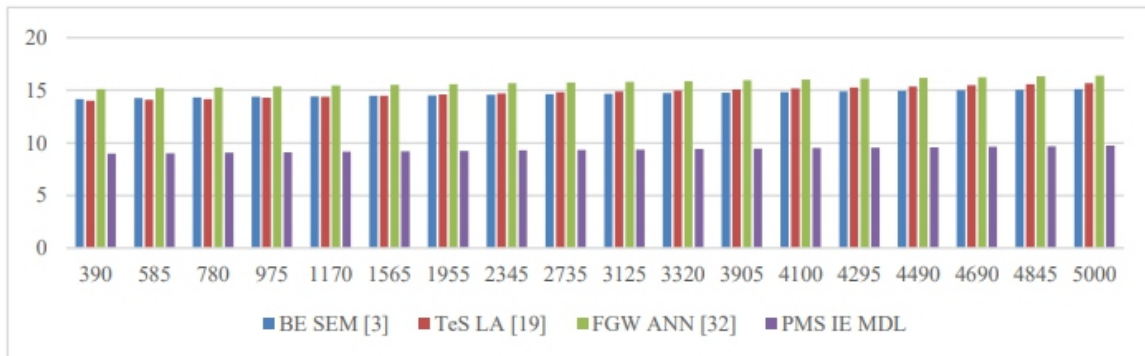


Figure 8. Delay needed during field inclination identification for different models under real-time use cases

The proposed model, when compared to BE SEM [3], TeS LA [19], and FGW ANN [32], performs 14.5% better than the original, 18.5% better than the second-best, and 15.2% better than the third-best models under variety of use cases. The use of a pre-trained CNN model and RNN helps reduce computational redundancy when evaluating inclination classes, which is why this improvement has been obtained for multiple use cases. Since the proposed model outperforms many existing models in terms of classification and inclination detection, it can be put to good use in a wide range of real-time application deployment scenarios.

V. CONCLUSION & FUTURE SCOPE

The proposed model is capable of improving accuracy of inclination classification for multiple users due to integration of high-density feature extraction & selection layers. These layers comprise of LSTM & GRU Models, and are used to support classification operations. These operations showcase further enhancements due to integration of psychological questionnaires which assist in identification of instantaneous student behaviour & inclination levels. Due to which, the model is capable of achieving 16.3 percent better accuracy than BE SEM [3], 14.5 percent better than TeS LA [19], and 10.4 percent better than FGW ANN [32] for multiple types of comparisons. It also showcased 18.5% better precision than that of BE SEM [3], 14.5% better than that of TeS LA [19], and 8.3% better than that of FGW ANN [32] based on these evaluations. The model also showcased 14.8 percent better recall than BE SEM [3], 15.4 percent better than TeS LA [19], and 18.3 percent better than FGW ANN [32] for RFI performance under multiple evaluations. In terms of speed, when compared to BE SEM [3], TeS LA [19], and FGW ANN [32], performs 14.5% better than the original, 18.5% better than the second-best, and 15.2% better than the third-best models under variety of use cases. The use of a pre-trained CNN model and RNN helps reduce computational redundancy when evaluating inclination classes, which is why this improvement has been obtained for multiple use cases. Since the proposed model outperforms many existing models in terms of classification and inclination detection, it can be put to good use in a wide range of real-time application deployment scenarios. In future, the proposed model's performance can be improved via integration of multiple deep learning models including Generative Adversarial Networks (GANs), Q-Learning, and Auto encoders. It must also be validated on larger sets, and can be optimized via integration of multiple bioinspired models that can optimize accuracy and integrate large set of analysis parameters for clinical use cases.

