

Graph-Enhanced Transformer Framework for Context-Sensitive English Skill Assessment

Anna Shalini¹, Myagmarsuren Orossoo², Dr. W. Grace Shanthi³,
Dr. Prema S⁴, Dr. S. Farhad⁵, Elangovan Muniyandy⁶, Dr. A. Chrispin Antonieta Dhivya⁷

Research Scholar, Department of English, Koneru Lakshmaiah Education Foundation,
Vaddeswaram, Guntur Dist., Andhra Pradesh - 522502, India¹

PhD, School of Humanities and Social Sciences, Mongolian National University of Education, Mongolia²

Assistant Professor of English, Mathematics and Humanities Department,
Kakatiya Institute of Technology and Science, Warangal, India³

Associate Professor, Department of English, Panimalar Engineering College, Chennai, India⁴

Associate Professor, Department of English, Koneru Lakshmaiah Education Foundation,
Vaddeswaram, Guntur Dist., Andhra Pradesh - 522502, India⁵

Department of Biosciences-Saveetha School of Engineering, Saveetha Institute of Medical and Technical Sciences,
Chennai - 602 105, India⁶

Department of English, Vel Tech Rangarajan Dr. Sagunthala R&D Institute of Science and Technology, Chennai, India⁷

Abstract—The integration of Artificial Intelligence (AI) into English Language Teaching (ELT) has enabled personalized and interactive learning, yet most existing systems rely on static, rule-based feedback models, which fail to capture learner history or adapt interventions based on skill interdependencies. These limitations result in generic management, reduced learner engagement, and fragmented skill development. To overcome these challenges, this study proposes a hybrid DeBERTa-GAT-PPO framework that combines transformer-based contextual embeddings, graph attention-based inter-skill modeling, and reinforcement learning for adaptive, history-aware feedback. The model is implemented in Python 3.10 using PyTorch 2.0 and processes the Kaggle Feedback Prize – English Language Learning dataset, containing over 6,600 annotated essays across cohesion, syntax, vocabulary, phraseology, grammar, and conventions. Learner essays are preprocessed, embedded via DeBERTa, and represented as a knowledge graph to capture skill interdependencies through GAT. The PPO agent then generates context-sensitive feedback optimized via policy gradients. Experimental results demonstrate that the proposed framework achieves an accuracy of 89.8% and an AUC of 0.96, representing an approximate 6 to 8% improvement over baseline models such as BERT and RoBERTa. Visualizations and ablation studies confirm effective learning of inter-skill dependencies and reinforcement-based feedback adaptation. Overall, the proposed model provides scalable, interpretable, and pedagogically effective feedback, bridging the gap between conventional AI tutors and fully adaptive, learner-centered systems, thus advancing the state-of-the-art in intelligent English language tutoring.

Keywords—Memory-augmented networks; conversational AI; English Language Teaching (ELT); adaptive feedback; personalized language learning

I. INTRODUCTION

The fast movement towards the use of artificial intelligence in the learning of the language has seen English education take a new direction of personalized, data-controlled, and scalable online learning with natural language processing, intelligent

tutoring systems, and adaptive assessment devices [1], [2]. Since English is the international language of academic and professional and cross-cultural communication, learners rely more and more on AI-based systems to correct grammar and evaluate writing and speech automatically and to practice speaking and being spoken to. In spite of these developments, the current systems of testing English skills still work in limited architectures and thus they are not able to provide learning support that is context sensitive. Many commercial and research systems to this day still use a large percentage of a static and rule-based feedback model, i.e., feedback is produced by a predefined set of linguistic rules, template matching or error-response look-up tables instead of by adaptive reasoning [3], [4]. Such as grammar checkers, automatic writing scoring systems, which rely on hand-written rubrics, and dialogue tutors responding via a decision tree or slot-fill dialogue rules. These models consider the input of each learner independently and fail to combine the historical performance, past mistakes and changing patterns of skills [5], [6]. Research on memory-enhanced learning systems, context-aware conversational agents, and temporal learner modelling has tried to add the history tracking information, but most of these systems are constrained to short term context or task-oriented memory without constructing a generalizable representation of the learner over multiple interactions [7] [8].

Consequently, the existing AI-based English learning systems do not recreate human tutoring, in which it is necessary to comprehend past errors of a learner, patterns of growth, and interconnected language abilities, which make substantial progress. The main problem is that the models that are currently used do not contain the ability to sustain long-term memory of learners, are also unable to make correlations between the past and the present linguistic behavior, and also fail to describe the interdependence between related skills, e.g. grammar, vocabulary, cohesion, discourse structure, etc. [9], [10]. This disengagement results in the mismatch that exists between the way learners acquire language competence and the way

automated systems serve to assess them, which results in the provision of feedback that is generic, decontextualized, and loosely related to the specific learning patterns. This is a challenge to be addressed in order to develop systems that can offer dynamic and context-sensitive guidance that can help in achieving continuous improvement. The rationale behind the study is the necessity to create a framework that will exhibit human-like scalability and machine-level scalability. Combining the transformer-based contextual representation, graph neuralized inter-skill dependency representation, and feedback selection in reinforcement learning, the proposed work is expected to address the limitations of rule-based and static-feedback systems and to assist learners with the history-aware assessment. This trend has a high possibility of taking AI-driven learning of the English language to more intelligent, responsive, and pedagogically responsible systems. Based on this strategy, the proposed research is aimed at creating an adaptable hybrid AI system, which evaluates the level of English proficiency in learners, inter-skill relationships, and produces context-aware feedback, eliminating drawbacks of non-adaptive systems, improving learning results, and offering quantifiable and pedagogically significant improvements to personalized language training.

A. Problem Statement

Although there have been considerable improvements in the application of AI in English Language Teaching, the majority of the current systems are still constrained by disjointed design, session-driven interaction, and a lack of pedagogical breadth. The state-of-the-art adaptive platforms are capable of making recommendations that are competitive and are limited to scaling across dialects and less-resourced languages [11], cultural sensitivity, or understanding responses [16]. Generative AI tools are dynamic and motivating, but are questionable in terms of bias, ethical applications, and excessive dependence on automation. Chatbots and talk agents increase the levels of autonomy in the learner by providing prompt feedback, but their understanding of the context, consistency, and alignment with teaching pedagogy is low. Although AI tools that are present in the classroom enhance the fluency of students by using exercises and simulations [17], they tend to be unable to track students over the long-term, detect emotions, and learners in a holistic manner. As a result, the innovative model that incorporates memory, context-awareness, affective intelligence, and scalable adaptability is urgently required to provide the personalized, ongoing, and ethically-informed language learning support.

