

Automatic Essay Score Prediction in Bengali

Piyal Roy¹, Rajat Pandit^{1*} and Sudip Kumar Naskar²

¹Computer Science, West Bengal State University, West Bengal, India.

²Computer Science and Engineering, Jadavpur University, West Bengal, India.

*rajatpandit123@gmail.com

Abstract:

Automatic essay scoring (AES) is a crucial task in educational technology that aims to evaluate student essays objectively and accurately. While numerous AES models have been proposed over the years for languages with large spoken populations, there has been very little research on this topic for low-resource and morphologically complex languages, such as Bengali. To address this gap, we developed a custom dataset in Bengali designed explicitly for this task, on which we carried out comprehensive experiments using various machine learning and language processing techniques. The essays were preprocessed and analyzed to extract features relevant to the quality of the essays. We employed multiple regression models, including simple linear regression, Decision Tree Regression, Ridge Regression, LASSO Regression, Elastic Net Regression, and Support Vector Regression (SVR) models, to predict the scores for Bengali essays. Furthermore, we also conducted experiments with state-of-the-art Large Language Models (LLMs), such as ChatGPT and Gemini, to assess their performance on this task. They were evaluated using the Pearson Coefficient. In our experiments, the Decision Tree Regressor outperforms other models, achieving a Pearson correlation coefficient of 0.905. However, we found that LLMs are not well-suited for this task, as they do not achieve the same performance as traditional regression models. Only ChatGPT performed well, with a Pearson correlation of 0.605. Our findings suggest that integrating LLMs into the system's regression-based AES may increase the system's overall accuracy for Bengali.

Keywords: AES, Regression, LLM, Perplexity, NLP.

1. INTRODUCTION

1.1. Background and Motivation

Essays play a significant role in the curricula of schools and universities worldwide. It serves as a good instrument to evaluate students' understanding, competence, and critical thinking abilities. Essay evaluation is a problem that has been machine-done since it has manual conduct. It is time-consuming, expensive, and prone to human error as the manual essay scoring process is comprehensive but is still in use today. Believing that there is a growing demand for essays that are to be evaluated, it becomes necessary to automate the scoring system. There are also systems, such as AES, based on machine learning algorithms that evaluate and score human-written texts; an example of the contribution the field of Natural Language Processing brings. There are several benefits to Automated Essay Scoring: it has high reliability, is scalable, scoring is consistent, and this translates to reduced time spent scoring. For both theoretical and experimental works on AES for English and other overt languages, research has been conducted in the literature. However, research on low-resource languages, such as Bengali, is still in its nascent stages and needs to be explored further.

1.2. Importance of Essay Scoring in Educational Domain

Essays are an essential part of the educational system, as they test a student's ability to write, identify ideas, analyze, and present a coherent argument. They are extensively used in tests such as standardized tests, classroom evaluations, and admissions processes. Particularly in the widely spoken language of Bengali, a system incorporating the use of an automated essay scoring system is of the utmost importance. These systems also enhance fairness and affordability and allow wider audience to evaluate. The AES is needed in countries where Bengali is one of the official or regional languages. The reason for this is that many students will be used for formal tests and standard assessments. A Bengali essay would be scored with an automated system because this would not only reduce manual load but would also provide students with timely performance feedback, and as a result promote a culture of continuous learning and development.

1.3. Challenges Specific to Bengali Essay Scoring

The presentation of the development challenges of an AES system for Bengali, due to the scarcity of annotated resources and linguistic properties of Bengali are given. Bengali is an Indo-Aryan, or Indo-European, language. It has a great deal of inflection: several suffixes that can change the meaning of words, and complex verb conjugation. The feature extraction and model training become much more complicated due to this linguistic complexity. In addition, we face an unsurmountable hurdle – there are no well-annotated Bengali essay datasets that we need to map and evaluate in order to match the student's work.

