

Advancing Dynamic Learning Pattern Recognition with Enhanced Prefix Span Algorithm Model and LSTM

¹Ramesh M and ²Jayashree R

¹Department of Computer Science and Applications, Faculty of Science and Humanities,
SRM Institute of Science and Technology, Vadapalani, Chennai, Tamil Nadu, India.

²Department of Computer Applications, Faculty of Science and Humanities, SRM Institute of Science and Technology,
Kattankulathur, Chengalpattu, Tamil Nadu, India.

¹rameshm@srmist.edu.in, ²jayashrr@srmist.edu.in

Correspondence should be addressed to Jayashree R : jayashrr@srmist.edu.in

Article Info

Journal of Machine and Computing (<https://anapub.co.ke/journals/jmc/jmc.html>)

Doi: <https://doi.org/10.53759/7669/jmc202606002>

Received 23 May 2025; Revised from 30 August 2025; Accepted 20 September 2025

Available online 13 October 2025.

©2026 The Authors. Published by AnaPub Publications.

This is an open access article under the CC BY-NC-ND license. (<https://creativecommons.org/licenses/by-nc-nd/4.0/>)

Abstract – In Dynamic Learning Pattern Recognition, the progress needs new ways of reliably detecting and describing dynamic patterns in big data. Additionally, based on the above research, this thesis proposes a novel cycle of the development of the dynamic learning pattern recognition through the combination of Enhanced Prefix Span Algorithm Model, as well as Long Short-Term Memory (LSTM) networks. An important part of our work in solving the issue of identifying dynamic patterns in textual data in online caselaw finding applications of sequential pattern mining methods and deep learning methods. On that basis, this study will then accelerate the traditional PrefixSpan algorithm to extract the sequential patterns of the data set taking into consideration the dynamic nature of the pattern. Then add LSTM networks to model and learn such sequential patterns and thus identify long term dependencies and time information. The suggested methodology has been researched and the performance of the designed methodology compared to other classifiers such as MultinomialNB, Logistical Regression. The findings demonstrate the high performance of the proposed model whose accuracy, recall, and F1-score are 98.45, 98.62, and 97.3, respectively. This study is to combine sequential pattern mining with the deep learning, which serves to the advancement of dynamic learning pattern recognition methods and to give certain promising suggestions in the cases such as natural language processing and predictive modeling. It is an effort by python implementation to promote the efficiency and precision of recognition of learning patterns in different modalities. Here, we apply it in assembling a sound framework of the effective recognition and adaptation to changing patterns within textual data to assist in enhancing and more efficient learning on a variety of modalities in different situations.

Keywords – Dynamic Learning Pattern, Prefix-Span Algorithm, Predictive Modeling, Textual Data, Temporal Information.

I. INTRODUCTION

Pattern recognition is one of the most important issues related to machine learning and artificial intelligence as the systems can identify the information related to the perceived semantics and make decisions based on such patterns [1]. The conventional methods of pattern recognition are based on the static models in the majority of aspects. Finite data sets were typically used to train these models. Even though the efficient terms were used in the case of the methods applicable to the static techniques in a stable environment, the methods have not been effective as the data distribution underlying them changes with time. Learning pattern recognitions that are new dynamic are approaches of updating models with new features of changes in data patterns as they come. Some of the new systems that include adaptive algorithms, online learning and incremental learning include techniques that enable the models to be dynamic and adapt to new data [2]. Dynamic learning methods are contrasted with methods of batch learning- most methods of retraining need a complete retraining on the full dataset. New information is inserted in the Lake but without forgetting what one has already learned before. Concept drift recognition and management are also one of the most characteristic and distinctive strengths of dynamic learning. Concept drift is the scenario in which there are gradual or drastic alterations in the underlying distributions of data. Otherwise, concept drift may significantly deteriorate with time the model performance [3].

Nevertheless, dynamic learning systems attempt to address this issue by tracking data streams and replacing model parameters or structures in response to observed changes. A fundamental course under the dynamic learning is also online learning as it maintains the model with the newly available information on-the-fly [4]. Online algorithms can also implement model incremental learning and the information can be added by them without forgetting the prior knowledge; under leverage to batch-wise computations. This processing has, consequently, enhanced the flexibility of the changing trends of the data in addition to reducing computation consuming the total retraining occurs [5].

Pattern recognition is a feature that is driven by the flexibility of algorithms that evolve on-the-fly with the dynamic changes in input data and the environment. Algorithms operate on varying parameters which keep on adjusting their inner mechanism to reach performance in the environment on which the environment variables are changing. Such versatile flexibility is the characteristic feature of a wide range of applications. In financial analytics, e.g. in dynamic learning models, real-time market data is processed by dynamic learning models to modify trade strategies that react to market volatility and the economy [6]. Streaming of patient data in healthcare is useful in identifying abnormalities at the earliest stage of their occurrence to aid in the early diagnosis and treatment. These systems do not only track the dynamics of individual patients but also track the dynamic paradigms of medical research. One area where dynamic recognition models can do more than good is in the field of cybersecurity, where it is necessary to trace the current network traffic and behavioral patterns to identify and counter any emerging threats. As each interaction takes place, these models become more skilled at defining what potential risks appear like, and thereby becoming more resistant to new types of attack [7]. Dynamic models of Natural language processing are therefore typified by the constant learning between the interaction with the user and the evolving lingual landscape to offer real time contextual deviance and monitor the fluctuating linguistic fashions; thereby, enhancement of understanding and production of natural lingo [8].

The dynamic learning pattern recognition has been at the forefront of most of the processing to change natural languages parsing systems into the actual fact that they can indeed be made to learn dynamically in changing language patterns and contextual factors. The traditional NLP methods are often based on the utilization of the eponymous static models that are trained on concrete datasets, and which embody the slow separation of the language and the world of the language traditions. It allows NLP models to continue learning real-time textual data and make prompts and internal alterations as necessitated by them [9]. This is one of the most common applications of such a self-adaptive method: sentiment analysis, in which models interpret and categorize emotions written in text. The dynamic models will be capable of evolving with those variations in which the sentiment is manifested: new slang, social narratives or just what folks think or say in the street with regard to something different than previously [10]. Also important is the role of dynamic pattern recognition in the NLP is data companion to the generative language end task, such as machine translation and summary construction; such tasks include real-time feedback by users or recently available parallel corpora, and therefore facilitate the improvement of their results, as well as to allow the contextual accurateness of translations and conciseness and coherence of summaries. This ensures procedural continuing experiences of language models with the natural flow of the human language development [11].

The flexibility of NER frameworks is provided by the rules which permit networks to adjust their pattern of entity recognition and thus permit the networks to adapt their pattern to new entities or in cases where the standards used to name things vary between languages or domains. This will ensure that the algorithms of NER are accurate and up to date in a large number of dynamic textual situations. Furthermore, the NLP algorithms NLP algorithms are also made more resilient and flexible through dynamic learning identification of patterns [12]. The use of languages and customs are the opposite to general purpose databases about such disciplines like medical care, law or technical literature. Since dynamic NLP algorithms can keep on learning domain specific linguistic patterns and adjusting their projections and interpretations to suit the requirements of various applications or sectors, they are capable of possessing the properties of distributing algorithm in constructing lexical semantics [13]. One of the dynamic learning pattern recognition challenges is the streams of information of varying numbers and speeds that handle security and privacy issues in real time processing and balancing complexity and computing efficiency. To achieve dynamic pattern recognition, the most significant goal of research is to ensure the reliability of algorithms that are able to handle large and complex streams of information and improve the model transparency as well as the ethical issues related to the algorithmic decision-making in dynamic conditions [15].