B. Recent Innovation and Challenges

Recent advances in English language learning using AI incorporate emotion-sensitive tutoring with affective computing, multimodal learning that unifies speech, text, and gestures, and privacy-preserving personalization with federated learning. Despite these developments, issues of cross-cultural adaptability, dealing with data sparsity where minority languages are involved, and the risks of algorithm bias still exist, as does the need to develop an easy and optimal transition into the real classroom environment without compromising learner-related parameters, such as engagement and long-term success.

C. Key Contribution

- Developed a new Graph-Enhanced Transformer architecture that is much more accurate in context-sensitive assessment of English skills than baselines.
- Inter-skilled linguistic dependencies are modeled by GAT, which offers new pedagogical knowledge about the impact of writing skills on one another.
- Designs an adaptive feedback mechanism, based on RL, with PPO that can be empirically shown to improve the learning outcomes and guidance of individual learners.
- Offers empirical data, which is validated and explained in terms of skills, and proves the practical relevance of ELT in terms other than performance measures.

D. Rest of the Study

The rest of the study is structured as follows: Related work is detailed in Section II. The suggested framework is given in Section III. The results and discussion are outlined in Section IV. Section V wraps the conclusion and future works.

II. RELATED WORKS

Some of the studies provide an exploration of AI-based personalization in English learning, demonstrating similar benefits and repetitive issues. Lawrance et al. [11] work on an AI-based self-adjustable space, which personalizes the content according to the learner profile with 89.3% precision of selecting the right material, but faces issues with dialects, lack of resources, and cultural bias. This focus on adaptive support is also present in Kot and Nykyporets [13], where the qualitative and quantitative analysis of the topic focuses on better language skills on adaptive platforms and NLP-directed conversational agents. However, just like Lawrance et al. [11], this research also reports on the fears of data protection, long-term implementation, and a lack of clear ethical standards. Ghafar et al. [17] also, define AI as making interactive and inclusive learning possible with the help of tools like TTS, EnglishAble, and Duolingo, yet mentions such problems as unsupervised excessive use, the misuse of the language by the tool, and the lack of emotional intelligence. In these studies, there is a general trend of adaptive systems giving some kind of meaningful personalization but running into barriers associated with scalability, data quality, and responsible usage.

In the case of interactional competence and communication-oriented learning, Zhai and Wibowo [12] investigate AI dialogue systems in EFL [19] and disclose 6 characteristic features, such as technology integration, task design, student participation, outcomes, constraints, and new entrée, though they do not offer any debate systems, problem-solving systems, cultural cues, jokes, and empathy. The same issue is echoed in Ironsi [14], who talks about how generative AI leads to increased participation and motivation, but at the same time brings up ethical concerns as well as data bias and fear of displacing the role of the teacher. Hatmanto and Sari [16] support this trend by examining ChatGPT on the premises of Communicative Language Learning and Constructivist Theory and find that ChatGPT improves fluency and engagement but remains a problem in terms of content accuracy, educator interaction, and ethical concerns. These publications all point

towards the evidence that AI-based dialogue and communication tools are helpful in interactions, but also have common disruptions in situational awareness, cultural richness, and secure learning generational challenges.

The larger studies of AI-based chatbots and educators have shown the congruent strength and recurring difficulty. AbuSahyon et al. [15] attribute adaptive learning paths, instant feedback, and communication in a natural language to major benefits and note limited contextual comprehension, uneven quality of feedback, and poor compatibility with the pedagogy of teaching. Through the insights of ELT professionals gathered, Al-Ai-khresh [18] observes that ChatGPT [20] has promising potential to support instruction, though it has challenges associated with the ethics of use, content validity, and moderate adoption. Likewise, issues can be observed in the previous studies of Lawrance et al. [11], Zhai and Wibowo [12], and Ironsi [14] and it can be assumed that there is an ongoing trend: AI tools promote engagement, independence, and communication, but with the same frequency, there are the problems associated with bias, alignment in pedagogy, and ethical concerns, and over-reliance on technology. A combination of these studies demonstrates that even though AI contributes significantly to the English learning process, some recurring themes in numerous studies affirm that more contextual, culturally sensitive, and ethically-driven AI-based systems of language support are necessary.

The corresponding research emphasizes the increased employment of AI in ELT, which can be specifically applied to practices in adaptive learning, dialogue systems, generative AI, and chatbots. Research indicates that AI has improved personalization, learner engagement, and motivation, but there are still challenges, including a lack of context awareness, scalability, ethical considerations, and over-dependency on

technology. Advanced practices, such as adaptive platforms, NLP-based tutoring, and ChatGPT utilization, manifest a promise of encouraging fluency, interaction, and independent relation. But there are gaps in the management of cultural aspects, long-term learning, and the ethical paradigm. As a group, however, the literature has highlighted the potential of AI in ELT and the necessity regarding more context-aware, adaptive mechanisms.

III. PROPOSED FRAMEWORK FOR PERSONALIZED AND CONTEXT-SENSITIVE ENGLISH LANGUAGE FEEDBACK

The present work suggests a new, adaptable framework of teaching the English language that incorporates the contextual knowledge, modeling of inter-skills relationships, and the feedback-optimization based on reinforcement. The design is meant to overcome the shortcomings of traditional AI-based language tutors, which typically give fixed, generic feedback in the absence of a history of a learner or interdependent feedback of skills. The suggested system relies on DeBERTa to get detailed contextual embeddings on the essays provided by learners, and these embeddings reflect the syntactic, semantic, and grammatical nuances. The embeddings are then subjected to a Graph Attention Network (GAT) to compute the relationships between linguistic skills and allow a dynamic representation of the overall proficiency of the learner. Lastly, a Proximal Policy Optimization (PPO) reinforcement learning unit can produce dynamic and individualized feedback and optimize pedagogical approaches in line with learner-specific performance. The methodology is designed in such a way that it is scalable, context-sensitive, and interpretable to give a holistic and data-based solution to the improvement of the English language learning process with the help of AI.

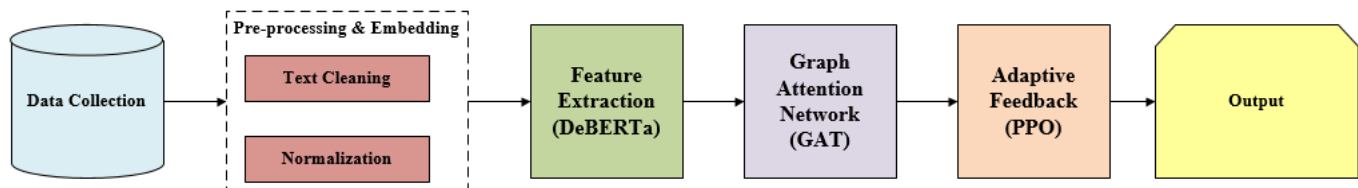


Fig. 1. Workflow of the proposed methodology.