1.4. Objectives and Scope of the Research

The purpose of this research is to develop an efficient and secure AES system for Bengali essays. Compiling a comprehensive dataset of Bengali essays with standardized scorers' evaluations is what this amount to. We prepared an annotated essay dataset to train and evaluate the AES models. The task in this work is to identify and extract meaningful linguistic, semantic, and syntactical features of the Bengali essays and then train models to score essays being evaluated by essay scoring. We will utilize various deep learning and machine learning algorithms to score an essay. The standard metrics will be used to train and test the models for selection and validation. However, we hope that our study can contribute to the educational evaluation of low-resource languages and should be relevant to both the educational technology community and natural language processing.

2. LITERATURE REVIEW

Because of its ability to provide easily comparable, accurate, and exhaustive assessment of written language, the task of AES has attracted a great deal of interest within NLP. The AES roots can be traced back to the early 1960s when Page introduced “Project Essay Grade” (PEG), which used simple statistical indices to serve as an indicator of essay quality (Page, 1966). The centuries of AES have progressed from initial methods of data handling over the 1990s to today's learning systems. Despite the evolution of computational linguistics over the years, there have been many techniques and algorithms offered in literature, including latent semantic analysis (LSA) (Landauer, Foltz, & Laham, 1998), support vector machine (SVM) (Dikli, 2006), and the recent one such as deep learning such as convolutional neural networks (CNN) and recursive neural networks (RNN) (Dong & Zhang, 2016).

While a considerable amount of research has been conducted on AES in English, a problem arises when implementing AES in other languages, specifically the Bengali language. To cite an instance, Islam, 2011 acknowledge that Bengali is characterized by high linguistic density and morphological variety, which necessitates personalization at the essay grader level for accurate calculation. Training accurate models is made challenging by the lack of Bengali corpora with annotations, as noted by Chakraborty et al. (2021). Although feature engineering has always been an important part of constructing AES systems, it remains an integral component of these systems to this date. Shermis and Burstein (2013) refer to them as the first generation, and they were primarily based on basic count statistics, such as word frequencies, syntactic density, and lexical variety. However, due to the latest advancements in neural networks, embeddings, and attention mechanisms (Vaswani et al., 2017), feature extraction has mainly been replaced. Recently, some works using Bengali have employed character-level embedding in conjunction with bidirectional LSTMs (Karim et al., 2020). It has been proven that these methods are fruitful, and in these evaluations, there is potential to conduct further research on these models to optimize them for the required characteristics of Bengali essays.

The results generated by the models (whether correct or not) are then compared to the assessment scores provided by human annotators for the AES systems. As suggested by Williamson (2009), the commonly used assessment tools are Pearson correlation, quadratic weighted kappa, and accuracy. However, the debate on the use of Bengali has arisen due to similar use caused by the absence of a standard evaluation dataset. This has led to variations in the outcomes of different studies (Islam & Hoque, 2013). This will help avoid these discrepancies and determine the models that would be useful for rating. Uniform feedback provided to students from remote and low-income areas can be equally effective in promoting equity among Bengali learners through an effective AES system, as noted by Saha and Pal (2020).

Advanced limes have already been developed for the AES systems; however, their implementation in Bengali is a significant issue, both in terms of development and prospects. The Bengali language is

3. METHODOLOGY

3.1. Description of the Dataset

3.1.1. Source of Essays

Bengali Data Collection			
Name: _____			
Age: _____	Gender: Male ☐ Female ☐ Rather not say ☐		
Class: Up to X ☐ XI-XII ☐ UG ☐ PG ☐ Department (In case of UG/PG): _____			
Write a short Bengali essay on “একটি ব্রহ্মের অভিজ্ঞতা”:			

Figure 1: Blank essay collection form.

