

Application of Artificial Intelligence Technology in Enhancing the Effectiveness of Production-oriented Approach to Literacy Instruction and Instructional Design

Ru Zhao¹, 

¹ Department of Management Engineering, Anhui Communications Vocational & Technical College, Hefei, Anhui, 230051, China

ABSTRACT

In the era of artificial intelligence, human-computer collaborative teaching has become a new picture of future development in the field of education. Based on the theory of human-computer collaboration and the theory of production-oriented approach (POA), this paper constructs a university English POA teaching model based on human-computer collaboration. It also combines the speech recognition algorithm, S-T behavioural analysis method and social network analysis method to conduct a case study on the current situation of college English classroom teaching under this instructional design model. Meanwhile, a teaching experiment is designed to verify the effectiveness of the constructed POA teaching model. The results of the case study show that most of the university English courses favour the lecture mode, with less interaction between students, and the classroom is dominated by teacher lectures and teacher-student interactions, but at the same time, many teachers begin to experiment with the discussion mode, which increases teacher-student interactions and student-student interactions in the classroom. In addition, the experimental group adopts the POA teaching mode and the control group adopts the traditional lecture mode, and its independent samples t-test results show that the experimental group is significantly better than the homogeneous control group in the dimensions of interest, ability, attitude, and test scores in English literacy after the experiment ($P < 0.05$), which suggests that the combination of AI technology and the production-oriented method can effectively improve the effectiveness of the design of university English literacy teaching and achieve better teaching effectiveness and has potential application value.

Keywords: speech recognition, S-T behavioural analysis, social network analysis, production-oriented method of reading and writing instruction

1. Introduction

In language teaching, combining reading and writing is an important way to cultivate students' comprehensive language use ability. With the continuous updating of education and teaching theories,

 Corresponding author.

E-mail address: 13625513712@163.com (R. Zhao).

Received 20 February 2024; Revised 15 May 2024; Accepted 20 November 2024; Published Online 15 April 2025.

DOI: [10.61091/jcmcc127a-326](https://doi.org/10.61091/jcmcc127a-326)

© 2025 The Author(s). Published by Combinatorial Press. This is an open access article under the CC BY license (<https://creativecommons.org/licenses/by/4.0/>).

the production-oriented method, as a new teaching concept, has shown its unique advantages in the teaching practice of many subject areas, especially in the educational background that emphasizes student-centeredness and practice-orientedness, applying the production-oriented method to teaching has become a teaching reform trend worth exploring [1-2]. Second language teaching is more difficult and will be affected by the first language, increasing the difficulty of teaching. Students in the learning process, will also feel more difficult, and will produce boredom or even resistance, which affects the learning effect, and can not guarantee the learning efficiency and quality [3]. Production-oriented method is a good solution to this problem.

Production-oriented method is a second language teaching method proposed by Professor Wen Qiufang, which bridges the gap between receptive language learning and production language application, solves the problem of separation of students' learning and use, and realizes the use of learning to promote learning and learning to use [4]. The theoretical system of production-oriented method includes three parts: teaching philosophy, teaching assumption and teaching process. Educational philosophy includes the theory of whole-person education, the theory of learning center, and the theory of learning-use integration. Whole-person education means that the ultimate goal is to promote the development of students, and teaching should be student-centered; learning-centeredness emphasizes that classroom teaching is a tool for effective learning and meets the needs of students; learning-use-integratedness emphasizes that students learn from the textbooks and apply the knowledge to practice, and puts forward a new educational concept of learning by using, which helps students to apply the knowledge to practice[5-7]. Teaching assumptions include production-motivated assumption, input-enabling assumption, and choice learning assumption. Production-motivated hypothesis means that when teachers ask students to complete language production tasks, students will find their own deficiencies, which will stimulate the desire to learn in order to obtain better academic performance. Input-enabled hypothesis is based on the premise of production-motivated, teachers in the classroom teaching, for students to reasonably select the input material and explanation, to help students to achieve better academic performance. The hypothesis emphasizes the selective learning, in the input-enabled concept, from the input emphasizing selective learning, under the concept of input facilitation, from the input materials selecting useful information for [8-9]. The teaching process includes motivating, enabling and assessing. In the motivating teaching process, teachers allow students to experiment with productions so as to stimulate the enthusiasm for learning. In the enabling teaching process, teachers provide learning materials and give help to students. In this process, it is necessary to strengthen the teacher-student evaluation, and to strengthen the mutual critique and modification between students and students [10-11]. Through the theoretical system of production-oriented method, the teaching design is also applied, in the process of implementation, the teachers according to the actual situation of the students, choose appropriate reading materials and tasks, and reasonably arrange teaching steps, focusing on students' cooperation and communication and personality development [12-13]. Through the design of production-oriented reading teaching, students can obtain more specific and practical knowledge and ability in reading.

With the wide application of artificial intelligence (AI), the teaching process of production-oriented method also has the figure of AI [14]. However, the potential of artificial intelligence must be more than that, and its various applications in the teaching of production-oriented method have important research significance.

Based on artificial intelligence technology, this paper realizes a human-computer collaborative university English POA reading and writing teaching model by introducing the theory of educational goal taxonomy and the theory of production-oriented approach (POA). In order to analyse the application environment of the model, speech recognition algorithm, S-T behavioural analysis method and social network analysis are comprehensively used to deeply explore the model of college English classroom teaching. And in the end, a teaching control experiment is designed to analyse the superiority of the constructed POA teaching model compared with the traditional teaching model by

using independent samples t-test.

2. Teaching University English by Production-oriented Method Based on Artificial Intelligence Technology

2.1. The Construction of University English Teaching Mode Supported by Human-Computer Collaboration

This paper takes college English reading and writing teaching as a specific research object, combines the theory of taxonomy of educational objectives and the theory of production-oriented approach (POA), and constructs a human-machine synergistic college English POA teaching model as shown in Figure 1. The model is divided into the human-computer synergy layer, the POA teaching layer and the teaching goal layer, and the core content is the human-computer synergy layer, including human-computer diagnosis, human-computer promotion and human-computer optimisation, aiming at the precise implementation of the POA teaching concept through the design of teaching activities in these three human-computer synergistic links, the realisation of personalised college English teaching goals, and the promotion of students' English proficiency in a comprehensive manner.

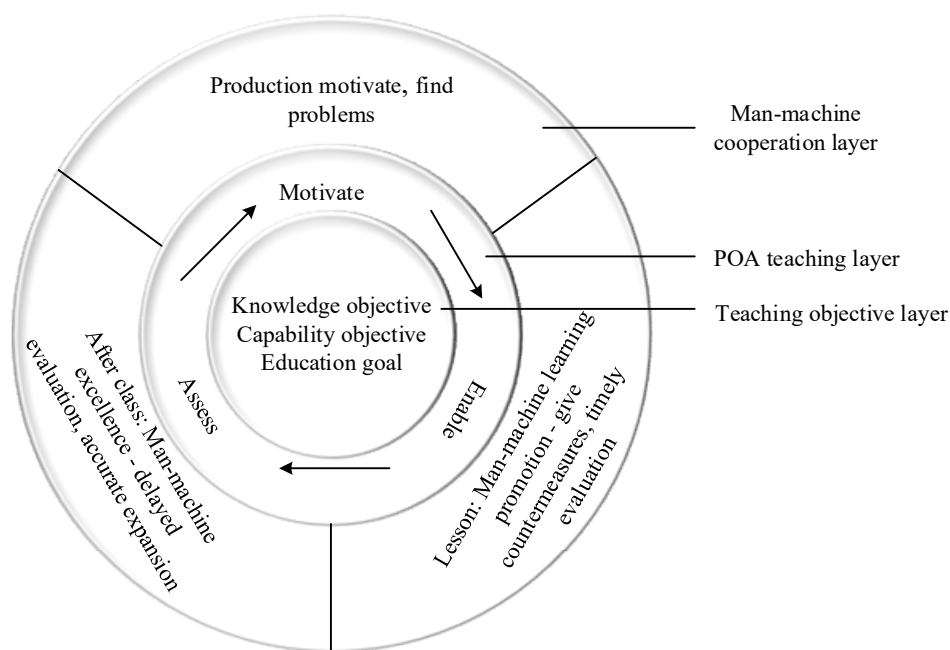


Fig. 1. College English POA teaching mode supported by man-machine collaboration

2.1.1. Teaching objective level. In the layer of teaching objectives, the knowledge objective refers to the knowledge of language forms, including phonetics, vocabulary, chapter, syntax and pragmatics. The capability objective is the higher-order competence of developing students' communicative negotiation skills, analytical and problem-solving skills, intercultural communication skills, and critical thinking skills in the use of English for communication. The education objective aims to integrate the elements of ideology and politics, such as values, workplace ethics, current affairs and politics, national policies, economic development and excellent traditional culture, into the teaching of college English courses, so as to achieve the organic unity of knowledge transmission, ability cultivation and value shaping. The teaching goal layer is production-oriented, shifting from the traditional single goal of language knowledge acquisition to the multi-dimensional goals of language, competence and values, in which the input of lower-order knowledge mainly relies on intelligent technology, while the higher-order competence goals need to be undertaken by teachers.

2.1.2. POA Teaching layer. The POA teaching process consists of three important links: “motivating, enabling and assessing”. Through the motivating process before the class, we can find out the problems in the students' production, clarify the objectives and key points of the unit production, stimulate the desire and enthusiasm for learning, and make the enabling process in the class more accurate and efficient. Through the classroom enabling, the teacher provides language learning materials around the unit production tasks, provides timely feedback on students' problems and gives countermeasures, and gives key explanations and guidance on the language content, task forms and communication strategies that students do not understand. On this basis, the teacher organizes production activities such as group discussions and group presentations, and the teacher completes immediate evaluations by means of teacher-student and student-student mutual evaluations, so as to help students recognize their own shortcomings, thus stimulating their desire to improve the quality of production and forming a learning enabling effect. After the lesson, students improve their production tasks and submit them again, and the teacher completes the delayed evaluation by summarizing and summarizing, consolidates and improves the existing achievements, and pushes supplementary resources on the platform according to the students' completion situation, and formulates the motivating tasks, teaching objectives and teaching contents of the next lesson. The three teaching links of the POA teaching layer are closely connected and interlocked, which provide important theoretical support for the human-computer collaborative teaching.