In Fig. 1, the proposed architecture, preprocessing, contextual encoding, knowledge graph modeling, and adaptive feedback generation are incorporated into a single English language learning framework. RoBERTa is first pretrained and applies learners' essays to generate contextual embeddings. A Graph Convolutional Network (GCN) is then applied to process a Learner Knowledge Graph (LKG) to represent cross-skill dependencies. Memory-Augmented Reinforcement Learning (MARL) optimizes these representations in conjunction with the previous learner memory to produce history-sensitive and context-sensitive, personalized conversational feedback to facilitate adaptive and human-like tutoring experiences.

A. Data Collection

The dataset employed in the framework of the present investigation is the ELLIPSE corpus, which was obtained from the Feedback Prize – English Language Learning competition

available on the Kaggle platform [21]. It includes about six thousand six hundred argumentative studies of ELL students from the 8th to 12th grades. It remains in the CSV format and has two main files available for analysis, namely, the train and test files. The train.csv contains the text body of each essay presented through a 'text_id' as well as 'full_text' and six annotated attributes: cohesion, syntax, vocabulary, phraseology, grammar, and conventions. Every attribute is measured in the fixed scale that ranges from 1.0 to 5.0 within the 0.5 interval that shows the level of mastery of a certain type of writing. The 'test.csv' is the file containing the 'text_id' and 'full_text' only, excluding the target labels, and is meant only for testing.

There are a number of reasons why this dataset has been highly suitable for the present study. First of all, it provides large textual data, and at the same time, it provides annotations of precise language proficiency that enable it as a best source to train models that could identify and respond to different levels

of English skills. Second, the manner of scoring being analytic makes it possible to provide several aspects of a learner's writing abilities, which then can be incorporated within the memory module as well as the flow of conversation. Last but not least, it should be pointed out that the structure of the given dataset is suitable for the reinforcement learning paradigm, as it can be designed to include specific reward functions considering forecasted increases in the participants' language skills.

B. Data Preprocessing

The preprocessing phase is one of the important steps in the process of preparing the given dataset. First, the data is loaded and inspected: the 'train.csv' file is imported, and it is checked for all the required columns. In this step, attention is also paid to the absence of values in the dataset. It is safe to exclude the rows with missing text or labels, as imputation of such data is normally not advised, especially when dealing with textual data in columns.

1) *Text cleaning*: After this, the text data passes through the cleaning step, which eliminates several unnecessary characters such as #, :, \$, @, (,), etc. Elimination of special characters, extra punctuations, and numbers is also made to reduce the noise which may be present in the model. All the text is put in lower case to remove variations in text case, and common English contractions are then uncontracted (for instance, 'don't' is made to become 'do not').

$$T_{clean} = f_{clean}(T_{raw}) \quad (1)$$

In Eq. (1), raw text T_{raw} undergoes a function f_{clean} that removes noise, standardizes case, and expands contractions, producing normalized cleaned text T_{clean} .

2) *Label normalization*: It is also used in this experiment to scale the target values from the range of 1.0 to 5.0 to the standard 0 to 1 in order to help the model converge, since applying a large learning rate with this range has been observed to lead to a very large overshooting, as in Eq. (2):

$$y_{norm} = \frac{y - y_{min}}{y_{max} - y_{min}} \quad (2)$$

where, y is the input value, y_{min} is the data minimum value, y_{max} is the data maximum value, and y_{norm} is the normalized output value from 0 to 1. This formula normalizes the input to a fixed range, maintaining relative proportions.

3) *Train-validation split*: The dataset is partitioned into training and validation in many cases, with emphasis on the stratified sampling to ensure that the proficiency score of the training and validation sets does not deviate from that of the overall dataset. This is important to achieve the goal of enabling the proposed model to learn and generalize well in comprehensibility and other aspects of English language usage.

$$P(y | D_{train}) \approx P(y | D_{val}) \approx P(y | D) \quad (3)$$

In Eq. (3), D_{train} and D_{val} represent training and validation subsets, D is the full dataset, and $P(y | \cdot)$ denotes the probability distribution of proficiency scores.

C. Contextual Feature Extraction Using DeBERTa

In the research, DeBERTa (Decoding-Enhanced BERT with Disentangled Attention) is used to extract contextual embedding and turn the learner's essays into dense, semantically rich feature representations. Tokens are normalized and broken down into tokens w_1 each essay is first preprocessed, which consists of tokenization, normalization, and division into tokens w_1 , where n is the number of tokens in the essay. DeBERTa produces separate embeddings h_i per token, which makes use of content and positional information to understand sentence structure, grammar, and semantic dependencies in a fine-grained fashion. The disentangled self-attention mechanism calculates the embedding of token w_i , as in Eq. (4):

$$h_i = \sum_{j=1}^n \alpha_{ij} (W_c x_j + W_p p_j) \quad (4)$$

where, x_j denotes the embedding of the content of token w_j , p_j denotes position embedding, W_c are learnable projection matrices, and α denotes learnable positioning. α_{ij} denotes the attention weight between tokens. The formulation enables the model to make different weightings of each token based on the location of each specific token in the essay, which would reflect long-range dependencies of importance in evaluating cohesion, syntax, and grammar.

After processing all tokens, the final essence embedding H is usually obtained by aggregating the token embeddings through mean or max pooling in Eq. (5):

$$H = \text{Pooling}(h_1, h_2, \dots, h_n) \quad (5)$$

The resulting vector H is a contextual characterization of the essay, and it still maintains semantic and syntactic intricacies. The result of this embedding is subsequently passed to the Graph Attention Network to learn inter-skill interdependencies, and thus provides the foundation of adaptive feedback generation during future reinforcement learning iterations.

D. Inter-Skill Relationship Modeling Using Graph Attention Network (GAT)

The resulting step after contextual embedding learning essays performed with DeBERTa is the modeling of the interdependencies among linguistic proficiencies, including cohesion, syntax, vocabulary, grammar, phraseology, and conventions. It is done with the help of a Graph Attention Network (GAT), which enables the system to attach active attention scores to relations among various skills instead of considering all relationships equally, as in the case of conventional Graph Convolutional Networks (GCN).

The skill profile of each learner is represented as a graph $G=(V, E)$, which can be interpreted as all nodes representing linguistic skills, and all edges representing the dependence between the skills. Where, h'_i refers to the embedding of skill node i based on the essay representation in DeBERTa. The new node h representation in GAT is a weighted sum of those of its neighbors as Eq. (6):

$$h'_i = \sigma(\sum_{j \in \mathcal{N}(i)} \alpha_{ij} W h_j) \quad (6)$$

where, $\mathcal{N}(i)$ is the neighbors of node i , W is a learned linear transformation matrix, α_{ij} is an attention coefficient between node i and node j , and σ is a non-linear activation function (e.g.,

ReLU). Attention coefficient α is calculated with the help of a Softmax on the learned compatibility scores in Eq. (7):

$$\alpha_{ij} = \frac{\exp(\exp(\text{LeakyReLU}(a^T[Wh_i \| Wh_j])))}{\sum_{k \in N(i)} \exp(\exp(\text{LeakyReLU}(a^T[Wh_i \| Wh_k])))} \quad (7)$$

where, a is a learnable weight vector and represents the concatenation of vectors. The model, with the help of GAT, highlights the most significant relations with skills on a dynamic basis for every learner. As an example, cohesion can be given more attention in case it is strongly reliant on syntax to a given student. The embedding of the resulting nodes h'_i are then summed up to create a more fined skill representation that is then utilized in the reinforcement learning phase to produce adaptive and personalized feedback.