References

- [1] Altuwairqi, K., Jarraya, S.K., Allinjaw, A. et al. Student behavior analysis to measure engagement levels in online learning environments. *SIVIP* 15, 1387–1395 (2021). <https://doi.org/10.1007/s11760-021-01869-7>
- [2] Bao, W. COVID-19 and online teaching in higher education: A case study of Peking University. *Hum Behav & Emerg Tech.* 2020; 2: 113–115. <https://doi.org/10.1002/hbe2.191>
- [3] Hoi, VN, Le Hang, H. The structure of student engagement in online learning: A bi-factor exploratory structural equation modelling approach. *J Comput Assist Learn.* 2021; 37: 1141–1153. <https://doi.org/10.1111/jcal.12551>
- [4] Xiaoying Shen, Chao Yuan, "A College Student Behavior Analysis and Management Method Based on Machine Learning Technology", *Wireless Communications and Mobile Computing*, vol. 2021, Article ID 3126347, 10 pages, 2021. <https://doi.org/10.1155/2021/3126347>
- [5] Hellings, J., Haelermans, C. The effect of providing learning analytics on student behaviour and performance in programming: a randomised controlled experiment. *High Educ* (2020). <https://doi.org/10.1007/s10734-020-00560-z>
- [6] Lin, CL., Jin, Y.Q., Zhao, Q. et al. Factors Influence Students' Switching Behavior to Online Learning under COVID-19 Pandemic: A Push–Pull–Mooring Model Perspective. *Asia-Pacific Edu Res* 30, 229–245 (2021). <https://doi.org/10.1007/s40299-021-00570-0>
- [7] Mailizar, M., Burg, D. & Maulina, S. Examining university students' behavioural intention to use e-learning during the COVID-19 pandemic: An extended TAM model. *Educ Inf Technol* (2021). <https://doi.org/10.1007/s10639-021-10557-5>
- [8] Morris, J., Slater, E., Fitzgerald, M.T. et al. Using Local Rural Knowledge to Enhance STEM Learning for Gifted and Talented Students in Australia. *Res Sci Educ* 51, 61–79 (2021). <https://doi.org/10.1007/s11165-019-9823-2>
- [9] Nandi, R., Nedumaran, S. Understanding the Aspirations of Farming Communities in Developing Countries: A Systematic Review of the Literature. *Eur J Dev Res* 33, 809–832 (2021). <https://doi.org/10.1057/s41287-021-00413-0>
- [10] Khan, Sharifullah & Hwang, Gwo-Jen & Abbas, Muhammad & Rehman, Arshia. (2018). Mitigating the urban–rural educational gap in developing countries through mobile technology-supported learning. *British Journal of Educational Technology*. 50. 10.1111/bjet.12692.
- [11] Nelson, I.A. (2016), Rural Students' Social Capital in the College Search and Application Process. *Rural Sociology*, 81: 249–281. <https://doi.org/10.1111/ruso.12095>
- [12] L. Wenyuan, W. Siyang and W. Lin, "A Multiperson Behavior Feature Generation Model Based on Noise Reduction Using WiFi," in *IEEE Transactions on Industrial Electronics*, vol. 67, no. 6, pp. 5179–5186, June 2020, doi: 10.1109/TIE.2019.2927179.
- [13] R. Zheng, F. Jiang and R. Shen, "Intelligent Student Behavior Analysis System for Real Classrooms," *ICASSP 2020 -2020 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, 2020, pp. 9244–9248, doi: 10.1109/ICASSP40776.2020.9053457.
- [14] Y. Ikawati, M. U. H. Al Rasyid and I. Winarno, "Student Behavior Analysis to Detect Learning Styles in Moodle Learning Management System," *2020 International Electronics Symposium (IES)*, 2020, pp. 501–506, doi: 10.1109/IES50839.2020.9231567.
- [15] A. Fekry, G. Dafoulas and M. Ismail, "Automatic detection for students behaviors in a group presentation," *2019 14th International Conference on Computer Engineering and Systems (ICCES)*, 2019, pp. 11–15, doi: 10.1109/ICCES48960.2019.9068128.
- [16] L. Wang, Z. Zhen, T. Wo, B. Jiang, H. Sun and X. Long, "A Scalable Operating System Experiment Platform Supporting Learning Behavior Analysis," in *IEEE Transactions on Education*, vol. 63, no. 3, pp. 232–239, Aug. 2020, doi: 10.1109/TE.2020.2975556.
- [17] S. Delgado, F. Morán, J. C. S. José and D. Burgos, "Analysis of Students' Behavior Through User Clustering in Online Learning Settings, Based on Self Organizing Maps Neural Networks," in *IEEE Access*, vol. 9, pp. 132592–132608, 2021, doi: 10.1109/ACCESS.2021.3115024.
- [18] H. Wan, Q. Yu, J. Ding and K. Liu, "Students' behavior analysis under the Sakai LMS," *2017 IEEE 6th International Conference on Teaching, Assessment, and Learning for Engineering (TALE)*, 2017, pp. 250–255, doi: 10.1109/TALE.2017.8252342.
- [19] J. Deng, "College Students' Online Behavior Analysis During Epidemic Prevention and Control," *2021 13th International Conference on Measuring Technology and Mechatronics Automation (ICMTMA)*, 2021, pp. 796–799, doi: 10.1109/ICMTMA52658.2021.00183.
- [20] P. Sun, Z. Wu, Y. Zhang and Y. Yang, "Analysis and Modeling of Learning Behaviors on Intelligent Tutoring Website," *2014 Tenth International Conference on Computational Intelligence and Security*, 2014, pp. 729–731,