2.1.3. Man-machine cooperative layer. Intelligent technology has become the cognitive outsourcing of teachers by virtue of its powerful computing, identification, sensing, acquisition, tracking, convergence, analysis and monitoring capabilities, making “technology + teacher” collaborative teaching the norm and the main form of human-computer collaboration. The data-driven intelligent technology favors quantitative cultivation, which is able to deal with the repetitive affairs of large-scale and fine computation in teaching, but it lacks due humanistic care and moral enhancement, and is unable to realize the goal of human education. In contrast, the teacher's qualitative-focused cultivation approach can make up for the lack of humanity in big data, with teachers making teaching decisions based on data collected by AI, while taking on tasks such as inspirational guidance, lesson plan design, moral cultivation, value shaping, emotional communication, and the cultivation of critical thinking and imagination. Therefore, teachers and smart technologies are needed to synergize to promote the overall development of students.

2.2. *University English Teaching Design Supported by Human-Computer Collaboration*

Based on the college English POA teaching model under the human-computer collaboration constructed in the previous papers, this paper designs the teaching program and content of college English relying on the Super Star Learning Pass and iwrite intelligent composition reviewing platform.

2.2.1. Pre-course: human-computer diagnostics - production-motivated localization problems. Before the lesson, the teacher presents the video of the communicative scene of this unit on Learning Access, and students view it independently and then complete the pre-class language test to check their mastery of relevant vocabulary and phrases. At the same time, the teacher releases the production tasks of the unit on iwrite. Upon completion of the production tasks, students recognize their own deficiencies in language skills and strategies, creating a sense of hunger for learning and stimulating their desire for further learning. Teachers use the platform's feedback data to understand the learning situation and pinpoint students' problems in terms of language knowledge, structure, skills and strategies, so as to design teaching programs.

2.2.2. During the lesson: human-computer enables learning - giving countermeasures, timely evaluation. The human-computer facilitation of the in-lesson session includes 4 aspects: language, structure, skills and strategies.

Language enabling: Teachers divided cooperative groups according to the pre-course test. With the help of the multimodal language input materials in the textbook, the group discussion tasks are pushed on Learning Link, and the students explore on their own first, then discuss in groups and upload the exploration results. Based on the statistical visualization results and high-frequency mistakes, the teacher provides targeted guidance to students through language point explanation.

Structure for learning: Teachers provide sample texts, give structural explanations, explain the structure of the texts, and build scaffolding for students with the help of mind maps. On this basis, the teacher pushes the structure selection task, and the students internalize the chapter structure of the sample discourse and complete the structure-enabling task. Based on the visualization scores, teachers provide precise explanations and reinforcement guidance for weak links.

Skills for learning: Teachers push fill-in-the-blank exercises on StudyPass, and students explore on their own. Teachers analyze students' shortcomings in descriptive skills based on the instant feedback from the platform and provide accurate corrections. The teacher then pushes the group presentation task to build "scaffolding" for the students. Students negotiate and produced group results through knowledge sharing and group discussion, uploaded the results to the Learning Channel and make group presentations, reflect on and evaluate the group learning results through inter-group evaluation and teacher evaluation, and provide immediate guidance and feedback during inter-group evaluation to provide solutions to students' doubts, prompting students to check, summarize and precisely improve their expression skills.

Strategies for learning: Teachers first analyze the communication strategies in the discourse. Secondly, they set up multiple-choice exercises and fill-in-the-blank exercises to test the effect of enabling learning. After the students complete the exercises on their own, the teacher shares the excellent samples of the students according to the instant diagnosis of the platform, and promotes the precise improvement and reflection of the students' language strategies through peer assessment.

2.2.3 After class: human-computer optimized learning-extended assessment, accurate expansion. The human-machine e-learning session after class mainly consists of the intelligent platform and the wisdom of teachers and students to jointly assess the unit production tasks submitted by students. Firstly, the teacher pushes the assessment criteria of the assessment tasks, students submit the assessment tasks on iwrite after class, and after the machine automatically grades them, students complete self-assessment and peer assessment according to the assessment criteria. Students revise the production task again and submit it. The teacher selects the outstanding samples and marks them as excellent model essays to share with students, and gives feedback on the common problems that students have. Students reflect and summarize and revise their essays and submit them again. The teacher's delayed assessment helps consolidate the learning effect in a timely manner, accurately expands students' existing knowledge and ability structure, and helps students improve the processing depth of the production task.

3. An Analytical Model of the Effectiveness of the Production-oriented Approach to Teaching English at the University Level

In order to assess the effectiveness of the proposed production-oriented learning teaching model and instructional design of college English, this paper introduces the speech recognition algorithm, the S-T classroom teaching analysis method, and the social network method by analyzing and assessing the classroom interaction behaviors of teachers and students.

3.1. Speech Recognition Algorithm

Speech recognition is a technology that uses machines to recognize and understand speech signals and convert them into appropriate text and commands. Its core is to find the best match by feature

extraction and parsing of unknown speech and then comparing it with known speech. The speech recognition process is shown in Figure 2.

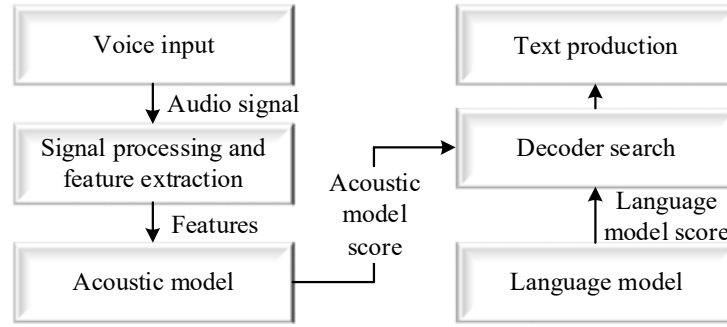


Fig. 2. Speech recognition process

3.1.1. Speech signal acquisition. Voice signal acquisition is the collection of the user's voice signal through a physical device such as a microphone or answering machine. At this point the sound signal is analog and still needs to be converted to a digital signal that the computer can understand.

3.1.2. Speech signal preprocessing. The acquired speech signal is a waveform graph, which firstly needs to be filtered and processed to eliminate the noise, by adding windows and sub-frames to the sound wave signal, as a way to get a short-time smooth signal. Adding window is the process of cutting the signal, and the commonly used window functions are Hamming window and rectangular window. The window length is denoted by N . The Hamming window function is shown in Equation (1).

The Hamming window function is shown in equation (1):

$$\omega(n) = \begin{cases} 0.54 - 0.46 \cos\left(\frac{2\pi n}{N-1}\right), & 0 \leq n \leq N-1 \\ 0, & \text{others} \end{cases} \quad (1)$$

The rectangular window function is shown in equation (2):

$$\omega(n) = \begin{cases} 1, & 0 \leq n \leq N-1 \\ 0, & \text{others} \end{cases} \quad (2)$$

After obtaining the frame signal, it is also necessary to carry out endpoint detection on the frame signal to improve the authenticity and validity of the frame signal, so as to enable the feature extraction of the real frame signal afterwards.

3.1.3. Extraction of characteristic parameters of the speech signal. The Mel frequency cepstrum coefficient (MFCC) is one of the most commonly used methods to parameterize the speech signal by extracting the eigenvalue coefficients of the eigenvalues. In MFCC, for the convenience of calculation, it is necessary to convert the actual sound frequency with Mel frequency, so that the sound frequency and Mel frequency are linearly related, and the conversion formula is:

$$Mel(f) = 2595 \times \log\left(1 + \frac{f}{700}\right) \quad (3)$$

where f is the actual frequency value of the speech signal.

The extraction process of MFCC is as follows: first of all it is necessary to perform a Fast Fourier Transform (FFT) on the signal obtained by preprocessing $x(n)$. The formula of the FFT is shown in Eq. (4):

$$X_{\alpha}(n) = \sum_{n=0}^{N-1} x(n) e^{\frac{-2\pi jnk}{N}}, 0 \leq k \leq N \quad (4)$$

where $x(n)$ is the signal.

After obtaining the energy distribution function of the signal, in order to filter out some non-essential waveforms, to simplify the calculation and to improve the resonance of the waveforms, it is necessary to introduce a triangular filter for filtering operation. The frequency response function of the triple filter is shown in equation (5):

$$H_m = \begin{cases} 0 & k < f(m-1) \\ \frac{2(k-f(m-1))}{(f(m+1)-f(m-1))(f(m)-f(m-1))} & f(m-1) \leq k \leq f(m) \\ \frac{2(f(m+1)-k)}{(f(m+1)-f(m-1))(f(m)-f(m-1))} & f(m) \leq k \leq f(m+1) \\ 0 & k > f(m+1) \end{cases} \quad (5)$$

where $f(m)$ is the Mel frequency and $\sum_{k=0}^{N-1} H_m(k) = 1$.

The corresponding logarithmic energy E^q is obtained by performing a logarithmic operation on the production value after the triangular filter, after which a discrete cosine transform is required as a means of obtaining the MFCC features.