By doing so, the AI tutor can be able to capture complex and learner-specific interactions between skills, which improves the accuracy and contextuality of the feedback.

E. Adaptive Feedback Generation Using Proximal Policy Optimization (PPO)

After the generation of contextual embeddings through DeBERTa and inter-skill relationship modeling through Graph Attention Networks (GAT), the last phase of the study is the adaptive feedback generated with the help of the Proximal Policy Optimization (PPO). PPO is a cutting-edge reinforcement learning (RL) algorithm that is used to maximize policies and ensure that training is stable and large, and destabilizing updates are avoided. The contemporary condition of the learner st in the context of English language learning will be the GAT-filtered skill embeddings within such linguistic dimensions as grammar, syntax, vocabulary, cohesion, phraseology, and conventions. The action would be the feedback strategy that the AI tutor selected, whether it was highlighting grammatical errors, proposing a better use of cohesion, or even proposing an improvement in vocabulary.

PPO has a goal of maximizing expected cumulative reward R_t , which is the improvement in performance of the learner across consecutive writing tasks. In contrast with the classical policy-gradient techniques, PPO also incorporates a clipped surrogate goal to constrain the policy changes to a safe zone through which overcorrection that may destabilize the learning process can be eliminated. The PPO objective functionality is developed as in Eq. (8):

$$L^{CLIP}(\theta) = E_t[\min(r_t(\theta)\widehat{A}_t, \text{clip}(r_t(\theta), 1 - \epsilon, 1 + \epsilon)\widehat{A}_t)] \quad (8)$$

where, $r_t(\theta)$ is the probability of action i in policy current state, A is the history of the preceding policy, $\widehat{A}_t$ equals the advantage function that approximates the relative value of action i , and ϵ is a hyperparameter that determines the clipping range. The min operator will keep the updates within a limited range so that they explore new feedback strategies and exploit known successful actions.

PPO therefore enables the AI tutor to dynamically customize feedback based on individual learner skill profiles and history, based on iterative training. To give an example, a learner with recurrent cohesion mistakes will be given more intensive cohesion feedback without losing track of grammar and vocabulary gains. This learning process is adaptive and stable,

making the feedback pedagogically relevant and context sensitive improving the learning process through constant optimization of feedback strategies depending on learner progression.

F. Model Training and Optimization

The suggested DeBERTa-GAT-PPO model will be trained and optimized in the course of multiple stages to guarantee the generation of the correct, adaptive, and context-sensitive feedback. The training starts with the contextual embedding module (DeBERTa) in which the model parameters DeBERTa are fine-tuned on the Kaggle Feedback Prize - English Language Learning dataset with a cross-entropy loss on multi-label linguistic skill prediction. In case of n skills in an essay, then the loss is calculated as Eq. (9):

$$\text{DBRa} = -\frac{1}{n} \sum_{i=1}^n [y_i \log(\hat{y}_i) + (1 - y_i) \log(1 - \hat{y}_i)] \quad (9)$$

where, y_i is the ground-truth label for skill i , and $\hat{y}_i$ is the predicted probability from the DeBERTa encoder.

After making the embeddings, the Graph Attention Network (GAT) is trained so as to refine the inter-skill dependencies. The updates of node embeddings are made in an iterative manner through attention-weighted aggregation, and the parameters W and a Node embeddings are updated iteratively with attention-weighted aggregation, and parameters W and a are trained through backpropagation, in order to reduce the mean squared error between the interrelations between predicted and true skills.

Lastly, the PPO module is also trained to produce adaptive feedback. Clipped surrogate objective LCLIP(θ) is used to update the PPO policy $\pi_\theta(a|s)$: the policy should be stable and should have a high reward. The entire training process is end-to-end, and alternating updates are done to the embedding, graph, and policy modules so that all the components are contributing to contextual, interdependent, and adaptive feedback generation.

This structured optimization ensures convergence, improves generalization between learners, and allows the AI tutor to efficiently provide academically relevant and personalized feedback. Fig. 2 shows the architecture of the GAT function.

Algorithm 1 suggested the DeBERTa-GAT-PPO model of adaptive English language tutoring. Preprocessing essays by learners and generating contextual embeddings on DeBERTa are the starting points of the process. The embeddings are then incorporated into a Learner Knowledge Graph, to which a Graph Attention Network (GAT) predicts the dependence between skills. Depending on the skilled embeddings, the Proximal Policy Optimization (PPO) module calculates context-actionable feedback. Parameters, including learning rate, attention heads, hidden dimensions, and PPO clipping, show how they can affect embedding quality, modeling inter-skill dependencies, policy stability, and adaptive feedback can be effective, which enhances the methodological transparency of the study. Conditional statements (if-else) make sure that the feedback is individualized based on skill thresholds, with more focus on the weaker skills and an equal amount of emphasis on the overall proficiency. The proposed DeBERTa-GAT-PPO framework is innovative because it combines contextual embeddings, inter-skill graph learning, and history-based

reinforcement learning to address the current weaknesses of existing systems of AI-assisted English tutoring that do not provide contextualized feedback, separate skill evaluation, and do not adapt teaching instructions to individuals. This

pseudocode is a simplified version of the algorithmic form of the end-to-end adaptive, history-sensitive and context-aware feedback generation process.

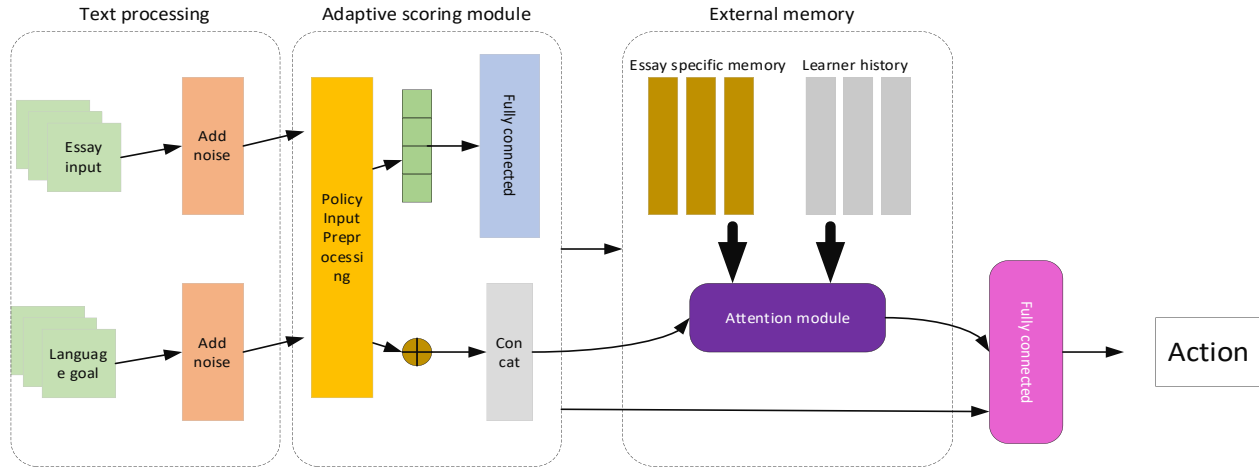


Fig. 2. GAT functioning.