This is a significant move towards the effective collection and adaptation of dynamic learning pattern through the improvement of Long Short-Memory (LSTM) networks, and to the enhancement of the Prefix Span algorithm model. The Prefix Span is a technique of information searching on streamed recurrent patterns and it was initially created to search sequential patterns in mining. A application of dynamic learning capabilities to this strategy assists the algorithm to learn faster with minute variation and pattern of ever-shifting streams. The LSTM networks also improve the way in which the framework may treat the temporal aspects and long-term correlations of sequential information. LSTM networks have the strength to memorize and study patterns over an extended duration and therefore can be applied in the time series or sequence data stream analysis. The addition of the LSTM layers assists the dynamic learning pattern recognition system to acquire the capacity to unravel untangled intertwinement and couplings in developing patterns of data, escalating model flexibility and forecast precision. It builds a robust dynamic learning framework on the basis of the enhanced Prefix Span algorithm framework and LSTM networks that are capable of continually learning and adapting to the evolving trend of information. Therefore, with these two approaches put together, it will be not only

easier to mine streaming data to extract recurring pattern using an efficient mining approach, but it will also enable the model to extract and utilize temporal dependencies to enable it to make more precise predictions and decision-making. Due to this, this progress can be implemented in the field of natural language processing, banking, medicine, and information security where the rapidity of analyzing and reacting to the ever-evolving info is so of crucial importance. The key contributions of the research study is described as,

- The improved Prefix Span Algorithm Model is a powerful model to perform the sequential pattern mining, where one can detect recurring patterns in dynamic learning process. The model is an effective representation of temporal dependencies and sequential relations among the textual data that enables the recognition of the emerging patterns accurately.
- Deep learning frameworks, such as LSTM, the effective characteristic of the model, increase the ability to recognize long-term relationships and contextual data in sequential data. The proposed solution, which makes use of LSTM, demonstrates high performance in identifying complex and changing patterns which makes the dynamic learning pattern recognition systems more accurate and reliable.
- The synergistic integration of the improved Prefix Span Algorithm Model, as well as LSTM, leads to the significant improvement in the dynamic learning patterns recognition, which can be considered in a wide range of applications that involve the real-time analysis of patterns and decision-making.

The remainder of the paper is organized as follows: Section 2 presents the literature review; Section 3 presents Problem statement; and Section 4 presents the proposed methodology. Section 5 provides an overview of the results. Section 6 contains the work's conclusion.

II. RELATED WORKS

Although the attention-based technologies have shown potential in recognition of scenario text, they may lead to catastrophic out-of-vocabulary performance because of overreliance on contextual information. Despite the excellent results of general expectations research, later studies indicate that despite excellent imagery, it is still difficult to identify invisible text reliably. The proposed system, Context-based contrastive learning, is supposed to deal with this problem, creating symbols in a range of circumstances and minimizing contrastive loss on their embeddings. The complexity of scene textual information and its variety may also be factors that restrain the efficacy of ConCLR because of its generalization efficiency in real-world situations. The complexity of computing, which is associated with developing characters and contrastive optimization strategies, may also limit the ability of ConCLR to scale to large data sets. Representation learning is also dependent on embeddings and this could lead to interpretation and modelling comprehensibility problems. Furthermore, the quality and the variety of information to train on could affect the efficiency of ConCLR that could potentially result in suboptimal results on certain text recognition tasks. ConCLR could be useful in dealing with the limitations of attention-based methods; however, further research is necessary to increase its ability to withstand and scale to realistic STR implementations [16].

Although self-supervised word recognition schemes offer encouragement due to exploitation of unlabelled natural images, they also have challenges that occur as a result of domain difference between synthetic and real information. Trying to cover the gap partially, through contrastive learning, the dependence upon discriminating learning per se might be inadequate to simulate the peculiarities of text perception, particularly in the conditions of the highly varying scene texts. This is more sustainable method and thus can be defeated by the tremendous computation and high complexities of masked image modeling combination. Therefore, the approach corrects both comprehension and explainability problems in an exciting way in that it heavily relies on unmasked image modeling to be used in context building and contrastive instruction in the process of discrimination. Major comments may be inserted on the natural volume and nature of training data in relation to performance enhancement as a restraining element to the generalized capabilities of the strategy. Nevertheless, this performance improvement might not be generalized across all text-recognition tasks and datasets, and additional development and testing of the approach might be needed before it can be applicable to industrial use. In conclusion, the suggested method can also be unavailable to the less equipped or more ignorant investigators and practitioners because it relies on the initial training and fine-tuning guidelines that can be accompanied by considerable computing resources and skills. Finally, the proposed method can be offered to the future development of self-supervised text recognition. Its weaknesses and the need to generalize it onto the real-life situations have not been addressed in the further research [17].

Total freedom of computer vision system to identify handwritten text is still a challenge however. The line segmentation and textual line detection are commonly considered as two distinct methods and this inefficiency could lead to the further deterioration of efficiency. The wide-ranging integrative architecture that takes an integrative approach as well as mixed consideration might appear to be a viable alternative, but the architecture might be lacking in its approach to the variations of handwritings as well as configuration complexities. This analysis of the image of one sentence, starting at the beginning to the end, particularly of long articles or pictures containing a lot of text, can be very computationally costly and also inference-time-sensitive. The attention modules can also implicitly sub divide the lines in terms of text readability and consistency as well. The lack of certainty in handwriting, sound, and the flaws might make it difficult in identifying such patterns in decoding component. This performance degradation could also occur when it is used to work with handwritten text to learn language with complex character systems or writing variations. Also, storage

and computational constraints might limit the flexibility of the proposed algorithm to work with large data sets or algorithms in the real world. Perhaps further studies and validation will be required to identify the generalization abilities of the algorithm with different handwriting methods, typefaces and dialects. Moreover, the algorithm still has the potential to be developed regarding its cognition of the projections and resistance to visual distortion like ambient cluttered and obstructions. All in all, although the proposed integrated end-to-end system is seemingly a viable solution to the issues with sentence text identification, there is still more research to undertake to overcome the above-mentioned problems and increase the utility of the model [18].

Cross-modal retrieving is very challenging between the texts and images because their semantic difference between the visible and verbal interpretations is extreme. However, the proposed solution can attempt to fill this gap by finding a common embedding space, but it may fail to encounter nuances and more complex semantic connections in image-text communication. The model may be impaired in terms of the degree and quality of provided training set, leading to the less-than-perfect alignments between the optical and textual properties, despite the fact that the semantic connection-based information is also provided and the global reasoning process is also used. Moreover, issues with scalability and computational cost can also be encountered due to the use of gating and memory-based operations in semantic reasoning as well as when processing large quantities of data or applications in actual applications. The algorithm proves to be efficient and easy compared to more complex local matching algorithms but could sacrifice a degree of very fine corresponding accuracy, particularly in situations where there are complex semantic relations or unreliable image-text relations. Though the technique proves to be efficient and productive concerning the accommodation learning and general representation learning, it might not be capable of delivering to various sets of text and image information with various features and forms of representation. Despite such disadvantages, the method offers a promising new avenue of cross-modal extraction, focusing on the potential of delivering excellent results with simplified global combining methods [19].