doi: 10.1109/CIS.2014.170.

- [21] M. E. Rodríguez, A. -E. Guerrero-Roldán, D. Baneres and I. Noguera, "Students' Perceptions of and Behaviors Toward Cheating in Online Education," in *IEEE Revista Iberoamericana de Tecnologías del Aprendizaje*, vol. 16, no. 2, pp. 134-142, May 2021, doi: 10.1109/RITA.2021.3089925.
- [22] S. Lai, B. Sun, F. Wu and R. Xiao, "Automatic Personality Identification Using Students' Online Learning Behavior," in *IEEE Transactions on Learning Technologies*, vol. 13, no. 1, pp. 26-37, 1 Jan.-March 2020, doi: 10.1109/TLT.2019.2924223.
- [23] F. B. Shaikh, M. Rehman, A. Amin, A. Shamim and M. A. Hashmani, "Cyberbullying Behaviour: A Study of Undergraduate University Students," in *IEEE Access*, vol. 9, pp. 92715-92734, 2021, doi: 10.1109/ACCESS.2021.3086679.
- [24] D. Aviano, B. L. Putro, E. P. Nugroho and H. Siregar, "Behavioral tracking analysis on learning management system with apriori association rules algorithm," 2017 3rd International Conference on Science in Information Technology (ICSITech), 2017, pp. 372-377, doi: 10.1109/ICSITech.2017.8257141.
- [25] M. Awais Hassan, U. Habiba, H. Khalid, M. Shoaib and S. Arshad, "An Adaptive Feedback System to Improve Student Performance Based on Collaborative Behavior," in *IEEE Access*, vol. 7, pp. 107171-107178, 2019, doi: 10.1109/ACCESS.2019.2931565.
- [26] A. M. Villa, C. P. Molina and M. Castro, "Students' Behavior When Connecting to the LMS: A Case Study at UNED," in *IEEE Revista Iberoamericana de Tecnologías del Aprendizaje*, vol. 14, no. 3, pp. 87-94, Aug. 2019, doi: 10.1109/RITA.2019.2942253.
- [27] D. Liu, Y. Zhang, J. Zhang, Q. Li, C. Zhang and Y. Yin, "Multiple Features Fusion Attention Mechanism Enhanced Deep Knowledge Tracing for Student Performance Prediction," in *IEEE Access*, vol. 8, pp. 194894-194903, 2020, doi: 10.1109/ACCESS.2020.3033200.
- [28] Z. Xu, H. Yuan and Q. Liu, "Student Performance Prediction Based on Blended Learning," in *IEEE Transactions on Education*, vol. 64, no. 1, pp. 66-73, Feb. 2021, doi: 10.1109/TE.2020.3008751.
- [29] P. Jiang and X. Wang, "Preference Cognitive Diagnosis for Student Performance Prediction," in *IEEE Access*, vol. 8, pp. 219775-219787, 2020, doi: 10.1109/ACCESS.2020.3042775.
- [30] H. Tang et al., "Predicting Green Consumption Behaviors of Students Using Efficient Firefly Grey Wolf-Assisted K Nearest Neighbor Classifiers," in *IEEE Access*, vol. 8, pp. 35546-35562, 2020, doi: 10.1109/ACCESS.2020.2973763.