Bengali Data Collection

Name: <u>Anupam Mondal</u>				
Age: <u>22</u>	Gender: <u>Male</u>	Female	Rather not say	
Class: <u>Upto X</u>	<u>XI-XII</u>	<u>UG</u>	<u>PG</u>	Department (In case of <u>UG/PG</u>) <u>Bengali</u>

Write a short Bengali essay on "যাকি কালের অধিকাৰ":

06

10

শিৰোনাম?

ইতিহাসিক কালজ্ঞানৰ ব্যৱহাৰ আৰু অৱস্থা
 মানৱে এনে ভাৱকেই হিচাপে অৱস্থাত
 বিকাশ আৰম্ভ কৰিছিল যিটোকে কালজ্ঞান বোলে।
 তেওঁৰে সাধুৰূপে ও বিহীনকৈ হৈছে আৰম্ভণি
 ভাৱমূলক জ্ঞান, এই হৈছে ভাৱমূলক জ্ঞান
 কালজ্ঞান- অৱস্থা।

-মাধৱী ও বিহীনকৈ হৈছে আৰম্ভণি
 জীৱনৰ জ্ঞান (জ্ঞান) হৈছে, মাধৱী আৰু
 হৈছে এনে ভাৱমূলক জ্ঞান আৰম্ভণি
 কালজ্ঞানৰ পৰা আৰম্ভ হৈছে জ্ঞান, হৈছে
 ৪৩০ৰ পৰা জ্ঞান হৈছে জ্ঞান, এনে হৈছে
 কালজ্ঞান হৈছে জ্ঞান হৈছে জ্ঞান হৈছে
 কাল জ্ঞান হৈছে হৈছে হৈছে হৈছে
 কাল জ্ঞান, জ্ঞানৰ জ্ঞানৰ জ্ঞানৰ জ্ঞান
 জ্ঞান জ্ঞান ও জ্ঞানৰ জ্ঞান হৈছে
 জ্ঞানৰ জ্ঞানৰ জ্ঞানৰ জ্ঞান
 জ্ঞানৰ জ্ঞানৰ জ্ঞানৰ জ্ঞান

Figure 2: Filled essay collection form.

3.1.2. Data Preprocessing Steps

Piyal Roy et al 536-543

Verification of Transcriptions: All the transcribed text was carefully compared to handwritten responses to confirm the accuracy of the digitization. Hence, this effort was made to obtain a 'careful' comparison, intended to minimize transcription errors (such as misreading handwriting) and, therefore, ensure high reliability in the data.

Formatting Consistency: All the essays were thoroughly formatted to have correct punctuation, spacing, and overall structure. This process was designed to address some of the inevitable issues that could arise during the transcription process and maintain a consistent formatting style. The accuracy and integrity of the data were to be maintained by ensuring uniformity.

Data organization: The transcribed essay was matched with the score that was assigned to it by the human grader and placed into a structured CSV file. The whole essay text file was stored in one column of each row of the file, while the associated score was stored in another adjacent column. The organization ensured that, in the later stages of model training, the setup provided seamless access and analysis.

The preprocessing steps helped ensure that the essays comprising the data were adequately digitized and formatted for the development of this essay scoring model.

3.1.3. Annotation and Scoring

The digitized and preprocessed essays were evaluated by a domain expert with extensive knowledge of the Bengali language and literature. The expert graded each essay based on a set of predefined criteria, which included coherence and organization, assessing the essay's structural flow and logical presentation of ideas; language and grammar, focusing on accuracy, complexity, and proper use of grammar, syntax, and vocabulary; content relevance, evaluating the alignment of the essay with the given topic and the depth of the ideas presented; and creativity and originality, considering the uniqueness and creativity in expressing ideas. It was these criteria that guaranteed a thorough and uniform evaluation of the essay.

Each essay was scored by an expert on a scale of 1 to 10, with 1 indicating a poorly written essay with little or no relevance or coherence and 10 indicating an excellently written essay that met all the parameters to the highest extent. In the next column of the corresponding CSV file, the essay text was aligned to the given score for each essay.