Current measurements on Handwriting Text Identification algorithms have focused on the editing lengths at characters and word-levels; but changing to page long transcribed presents it with further challenges. The quality of the transcribing and the precision of the order of reading should be considered when evaluating the page-level HTR technologies. However, inconsistencies which are encountered at the page level transcribed could not be adequately reflected in the existing measures. Although the two-pronged evaluation procedure suggested by us fits this issue perfectly, it may still experience certain issues with complex document design and handwriting [20]. Moreover, even the slightest inaccuracy in the transcription and reading order might go unnoticed in case such simple measures as Word Error Rate and Bag of Words Word Error Rate are used. Even though the Hungarian Word Rate measure was recently announced with confidence, additional testing and refinement might be required to ensure the measure remains consistent across different information and circumstances. Despite these disadvantages, our findings indicate that the evaluation of HTR systems on the scale of a page is necessary and provide the methods of performance evaluation improvement in this case. These limitations can be surmounted and the reliability and relevance of page-level assessment measures to HTR systems enhanced through further research.

Even though heterogeneity of appearances contributes to the challenge of word recognition of actual scenery image, the proposed solution enhances the adaptation of Twin Support Vector Equipment through normalization of manifolds. Nevertheless, the complexity and heterogeneity of images of scenarios might limit the effectiveness of multidimensional normalization. Although ambience and the regularize properties have been added to T-SVM, the algorithm may fail to provide a large tolerance of text/background styles. Moreover, with an addition of a revalidation module, there is an attempt to reduce the false positives but it may not be entirely effective to reduce noise and complexities in natural scene images. Depending on step-by-step procedures unfortunate to locate, identify and put together text objects may lead to inefficiencies and in particular, text deformation or overlapping. Although the proposed model has demonstrated better performance compared to other state of the art methods in terms of the accuracy levels it might not be quite competent when working with different sets of data or even real-world cases. Further research is required in order to make the model more efficient in the changing environmental settings, and in the process of handling the hard text recognition tasks. Nevertheless, despite the few weaknesses, the given approach demonstrates that it is effective in the context of text recognition of image natural scenes, which provides another step in the evolution of this field [21].

III. PROBLEM STATEMENT

Past studies on pattern recognition particularly in dynamic learning environments have not met considerable challenges when extended to these events as written text, self-learning, handwritten text, cross-modality text information retrieval and document level performance evaluation of handwriting recognition systems. Issues that have been faced have been in the flawed handling of out-of-vocabulary words, mismatch between synthetic and real-life data, inability to measure the performance of the questionable document layout settings and alternative handwriting, discrepancy between visualization and textual data resulting in the performance of the system, and absence of an extensive evaluation measure of the complete-page handwritten text recognition. The following paper introduces a new model that combines Enhanced Prefix Span Learning and Long Short-Term Memory (LSTM) networks in combating these drawbacks. Prefix Span Algorithm is used to work on sequential pattern mining in an efficient manner, and LSTM models are used to work on dependencies among sequential data in their long term. The joint efforts are planned to assist with more generalized recognition ability,

strength, and correctness in dynamic learning pattern recognition [20]. Through the integrated methodology in investigating the field of these problematic matters, the study will cover the major constraints of the existing models. The development of our integrated system is fairly proven by the fact that, experimental validation and performance analysis with certain contemporary solutions are able to validate the progress of our system in moving dynamic learning forward.

IV. PROPOSED DYNAMIC LEARNING PATTERN RECOGNITION USINGPSA AND LSTM

The dynamic learning pattern recognition is highly developed and has a number of important steps. The data on which that is done is approximately 15000 sentences in which each sentence is labeled as either Visual or Auditory style of learning. Moreover, the above data is purged using wrangling and exploratory data analysis (EDA) and a holistic process of preprocessing text that could involve the release of tokenization, lemmatization, removal of stop words, removal of punctuation, and text normalization. These are pre-prepared to conduct deeper analysis. The pattern explorations and outlier detection in the cleaned text then give meaning into the structure and trends in work. The improved algorithm model of prefix span operates the sequential pattern mining in the sequences to segregate the repetitive sequences of dynamic environment of learning. The long-term memory (LSTM) networks are then added so that the objective can be enhanced to attain superior recognition between the learned patterns but in long terms. In this way, this structure is offering a very solid solution towards identify the dynamic patterns in the texts because the comprehension of various learning styles and behaviors is more integrated in this case. **Fig 1** illustrates the overall process.

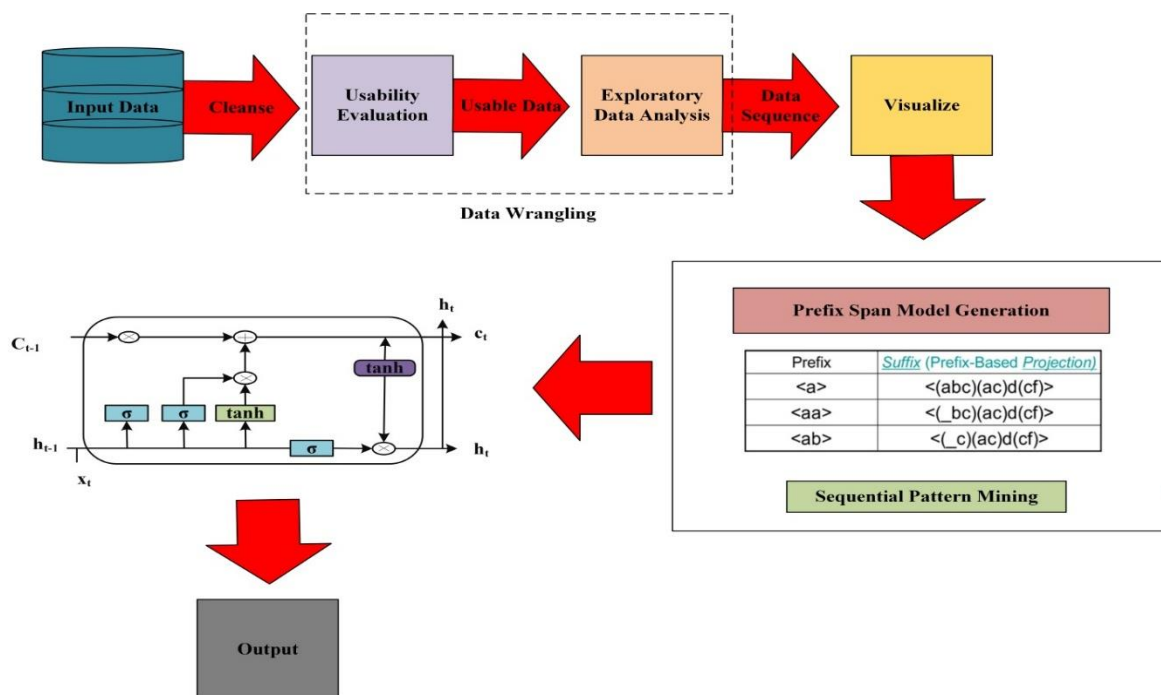


Fig 1. Workflow of Proposed Methodology.