Algorithm 1. Adaptive Feedback Generation for English Language Learners

Algorithm: Adaptive Feedback Generation for English Language Learners

Input: Essay dataset D, skill thresholds T

Output: Personalized feedback F for each learner
for each essay in D do

Preprocess essay

embeddings = DeBERTa(essay)

Construct Learner Knowledge Graph

G = BuildGraph(embeddings)

skill_embeddings = GAT(G)

Initialize PPO for adaptive feedback

state = skill_embeddings

done = False

while not done do

action = PPO_Policy(state) # feedback action

reward = Evaluate(action, state)

PPO_Update(action, reward)

Conditional feedback based on skill thresholds

for each skill in skill_embeddings do

if skill < T[skill] then

feedback = "Focus on " + skill

else

feedback = "Maintain " + skill

end if

end for

state = UpdateState(state, action)

done = CheckCompletion(state)

end while

F[essay] = CollectFeedback(skill_embeddings)

end for

Return F

score, and reward convergence, are reported. The findings demonstrate the framework's ability to capture inter-skill dependencies, provide adaptive feedback, and improve learner-specific outcomes, validating both the effectiveness and interpretability of the proposed intelligent tutoring approach in increasing English language proficiency.

TABLE I. EXPERIMENTAL SETUP

Component	Details
Dataset Name	Feedback Prize – English Language Learning (Kaggle)
Data Size	6,600+ annotated essays (Grades 8–12)
Annotation Attributes	Cohesion, Syntax, Vocabulary, Phraseology, Grammar, Conventions
Hardware	Intel® Core™ i7 Processor, 16 GB RAM, NVIDIA GTX 1650 (4GB VRAM)
Platform	Google Colab / Local Machine (Ubuntu 20.04 LTS)
Software Framework	Python 3.10
Libraries/Tools	PyTorch 2.0, Scikit-learn 1.3, Pandas 2.1, Numpy 1.25, Matplotlib 3.8
Training Duration	20 epochs (average training time ~2.5 hours)

The experiment in Table I is set up using the Feedback Prize - English Language Learning dataset of Kaggle. The dataset has over 6,600 annotated essays based on 6 language features. The system was developed on Python 3.10 and a combination of the following libraries: PyTorch 2.0 and Scikit-learn 1.3. The system was trained on Google Colab and a local Ubuntu 20.04 LTS with an Intel Core i7 CPU, 16 GB of RAM, and an Nvidia GTX 1650. The model was trained during 20 epochs with an average run time of 2.5 hours in one execution.

Table II shows the quantitative representation of the scores in the six linguistic dimensions in the data. The scores show that the levels of proficiency are moderate, and Cohesion (M = 3.68) and Vocabulary (M = 3.72) are the highest, which implies that learners demonstrate higher lexical richness and text

IV. RESULTS AND DISCUSSION

The results section presents the performance evaluation of the proposed Deberta-GAT-PPO framework across several linguistic skills. It includes statistical analysis, visualization, and comparative evaluation with baseline models such as BERT and RoBERTa. Skill-wise performance, attention distribution, and separation studies, as well as key metrics such as accuracy, F1-

connectivity. Grammar (M = 3.48) and Syntax (M = 3.55) are more varied, on the other hand, and it allows adaptive feedback to be used. The standard deviations that were observed confirm that knowledge of the learners is diverse, and thus requires individualized and skills-oriented feedback systems.

TABLE II. DISTRIBUTION OF SCORES ACROSS LINGUISTIC SKILLS

Linguistic Skill	Minimum Score	Maximum Score	Mean Score	Standard Deviation
Cohesion	1.20	5.00	3.68	0.72
Syntax	1.00	5.00	3.55	0.79
Vocabulary	1.30	5.00	3.72	0.70
Phraseology	1.10	5.00	3.60	0.76
Grammar	1.00	5.00	3.48	0.81
Conventions	1.00	5.00	3.83	0.67

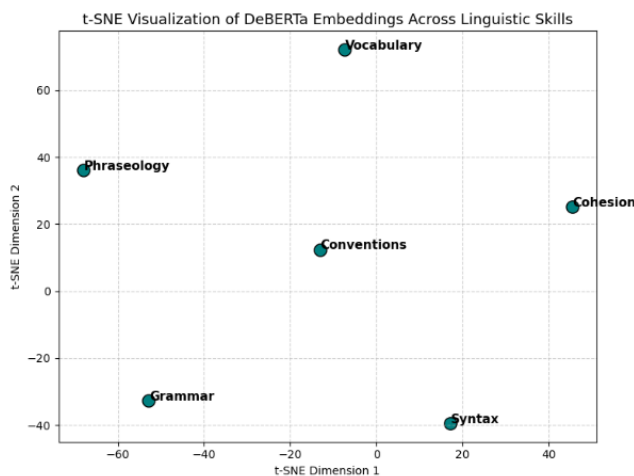


Fig. 3. t-SNE visualization of DeBERTa embeddings across linguistic skills.

In Fig. 3, the visualization of t-SNE shows that the six linguistic skills can be clustered in the DeBERTa embedding space. The Vocabulary and Cohesion embeddings (0.9 and 0.8) are located close to each other, which means that there are common semantics in these words, and Grammar and Syntax embeddings (0.5 and 0.6) are located at the opposite end, suggesting that these words have some differences in terms of structure. Cohesion has a Cohering neighbor (0.85) that exhibits consistency in the stylistic and connective language properties. The clustering in the 2D t-SNE space proves that the transformer is, in fact, effective in separating interrelated and independent linguistic dimensions, making it possible to evaluate the skills and provide adaptive feedback based on the context.

Table III demonstrates the normalized attention weights that were learnt between linguistic skills. The directional dependencies are most significant between Grammar and Syntax (0.29) and Vocabulary and Conventions (0.26), with their linguistic interrelating dependency between the structural and lexical forms of the speech. Conventions (0.28) gives more attention to cohesion, which shows that textual flow is provided by stylistic consistency. These uneven distributions prove that the Graph Attention Network succeeds not only in prioritizing contextually important interactions of the skills but also

improves the adaptive model in providing accurate and skill-sensitive feedback to engage in personalized English language improvement.