	A	B
1	essay	marks
1	দীঘা ভ্রমণ দীঘা একটি পর্যটন কেন্দ্র। এটি বঙ্গোপ সাগরের তীরে অবস্থিত। দীঘা একটি বিখ্যাত পর্যটন কেন্দ্রের মধ্যে একটি। এখানে বহু মানুষ ছুটি পেলে সময় কাটাতে আসে। সমুদ্র দেখতে আসে এবং স্বাস্থ্যকর পরিবেশে কিছুদিন সময় কাটায়। এখানে অনেক নামিদানি হোটেল আছে। অনেক পর্যটকরা এসে হোটলে থাকে তাদের হোটেল থেকে আয় করে জীবিকা সংগ্রহ করে। এছাড়া অনেক ছোট ছোট বোট, লঞ্চ, স্টিমার দেখা যায়। ভারতবর্ষের বহু মানুষ এখানে বসবাস করে। আমি সকলের মত আমিও বেড়াতে ভালোবাসি। দীঘা সমুদ্র দেখার আমার খুব ইচ্ছা ছিল। গত সরস্বতী পূজার ছুটিতে আমার মা, বাবা, দিদি আর আমি দীঘা বেড়াতে গিয়েছিলাম। হাওড়া থেকে কালকা মেলে চেপে দিল্লি বেড়াতে গিয়েছিলাম। সেই প্রথমবার আমি দূরপাল্লার ট্রেনের চেপে যেতে মনে দারুন আনন্দ হচ্ছিল। বেকাল পাঁচটায় ট্রেনে চেপে সারারাত ট্রেনে কাটিয়ে ভোর তিনটায় দীঘা পেছনে পৌঁছালো। নিউ দীঘা স্টেশনে। বাকবাকে একটা সুন্দর স্টেশন দেখলাম। সেখান থেকে আমরা একটি অটো নিয়ে একটা হোটলে পৌঁছলাম। প্রাকৃতিক সৌন্দর্য- খুব সুন্দর এই হোটেল। অনেক ছোট ছোট ঘর। হোটেলের দুতলা ঘরে আমাদের জায়গা হল। ভোরের আবহা আলোয় সমুদ্র দেখতে পেয়ে মন আনন্দে ভরে গেল। বেরিয়ে পরলাম হোটেল থেকে সূর্য দেখতে। অদৃশ্য এক দৃশ্য দেখলাম। প্রথমে মনে হল সূর্য সমুদ্রের জল থেকে উঠছে। আমার সামনে পিছনে দুদিকে দূরবাম জলের রং। সূর্য যেন খুব তাড়াতাড়ি উঠে পড়ল। এখানে শুধু সেই স্মৃতি চিহ্নটা মনে পড়ে। মনে হয় আবার যাবো সমুদ্র সৈকত দেখতে।	6
2	বকখালি ভ্রমণ আমার প্রিয় উইকেন্ড ডেস্টিনেশন গুলির মধ্যে পশ্চিমবঙ্গের বকখালি অন্যতম। যদিও এখানকার সমুদ্র দীঘার মত যদিও খুব উত্তাল নয় এবং সব সময় তেউ পাওয়া যায় না, জোয়ারের সময় ছাড়া। বিচ এর সাইড এ অনেক দোকান আছে ইচ্ছা হলে টুকি টাকি ঘর সাজানোর জিনিস কিনতে পারা যায়। ট্রেনের টিকিট বুকিং নিয়ে কোন চিন্তা নেই। এবার আমার বকখালি যাত্রা শুরু করছি। তখন সময় সকাল 6:40। হাসনাবাদ শিয়ালদহ লোকাল ঘরে আমি শিয়ালদহ স্টেশনে পৌঁছলাম। যেখান থেকে নামখানে লোকাল ঘরে বকখালি গিয়ে পৌঁছলাম। মোটামুটি সময় হবে 11:30 মতো। সেখানে গিয়েই সমুদ্রের দৃশ্যটা অতি সুন্দর। তো আমরা কয়েকজন বন্ধু মিলে গিয়েছিলাম। ওখানে গিয়ে পরে আমরা আগে একটা Hotel খুঁজলাম কারণ আমাদের ইচ্ছা ছিল যে দুই রাত আমরা ওখানে কাটাবো। যথারীতি সেই দিন আমরা সকালে দোকান থেকে খাওয়া দাওয়া করে ঘুরতে বেরিয়ে পড়লাম। বকখালি তে সন্ধ্যার থেকে দৃশ্যটা আরো সুন্দর লাগে তার কারণ সমস্ত দোকান শারি সারি লাইন দিয়ে থাকে। সমস্ত রকমের ঘর সাজানোর জিনিস পাওয়া যায় এখানে। এক হিসেবে ওখানে যারা নিত্য যাতায়াত করে বা ওখানকার বাসিন্দা তারা খুবই ভালো মনের মানুষ তার কারণ হল আমার একটা বন্ধুর শরীর খারাপ হওয়ায় আমরা কি করব সেটা বুঝতে না পেরে একটা দাদাকে ডেকেছিলাম সেই দাদাটা শুধু আমাদের কে পরামর্শই দেয়নি, আমাদেরকে নিজে Doctor এর কাছে নিয়ে গিয়েছিলেন। এই থেকে তাদের আচার-আচরণের অনেকটা পরিচয় মিলে এবং সেখানকার পরিবেশ শুধু মৌসুমী প্রকৃতির হয়ে থাকে। তাই খুব ভালোভাবেই মানুষ সেখানে মিশে যেতে পারে। শীতকালেই বেশি যেতে ভালো লাগে। সেখানকার লোক সংস্কৃতির দিক দিয়ে দেখতে গেলে দেখা যায় যে কোন রকম কোলা হল নেই। শান্ত পরিবেশ হয়ে থাকে। অনেক ভালো গাছ, নারকেল গাছ দেখতে পাওয়া যায় এবং নানা ধরনের ফুল গাছ লাগানো থাকে যা একটা গ্রাম্য সুন্দর পরিবেশের বার্তা বহন করে আর আমি ওখানে গিয়ে অনেক খুশি।	6