3.1.4. Speech recognition models:

1) Hidden Markov Models. A Markov chain is a special kind of Markov process, which can be described by the following mathematical process since the states and times of the nodes are discrete:

Suppose there exists a random sequence S_t with $\theta_1, \theta_2, \theta_3, \dots, \theta_t$ states, which shifts from one state to another over time, remembering that the probability of being in a state at a $(t+k)$ moment is q_{t+k} and that this probability is only related to the probability q_t of being in a state at its t moment, i.e., the following equation:

$$P(S_{t+k} = q_{t+k} | S_t = q_t, \dots, S_1 = q_1) = P(S_{t+k} = q_{t+k} | S_t = q_t) \quad (6)$$

where $q_1, q_2, q_3, \dots, q_n \in (\theta_1, \theta_2, \theta_3, \dots, \theta_n)$.

Then call S_t a Markov chain and P_{ij} the transfer probability from state i to state j , denoted by:

$$P_{ij}(t, t+k) = P(q_{t+k} = \theta_j | q_t = \theta_i) \quad (7)$$

where $1 \leq i \leq j \leq N$, i, j, t are positive integers.

When $P_{ij}(t, t+k)$ and moment t are independent, this Markov chain is called a sub-Markov chain, i.e., it has the following equation:

$$P_{ij}(t, t+k) = P_{ij}(k) \quad (8)$$

When the value of k takes 1, $p_{ij}(1)$ is written as the one-step transfer probability, referred to as transfer probability a_{ij} . The transfer probability matrix $A = [a_{ij}]$ can be obtained, where $a_{ij} \geq 0, \forall i, j > 0$ and $\sum_{j=1}^N a_{ij} = 1$.

The k -step transfer probability $a_{ij}(k)$ can be obtained by the one-step transfer probability a_{ij} . At this point, it is also necessary to know the distribution of the model's first trial, which is usually denoted as $\pi = (\pi_1, \pi_2, \dots, \pi_N)$, and π_i denotes the probability that the initial state of the model is θ_i . Which $\sum \pi_i = 1, \forall \pi_i \geq 0$.

In practice, some problems can not be described completely using Markov chain, so the Hidden Markov (HMM) model is introduced on this basis, relative to the Markov model, the HMM model introduces a stochastic process to represent the state transition of the event, and the change of the

state is perceived through the stochastic process, which can be described as:

$$\lambda = (N, M, \pi, A, B) \quad (9)$$

where N denotes the number of states in the Markov in the model, denoted as $\theta_1, \theta_2, \dots, \theta_N$. M denotes the number of observations per state, denoted as $V_1, V_2, \dots, V_M$. π denotes the initial trial probability, denoted as $\pi = (\pi_1, \pi_2, \dots, \pi_N)$. A denotes the state transfer probability matrix, denoted as $A = [a_{ij}]$, and B denotes the observation probability matrix, denoted as $B = [b_{ij}]$.

2) Neural network model. In ANN (Approximate Nearest Neighbor) artificial neurons are its basic building blocks. There are multiple inputs in the network model, which are weighted and calculated by the neural units in each layer to obtain the production values. A typical neuron is described as follows:

(1) Input vectors $X_j = \{x_1, x_2, \dots, x_n\}^T$, $1 \leq i \leq n$, x_i denote the input of the i th neuron.

(2) Connection strength W_{ij} denotes the connection strength between neuron node i and neuron node j .

(3) Threshold b_j of neuron j , where the threshold node is denoted using $x_0 = 1$ fixed bias input nodes, then the connection strength $w_{0j} = -1$ of the neuron adjacent to it.

The production of neuron j is weighted and summed:

$$z_j = \sum_{i=0}^n x_i w_{ij} = \sum_{i=1}^n x_i w_{ij} - b_j \quad (10)$$

The production state of neuron j is:

$$y_j = f(z_j) = f\left(\sum_{i=1}^n x_i w_{ij} - b_j\right) \quad (11)$$

Where $f(\cdot)$ function denotes the activation function, a Sigmoid function is commonly used to represent the relationship between the input and production of a neuron. It can fix the value of the element at $[0, 1]$. Its expression is as follows:

$$\text{sigmoid}(x) = \frac{1}{1 + \exp(-x)} = \frac{e^x - e^{-x}}{e^x + e^{-x}} \quad (12)$$

Only multi-layer neural networks can describe the complex node dependencies, so deep neural networks (DNN) are introduced.

Unlike the traditional forward computation method used in Hidden Markov Models, the backward computation method is mainly used in DNN-HMM. The backward computation method is shown in Fig. 3. V denotes the feature vector; V^0 denotes the input feature vector, and V^L denotes the production feature vector, where L denotes the total number of layers in the DNN model. If the current number of computed layers l is less than the total number of layers L , it is necessary to calculate the current excitation vector value to determine the state of the current layer Z^l , where B is the deviation coefficient matrix of the current layer and W is the deviation weight matrix of the current layer. Through the processing of the weighted value of Z , and then after the activation function can be obtained after the operation of the state value of the current layer, and this forward projection, you can get the probability value of each state at a certain moment. After getting the initial calculation value, we need to do regression operation. The commonly used regression function is Softmax function, the input unit of Softmax regression has one to many, through the normalization operation, to determine the characteristics of the production vector, so as to determine the final value of the backward calculation. The advantage of backward computation over forward computation is that the assessment is accurate and focuses on analyzing the posterior state of the signal.

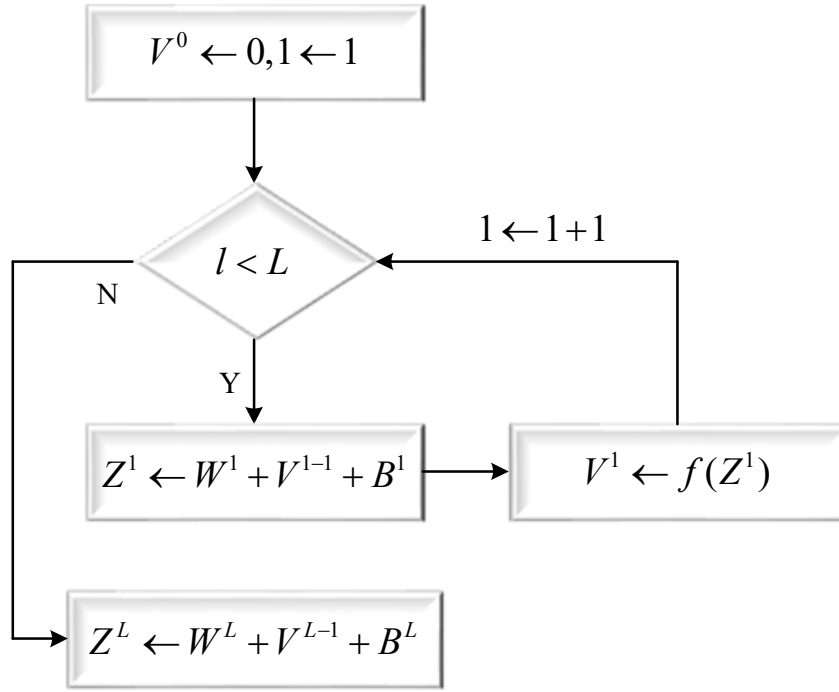


Fig. 3. Backward calculation method

Acoustic modeling based on DNN enables the machine to learn the association between feature values very well and make full use of the correlation between the frames of speech signals, so that the machine can think like a human being and strongly classify the received signals, so in this paper, we combine HMM and DNN for acoustic modeling in order to perform the task of speech recognition. The structure of the DNN-HMM model is shown in Fig. 4.

In this model, each node in the HMM is used to describe the dynamics of the speech signal, and the production of each node in the DNN is used to estimate the probability of an observation feature for a state in the HMM. Let the sequence of words be $W = (w_1, w_2, \dots, w_T)$ and the sequence of observations be $O = (O_1, O_2, \dots, O_T)$. The speech recognition problem can be represented by finding the sentence W that maximizes $P(W | O)$. Namely:

$$W = \max(P(W | O)) = \arg \max(O | W)P(W) \quad (13)$$

where $P(O | W)$ is the acoustic model and $P(W)$ is the language model.

In the acoustic model, the phoneme is usually taken as the modeling unit, and the acoustic model is further simplified by setting the sequence of articulatory states corresponding to the word sequence W as $Q = (q_1, q_2, \dots, q_T)$:

$$\begin{aligned} P(O | W) &= \sum P(O, Q | W)P(Q | W) \\ &\approx \text{mdx} \pi(q_0) \prod_{t=1}^T \alpha_{q_{t-1}q_t} \prod_{t=1}^T P(O_t | q_t) \end{aligned} \quad (14)$$

where $\alpha_{q_{t-1}q_t}$ denotes the transfer probability between states q_{t-1} and q_t .