Data Wrangling

The dataset is included in approximately 15,000 sentences and marked by learning styles Visual, Auditory or Kinesthetic. All sentences belong to one of the categories, which implies the most desirable mode of learning to interpret the text. The resources in the visual form of images, diagrams or charts are the most beneficial to the visual learners, and the opposite is true because auditory learners would absorb the information more effectively when it is presented in form of the verbal presentation, in form of lectures or audio tapes. The fact that it is exhaustive and covers all the learning differentials makes this dataset significant to the study of the way people process and internalize information in different modalities [22].

Table 1. Dataset

Type	Count
Visual	743
Auditory	257

Similar to **Table 1**, the number of distinct values of the column Type demonstrates a distribution across the three styles of learning: Visual, Auditory. The 743 instances that are related to the Visual learning style indicate that in the dataset, people are the ones who like visual aids and materials most compared to other types. Conversely, the Auditory learning style is deemed to be less common and only a handful of 257 instances imply a far lesser number of people, who

like to learn through the Auditory medium like lectures or discussions. The distribution thus captures diversity in the learning preferences of the samples of the dataset whereby the most preferred would be Visual, which is then followed by Auditory.

Table 2. Sample Data

Sentence	Type
895 If an important concept is hazy or difficult, ...	Visual
35 Mam loved people, adored God and family and al...	Auditory
918 We have condensed the ideas into our own words...	Visual
169 The routines at Lochranza Court suited her dow...	Auditory
579 To study a distribution, take random samples f...	Visual

Table 2 represents a sample dataset of paired sentences and categories with a visual and auditory representation of sentences in each row. This is the structured format in which the contents in the linguistic data are described differently in that, auditory and visual categories are processed differently. These kinds of reference will initiate a superior comprehension of the inner semantics and contextual characteristics with submerging into more analysis.

Text Pre-Processing

Tokenization

The first stage of a given text is termed tokenization and it is divided into the individual words or tokens that make up the text. This is disintegration of information which can be easily processed further when converted into tokens. In this the token can take the form of words in the sentence or any other form of the word separated by a white space. This is the most simplistic form of tokenization of NLP tasks. In this method, a token has a tendency of associating with a sentence. Tokens might however be in the form of words with extra characters or punctuations that have been interspersed with them since the text is not processed. This paper provides an in-depth analysis of the influence of the evolving pattern identification problem of different tokenization algorithms as displayed below. Phrase or document preprocessing can be a preliminary step to tokenization. We employ the following preprocessing procedures in this method,

- Convert the sentence to lowercase.
- Replace instances of "n't" with "not."
- Remove any occurrences of "@name."
- Eliminate all punctuation marks except for question marks.
- Remove any other non-alphanumeric characters from the sentence.

Lemmatization

Normalisation is performed in the form of lemmatization to normalise the terms to their base or vocabulary form following tokenization. Ensuring that all the word variants are always filled in, will contribute to lessening the level of dimensionality of writing data and to the quality of the future researches. Lemmatization and stemming are similar in the fact that the basic form of the word is restored. The difference is that the irregular cases and morphophonemic norms also are considered. The lemma representations of the phrases in the models are used in a similar way as stemming. Lemmatization refers to the replacement of a word by its lexicon. This approach forms the lemma because it can be found in a dictionary through analysis of whether a word is placed in the middle or end of a sentence and also by removing the inflexional parts of a word that are ending (e.g. Performance is considerably enhanced, replaced with Performer be greatly improved). Lemmatization enhances the effectiveness of user extraction, where the various forms of words are reduced into one lemma.

Stopword Removal

Stopwords are common expressions that do not necessarily contribute to any linguistic meaning of the text. They are eliminated. This phase can help in the reduction of background noise and the focus on the main ideas of the text. Stopwords are the word that are used frequently in any type of text in a certain language because they are unable to differentiate. We drop stopwords in a statement with help of the stopwords list provided by python NLP library with the help of this tokenization method. Stop words, which are usually frequent in a language, are believed to be amongst the least informative words (i.e. the stop words alone do not contribute any meaning to a document). Stop words can also be not effective to retrieve information as they are language-specific and cannot be used as keywords in text-extraction algorithms. Stop words, which are articles, adjectives, conjunctions, and prepositions, are very common in the texts and are not related to a specific subject. TC work is normally done after the removal of stop words. Removal of stop words in a set of data minimizes its size. The words such as off, a, the, in, an, with, and, and to are called stop words. Depending on the list that is used there are usually about 400 to 500 stop words in one language.

Punctuation Removal

Punctuation is fundamentally removed in the text; this is to simplify text, and simplicity adds to uniformity in the word presentation.

Text Cleaning

All special characters are removed, and all words are reduced to lowercase; in other words, these texts are standardized. The stage preconditions the necessity of providing the text with the consistency between each other, which, in its turn, results into the quality of the further analyst and modeling.

The proposed methodology divides the text into tokens or single words or tokens. A then attempt is made to cut down the words to their base or dictionary forms hence the words are represented uniformly. To eliminate on such a basis, as counterparts to whom, either in the cognitive stream or in the highest grammatical relation of the phrase, which is concerned in ruling its significance, the least grammatical relation, the argument which attempts to assert significance, is to decrease auditory noise and emphasize on arguments which attempt to proclaim significance. In this way, once they took away them, they would take away the punctuations, thereby completing the cleaning process. The words were then converted into lower case and all non-alphabetical characters were edited out. This would cement and knot the strings and thus present much more challenging duty ahead to any pattern detecting applications that would attempt to work on it. Thus, now, textual data, having undergone all these preprocessing methods, is now about to fly off to the subsequent stage of analysis and modelling-on-the-Enhanced Prefix Span Algorithm Model structure.

Exploratory Data Analysis for Identifying Outliers

Exploratory Data Analysis (EDA) entails important processes into interpreting structures, determination of patterns or tendencies that were unnoticed in the data. Once the text pre-processing methods are completed, we do EDA in which we move on to further analysis of the processed textual data. The former is conducted by carrying out an analysis of the related distribution of token lengths and subsequently discovering how various words differ in size throughout the whole corpus. Such a process gives a concept of text granularity, and the presence of any anomalies or outliers is also revealed. Words frequency distribution should be created to be further studied, a great deal along the same lines, however, on the visualization tools such as word clouds or bar graphs to produce the most frequent words. This graphical representation completes the general conversation and begins to bring up common themes/topics in the text, and this also assists in giving an overview of the direction that one should take in the future analysis. Moreover, we will examine general learning patterns in the text that will give us the information about some of the prevailing attitudes or emotions there. Therefore, we should be able to extract useful meaningful thoughts out of pure EDA on the refined text to simplify the following model and analysis processes, with the Enhanced Prefix Span Algorithm Model, to a great extent.