- [31] R. A. Rasheed, A. Kamsin and N. A. Abdullah, "An Approach for Scaffolding Students Peer-Learning Self-Regulation Strategy in the Online Component of Blended Learning," in *IEEE Access*, vol. 9, pp. 30721-30738, 2021, doi: 10.1109/ACCESS.2021.3059916.
- [32] X. Li, Y. Zhang, H. Cheng, F. Zhou and B. Yin, "An Unsupervised Ensemble Clustering Approach for the Analysis of Student Behavioral Patterns," in *IEEE Access*, vol. 9, pp. 7076-7091, 2021, doi: 10.1109/ACCESS.2021.3049157.
- [33] S. Vhaduri, S. V. Dibbo and Y. Kim, "Deriving College Students' Phone Call Patterns to Improve Student Life," in *IEEE Access*, vol. 9, pp. 96453-96465, 2021, doi: 10.1109/ACCESS.2021.3093493.
- [34] P. Song, Y. Gao, Y. Xue, J. Jia, W. Luo and W. Li, "Human Behavior Modeling for Evacuation From Classroom Using Cellular Automata," in *IEEE Access*, vol. 7, pp. 98694-98701, 2019, doi: 10.1109/ACCESS.2019.2930251.
- [35] W. M. Al-Rahmi, N. Yahaya, M. M. Alamri, N. A. Aljarboa, Y. B. Kamin and F. A. Moafa, "A Model of Factors Affecting Cyber Bullying Behaviors Among University Students," in *IEEE Access*, vol. 7, pp. 2978-2985, 2019, doi: 10.1109/ACCESS.2018.2881292.
- [36] F. D. Pereira et al., "Explaining Individual and Collective Programming Students' Behavior by Interpreting a Black-Box Predictive Model," in *IEEE Access*, vol. 9, pp. 117097-117119, 2021, doi: 10.1109/ACCESS.2021.3105956.
- [37] S. A. Priyambada, M. Er, B. N. Yahya and T. Usagawa, "Profile-Based Cluster Evolution Analysis: Identification of Migration Patterns for Understanding Student Learning Behavior," in *IEEE Access*, vol. 9, pp. 101718-101728, 2021, doi: 10.1109/ACCESS.2021.3095958.
- [38] S. Nam, G. Frishkoff and K. Collins-Thompson, "Predicting Students Disengaged Behaviors in an Online Meaning Generation Task," in *IEEE Transactions on Learning Technologies*, vol. 11, no. 3, pp. 362-375, 1 July-Sept. 2018, doi: 10.1109/TLT.2017.2720738.
- [39] W. -Y. Hwang, I. Q. Utami, S. W. D. Purba and H. S. L. Chen, "Effect of Ubiquitous Fraction App on Mathematics Learning Achievements and Learning Behaviors of Taiwanese Students in Authentic Contexts," in *IEEE Transactions on Learning Technologies*, vol. 13, no. 3, pp. 530-539, 1 July-Sept. 2020, doi: 10.1109/TLT.2019.2930045.
- [40] R. Lavi, Y. J. Dori, N. Wengrowicz and D. Dori, "Model-Based Systems Thinking: Assessing Engineering