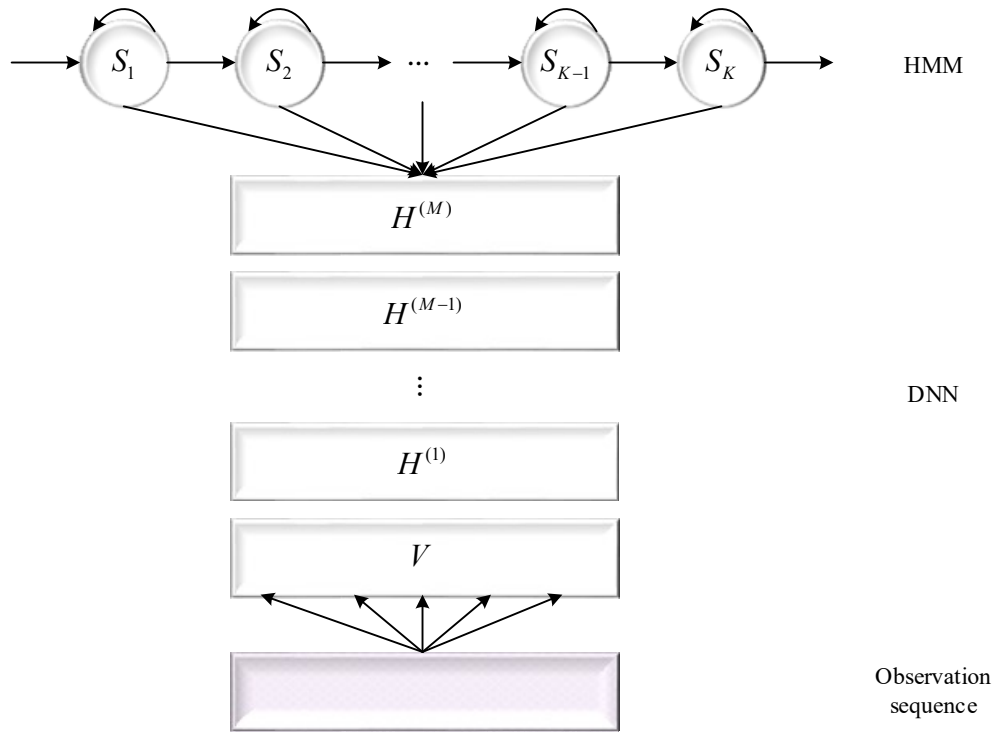


Fig. 4. DNM-HMM model structure

3.2. S-T Behavioral Analysis

S-T behavior analysis method is a more mature classroom teaching analysis method[15]. The method first calculates the proportion of teacher and student behavioral hours in the classroom, and then conducts classroom teaching analysis through the change of teacher and student behavioral hours in the class check. S-T behavioral analysis method divides the behavior in the process of classroom teaching into teacher's behavior (S) and student's behavior (T), and the definitions of the two kinds of behaviors are shown as follows:

T behavior: i.e., teacher information behavior, which is manifested in the teacher's explanation, question and answer, board writing, multimedia presentation, etc.

S behavior: i.e. student information behavior, which is manifested as students asking questions, answering questions, silent thinking, taking notes, etc.

The S-T behavior analysis method usually combines the teaching video with a fixed sampling length to code the teacher-student behaviors, i.e., the whole teaching video is segmented by a fixed length, and the teacher-led T behaviors and student-led S behaviors shown in each segmented segment are labeled and coded. After encoding, the data obtained are then calculated to draw the corresponding data graphs.

The traditional S-T behavioral coding method for coding usually adopts a fixed duration segmentation method, taking a fixed segmentation duration from 3-30s to segment the audio. Assuming that the total length of a lesson is T , the number of codes is N , and the number of teacher-student interactions is $N-1$, the first segment, if it is the teacher's audio, is divided into the teacher behavior segment, coded as i , i.e., $C_i = 1$, and the other audios are divided into the student behavior segments, coded as 0 , i.e., $C_i = 0$. Then, the proportion of the teacher-dominated time in the classroom, Rt , and the rate of teacher-student interactions, Ch , can be found in the following equations:

$$Rt = \frac{1}{N} \sum_{i=1}^N C_i \quad (15)$$

$$Ch = \frac{1}{N} \sum_{i=1}^{N-1} \{ (1.0 - C_i) * C_{i+1} + C_i^* (1.0 - C_{i+1}) \} \quad (16)$$

As can be seen from the above formula, traditional S-T behavior coding is calculated and evaluated by coding number N . The number of codes N is the same as the total time duration. The length of the segmentation time affects the size of N when the total duration is constant, i.e., the Rt and Ch values have an obvious relationship with the segmentation duration and lack time invariance. In order to address the effect of segmentation duration on the coding analysis of S-T behaviors, a variable-duration S-T behavior coding method has been proposed, i.e., when the instructional behavior transitions, the data at this point are recorded and coded in this way. However, this method has the situation that when the difference in coding length is too large, the analysis using the S-T behavior method is prone to large errors.

In this paper, we use a method of coding S-T behaviors by subdividing the average length of time into variable lengths. Based on the teacher-student behavioral time segments for variable time length segmentation, that is, the teacher-student behavioral jump point as the segmentation point, the audio is divided into teacher behavioral segments and student behavioral segments for S-T behavioral coding, take the average of all time segments as the average length, use the average length to subdivide the S-T behavioral coding, and then take the first half of the first half and the second half of each jump point as the interactive behavior, and the remaining segments are all the silence behavior.

The total duration of a lesson is T . The number of variable duration coded segments is N . For the i st segment, the duration is t_i . If it is teacher audio, it is divided into teacher behavior segments coded as 1, i.e., $C_i = 1$, and the other audio is divided into student behavior segments coded as 0, i.e., $C_i = 0$. The number of interaction coded segments is M . For the j th segment, the duration is t_j .

If it is interaction behavior, it is coded as 1, i.e., $I_j = 1$, and if it is silence behavior, it is coded as 0, i.e., $I_j = 0$. Variable duration, the coded teacher-led percentage of time Rt and teacher-student interaction rate Ch can be solved for in the following equation:

$$Rt = \frac{1}{T} \sum_{i=1}^N (C_i * t_i) \quad (17)$$

$$Ch = \frac{1}{T} \sum_{j=1}^M (I_j * t_j) \quad (18)$$

3.3. Social Network Analysis

Currently, social network analysis in education is mostly used to analyze the interaction behaviors of students and teachers in online course forums, while the interaction behaviors in offline classroom environments are rarely analyzed because offline classroom data are not as concentrated and easy to be extracted as in online forums[16]. Therefore, in this paper, on the premise of analyzing classroom audio data and solving the problem of offline classroom data collection, social network analysis is innovatively applied to the analysis of offline classroom interaction behavior. The speakers in the classroom are corresponded to a single node in the social network, and the weight of the node is used to represent the total time of the speaker's speech: the line between two nodes indicates that there is an interaction between the speakers, and the weight of the line between the two nodes is used to represent the sum of the interaction time, and the classroom interaction behavior in the classroom is analyzed in this way, and a social network graph is constructed.

Next, the teacher-student behavior conversion rate, classroom interaction distance, classroom interaction density, social network density, social network diameter, individual node degree and node average degree in classroom interaction behavior are calculated.

1) Teacher-student behavior conversion rate reflects the frequency of interaction in the course of lectures. The higher the frequency of teacher-student interaction in the classroom, the greater the

resulting teacher-student behavior conversion rate, which is calculated by formula (19):

$$Rt = \frac{f_{st}}{T \cdot O_c} \quad (19)$$

Where, f_{st} represents the number of times the teacher-student interaction behavior is transformed; T represents the classroom audio duration, and O_c is the sampling frequency of the audio.

2) Classroom interaction distance reflects how students interact with the teacher and students interact with each other in the classroom. The larger the classroom distance is, the smaller the interaction intensity is. The classroom interaction distance is calculated from equation (20):

$$d_i = \frac{1}{n-1} \sum_{j \neq i} d_{ij} \quad (20)$$

3) Classroom interaction density reflects the overall classroom interaction behavior. The interaction density formula is:

$$I_n = \frac{\sum(l \times \omega_l)}{\sum(n \times \omega_n)} \quad (21)$$

Where, l denotes the sum of the number of edges between all nodes; ω_l denotes the weight of the edges; n denotes the total number of nodes, and ω_n denotes the weight of a single node.

4) Social network density reflects the closeness of interactions among classroom members. The closer the social network density is to 1, the closer the interaction behavior between members. The social network density is calculated as shown in equation (22):

$$D_n = 2 \frac{l}{n(n-1)} \quad (22)$$

5) The social network diameter reflects the autonomy of student-to-student interactions. The larger the network diameter, the greater the autonomy of the students in that classroom. The network diameter is calculated by the formula:

$$d_n = \max(dia) \quad (23)$$

In the formula, dia is used to represent the distance between two nodes.

6) The average degree of the network can visualize the number of interactions in the classroom, and the higher the average degree, the more frequent the classroom interactions and the more positive the overall classroom atmosphere. The formula for calculating the average degree is:

$$D_{ave} = \frac{\sum D}{n} \quad (24)$$

7) By comparing the node degree values of each node, we can find out the interaction differences between students with different attributes in the classroom, and the larger the node degree value, the more active the students are in the classroom. Therefore it can reflect the degree of activity in the classroom to a certain extent. The formula for calculating the node degree number is:

$$D = ID + OD \quad (25)$$

where ID represents the in-degree of the node and OD represents the out-degree of the node.

4. Case Studies

In order to validate the usefulness of the constructed model, this chapter analyzes the classroom

interaction behaviors from two perspectives: single case and case comparison, respectively.

4.1. Single Case Analysis

4.1.1. Vocal recognition and clustering of classroom speech. In this paper, the course of college English is selected as a case study for analysis. Firstly, the classroom video is processed and converted into an audio file in wave format with a bit rate of 256kbps and a mono audio file with a sampling rate of 16k. To achieve speaker differentiation in classroom audio, the process of feature extraction, active speech detection, speaker change detection, speaker clustering, etc. is required, and then speech recognition is performed on the audio file based on the speech recognition algorithm, and finally S-T behavioral analysis and social network analysis are performed in turn.

1) Speech activity detection. There are four parameters, namely, Minimum Speech Duration, Minimum Non-Language Duration, Seg. Before Expansion, and Seg. After Expansion, which have an impact on the active speech detection process. Among them, the first two parameters have a greater impact on the final result, i.e., the minimum speech that can constitute a speaker change and the minimum non-speech that will separate two consecutive turns. The last two represent the feature padding between each speech segment and have a smaller impact on the results.

In order to compare the performance of different values assigned to these parameters, accuracy is measured using VER, with smaller values of VER indicating better performance of the algorithm. Where VER is defined as:

$$VER = \frac{SaN + NaS}{TotalTime} \quad (26)$$

SaN denotes misidentification of speech segments as non-speech segments and NaS denotes misdetection of non-speech segments as speech segments.