Figure 3: Sample of the CSV file.

The annotation of the dataset was performed thoroughly and also matched the standard of the Bengali Language assessment. The scores served as a benchmark for training and testing the AES system, enabling it to learn and predict scores for new essays.

3.2. Feature Extraction

The development of an AES system requires a feature extraction step. In this case, it involves identifying and quantifying some attributes of the essays that can help predict the overall quality of the essays and the score for a uniform rating scale.

The word count feature keeps track of the total number of words in each essay. It provides a straightforward yet highly effective method of measuring essay length – something that usually correlates with quality. Relevance and coherence, as well as having more extended essays, can suggest more extensive responses. The number of words in each digitized essay was counted, serving as a measure of this feature.

Average word length is calculated by dividing the total number of words by the total number of essays. Its usefulness is to see the amount of work put into the language the writer used. Essays with above-average word lengths may suggest a higher level of writing proficiency.

Sentence complexity is a measure of how complex the sentences in an essay are. For this experiment, two metrics were used, namely, average words per sentence and maximum words per sentence. The average number of words per sentence provides a general indication of the number of words used in each sentence. In contrast, the maximum number of words in a sentence indicates sentences that are loaded with words that are either complex or compound. Higher sentence complexity does not necessarily indicate a more sophisticated writing style; it must also be readable.

The spelling error feature intercepts the number of spelling mistakes made in each essay. A spelling error can interrupt the attractiveness of an essay as it implies the quality of the writer's attention to detail and language proficiency. This feature was computed across every word from an essay by counting the number of misspelled words against a list of correctly spelled Bengali words.