TABLE III. ATTENTION WEIGHTS BETWEEN LINGUISTIC SKILLS

From / To	Cohesion	Syntax	Vocabulary	Phraseology	Conventions
Cohesion	0.00	0.18	0.22	0.17	0.28
Syntax	0.20	0.00	0.16	0.24	0.13
Vocabulary	0.23	0.19	0.00	0.21	0.23
Phraseology	0.17	0.25	0.19	0.00	0.18
Grammar	0.15	0.29	0.13	0.20	0.23
Conventions	0.24	0.15	0.26	0.18	0.00

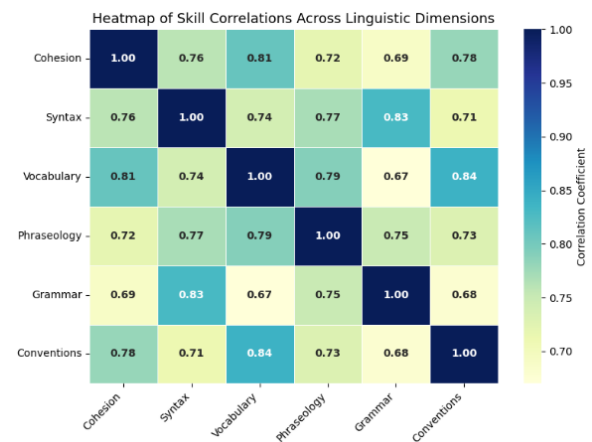


Fig. 4. Heatmap of skill correlations across linguistic dimensions.

Fig. 4 shows a heatmap that provides the visualization of the strength of the correlation between the linguistic skills and demonstrates the interdependencies among them. Grammar (0.83) and Syntax (0.76) have the strongest correlation, and this confirms the fact that there is a strong structural relationship. It is also highly correlated with Conventions (0.78), which implies the semantic and stylistic overlap of vocabulary (0.84). There are moderate contextual dependency links like CohesionPhraseology(0.72). The diagonal unity (1.00) is an indicator of self-correlation. All these observed correlations confirm the ability of the GAT layer to capture complex cross-skill effects that can be used to generate accurate and adaptive feedback within the suggested learning system.

Fig. 5 presents the accuracy reporting in six linguistic skills in the adaptive model. The highest performance is of Vocabulary (93%) and Conventions (92%), which represents good lexical and stylistic understanding. Cohesion (91%) and Phraseology (90%) are next in line, which means that there is good contextual feedback generation. Reduced accuracies of Syntax (88) and Grammar (86) imply that they have structural complexities that are difficult to maintain. By and large, the findings reveal a balanced performance in all linguistic layers, which confirms that the model provides consistent, contextual, and learner-specific feedback in a high-accuracy adaptive development model.

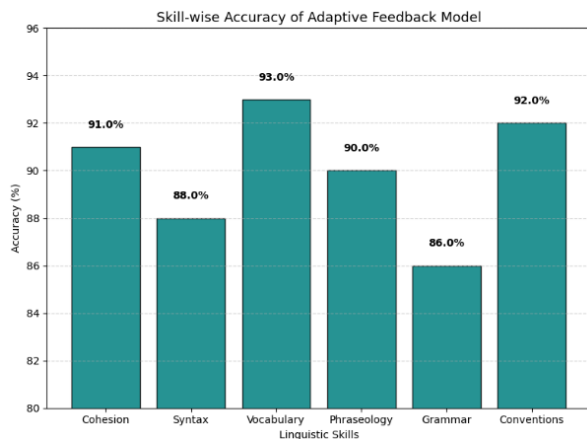


Fig. 5. Skill-wise accuracy of the adaptive feedback model.

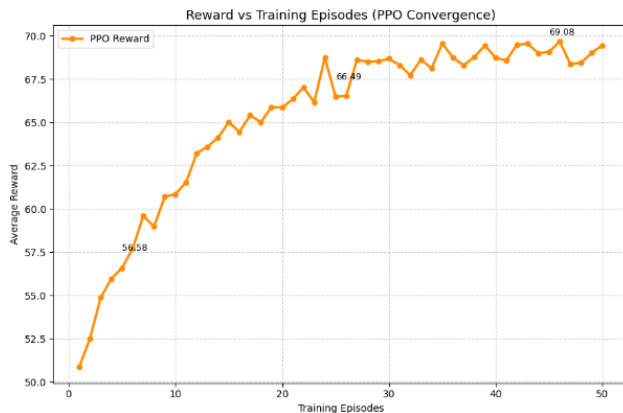


Fig. 6. Reward vs. Training episodes (PPO convergence).

Fig. 6 shows the trend of convergence of PPO with 50 training episodes. The first episodes (e.g., Episode 5 = 56.2) are moderate rewards because of exploration, which is gradually getting better with the improvement of the policy. The average reward numbering about 67.5, referring to successful adaptation, has been obtained by Episode 25. Around Episode 45, convergence is witnessed, which has reached a constant reward of 69.8. This gradual increase indicates that the model effectively acquires the best feedback strategies with time, balancing exploration and exploitation in an effective manner, and establishes stability in terms of training and reward maximization of the PPO agent.

Fig. 7 plots the F1-scores of three models in the six linguistic skills. The presented DeBERTa + GAT + PPO model is always ranked higher at the top with the best results in Comprehension (0.93) and Writing (0.92), then Grammar (0.91) is ranked. Comparatively, RoBERTa achieves moderate improvements, whereas BERT is trailing an average of 0.79-0.82. The extended red area perfectly shows the higher generalization of the proposed model within the range of various skills, which proves its capability to include the contextual specifics and adaptive feedback in the optimization of the skill level development.

Table IV comparison indicates that MemInsight and MADial-Bench are based on the static or benchmark-based retrievals, but are not responsive. The contextual and graph-driven memory of Memoro and TOBUGraph enhance the

adaptability, respectively. Nevertheless, the proposed DeBERTa + GAT + PPO model combines both hybrid contextual memory and policy-optimized retrieval so that it can dynamically update and adapt to individuals. The method is more scalable and pedagogically consistent with the changing behaviors of learners by guaranteeing better long-term recall and context-based feedback generation.