Sequential Pattern Mining by Employing Prefix-Span Algorithm

PrefixSpan proposes hierarchies of the projected databases by analyzing solely pieces that are locally frequent and recursively projecting a sequence dataset into a descending chain of smaller projected database of sequences. The next task is to reduce the amount of work to produce candidate subsequences and search through the entirety of sequential structures. Reducing the size of the projected database as well as the cost of testing a likely candidate sequencing at all possible locations is the next task. Each element may have the pieces rearranged in another way before they are checked against a sequence of the prospective candidates. It can be assumed that the elements of a part of a sequence are always sorted alphabetically since they can be placed in any order, and it does not cause them to lose their generality. E.g. instead of being represented as $(a(bac)(ca)d(f\ c))$, the sequences in S with Sequence_id 10 representing our case are represented as $h(a(abc)(ac)d(cf))$. A sequence represented in this pattern is different. Then, we explore whether there is an option to change the order of projection of items when creating a projected dataset. It can be rationalized to consider every possible sequence and projections of the database that accompany the sequence in an orderly way in instances when one projects only the suffix of the sequence and keeps the sequence of the prefix.

Algorithm 1: Prefix Span Model

Input: A sequence database S and a minimum support threshold $min_support$.

Output: The complete set of sequential patterns

Initialize an empty set of sequential patterns.

Call the function Prefix-Span with an empty prefix α , length $l=0$, and the sequence database S .

Parameters; α as a sequential pattern, l as the length of α , and $S|\alpha$ as the α -projected database.

Define the function Prefix-Span ($\alpha, l, S | \alpha$) to recursively mine sequential patterns:

If α is empty, create a sequence database $S | \alpha$ from the original sequence database S .

Scan $S|\alpha$ once to find frequent items b that:

Can be appended to α to form a sequential pattern.

For each frequent item b , append it to α to form a new sequential pattern α' and output α' .

For each α' , construct the α' -projected database $S | \alpha'$ and recursively call Prefix-Span ($\alpha', l + 1, S | \alpha'$).

Execute the Prefix-Span subroutine recursively until all sequential patterns are discovered and outputted.

The Algorithm 1 improves the PrefixSpan algorithm by applying dynamic learning pattern recognition concepts. It is an effective tool to extract sequential patterns of the dataset S keeping in view the changing patterns of patterns in dynamic environments. The algorithm allows improving the pattern recognition methods in the framework of the given study by modifying PrefixSpan in order to fit dynamic learning scenarios.

Pattern Recognition Using LSTM

The more sophisticated LSTM operation learns to regulate information flow to prevent the disappearance of the gradient and simplify the process of the recurrent layers to identify relationships in the long term. The LSTM structure with unit representation is illustrated in **Fig 2**. RNN is also prone to the explosion or disappearance of gradients. The gradient following a phrase may not be able to back-propagate to the original point of a sentence due to a variety of nonlinearity changes, but, RNNs consider deep neural networks that traverse multiple time instances. Such issues are the main sources of motion that contributed to the creation of the LSTM framework the innovative architecture involves the introduction of a memory cell. The four major components of the memory cell include input, output, forget gates and a suitable cell to store the memory. In an LSTM neural network, the four gates can be represented by,

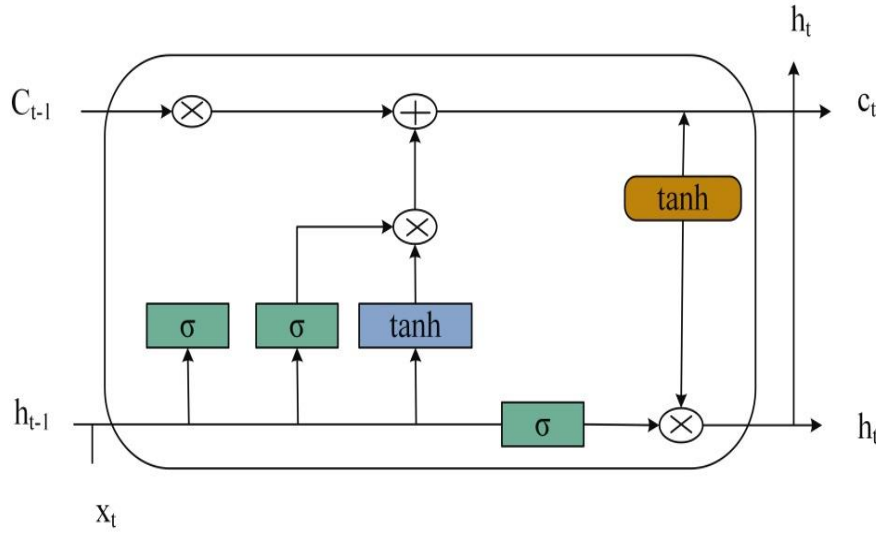


Fig 2. LSTM Structure.

The forget gate f_t is determined in equation (1), how much to forget in terms of the previous cell state. The sigmoid activation function for L_f calculates the sigmoid from the summed input x_t , and the weighted previous hidden state h_{t-1} .

$$f_t = \sigma(M_f x_t + L_f h_{t-1} + c_f) \quad (1)$$

The input modulation gate g_t corresponds to the candidate update to the cell state in this situation. $\tanh$ is the hyperbolic tangent function, which squashes the input sum, composed of the current input x_t , the previous hidden state h_{t-1} (weighted by L_g),

$$g_t = \tanh(M_g x_t + L_g h_{t-1} + c_g) \quad (2)$$

The amount of new information to store in the cell state, i.e. i_t is given by Equation (3). The input sum for the sigmoid function is σ , the previous hidden state h_{t-1} (weighted by L_o), and a bias term c_i .

$$i_t = \sigma(M_i x_t + L_o h_{t-1} + c_i) \quad (3)$$

Equation (4) computes the output gate o_t , which it determines how much to output of the cell state x_t as hidden state. Then the sigmoid activation function σ works over the input sum resulting from h_{t-1} (weighted by L_o), and a bias term L_o .

$$o_t = \sigma(M_o x_t + L_o h_{t-1} + c_o) \quad (4)$$

V. RESULTS AND DISCUSSION

The results section presents a critical discussion of the measures and analyses used in the proposed approach towards the identification of dynamic learning patterns. Based on the interpretation of the details of the data distribution, it can be seen that visual modes of learning are extremely preferred in the data. Sentence length analysis, in its turn, illuminates on the markers of textual difficulty which, in its turn, are also effective in terms of model tuning. Frequencies visualizations can also play a role in the derivation of essential themes and discussions in the corpus. The fact that the classification outcomes and their comparison with the existing models prove the superiority of the Prefix-Span algorithm referred to here is supported by the values of extremely high accuracy, recall and F1-score. These findings corroborated the finding of evolving trends to effective decisions and exploration of knowledge in reading materials. Thus, the findings demonstrate the importance of innovative methods such as sequential pattern mining and deep learning in the process of revitalizing the process of identification of patterns in dynamic learning settings applicable to other domains.

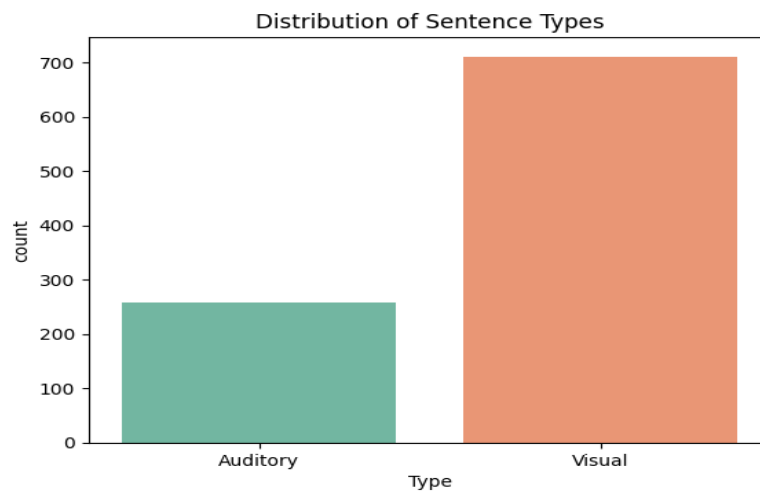


Fig 3. Distribution of Sentence Type.