Perplexity score is defined as the measure of how well a probabilistic language model predicts a piece of text. For this work, the perplexity score was calculated for each essay using a fine-tuned KenLM model (Heafield, 2011). Lower perplexity scores indicate that the essay effectively describes the pattern learned by the language model, generally associated with fluency and coherence. On the other hand, higher values in perplexity indicate more unexpected or unconventional use of language, which can affect readability and quality.

3.3. Model Selection and Training

Developing an effective AES system is indeed the critical choice of model and training process. Prediction of the essay scores was done through several different regression algorithms, and we compared with LLMs as well.

3.3.1. Choice of Algorithms

Several regression models were chosen to predict the essay scores based on the ability to deal with different data distributions and the effectiveness of these models in different predictive tasks. The following algorithms were utilized:

Simple Linear Regression: For its simplicity and interpretability, and due to the problem of the data, a simple linear regression was selected. It is to derive a linear relationship between the features extracted from the essays and the target scores of the essays.

Decision Tree Regression: In this case, we used decision tree regression to capture more complex relationships in data. The decision tree can also learn additional interactions that linear regression may miss by dividing the data into different branches based on feature values.

Ridge Regression: It utilizes ridge regression by implementing the linear regression model, adding a regularization term that makes the model less sensitive to feature correlations.

LASSO Regression: LASSO Regression is also based on regularization terms, which means it sets some feature coefficients to zero during feature selection. The model is simplified by focusing only on the most important features.

Elastic Net Regression: Elastic Net Regression combines the penalty of Ridge and LASSO regression, giving a trade-off between these two. Specifically, feature selection has the additional advantage of being a useful approach in either a robust feature selection method when multiple features are correlated or in model regularization.

We also used two LLMs ChatGPT and Gemini to predict essay scores and assess the ability of these LLMs and traditional regression models to do so. These models are trained using deep learning techniques that have been pre-trained on a huge corpus of text data that are capable of understanding and assessing essays in a better way.

3.3.2. Model Architecture

Each regression model was a simple architecture of the relationship between the extracted features and the essay scores. For LLM, the architecture involves generating essay text-based score predictions using pre-trained models without requiring any fine-tuning.

Simple Linear Regression, Ridge Regression, LASSO Regression, and Elastic Net Regression: The extracted features are used as the input variables in these models. The result was the continuous prediction of the score based on linear or regularized linear relations between the features and the target variable.

Decision Tree Regression: This model employs a tree-based structure, where each node represents a decision based on a feature value, ultimately leading to a final prediction at the leaves of the tree.

ChatGPT and Gemini: The essay scores were predicted by taking the raw text of each essay as input and using the LLMs, ChatGPT, and Gemini. They deployed their advanced deep learning architectures, based on layers of transformers, attention mechanisms, and other processing techniques, on the text to generate a score. For this specific task, we employed few-shot learning to improve the performance of these models, that is, to learn to produce essay scores given only a small sample of example essays with their corresponding accurate scores. This approach improved the models' learning from the pre-training task. It allowed them to build upon their existing skills and develop a deeper comprehension of the scoring criteria central to this research.

3.3.3. Training Process and Hyperparameter Tuning

This dataset was divided into training and test sets on an 80:20 basis. The models were trained on the training set and tested on the test set.

Fitting the regression models to the training data by minimizing the loss function was a key part of the training process. Based on this, hyperparameter tuning was performed on models such as Ridge, LASSO, and Elastic Net to determine the optimal regularization parameters, which define the strength of the penalty applied to the model coefficients.

ChatGPT and Gemini were not trained or tuned in this setting, as they were pre-trained models applied pointwise to the task. Finally, their performance was evaluated on the test set together with the regression models.