In order to ensure that there is enough speech for meaningful speech segmentation and clustering, this paper explores the minimum speech (MS) values of 0.5s, 0.8s and 1.0s. For Minimum Non-Speech (MNS), a range of values starting from the NIST standard of 0.2s and then using 0.2-1.2s values separated by 0.2s, where the 0.2s difference between manually annotated labels and automatically generated labeled scores is ignored. The VERs for cases with different MS and MNS values are shown in Table 1.

Through the above parameter tuning process, the best results are obtained when the minimum speech is 0.5s and the minimum non-speech is 1.2s, then this parameter will be used for active speech detection during the experiment.

Table 1. VER for different MS and MNS values

Voice Activity Detection					
MS	MNS	VER	MS	MNS	VER
0.5s	0.2s	0.397	0.8s	0.8s	0.166
0.5s	0.4s	0.208	0.8s	1.0s	0.118
0.5s	0.6s	0.165	0.8s	1.2s	0.122
0.5s	0.8s	0.094	1.0s	0.2s	0.547
0.5s	1.0s	0.075	1.0s	0.4s	0.483
0.5s	1.2s	0.071	1.0s	0.6s	0.359
0.8s	0.2s	0.529	1.0s	0.8s	0.254
0.8s	0.4s	0.404	1.0s	1.0s	0.208
0.8s	0.6s	0.196	1.0s	1.2s	0.178

2) Speaker change detection. For speaker change detection, three distance calculations, BIC growth window, GLR and KL2 fixed sliding window, are used for detection. In this case, the equation MTER is

used to compare the performance of the different parameters, and in order to differentiate from the VER for active speech detection, the parameters that have been calculated to give the best performance for active speech detection are selected when fine-tuning these parameters. Where MTER is defined as the ratio of miscalculated speaker change points to total speaker change points. A smaller value of MTER indicates better performance of the algorithm. Namely:

$$MTER = \frac{\#ofmissdturns}{\#oftotalturns} \quad (27)$$

(1) Fixed sliding window. For this method using the distance of GLR or KL2, the biggest parameter selection problem is to find a suitable threshold for the distance and this threshold is closely related to the data, the window size (wSz) and the window step size (WS) are parameters that need to be adjusted. It is reasonable to choose the size of 4s, 5s and 6s for the window size, and after some tests, the window step size is fixed to 0.5 seconds, because a larger step size will lead to performance degradation. For GLR, after some data testing, the MTER with fixed sliding window and GLR is shown in Table 2. The correct threshold for the classroom data used in the case should be between 2100 and 2600, values below 2100 will result in smaller MTER but more computationally intensive, if the problem is solved by clustering, it can also reduce the MTER again, but if more audio segments are cut out, it will result in a slow and incorrect clustering, so the threshold should be chosen between 2100 and 2600.

Table 2. MTER with fixed sliding window and GLR

Speech		
MWS	Threshold	MTER
4s	2100	0.219
4s	2350	0.408
4s	2600	0.591
5s	2100	0.294
5s	2350	0.352
5s	2600	0.548
6s	2100	0.329
6s	2350	0.361
6s	2600	0.442

The MTER with fixed sliding window and KL2 is shown in Table 3. The thresholds for KL2 were found to be very low and difficult to adjust, as small threshold changes resulted in large MTER differences. The optimal threshold for the case data seems to be between 12 and 24, however, the final performance is also poor, and further lowering the threshold introduces too many false alarms while improving the results. Therefore, for a fixed sliding window approach, this paper suggests that GLR is recommended on KL2 because GLR thresholds are easier to adjust to the data.

Table 3. MTER with fixed sliding window and KL2

Speech Segmentation		
MWS	Threshold	MTER
4s	12	0.309
4s	18	0.486
4s	24	0.635
5s	12	0.493
5s	18	0.632
5s	24	0.718
6s	12	0.629
6s	18	0.668
6s	24	0.727

(2) Growth window using BIC. For the growth window method using the BIC metric, the parameters to be tuned are the minimum window size (MWS), which defines the maximum window size, and the minimum window growth window step (WS), called δWS . The MTER with increasing window and BIC is shown in Table 4. It was found that it is preferable to set δWS to 0.1s, since larger values will only negatively affect the performance, even if it will make the algorithm slightly faster. Thus, in Table 4, δWS is fixed to 0.1s and the column is omitted. Alternatively, a smaller value of λ yields better performance, and similarly, it is fixed to 1.0 in Table 4.

From the results Table 4, it can be seen that a minimum window size (MWS) of 1.2s best suits the case data. Smaller values can lead to using too little data for decision making, while larger values may skip decision boundaries. As for the maximum window, defined by the window step size (WS), it has no weight in the final MTER when the MWS is set to 1.2s, but from the first two rows of the table we can see that it reaches the optimal position of 3.2s.

Therefore, in the case data, as the window increases and the BIC metric improves, the final speaker change detection is performed with better results. The final optimal solutions are 1.2s for MWS, 3.2s for WS. δWS for 0.1s and 1.0 for λ . It can be seen that the results using the BIC growth window are better than those obtained using the fixed sliding window method, and the final speaker segmentation

will be performed using these parameters.

Table 4. MTER with increased window and BIC

Speech Segmentation		
MWS	Threshold	MTER
0.6s	2.0s	0.189
0.6s	3.2s	0.182
0.6s	4.4s	0.208
1.2s	2.0s	0.167
1.2s	3.2s	0.167
1.2s	4.4s	0.167
1.8s	2.0s	0.164
1.8s	3.2s	0.164

(3) Speaker clustering. The last step is speaker clustering. There is only one parameter to be adjusted for this process, namely the λ of the BIC metric used. This is done by calculating the DER scores for different values of λ with the optimal parameter of the previous process fixed. A smaller value of DER indicates a lower error rate, and DER is defined as:

$$DER = \frac{\sum_{Segs} (ST * (\max(IS?, SD?) - CD?))}{\sum_{Segs} (ST * IS?)} \quad (28)$$

Where: *Segs* denotes a speech segment. *ST* denotes the duration (in seconds) of each speech fragment. *IS?* denotes whether it is a speaker or not. It takes the value of 1 if it is a speaker and 0 if it is not. *SD?* denotes whether it is detected as a speaker or not. It takes the value of 1 if it is detected as a speaker and 0 if it is not. *CD?* denotes whether it is detected accurately or not. It takes the value of 1 if the label is computed correctly and 0 if it is computed incorrectly.

The DERs for different values of λ are shown in Table 5. In Table 5, it can be seen that the optimal λ for clustering is in the range of 1.25, and based on this λ as well as the optimal parameters for each of the previous processes, the algorithm has a final DER of 0.142.

Table 5. DER for different input λ values

Speaker Clustering	
λ	DER
1	0.176
1.2	0.142
1.25	0.142
1.3	0.142
1.4	0.164
1.8	0.164

4.1.2. Analysis of classroom interactions. The results of the analysis and clustering of the speakers in this classroom can be obtained with the above optimal parameters, and the beginning and end times of each speaker's speech can be obtained by processing the data and calculating the values of Rt and Ch, and plotting the Rt-Ch diagram, and plotting the trajectory of classroom behaviors in the S-T diagram. There were 135 minutes in this classroom and a total of 800 behaviors (N) were produced, which contained 416 teacher behaviors (NT) and 384 student behaviors (NS), and a total of 162 continuous behaviors (g).

The values of Rt and Ch were calculated using Equation (15) and Equation (16) as 0.520 and 0.201,

respectively. Since $0.3 < Rt = 0.520 < 0.7$ and $Ch = 0.201 < 0.4$, this course is pedagogically modeled as a hybrid classroom.

The Rt-Ch plot obtained for this classroom is shown in Figure 5.

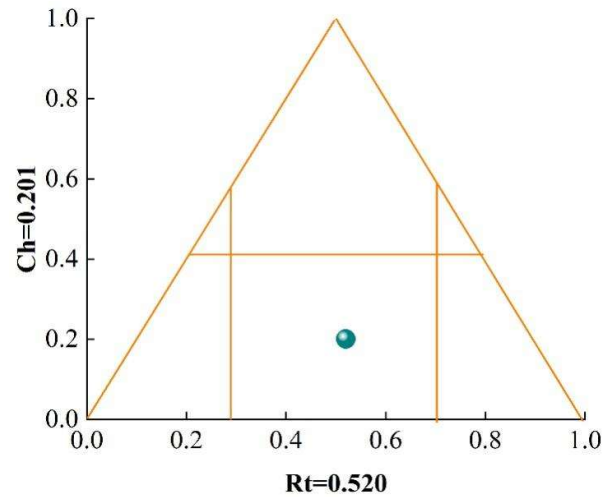


Fig. 5. Rt-Ch diagram of the class

The S-T diagram for this classroom is shown in Figure 6.

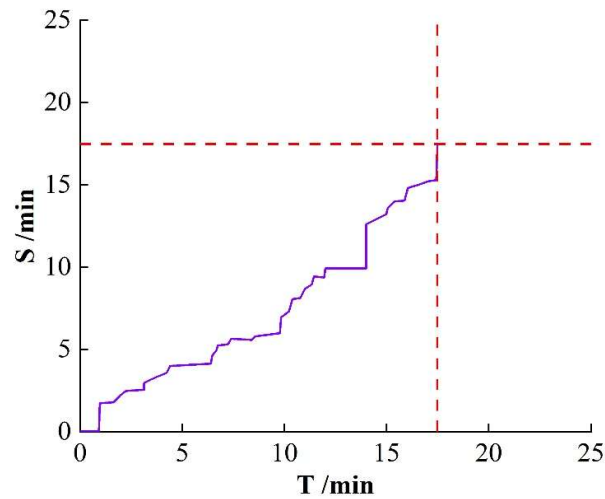


Fig. 6. S-T diagram of the class

Then the directed graph with nodes and edges weighted according to the result of voiceprint recognition is constructed as the social network graph, and the social network graph is shown in Fig. 7.