The **Fig 3** element explains the division of the types of sentences based on three varying styles of learning i.e. Visual, Auditory. Visual rates are the mode category, which has 743 instances, showing that there is a strong preference to the learning aids that were predominantly visual. Auditory is the second that comes with a considerably lower number at a mere 257 cases thus showing that this category is of a far less popular acceptability with the methods of learning such as lectures or discussions. The analysis is interesting in terms of the variety of types of learning preferences that represent the data set, particularly the presence of the visual methods.

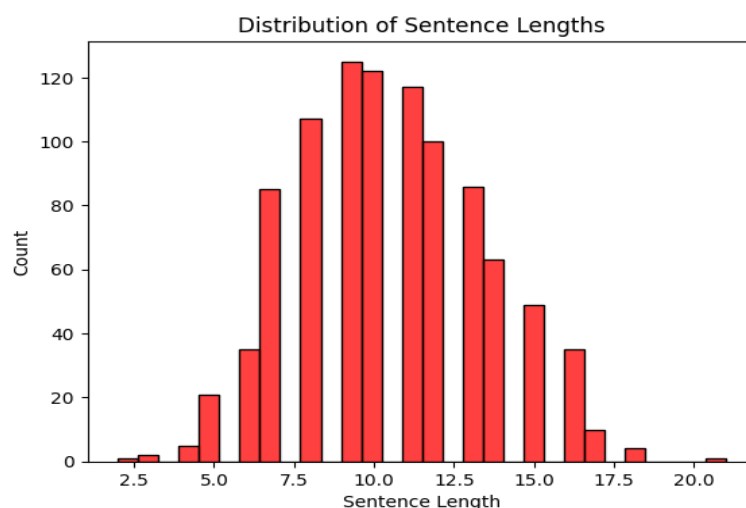


Fig 4. Distribution of Sentence Length.

The distribution of the length of sentences is depicted in the **Fig 4** that reveals the variation in the lengths of the sentences within a particular dataset or corpus. In that manner, it augments a main component of the expansion of dynamic learning patterns by the execution of a restructured Prefix Span Algorithm Model. The examination of sentence-

length distribution shows that the given dataset is highly complex structurally and linguistically, which are inherent characteristics of how well the model performs.

Fig 5 presented below is a bar graph indicating the frequency distribution of the 20 most commonly used words in the dataset or corpus. This graphic representation provides a helpful idea of lexical trends and language patterns that are present in the text. The analysis finds the most common terms and therefore can be useful to find key themes, and recurrent issues or discourse within the data.

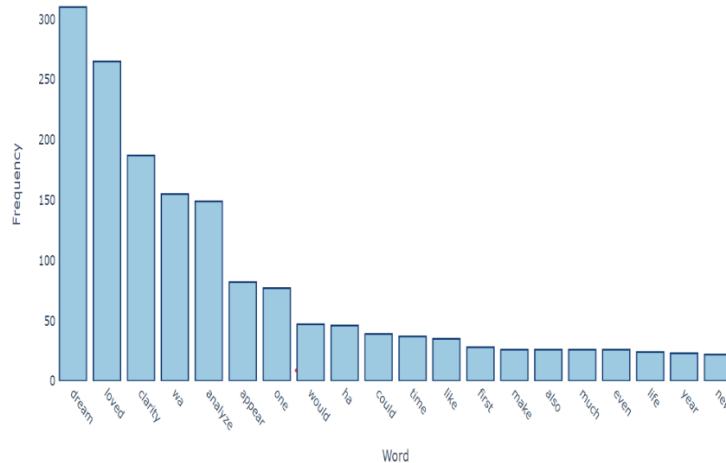


Fig 5. Bar-Plot of Top 20 Most Frequent Words.

Table 3. Text Before Pre-Processing

Sentence	Type	Cleaned Sentence
She loved children, and in default of them she...	Auditory	she loved children, and in default of them she...
This version is razor sharp, virtually flawles...	Visual	this version is razor sharp, virtually flawles...
Do you need two hours in the gym every day to ...	Visual	do you need two hours in the gym every day to ...

Table 3 presents the samples of text prior to pre-processing and each sentence is labeled auditory or visual, and then the cleaned text. Indicatively, the original sentence marked as auditory is left unchanged on cleaning. The visual category cases were standardized that is the conversion of the text to lowercase as well as the deletion of a punctuation mark. The impact of these preprocessing methods is that they create smoother versions of the original text, thereby making them easier to analyse, or process using computational methods, in the subsequent steps.

Table 4. Text After the Application of Word Net Lemmatizer

Sentence	Type	Tokens
For Derricke's final image is actually an idea...	Visual	[derricke, final, image, actually, idea, dream...
We derive our embarrassingly parallel VI algor...	Visual	[Derive, embarrassingly, parallel, vi, algorit...
She labored under the arduous burden of trying...	Visual	[Labored, arduous, burden, trying, achieve, cl...
His second major contribution is to analyze th...	Visual	[Second, major, contribution, analyze, underly...
The dream of many a French restaurateur is to ...	Visual	[Dream, many french, restaurateur, get, three...

Table 4 shows a sample of text that is undergone through the application of WordNetLemmatizer where each sentence is assigned with its type visual and presented in its lemmatized and cleaned appearance. The process of breaking words down into their root or dictionary form to facilitate the standardization and normalization of text in order to analyze it better is called lemmatization. This kind of lemmatization gives the sentences in the given case a homogenous and systematic way of rendering them. This step makes the text more interpretable and analytically more powerful, as it minimizes the variation in forms of words.

Fig 6 provides a confusion matrix that provides detailed analysis of the performance of the auditory and visual classes under classification. On the auditory category, 43 were categorized correctly and one was mistakenly categorized as

visual. Conversely, the visual group had 148 correct classifications with two of them being misclassified as auditory. This confusion matrix is essential in the establishment of the accuracy and efficiency of the model to distinguish auditory and visual learning data.

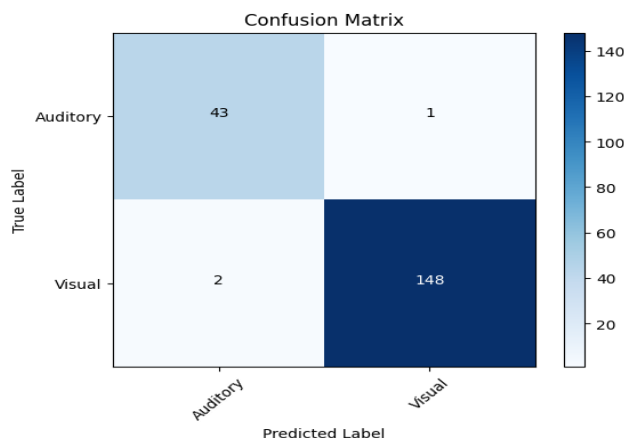


Fig 6. Confusion Matrix.