Student Teams," in IEEE Transactions on Education, vol. 63, no. 1, pp. 39-47, Feb. 2020, doi: 10.1109/TE.2019.2948807.

Instructions for Authors

Essentials for Publishing in this Journal

- 1 Submitted articles should not have been previously published or be currently under consideration for publication elsewhere.
- 2 Conference papers may only be submitted if the paper has been completely re-written (taken to mean more than 50%) and the author has cleared any necessary permission with the copyright owner if it has been previously copyrighted.
- 3 All our articles are refereed through a double-blind process.
- 4 All authors must declare they have read and agreed to the content of the submitted article and must sign a declaration correspond to the originality of the article.

Submission Process

All articles for this journal must be submitted using our online submissions system. <http://enrichedpub.com/> . Please use the Submit Your Article link in the Author Service area.

Manuscript Guidelines

The instructions to authors about the article preparation for publication in the Manuscripts are submitted online, through the e-Ur (Electronic editing) system, developed by **Enriched Publications Pvt. Ltd.** The article should contain the abstract with keywords, introduction, body, conclusion, references and the summary in English language (without heading and subheading enumeration). The article length should not exceed 16 pages of A4 paper format.

Title

The title should be informative. It is in both Journal's and author's best interest to use terms suitable. For indexing and word search. If there are no such terms in the title, the author is strongly advised to add a subtitle. The title should be given in English as well. The titles precede the abstract and the summary in an appropriate language.

Letterhead Title

The letterhead title is given at a top of each page for easier identification of article copies in an Electronic form in particular. It contains the author's surname and first name initial .article title, journal title and collation (year, volume, and issue, first and last page). The journal and article titles can be given in a shortened form.

Author's Name

Full name(s) of author(s) should be used. It is advisable to give the middle initial. Names are given in their original form.

Contact Details

The postal address or the e-mail address of the author (usually of the first one if there are more Authors) is given in the footnote at the bottom of the first page.

Type of Articles

Classification of articles is a duty of the editorial staff and is of special importance. Referees and the members of the editorial staff, or section editors, can propose a category, but the editor-in-chief has the sole responsibility for their classification. Journal articles are classified as follows:

Scientific articles:

1. Original scientific paper (giving the previously unpublished results of the author's own research based on management methods).
2. Survey paper (giving an original, detailed and critical view of a research problem or an area to which the author has made a contribution visible through his self-citation);
3. Short or preliminary communication (original management paper of full format but of a smaller extent or of a preliminary character);
4. Scientific critique or forum (discussion on a particular scientific topic, based exclusively on management argumentation) and commentaries. Exceptionally, in particular areas, a scientific paper in the Journal can be in a form of a monograph or a critical edition of scientific data (historical, archival, lexicographic, bibliographic, data survey, etc.) which were unknown or hardly accessible for scientific research.

Professional articles:

1. Professional paper (contribution offering experience useful for improvement of professional practice but not necessarily based on scientific methods);
2. Informative contribution (editorial, commentary, etc.);
3. Review (of a book, software, case study, scientific event, etc.)

Language

The article should be in English. The grammar and style of the article should be of good quality. The systematized text should be without abbreviations (except standard ones). All measurements must be in SI units. The sequence of formulae is denoted in Arabic numerals in parentheses on the right-hand side.

Abstract and Summary

An abstract is a concise informative presentation of the article content for fast and accurate Evaluation of its relevance. It is both in the Editorial Office's and the author's best interest for an abstract to contain terms often used for indexing and article search. The abstract describes the purpose of the study and the methods, outlines the findings and state the conclusions. A 100- to 250-Word abstract should be placed between the title and the keywords with the body text to follow. Besides an abstract are advised to have a summary in English, at the end of the article, after the Reference list. The summary should be structured and long up to 1/10 of the article length (it is more extensive than the abstract).

Keywords

Keywords are terms or phrases showing adequately the article content for indexing and search purposes. They should be allocated heaving in mind widely accepted international sources (index, dictionary or thesaurus), such as the Web of Science keyword list for science in general. The higher their usage frequency is the better. Up to 10 keywords immediately follow the abstract and the summary, in respective languages.

Acknowledgements

The name and the number of the project or programmed within which the article was realized is given in a separate note at the bottom of the first page together with the name of the institution which financially supported the project or programmed.

Tables and Illustrations

All the captions should be in the original language as well as in English, together with the texts in illustrations if possible. Tables are typed in the same style as the text and are denoted by numerals at the top. Photographs and drawings, placed appropriately in the text, should be clear, precise and suitable for reproduction. Drawings should be created in Word or Corel.

Citation in the Text

Citation in the text must be uniform. When citing references in the text, use the reference number set in square brackets from the Reference list at the end of the article.

Footnotes

Footnotes are given at the bottom of the page with the text they refer to. They can contain less relevant details, additional explanations or used sources (e.g. scientific material, manuals). They cannot replace the cited literature.

The article should be accompanied with a cover letter with the information about the author(s): surname, middle initial, first name, and citizen personal number, rank, title, e-mail address, and affiliation address, home address including municipality, phone number in the office and at home (or a mobile phone number). The cover letter should state the type of the article and tell which illustrations are original and which are not.