3.4. Evaluation Metrics

The Pearson correlation coefficient (r) is a statistical measure that indicates the strength and direction of the linear relation between the predicted essay scores and the human grader-assigned scores (Rodgers & Nicewander, 1988). The value of the Pearson correlation ranges from -1 to +1, where:

- **+1** indicates a perfect positive linear relationship, meaning that as the predicted scores increase, the actual scores also increase proportionally.
- **0** indicates no linear relationship between the predicted and actual scores.
- **-1** indicates a perfect negative linear relationship, meaning that as the predicted scores increase, the actual scores decrease proportionally.

In our work, we prefer a higher positive Pearson correlation coefficient because it indicates that our model is more closely aligned with human measures. This is a simple, interpretable metric for which a scoring system is commonly evaluated (Cohen et al., 2002). We calculated the predictions for each model on the test set and computed the Pearson correlation for each model's predictions to compare their direct performance. Important for the deployment of an AES was a metric that was both transparent and interpretable of the model's ability to reproduce the scoring patterns of the human expert.

4. EXPERIMENTS AND RESULT

4.1. Experimental Setup

This research relates to experiments running to investigate the effectiveness of different regression models and LLMs in predicting essay scores. We used Simple Linear Regression, Decision Tree Regression, Ridge

Regression, LASSO Regression, Elastic Net Regression, Support Vector Regression models. Also, with these new LLMs, like ChatGPT and Gemini, cutting edge language processing techniques were applied to discover new use cases to be used for essay scoring with the technologies. All models were built using Scikit-learn, one of the popular machine learning tools of the language Python (Pedregosa et al., 2011). The experimentation was carried on the Google Colab platform with sufficient computational power available in hand, and it was the best way to do machine learning on a large scale. We used perplexity score as another metric to evaluate how well model performs. Perplexity gave deeper insights into how the models could understand and process language.

4.2. Results of Different Models

We evaluated the performance of each model using the Pearson correlation coefficient, a standard metric for determining the linear relationship between the actual essay scores and the predicted ones. To assess the impact of the perplexity score on the models' performance, the raw Pearson correlation coefficients are compared to those obtained after the metric is introduced (Table 1).

Model	Without Perplexity	With Perplexity
Simple Linear Regression	0.507	0.598
Decision Tree Regression	0.875	0.905
Ridge Regression	0.514	0.536
LASSO Regression	0.514	0.588
Elastic Net Regression	0.536	0.530
Support Vector Regression	0.488	0.250

Table 1: Pearson Correlation with and without Perplexity

Table 2 presents the Pearson correlation coefficients for each model, including the LLMs, reflecting their accuracy in predicting essay scores.

Model	Pearson Correlation
Simple Linear Regression	0.598
Decision Tree Regression	0.905
Ridge Regression	0.536
LASSO Regression	0.588
Elastic Net Regression	0.530
Support Vector Regression	0.250
ChatGPT	0.605
Gemini	0.266

Table 2: Pearson Correlation of Different Models.

4.3. Result Analysis

The results of the experiments demonstrate several important facts. The best model resembles Decision Tree Regression, with a Pearson correlation coefficient of 0.905, indicating a robust relationship between the predicted and actual essay scores. This model most likely outperformed its counterparts due to its ability to effectively capture nonlinear relationships in the data. Among all the LLMs, ChatGPT demonstrated superior performance compared to other models, including Support Vector Regression and traditional regression methods such as Ridge and LASSO Regression, as indicated by a Pearson correlation coefficient of 0.605. These results indicate that ChatGPT and similar advanced language models achieve strong performance in essay scoring. Typically, adding the perplexity score to the list of metrics improved the Pearson correlation coefficients of all of these models, except for a few. In addition, Decision Tree Regression and Simple Linear Regression models also received an important contribution from perplexity, which means that their usage is indeed beneficial to predict the essay score. Overall, we found that the additional perplexity metric does little to improve the performance of our Support Vector Regression model, which is not a good indication that this metric is a good fit for this model. This means that when constructing automatic essay scoring systems, it is crucial to choose the type of models and evaluation metrics carefully, especially in cases where the nuance of the language matters. The integration of advanced

language processing techniques and expanded evaluation metrics led to virtually no decline in the accuracy or reliability of systems built on traditional regression models.