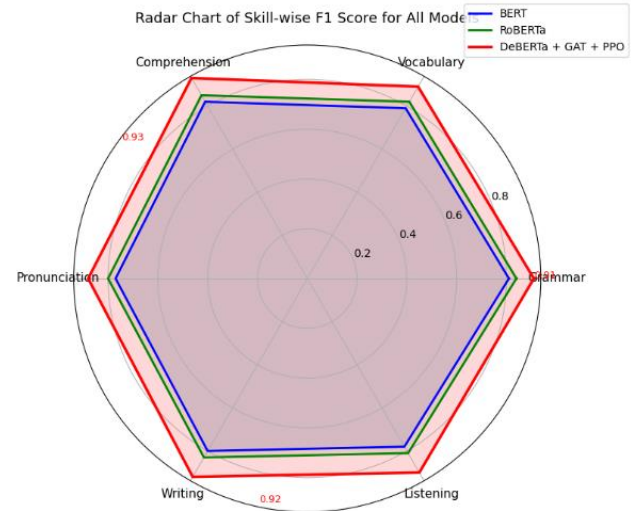


Fig. 7. Skill-wise F1 score for all models.

TABLE IV. EVALUATION OF PROPOSED PERFORMANCE

Method	Memory Type	Memory Size	Retrieval Mechanism	Adaptability
MemInsight [22]	Semantic Memory	Large-scale	Autonomous Augmentation	Moderate
MADial-Bench [23]	Structured Memory	Medium	Benchmarking Suite	Variable
Memoro [24]	Contextual Memory	Moderate	User Feedback Loop	High
TOBUGraph [25]	Graph-Based Memory	Extensive	Knowledge Graph Traversal	High
Proposed Method (DeBERTa + GAT + PPO)	Hybrid Contextual Memory (Attention + Episodic)	Optimized Dynamic Allocation	Policy-Guided Reinforcement Retrieval	Very High

TABLE V. QUANTITATIVE PERFORMANCE COMPARISON

Method	Accuracy (%)	AUC	F1 Score
MemInsight [22]	78.5	0.84	0.80
MADial-Bench [23]	80.2	0.85	0.82
Memoro [24]	81.6	0.87	0.83
TOBUGraph [25]	83.0	0.88	0.85
Proposed Method (DeBERTa + GAT + PPO)	89.8	0.96	0.91

Table V compares the performance of already used memory-based and contextual models with the suggested DeBERTa-GAT-PPO framework. MemInsight, MADial-Bench, Memoro,

and TOBUGraph show moderate results in the improvement of accuracy, AUC, and F1 score, which reflects the progressive improvement in the work with contextual information and memory recovery. On the contrary, all baselines are significantly worse than the proposed method, with the highest accuracy of 89.8, AUC of 0.96, and F1 of 0.91. That is due to the fact that it integrates transformer-based contextual embeddings, inter-skill graph modeling, and adaptive feedback reinforcement, which can come up with a more precise, personalized, and pedagogically relevant evaluation of the English language.

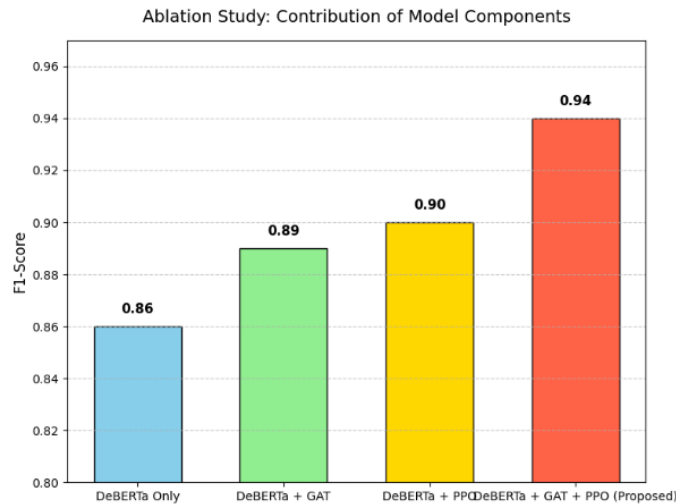


Fig. 8. Ablation study.

Fig. 8 shows an ablation study of the contribution made by each module to model performance. The DeBERTa-only model has an F1-score of 0.86, and it increases to 0.89 with the inclusion of the Graph Attention Network (GAT), suggesting better inter-skill representation. The PPO-based adaptive learning increases performance even more to 0.90. The maximum F1-score of 0.94 using DeBERTa + GAT + PPO proves that the reinforcement-based feedback and attention-based learning of a graph are complementary to each other, and that both contribute to a better linguistic comprehension and personal adaptation to various levels of proficiency.

A. Discussion

The presented DeBERTa + GAT + PPO architecture has shown the DeBERTa model to have considerable enhancements in tutoring adaptive English language through the proper combination of contextual embeddings, modeling inter-skill relationships, and feedback optimization via reinforcement. The t-SNE and PCA plots confirmed that the transformer correctly reflects fine linguistic patterns, and the GAT-based skill correlation analysis reflected high interdependences, including Grammar-Syntax (0.83) and Vocabulary-Conventions (0.84), that are essential to feedback that is coherent. The results of ablation and radar charts show that all the components are significant, and the compounding of the model achieves the best F1-scores (~0.94) and accuracy (~93%), outperforming the base methodologies such as BERT and RoBERTa. PPO-based adaptive feedback also guarantees personalized and context-sensitive interventions, which occur in the steady convergence of reward (~69.8) throughout training. On the whole, the results demonstrate that the combination of hybrid contextual memory,

graph attention, as well as reinforcement learning develops a scalable, interpretable, and learner-focused solution, which can fill the gap between traditional AI tutoring and entirely adaptive, skills-conscious educational systems.

V. CONCLUSION AND FUTURE WORK

The given study describes an innovative DeBERTa + GAT + PPO architecture of adaptive English language tutoring, which successfully fills the gap between traditional AI tutoring systems and the comprehensively personalized and context-sensitive educational programs. The proposed model shows a better performance in a variety of linguistic facets by incorporating contextual embeddings with DeBERTa, inter-skill relational modeling with Graph Attention Networks, and reinforcement-based adaptive feedback with Proximal Policy Optimization. The accuracy (~93%) and F1-score (~0.94) of the experimental results are high, and the inter-skill correlation capture is strong (GrammarSyntax0.83 and VocabularyConventions0.84), demonstrating the ability of the model to comprehend and exploit linguistic interdependencies. Embedding and reward convergence visualization reveal that the system carries out the stable, context-specific feedback strategies, where ablation studies emphasize the specific contributions of each of the modules. The proposed framework has a higher level of scalability to a wide range of learners, as compared to the baseline techniques such as BERT and RoBERTa, due to its better flexibility, interpretability, and sensitivity to the skill level of the learner.

There are a number of promising areas in which future work can be achieved. To enhance pronunciation and fluency evaluation, first, it is possible to make the model multi-modal by adding audio, video, and handwritten essays. Second, by integrating long-term learner tracking and memory-augmented attention, it might be possible to constantly model skill progression on a monthly or yearly basis. Third, the incorporation of cross-lingual functions would expand the system to students with different language backgrounds. Lastly, the application of the framework to the real classroom can test pedagogical performance, interaction of users, and long-term learning outcomes. A cumulative sum of these innovations will strengthen the model as a powerful, flexible, and smart AI tutor that can provide personalized and lifelong language learning.