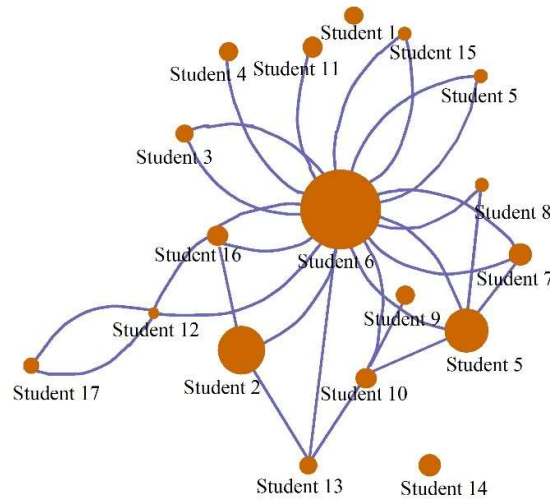


Fig. 7. Interactive network diagram of the class

The statistics of the number of all nodes in this classroom are shown in Table 6.

Table 6. Degrees of all nodes in the class

Teacher	286	Student 9	3
Student 1	2	Student 10	4
Student 2	4	Student 11	3
Student 3	3	Student 12	4
Student 4	3	Student 13	5
Student 5	3	Student 14	4
Student 6	4	Student 15	3
Student 7	4	Student 16	4
Student 8	6	Student 17	3

Based on Figure 7 and Table 1, the teacher-student behavior conversion rate, interaction density, network density, network diameter, node degree, and average degree can be calculated:

Calculate the teacher-student behavioral conversion rate:

$$Rt = \frac{f_{st}}{T \cdot O_c} = \frac{161}{800} = 0.201 \quad (29)$$

Calculate classroom interaction distance:

$$d_i = \frac{1}{n-1} \sum_{j \neq i} d_{ij} = 0.938 \quad (30)$$

Calculate the interaction density:

$$I_n = \frac{\sum (l \times w_l)}{\sum (n \times w_n)} = \frac{10769}{27948.74} = 0.385 \quad (31)$$

Calculate network density:

$$D_n = 2 \frac{l}{n(n-1)} = 0.183 \quad (32)$$

Calculate the network diameter:

$$d_n = \max(\text{diameter}) = 5 \quad (33)$$

Calculate the average degree:

$$D_{ave} = \frac{\sum D}{n} = 4.889 \quad (34)$$

Since $0.3 < Rt = 0.520 < 0.7$ and $Ch = 0.201 < 0.4$, this course is pedagogically modeled as a hybrid classroom, which is the type of course that makes up the largest percentage of current courses, and the vast majority of instructors' courses are currently hybrid. $d_n = 5 > 3$, $I_n = 0.385 > 0.08$, so the classroom belongs to the balanced structure; the interaction in this case classroom is more average; the interaction between teachers and students, students and students are well balanced; the length of students' speeches and the length of teachers' speeches also present a balanced state. $D_{ave} = 4.889 < 5$, $Rt = 0.201 > 0.1$, in the classroom mode this classroom belongs to the lecture mode. There is very little interaction between students, and the classroom is dominated by the teacher's lecture and teacher-student interaction. This case is dominated by the teacher's speech, which produces less interaction, but the interaction pattern is balanced, through the quantification and analysis of the course, it can be obtained that: this case belongs to the hybrid teaching model; the classroom interaction structure is balanced structure, and the classroom mode belongs to the lecture mode.

4.2. Comparative Analysis of Cases

The analysis method proposed in this paper solves the problem that the current classroom analysis relies on human manual analysis, so it can be realized to analyze a large amount of data. In the multi-case analysis, excellent university English open courses (Lessons 1- Lessons 20) were selected as the classroom data to be analyzed and processed, and after examining the quantity as well as the quality of the courses, a total of 20 courses were selected.

4.2.1. Vocal recognition and clustering of classroom speech. Because the selected courses in the multi-case analysis and the single-case analysis are university English courses and the quality of the data is similar, the parametric running method described in the previous section is used, and the specific algorithmic process will not be described in detail here.

4.2.2. Analysis of classroom interactions. The calculation results of the selected 20 courses were subjected to data analysis and parameter calculation, processed to obtain the beginning and end time of each speaker's speech, calculated to obtain the values of Rt and Ch for each course, and plotted the Rt - Ch diagram, and plotted the trajectory of classroom behavior on the S-T diagram. The S-T analysis metrics for each course are shown in Table 7.

In Table 7, it can be seen that the length of these courses is around 36 minutes to 45 minutes, the length of the video recording is in line with the length of a lesson, and the number of teacher-student behavioral transitions ranges from 60-172, indicating that the habits of different teachers in the classroom are different and varied.

Table 7. S-T analysis indicators for each course

Teacher	Duration /min	NS	NT	g	Rt	Ch
Course 1	42.77	616	230	90	0.272	0.105
Course 2	43.73	376	491	152	0.566	0.174
Course 3	44.14	431	417	128	0.492	0.150
Course 4	43.48	393	468	162	0.544	0.187
Course 5	39.90	294	503	113	0.631	0.141
Course 6	40.67	562	254	60	0.311	0.072
Course 7	39.45	569	271	123	0.323	0.145
Course 8	42.35	277	545	77	0.663	0.092
Course 9	41.37	353	468	143	0.570	0.173
Course 10	36.07	276	506	95	0.647	0.120
Course 11	43.80	250	615	100	0.711	0.114
Course 12	41.90	237	605	131	0.719	0.154
Course 13	39.63	459	371	91	0.447	0.108
Course 14	39.53	520	297	148	0.364	0.180
Course 15	39.85	321	476	131	0.597	0.163
Course 16	44.22	377	508	133	0.574	0.149
Course 17	43.94	384	486	88	0.559	0.100
Course 18	44.50	439	429	172	0.494	0.197
Course 19	39.17	335	482	101	0.590	0.122
Course 20	39.55	528	293	84	0.357	0.101

Rt-Ch plots were drawn for all courses as shown in Figure 8. The classroom mode can be determined from the plots drawn by calculating the Rt and Ch values for these courses: course 1 belongs to the practice mode, i.e., students speak more and the teacher-student behavior conversion rate is small. While course 11 and course 12 belong to the lecture mode, i.e., the teacher dominates in the interaction process. And most of the classes belong to the mixed mode, which is the main classroom mode in the current university English courses.

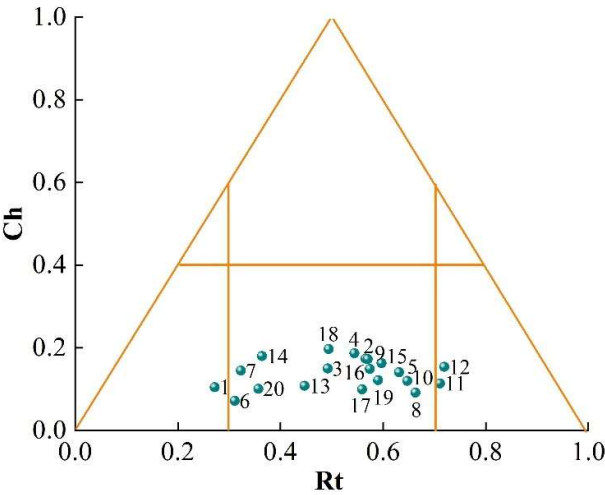


Fig. 8. Rt-Ch plots of 20 cases

Then based on the results of voice recognition to construct nodes and edges are weighted directed graph as a social network graph, in each classroom social network graph can be visualized in the classroom teachers accounted for the proportion of the classroom, the number of classroom speakers

as well as the interaction of the general situation, roughly grasp the classroom teacher-student interactions as well as student-student interactions of the distribution. Thus, it can be out of the calculation of teacher-student behavior conversion rate R_t , classroom interaction distance d_i , interaction density I_n , network density D_n , network diameter d_n , node degree D , average degree D_{ave} and other parameters, the social network analysis parameters of the 20 cases are shown in Table 8. The classroom interaction distance in the cases is between 0.73-1, the interaction density is 0.016-0.212, the network density is between 0.025-0.171, the minimum network diameter is 3, the maximum network diameter is 5, and the average degree is 3.631-6.553.

Table 8. Social network analysis parameters of 20 cases

Teacher	di	In	Dn	dn	Dave
Course 1	0.838	0.106	0.041	4	6.160
Course 2	0.945	0.098	0.056	4	4.255
Course 3	0.851	0.093	0.067	4	3.921
Course 4	0.839	0.092	0.064	4	5.088
Course 5	0.931	0.085	0.055	4	5.461
Course 6	0.776	0.096	0.054	5	5.335
Course 7	0.859	0.204	0.030	5	6.553
Course 8	0.847	0.057	0.072	5	4.296
Course 9	0.736	0.095	0.072	5	4.300
Course 10	1.000	0.016	0.171	3	3.837
Course 11	0.868	0.029	0.099	4	3.944
Course 12	0.820	0.065	0.082	4	3.631
Course 13	0.922	0.125	0.049	4	5.347
Course 14	0.856	0.212	0.025	5	6.184
Course 15	0.860	0.046	0.079	4	4.251
Course 16	1.000	0.072	0.052	3	4.789
Course 17	0.872	0.084	0.063	4	4.097
Course 18	0.873	0.105	0.050	4	4.163
Course 19	0.819	0.083	0.055	5	4.631
Course 20	0.880	0.108	0.044	4	5.525

A summary of the 20 case studies gives the classroom interaction structure for each course as:

Balanced structure: course 1, course 2, course 4, course 5, course 13, course 17, course 18, course 20.

Scattered structure: course 3, course 10, course 11, course 12, course 15, course 16.

Concentrated structure: course 6, course 7, course 8, course 9, course 14, course 19.

The 20 quality college English lessons are evenly distributed in terms of interaction structure, and it is speculated that the classroom interaction structure is more relevant to the individual teacher's teaching habits, and this speculation can be verified by analyzing different lessons of the same teacher when the amount of data is sufficiently supported. Analyze their classroom patterns as:

Indoctrination mode: course 8.