5. CONCLUSION & FUTURE WORK

In this work, we compare regression models and LLMs in this task, where learning to predict essay scores of essays in the Bengali language is considered. Also among the traditional models the results indicated that Decision Tree Regression better showed a high Pearson correlation with actual scores than the others. Furthermore, the grown up ChatGPT also outperformed many classic models and among LLMs. Overall, adding the perplexity score as an additional necessary metric may improve model performance for Decision Tree Regression and Simple Linear Regression. These results show that it is also possible to join the traditional machine learning techniques with the best of the class of language models to create robust and the accurate automatic essay scoring systems.

The use of regression models and LLMs for essay scoring in Bengali as presented by our work could also be of interest to the larger field. In the context of automated assessment, this language has been relatively little noticed. Finally, for future research, it is suggested that one study further validation of the models by fine-tuning their LLMs on much more extensive and diverse datasets.

REFERENCES

- [1] Islam, M. (2011). Automatic scoring of Bangla language essay using generalized latent semantic analysis.
- [2] Chakraborty, S., Nayeem, M. T., & Ahmad, W. U. (2021, May). Simple or complex? Learning to Predict Readability of Bengali Texts. In Proceedings of the AAAI Conference on Artificial Intelligence (Vol. 35, No. 14, pp. 12621-12629).
- [3] Islam, M. M., & Hoque, A. L. (2013, May). Automated Bangla essay scoring system: ABESS. In 2013 International Conference on Informatics, Electronics, and Vision (ICIEV) (pp. 1-5). IEEE.
- [4] Dikli, S. (2006). "An Overview of Automated Scoring of Essays." *The Journal of Technology, Learning, and Assessment*, 5(1), 4-28.
- [5] Dong, F., & Zhang, Y. (2016). "Automatic Features for Essay Scoring." *Proceedings of the 54th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, 1071-1080.
- [6] Landauer, T. K., Foltz, P. W., & Laham, D. (1998). "An Introduction to Latent Semantic Analysis." *Discourse Processes*, 25(2-3), 259-284.
- [7] Karim, M. R., Chakravarthi, B. R., McCrae, J. P., & Cochez, M. (2020, October). Classification benchmarks for the under-resourced Bengali language based on a multichannel convolutional LSTM network. In 2020, IEEE 7th International Conference on Data Science and Advanced Analytics (DSAA) (pp. 390-399). IEEE.
- [8] Page, E. B. (1966). "The Imminence of... Grading Essays by Computer." *The Phi Delta Kappan*, 47(5), 238-243.
- [9] Shermis, M. D., & Burstein, J. (2013). *Handbook of Automated Essay Evaluation: Current Applications and New Directions*. Routledge.
- [10] Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., ... & Polosukhin, I. (2017). "Attention is All You Need." *Advances in Neural Information Processing Systems*, 30, 5998-6008.
- [11] Williamson, D. M. (2009). "A Framework for Implementing Automated Scoring." *Journal of Educational Measurement*, 46(2), 135-150.
- [12] Rodgers, J. L., & Nicewander, W. A. (1988). Thirteen ways to look at the correlation coefficient. *The American Statistician*, 42(1), 59-66.
- [13] Cohen, L., Manion, L., & Morrison, K. (2002). *Research methods in education*. routledge.
- [14] Heafield, K. (2011, July). KenLM: Faster and more minor language model queries. In Proceedings of the sixth workshop on statistical machine translation (pp. 187-197).
- [15] Pedregosa, F., Varoquaux, G., Gramfort, A., Michel, V., Thirion, B., Grisel, O., ... & Duchesnay, E. (2011). Scikit-learn: Machine learning in Python. *Journal of Machine Learning Research*, 12, 2825-2830.