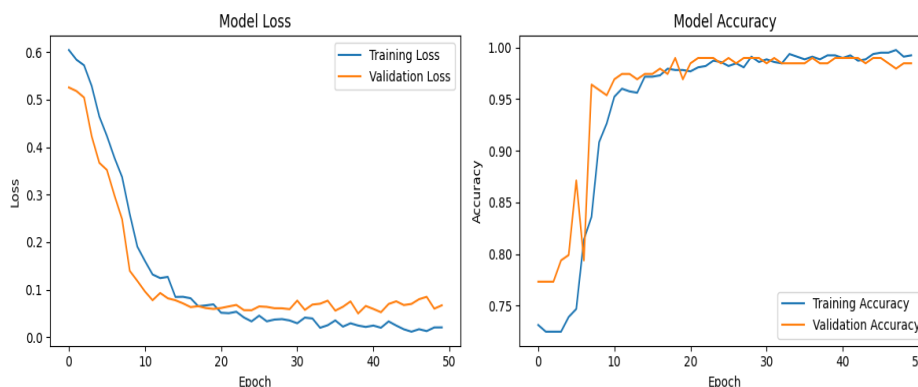


Fig 7. Model Accuracy and Loss.

In **Fig 7**, Model loss and model accuracy give the most significant indicators of determining the level of model or the created algorithm in machine learning. Model loss is usually the degree of well-performing of a model in its training stage usually expressed in a full numerical value indicating the variation between the actual and predicted outputs of a model. The lesser the loss the high the performance of the model, that is, the model is minimizing its errors in an excellent way. Model accuracy, conversely, displays the correct prediction of samples among all the ones assessed. It is, thus, a metric to the overall accuracy of the model in the classification of samples. When the values of the accuracy are high; it shows that the performance is better, however, when an accuracy score is one; it shows that all the predictions are correct. The test loss of the model is extraordinary at 0.067 and the test accuracy is at 98.45%. This value is representative of how well our method is able to capture the changing trends with minimal error.

Table 5. Classification Report of Proposed Method

	precision	recall	f1-score	support
0	0.96	0.98	0.97	44
1	0.99	0.99	0.99	150
Accuracy	-	-	0.98	194
Macro Average	0.97	0.98	0.98	194
Weighted Average	0.98	0.98	0.98	194

The classification report of the proposed method is provided in **Table 5**, and it shows the value of the assessment metrics as precision, recall, and F1-score of each class, and finally, the values of the overall accuracy, macro average, and weighted average. Precision is the degree to which the positive predictions are true and recall is the percentage of true positives that have been correctly identified and F1-score gives a harmonic mean of both the precision and recall. These findings suggest the great efficiency of the approach with accuracy ranging between 0.96 and 0.99; recall rates of 0.98 to 0.99; and F1 rates of 0.97 to 0.99. The total classification accuracy is 0.98, signifying the capacity of the model to correct prediction of all instances. Also, at the values of 0.97 and 0.98, respectively, the macro and the weighted averages of the precision, the recall and F1-score values show steady performance of the method in a scenario of class imbalance.

Table 6. Performance Metrics Comparison with Existing Classifiers

Methods	Accuracy	Recall	F1-score
MultinomialNB[23]	74.6%	70.1%	72.3%
Logistic Regression[23]	83.8%	82.3%	83%
MMCNet	92.4%	93.0%	92.7%
Proposed Prefix-Span	98.45%	98.62%	97.3%

Table 6 displays a general comparison of performance measures of the existing classifiers, Multinomial Naive Bayes (MultinomialNB) and Logistic Regression, alongside new MMCNet and the proposed Prefix-Span algorithm. The evaluation depends on the key performance indicators, i.e., Accuracy, Recall and F1-score. The results of both Multinomial Naive Bayes and Logistic Regression are average with an accuracy of between 74.6 and 83.8 per cent. Conversely, there is a major variance in MMCNet with a level of accuracy of 92.4. Based on all of the above, the Prefix-Span algorithm was determined to be the most efficacious in comparison to all other models with approximately an efficacious accuracy at 98.45, high recall, and also an equally satisfying F1-score. This actually underscores the excellence of the algorithm when it comes to identifying dynamic trends than the conventional way of classifying data.

Discussion

An analysis of the results in this part proves the validity of the suggested methodology of the recognition of dynamic learning patterns, particularly in the text-domain context. Distribution characteristic analysis shows a greater preference towards visual learning styles, on the basis of which we may deduce that the probability of users being extremely inclined to working with visual materials is likely to be high. Sentence length assessment is also the hallmark of gauging textual complexity—a significant aspect towards enhancing the accuracy and effectiveness of the model. Word frequency visualizations also revealed certain strong underlying themes and generally discussed items of the dataset. The outcomes of classification suggest that there is high performance motive in the results of commendable accuracy, recall and F1-scores; the Prefix-Span algorithm is good at this. Conclusively, one may infer that these findings refer to the clear value of sequential pattern mining being implemented together with a deep-learning approach to the successful isolation of dynamic tendencies in text information, which can critically contribute to decision-making and Knowledge discovery in most spheres of application.

VI. CONCLUSION AND FUTURE WORKS

The strong approach that is adopted in the research expands the understanding of dynamic learning behaviors of text material studied. The data set was systematically formatted to further analysis using the application of data cleaning process, exploratory data analysis and text processing. The effective sequential pattern mining of the dynamic educational contexts was accomplished by the implementation of Enhanced Prefix Span Algorithm Model. Moreover, the Long Short-Term Memory (LSTM) networks were also implemented to expand the model ability in tracking patterns and a narrower view of recognizing long term relationships in the dataset. It was a combination of such techniques that relevant insights into a diversity of learning behaviours and preferences and the intricate interaction between the learning method and the textual content of a specific individual were obtained. Essentially, this work contributes to the development of pattern recognition methods in adaptive learning environment and identifies a field that can potentially be researched further in this respect. Based on the findings and the methods of this study, there are a number of avenues that can most definitely be pursued as the future research directions. An example of optimization would be in regard to the sequential pattern mining algorithm such as the Enhanced Prefix Span Algorithm in terms of efficiency and scalability with respect to larger and more complicated datasets. Moreover, the development achieved in complex machine learning systems like within the domain of deep learning may also open new opportunities of pattern identification with respect to dynamic patterns. Other types that can be studied in the study analysis comprises of multimedia or user interaction logs that indicates a more complete and more valuable learning behaviour. Moreover, longitudinal studies can also be required to know how learning patterns evolve with time and how they evolve with respect to certain interventions. Secondly such research results in the educational technologies might be translated into practice by transforming theorization and progress into practice into actual effects on learners and educators in the framework of personalized learning systems, adaptive tutoring tools, and so on. Within this area, there remains the sense of determination to persist with that hope of expanding knowledge on the behavior, dynamic learning and developing the proper educational strategy and the systems.

CRedit Author Statement

The authors confirm contribution to the paper as follows:

Conceptualization: Ramesh M and Jayashree R; **Methodology:** Ramesh M; **Software:** Jayashree R; **Data Curation:** Ramesh M; **Writing- Original Draft Preparation:** Ramesh M and Jayashree R; **Visualization:** Ramesh M; **Investigation:** Jayashree R; **Supervision:** Ramesh M; **Validation:** Jayashree R; **Writing-Reviewing and Editing:** Ramesh M and Jayashree R; All authors reviewed the results and approved the final version of the manuscript.