Lecture mode: course 2, course 3, course 9, course 10, course 11, course 12, course 15, course 16, course 17, course 18, course 19, course 20.

Student-student discussion mode: course 5, course 13.

Faculty-student discussion mode: course 1, course 4, course 6, course 7, course 14.

Most of the courses favor lecture mode, which is the most common classroom mode of college English at present, with less interaction between students, and teacher lecture and teacher-student interaction dominating the classroom, and also many teachers begin to try the discussion mode, which

increases teacher-student and student-student interaction in the classroom, while the indoctrination mode is less and less being adopted under the requirements of the new curriculum standard nowadays.

5. Empirical Analysis of the Effectiveness of the Teaching Model

In order to verify the effectiveness of the designed college English production-oriented method teaching model, this chapter conducts a controlled teaching experiment, mainly comparing and analyzing the questionnaire survey and students' English reading and writing test scores in two dimensions, horizontally and vertically. The experimental class was taught college English reading and writing using the production-oriented method teaching model designed in this paper, while the control group was taught using the traditional teaching model.

5.1. Pre-laboratory

First, questionnaires and pre-tests on reading and writing were administered to the students in the experimental class and the control class, so as to understand the basic situation of the students in the two classes in terms of their English literacy skills, their interest in reading and writing and their attitudes towards reading and writing.

5.1.1. Questionnaires. In this paper, we now use SPSS28.0 software to carry out independent sample t-value test on the results of the questionnaire survey, which is analyzed in the following three aspects. 1) Students' interest in English reading and writing. A comparison of the English literacy interest of the experimental and control classes before the experiment is shown in Table 9. The questionnaire focuses on the following aspects of college students' interest in reading and writing: whether they like reading in English (question 1), whether they like writing in English (question 2), whether they like to actively do English reading questions (question 3) as well as actively write in English outside of class (question 4), and whether they actively participate in reading and writing classes for tasks they can complete (questions 5 and 6, which indicate the reading and writing courses' respectively active participation).

Firstly, we compare the mean values of the test items of the two classes. The mean values of students' liking of English reading in the control class and the experimental class are 3.31 and 3.28 respectively, and the mean value of the experimental class is 0.03 lower than that of the control class, and the mean value of students' liking of English writing in the control class is 3.28, which is 0.06 lower than the mean value of the experimental class (3.34), and the first two test items show that most of the students in the two classes don't like English reading as well as English writing very much. The mean values of students in the control class and the experimental class who would actively do English reading questions after school were 3.35 and 3.32 respectively, and the mean values of students who would actively write in English were 3.45 and 3.44 respectively. These two test items indicate that most of the students in the two classes did not utilize their time outside of school to actively practice reading and writing, and that the students' time for reading and writing was mainly concentrated in the classroom. In the English reading class, the mean values of students' active participation in the tasks they can accomplish in the control class and the experimental class are 3.52 and 3.51 respectively. In English writing class, on the other hand, the mean value of active participation of the students in the control class for the tasks that they can do by themselves is 3.41, which is 0.01 higher than the mean value of the experimental class (3.40). From the latter two test items, it can be seen that in English reading class as well as in writing class, the students of the two classes do not have enough self-confidence, and they are still reluctant to participate actively even in the tasks that they are able to do by themselves. Analyzing the data, it can be seen that the literacy interest of the students in both classes is roughly the same, the mean is basically around 3.4, and most of the students do not have a high interest in literacy, neither enjoying reading and writing in English, nor actively reading and writing in English. In English literacy classes, students in both classes are not very interested in

participating in classroom literacy tasks.

Table 9. Comparison of English reading and writing interest before the experiment

question	Class	Number of cases	Mean value	Standard deviation	Average standard error
1	Control class	54	3.31	1.159	0.172
	Experimental class	52	3.28	1.126	0.174
2	Control class	54	3.28	1.241	0.169
	Experimental class	52	3.34	1.235	0.176
3	Control class	54	3.35	1.304	0.178
	Experimental class	52	3.32	1.095	0.153
4	Control class	54	3.45	1.132	0.149
	Experimental class	52	3.44	1.101	0.152
5	Control class	54	3.52	1.113	0.146
	Experimental class	52	3.51	1.128	0.154
6	Control class	54	3.41	1.247	0.171
	Experimental class	52	3.40	1.106	0.162

Secondly, the results of the independent samples t-test for the literacy interest profiles of the experimental and control classes are shown in Table 10[17]. The p-value of the mean equivalence t-test significance for all question test items is greater than 0.05, indicating that there is no significant difference between the reading and writing interests of the students in the control class and the experimental class before the experiment.

Table 10. Independent sample T-test of reading and writing interest before the experiment

question	Levin's test for equality of variances				T-test		
	Assuming	F	Sig.	t	Sig.(2-tailed)	95% confidence interval of the difference	
						Lower limit	Upper limit
1	Equal variances	0.281	0.607	0.071	0.951	-0.467	0.502
	Unequal variances			0.071	0.951	-0.468	0.501
2	Equal variances	0.027	0.868	-0.758	0.464	-0.649	0.305
	Unequal variances			-0.758	0.464	-0.649	0.305
3	Equal variances	2.806	0.113	0.184	0.862	-0.424	0.512
	Unequal variances			0.185	0.861	-0.421	0.508
4	Equal variances	0.021	0.904	-0.039	0.984	-0.445	0.431
	Unequal variances			-0.039	0.984	0.445	0.430
5	Equal variances	0.026	0.885	0.052	0.969	-0.427	0.439
	Unequal variances			0.052	0.969	-0.427	0.439
6	Equal variances	1.309	0.281	0.027	0.978	-0.456	0.464
	Unequal variances			0.027	0.978	-0.453	0.461

2) Students' English literacy situation. A comparison of the English literacy situation of the experimental and control classes before the experiment is shown in Table 11. The questionnaire focuses on the following aspects of students' literacy: self-perception of reading and writing levels (questions 7 and 8), pre-reading prediction (question 9, mainly referring to the ability to predict the topic and content of an article based on its title), extracting information during reading (question 10, mainly referring to the ability to extract the main ideas and information of an article), reorganizing and reproducing the text based on key words after reading (question 11), constructing ideas before writing (question 12), the ability to choose appropriate vocabulary and structure according to the

needs of expression during writing (question 13), and the ability to self-correct after writing (question 14).

Before the experiment, the English literacy skills of the experimental class and the control class were generally the same, and the difference between the means of the two classes in question 7 to 14 was in the range of 0.01 to 0.09, and the mean values were all less than 4, which indicated that the students in the two classes were not confident in their reading and writing skills, and that there was not much difference in the reading and writing skills of the students in the two classes.

Table 11. Comparison of English reading and writing abilities before the experiment

question	Class	Number of cases	Mean value	Standard deviation	Average standard error
7	Control class	54	3.52	1.148	0.164
	Experimental class	52	3.46	1.079	0.162
8	Control class	54	3.62	1.045	0.141
	Experimental class	52	3.61	1.137	0.158
9	Control class	54	3.83	0.996	0.137
	Experimental class	52	3.75	0.973	0.141
10	Control class	54	3.84	0.964	0.130
	Experimental class	52	3.76	1.011	0.145
11	Control class	54	3.86	0.824	0.112
	Experimental class	52	3.79	0.815	0.116
12	Control class	54	3.67	0.943	0.129
	Experimental class	52	3.65	0.978	0.138
13	Control class	54	3.88	0.902	0.124
	Experimental class	52	3.79	0.797	0.115
14	Control class	54	3.84	0.856	0.116
	Experimental class	52	3.77	0.844	0.121

Meanwhile, the results of the independent samples t-test for literacy skills of the experimental and control classes are shown in Table 12. There is no significant difference between the means of all question test items ($P > 0.05$), indicating that there is no significant difference between the literacy skills of the students in the control class and the experimental class before the experiment.

Table 12. Independent sample T-test of reading and writing abilities before the experiment

question	Levin's test for equality of variances				T-test		
	Assuming	F	Sig.	t	Sig.(2-tailed)	95% confidence interval of the difference	
						Lower limit	Upper limit
7	Equal variances	0.141	0.728	0.452	0.664	-0.341	0.531
	Unequal variances			0.461	0.663	-0.338	0.529
8	Equal variances	0.153	0.716	0.037	0.983	-0.412	0.428
	Unequal variances			0.038	0.983	-0.415	0.431
9	Equal variances	0.218	0.661	1.134	0.307	-0.184	0.592
	Unequal variances			1.135	0.307	-0.181	0.590
10	Equal variances	0.174	0.692	0.468	0.652	-0.303	0.475
	Unequal variances			0.467	0.651	-0.305	0.478
11	Equal variances	0.515	0.486	1.194	0.243	-0.128	0.512
	Unequal variances			1.195	0.242	-0.128	0.509
12	Equal variances	1.563	0.227	1.874	0.068	-0.024	0.724
	Unequal variances			1.868	0.069	-0.25	0.725
13	Equal variances	0.148	0.722	0.702	0.497	-0.220	0.439
	Unequal variances			0.713	0.494	-0.217	0.435
14	Equal variances	1.723	0.205	1.161	0.263	-0.138	0.514
	Unequal variances			1.162	0.263	-0.138	0.513

3) Comparison of students' attitudes toward reading and writing in English before the experiment. A comparison of the situation of English literacy attitudes between the experimental and control classes before the experiment is shown in Table 13. The questionnaire focuses on the following aspects of students' attitudes toward reading and writing: opportunities to participate in the classroom (question 15), reading and writing selection (question 16), motivation to participate in classroom activities (question 17), the form of completing the task (questions 18 and 19, which indicate that reading is taught in the form of completing the task and English writing is learned, respectively), and the way of doing the activity (question 20, which indicates that a group activity is used to (complete teacher-created activities).