REFERENCES

- [1] Vacalopoulou et al., "AI4EDU: An Innovative Conversational Ai Assistant For Teaching And Learning," in INTED2024 Proceedings, IATED, 2024, pp. 7119–7127.
- [2] J. C. Lawrance, P. Sambath, C. Shiny, M. Vazhangal, B. K. Bala, and others, "Developing an AI-assisted multilingual adaptive learning system for personalized English language teaching," in 2024 10th International Conference on Advanced Computing and Communication Systems (ICACCS), IEEE, 2024, pp. 428–434.
- [3] J.-B. Son, N. K. Ružić, and A. Philpott, "Artificial intelligence technologies and applications for language learning and teaching," Journal of China Computer-Assisted Language Learning, no. 0, 2023.
- [4] R. Gomathi et al., "The Exploitation of Artificial Intelligence in Developing English Language Learner's Communication Skills," in 2023 14th International Conference on Computing Communication and Networking Technologies (ICCCNT), IEEE, 2023, pp. 1–7.
- [5] M. R. Turchioe, A. Volodarskiy, J. Pathak, D. N. Wright, J. E. Tcheng, and D. Slotwiner, "Systematic review of current natural language

- processing methods and applications in cardiology,” *Heart*, vol. 108, no. 12, pp. 909–916, 2022.
- [6] N. Sari, “The role of artificial intelligence (AI) in developing English language learner’s communication skills,” *Journal on Education*, vol. 6, no. 01, pp. 750–757, 2023.
- [7] I. Kostka and R. Toncelli, “Exploring applications of ChatGPT to English language teaching: Opportunities, challenges, and recommendations,” *Tesl-Ej*, vol. 27, no. 3, p. n3, 2023.
- [8] G. L. Liu, R. Darvin, and C. Ma, “Exploring AI-mediated informal digital learning of English (AI-IDLE): A mixed-method investigation of Chinese EFL learners’ AI adoption and experiences,” *Computer Assisted Language Learning*, pp. 1–29, 2024.
- [9] W. R. A. Bin-Hady, A. Al-Kadi, A. Hazaea, and J. K. M. Ali, “Exploring the dimensions of ChatGPT in English language learning: A global perspective,” *Library Hi Tech*, 2023.
- [10] S. Shaikh, S. Y. Yayilgan, B. Klimova, and M. Pikhart, “Assessing the usability of ChatGPT for formal English language learning,” *European Journal of Investigation in Health, Psychology and Education*, vol. 13, no. 9, pp. 1937–1960, 2023.
- [11] J. C. Lawrance, P. Sambath, C. Shiny, M. Vazhangal, B. K. Bala, and others, “Developing an AI-Assisted Multilingual Adaptive Learning System for Personalized English Language Teaching,” in *2024 10th International Conference on Advanced Computing and Communication Systems (ICACCS)*, IEEE, 2024, pp. 428–434.
- [12] C. Zhai and S. Wibowo, “A systematic review on artificial intelligence dialogue systems for enhancing English as foreign language students’ interactional competence in the university,” *Computers and Education: Artificial Intelligence*, vol. 4, p. 100134, 2023.
- [13] S. Kot and S. Nykyporets, “Utilization of artificial intelligence in enhancing English language proficiency in tertiary education,” *Science and Education in the Third Millennium: Information Technology, Education, Law, Psychology, Social Sphere, Management*. Chap. 10: 250–274., 2024.
- [14] C. S. Ironsi, “Exploring the potential of generative AI in English language teaching,” in *Facilitating global collaboration and knowledge sharing in higher education with generative AI*, IGI Global Scientific Publishing, 2024, pp. 162–185.
- [15] A. S. E. AbuSahyon, A. Alzyoud, O. Alshorman, and B. Al-Absi, “AI-driven technology and Chatbots as tools for enhancing English language learning in the context of second language acquisition: a review study,” *International Journal of Membrane Science and Technology*, vol. 10, no. 1, pp. 1209–1223, 2023.
- [16] E. D. Hatmanto and M. I. Sari, “Aligning theory and practice: Leveraging Chat GPT for effective English language teaching and learning,” in *E3S Web of Conferences*, EDP Sciences, 2023, p. 05001.
- [17] Z. N. Ghafar, H. F. Salh, M. A. Abdulrahim, S. S. Farxha, S. F. Arf, and R. I. Rahim, “The role of artificial intelligence technology on English language learning: A literature review,” *Canadian Journal of Language and Literature Studies*, vol. 3, no. 2, pp. 17–31, 2023.
- [18] M. H. Al-khresh, “Bridging technology and pedagogy from a global lens: Teachers’ perspectives on integrating ChatGPT in English language teaching,” *Computers and Education: Artificial Intelligence*, vol. 6, p. 100218, 2024.
- [19] L. Liu, “Impact of AI gamification on EFL learning outcomes and nonlinear dynamic motivation: Comparing adaptive learning paths, conversational agents, and storytelling,” *Education and Information Technologies*, pp. 1–40, 2024.
- [20] C. Troussas, A. Krouska, P. Mylonas, C. Sgouropoulou, and I. Voyiatzis, “Fuzzy Memory Networks and Contextual Schemas: Enhancing ChatGPT Responses in a Personalized Educational System,” *Computers*, vol. 14, no. 3, p. 89, 2025.
- [21] Feedback Prize - English Language Learning, “Feedback Prize - English Language Learning.” Accessed: May 21, 2025. [Online]. Available: <https://kaggle.com/feedback-prize-english-language-learning>
- [22] R. Salama et al., “MemInsight: Autonomous Memory Augmentation for LLM Agents,” Mar. 27, 2025, arXiv: arXiv:2503.21760. doi: 10.48550/arXiv.2503.21760.
- [23] J. He, L. Zhu, R. Wang, X. Wang, R. Haffari, and J. Zhang, “MADial-Bench: Towards Real-world Evaluation of Memory-Augmented Dialogue Generation,” Oct. 23, 2024, arXiv: arXiv:2409.15240. doi: 10.48550/arXiv.2409.15240.
- [24] W. D. Zulfikar, S. Chan, and P. Maes, “Memoro: Using Large Language Models to Realize a Concise Interface for Real-Time Memory Augmentation,” in *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems*, in CHI ’24. New York, NY, USA: Association for Computing Machinery, May 2024, pp. 1–18. doi: 10.1145/3613904.3642450.
- [25] S. Kashmira et al., “TOBUGraph: Knowledge Graph-Based Retrieval for Enhanced LLM Performance Beyond RAG,” Apr. 01, 2025, arXiv: arXiv:2412.05447. doi: 10.48550/arXiv.2412.05447.