Declaration of competing interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Data Availability Statement

The Datasets used and /or analysed during the current study available from the corresponding author on reasonable request.

Funding

No funding agency is associated with this research.

References

- [1]. M. Yang, B. Yang, M. Liao, Y. Zhu, and X. Bai, “Sequential visual and semantic consistency for semi-supervised text recognition,” *Pattern Recognition Letters*, vol. 178, pp. 174–180, Feb. 2024, doi: 10.1016/j.patrec.2024.01.008.
- [2]. G. Peng, “Segmentation-free Recognition Algorithm Based on Deep Learning for Handwritten Text Image,” *Journal of Artificial Intelligence and Technology*, Mar. 2024, doi: 10.37965/jait.2024.0473.
- [3]. X. Zhang et al., ‘Realcompo: Dynamic equilibrium between realism and compositionality improves text-to-image diffusion models’, arXiv preprint arXiv:2402.12908, 2024.
- [4]. K. Sai and R. Thokala, “Pattern Recognition in Medical Decision Support and Estimating Redundancy in Clinical Text,” *Proceedings of the 1st International Conference on Artificial Intelligence, Communication, IoT, Data Engineering and Security, IACIDS 2023*, 23–25 November 2023, Lavasa, Pune, India, 2024, doi: 10.4108/eai.23-11-2023.2343233.
- [5]. S. Venkatachalam and Rajiv Kannan, “Optimizing dynamic keystroke pattern recognition with hybrid deep learning technique and multiple soft biometric factors,” *International Journal of Computers Communications Control*, vol. 19, no. 2, Mar. 2024, doi: 10.15837/ijccc.2024.2.6097.
- [6]. D. Zhong et al., “NDOOrder: Exploring a novel decoding order for scene text recognition,” *Expert Systems with Applications*, vol. 249, p. 123771, Sep. 2024, doi: 10.1016/j.eswa.2024.123771.
- [7]. C. Thaokar, J. K. Rout, M. Rout, and N. K. Ray, “N-Gram Based Sarcasm Detection for News and Social Media Text Using Hybrid Deep Learning Models,” *SN Computer Science*, vol. 5, no. 1, Jan. 2024, doi: 10.1007/s42979-023-02506-5.
- [8]. X. Wu, K. Zhao, Q. Huang, Q. Wang, Z. Yang, and G. Hao, “MISL: Multi-grained image-text semantic learning for text-guided image inpainting,” *Pattern Recognition*, vol. 145, p. 109961, Jan. 2024, doi: 10.1016/j.patcog.2023.109961.
- [9]. S. Huang, W. Fu, Z. Zhang, and S. Liu, “Global-local fusion based on adversarial sample generation for image-text matching,” *Information Fusion*, vol. 103, p. 102084, Mar. 2024, doi: 10.1016/j.inffus.2023.102084.
- [10]. F. Yang, S. Yang, M. A. Butt, J. van de Weijer, and others, ‘Dynamic prompt learning: Addressing cross-attention leakage for text-based image editing’, *Advances in Neural Information Processing Systems*, vol. 36, 2024.
- [11]. J. Huo, J. Yu, M. Wang, Z. Yi, J. Leng, and Y. Liao, “Coexistence of Cyclic Sequential Pattern Recognition and Associative Memory in Neural Networks by Attractor Mechanisms,” *IEEE Transactions on Neural Networks and Learning Systems*, vol. 36, no. 3, pp. 4959–4970, Mar. 2025, doi: 10.1109/tnnls.2024.3368092.
- [12]. M. Yang, B. Yang, M. Liao, Y. Zhu, and X. Bai, “Class-Aware Mask-guided feature refinement for scene text recognition,” *Pattern Recognition*, vol. 149, p. 110244, May 2024, doi: 10.1016/j.patcog.2023.110244.
- [13]. S. Momeni and B. BabaAli, “A transformer-based approach for Arabic offline handwritten text recognition,” *Signal, Image and Video Processing*, vol. 18, no. 4, pp. 3053–3062, Jan. 2024, doi: 10.1007/s11760-023-02970-9.
- [14]. J. Zheng, L. Zhang, Y. Wu, and C. Zhao, “BPDO: Boundary Points Dynamic Optimization for Arbitrary Shape Scene Text Detection,” *ICASSP 2024 - 2024 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pp. 5345–5349, Apr. 2024, doi: 10.1109/icassp48485.2024.10447371.
- [15]. Z. Wei, Y. Gao, X. Zhang, X. Li, and Z. Han, “Adaptive marine traffic behaviour pattern recognition based on multidimensional dynamic time warping and DBSCAN algorithm,” *Expert Systems with Applications*, vol. 238, p. 122229, Mar. 2024, doi: 10.1016/j.eswa.2023.122229.
- [16]. X. Zhang, B. Zhu, X. Yao, Q. Sun, R. Li, and B. Yu, “Context-Based Contrastive Learning for Scene Text Recognition,” *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 36, no. 3, pp. 3353–3361, Jun. 2022, doi: 10.1609/aaai.v36i3.20245.
- [17]. M. Yang et al., “Reading and Writing: Discriminative and Generative Modeling for Self-Supervised Text Recognition,” *Proceedings of the 30th ACM International Conference on Multimedia*, pp. 4214–4223, Oct. 2022, doi: 10.1145/3503161.3547784.
- [18]. D. Coquenot, C. Chatelain, and T. Paquet, “End-to-End Handwritten Paragraph Text Recognition Using a Vertical Attention Network,” *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 45, no. 1, pp. 508–524, Jan. 2023, doi: 10.1109/tpami.2022.3144899.
- [19]. K. Li, Y. Zhang, K. Li, Y. Li, and Y. Fu, “Image-Text Embedding Learning via Visual and Textual Semantic Reasoning,” *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 45, no. 1, pp. 641–656, Jan. 2023, doi: 10.1109/tpami.2022.3148470.
- [20]. E. Vidal, A. H. Toselli, A. Ríos-Vila, and J. Calvo-Zaragoza, “End-to-End page-Level assessment of handwritten text recognition,” *Pattern Recognition*, vol. 142, p. 109695, Oct. 2023, doi: 10.1016/j.patcog.2023.109695.
- [21]. L. M. Francis and N. Sreenath, “Robust scene text recognition: Using manifold regularized Twin-Support Vector Machine,” *Journal of King Saud University - Computer and Information Sciences*, vol. 34, no. 3, pp. 589–604, Mar. 2022, doi: 10.1016/j.jksuci.2019.01.013.
- [22]. ‘Learning Style (VAK)’. Accessed: Mar. 26, 2024. [Online]. Available: <https://www.kaggle.com/datasets/zeyadkhalid/learning-style-vak>
- [23]. S.-S. Park and K. Chung, “MMCNet: deep learning-based multimodal classification model using dynamic knowledge,” *Personal and Ubiquitous Computing*, vol. 26, no. 2, pp. 355–364, Aug. 2019, doi: 10.1007/s00779-019-01261-w.

Publisher’s note: The publisher remains neutral with regard to jurisdictional claims in published maps and institutional affiliations. The content is solely the responsibility of the authors and does not necessarily reflect the views of the publisher.