Students' attitudes toward reading and writing instruction in the two classes were basically the same, as shown by the following two main points. First, the means of the six test items for both classes were observed. The difference between the mean values of the experimental class and the control class in each item is between 0.03 and 0.15, and the mean value is less than 4. This indicates that students in both classes have fewer opportunities to participate in the classroom, have basically the same views on literacy selection, are less willing to participate in classroom-organized activities, and basically have the same views on teaching English literacy through completing tasks and group activities.

Table 13. Comparison of English reading and writing attitude before the experiment

question	Class	Number of cases	Mean value	Standard deviation	Average standard error
15	Control class	54	3.92	1.108	0.148
	Experimental class	52	3.89	0.924	0.132
16	Control class	54	3.81	1.026	0.139
	Experimental class	52	3.74	0.947	0.134
17	Control class	54	3.82	0.879	0.121
	Experimental class	52	3.75	0.795	0.119
18	Control class	54	3.86	0.958	0.135

	Experimental class	52	3.71	1.063	0.151
19	Control class	54	3.68	0.862	0.118
	Experimental class	52	3.59	0.943	0.136
20	Control class	54	3.84	0.944	0.128
	Experimental class	52	3.81	0.875	0.124

Secondly, the results of the independent samples t-test of literacy attitudes of the experimental and control classes are shown in Table 14. From the T-test results, we know that the p-value of the mean equivalent T-test significance of all the test items is greater than 0.05, which means that there is no significant difference between the control class and the experimental class of students' attitudes toward reading and writing in English on these test items before the experiment.

Table 14. Independent sample T-test of reading and writing attitude before the experiment

question	Levin's test for equality of variances				T-test		
	Assuming	F	Sig.	t	Sig.(2-tailed)	95% confidence interval of the difference	
						Lower limit	Upper limit
15	Equal variances	0.014	0.924	0.501	0.634	-0.294	0.476
	Unequal variances			0.503	0.632	-0.292	0.474
16	Equal variances	0.056	0.827	1.114	0.283	-0.176	0.592
	Unequal variances			1.118	0.281	-0.174	0.591
17	Equal variances	0.192	0.674	1.175	0.253	-0.134	0.521
	Unequal variances			1.182	0.251	-0.132	0.519
18	Equal variances	1.915	0.182	0.854	0.404	-0.232	0.541
	Unequal variances			0.850	0.409	-0.234	0.543
19	Equal variances	1.428	0.248	1.825	0.076	-0.031	0.671
	Unequal variances			1.816	0.077	-0.034	0.675
20	Equal variances	0.093	0.785	0.116	0.923	-0.331	0.378
	Unequal variances			0.116	0.918	-0.336	0.376

5.1.2. English literacy pre-test scores. Comparison of the reading and writing scores of the two classes before the experiment is shown in Table 15. First, from the mean of the two classes' scores, the mean of the control and experimental classes are 3.674 and 3.581 respectively, with a difference of only 0.093, which indicates that there is not much difference between the two classes' reading and writing scores. Secondly, from the standard deviation of the two classes' performance, the control and experimental classes are 0.648 and 0.715 respectively, with a difference of 0.067, which indicates that there is not much difference in the performance of the students in the two classes, but the degree of dispersion is relatively large in both classes, which indicates that the performance of the students in both classes fluctuates more.

Table 15. Comparison of English reading and writing pre-test results

Finally, the independent samples t-test results of the English literacy pretest scores of the two

	Class	Number of cases	Mean value	Standard deviation	Average standard error
Grade	Control class	54	3.674	0.648	0.089
	Experimental class	52	3.581	0.715	0.102

classes are shown in Table 16. The p-value of all test items is 0.431, which is greater than 0.05, indicating that there is no significant difference between the pretest scores of the control class and the experimental class.

Table 16. Independent sample T-test of English reading and writing pre-test results

	Levin's test for equality of variances				T-test		
	Assuming	F	Sig.	t	Sig.(2-tailed)	95% confidence interval of the difference	
						Lower limit	Upper limit
Grade	Equal variances	1.281	0.283	0.815	0.431	-0.164	0.372
	Unequal variances			0.812	0.433	-0.165	0.374

5.2. Post-experiment

Questionnaires and literacy post-tests were again administered to students in the experimental and control classes, so as to find out whether there were any changes in English literacy, literacy interest and literacy attitudes of the students in the two classes.

5.2.1. Questionnaires:

1) Students' reading and writing interest situation. The results of the independent samples t-test of the reading interest situation of the students in the two classes after the experiment are shown in Table 17. From the T-test results, it is known that the mean equivalence T-test significance P-value of all test items is 0.000, which is less than 0.001, indicating that there is an extremely significant difference between the reading interest of the control class and the experimental class after the experiment.

Table 17. Independent sample T-test of reading and writing interest after the experiment

question	Levin's test for equality of variances				T-test		
	Assuming	F	Sig.	t	Sig.(2-tailed)	95% confidence interval of the difference	
						Lower limit	Upper limit
1	Equal variances	15.165	0.002	-7.914	0.000	-1.887	-1.128
	Unequal variances			-8.122	0.000	-1.876	-1.138
2	Equal variances	14.304	0.001	-6.347	0.000	-1.564	-0.817
	Unequal variances			-6.518	0.000	-1.559	-0.823
3	Equal variances	33.854	0.000	-5.926	0.000	-1.651	-0.817
	Unequal variances			-6.137	0.000	-1.642	-0.832
4	Equal variances	2.514	0.162	-5.818	0.000	-1.525	-0.708
	Unequal variances			-5.902	0.000	-1.526	-0.714
5	Equal variances	11.267	0.000	-4.825	0.000	-1.028	-0.482
	Unequal variances			-4.923	0.000	-1.059	-0.493
6	Equal variances	24.358	0.001	-6.295	0.000	-1.728	-0.905
	Unequal variances			-6.457	0.000	-1.719	-0.924

2) Situation of students' reading and writing skills. The independent samples T-test results of the reading interest situation of the two classes after the experiment are shown in Table 18. From the T-test results, it is known that the mean T-test Sig. (two-tailed) of all test items is 0.000, which is less than 0.05, indicating that there is a significant difference between the literacy skills of the control class and the experimental class after the production-oriented method of literacy teaching based on artificial intelligence technology.

Table 18. Independent sample T-test of reading and writing abilities after the experiment

question	Levin's test for equality of variances				T-test		
	Assuming	F	Sig.	t	Sig.(2-tailed)	95% confidence interval of the difference	
						Lower limit	Upper limit
7	Equal variances	15.817	0.253	-4.965	0.000	-1.261	-0.534
	Unequal variances			-5.003	0.000	-1.258	-0.541
8	Equal variances	1.362	0.004	-4.172	0.000	-0.957	-0.337
	Unequal variances			-4.185	0.000	-0.956	-0.338
9	Equal variances	0.874	0.000	-3.952	0.000	-0.996	-0.326
	Unequal variances			-3.976	0.000	-0.994	-0.331
10	Equal variances	8.318	0.341	-4.205	0.000	-1.048	-0.383
	Unequal variances			-4.356	0.000	-1.052	-0.388
11	Equal variances	7.356	0.006	-3.917	0.000	-0.989	-0.319
	Unequal variances			-3.964	0.000	-0.982	-0.324
12	Equal variances	1.419	0.264	-2.318	0.000	-0.689	-0.045
	Unequal variances			-2.349	0.000	-0.684	-0.048
13	Equal variances	4.472	0.035	-4.205	0.000	-1.023	-0.374
	Unequal variances			-4.324	0.000	-1.026	-0.378
14	Equal variances	9.044	0.003	-4.405	0.000	-1.007	-0.377
	Unequal variances			-4.468	0.000	-1.002	-0.395

3) Students' attitudes toward reading and writing. The results of the independent samples t-test of the English reading and writing attitudes of the experimental and control classes after the experiment are shown in Table 19. Except for the two test items of whether they have the opportunity to express their opinions and whether they want literacy materials to be relevant to their lives, the p-value of the test of homogeneity of variance for all variables is greater than 0.05, which means that the variances are equal, and the results of the t-test look at the results of the "hypothetical isotropy", and the results of the t-test tell us that the mean equivalence of all other test items is less than 0.05, indicating that there is a significant difference between these test items and the experimental class after the experiment. T-test results show that the mean equivalence of the other test items has a p-value of significance less than 0.05, which means that there is a significant difference between the attitudes of the control class and the experimental class towards reading and writing after the experiment for these test items.

Table 19. Independent sample T-test of reading and writing attitude after the experiment

question	Levin's test for equality of variances				T-test		
	Assuming	F	Sig.	t	Sig.(2-tailed)	95% confidence interval of the difference	
						Lower limit	Upper limit
15	Equal variances	3.004	0.089	-4.496	0.000	-1.063	-0.416
	Unequal variances			-4.448	0.000	-1.057	-0.424
16	Equal variances	7.852	0.007	-2.561	0.018	-0.759	-0.071
	Unequal variances			-2.515	0.016	-0.764	-0.082
17	Equal variances	8.456	0.006	-4.286	0.000	-0.945	-0.341
	Unequal variances			-4.367	0.000	-0.931	-0.352
18	Equal variances	1.224	0.304	-3.861	0.000	-0.934	-0.293
	Unequal variances			-3.922	0.000	-0.930	-0.304
19	Equal variances	0.627	0.483	-3.501	0.000	-0.862	-0.237
	Unequal variances			-3.554	0.000	-0.859	-0.242
20	Equal variances	3.804	0.065	-3.852	0.001	-0.904	-0.295
	Unequal variances			-3.943	0.001	-0.901	-0.304

5.2.2. English literacy post-test scores. The results of the independent samples t-test of the English literacy post-test scores of the experimental and control classes after the experiment are shown in Table 20. As can be seen from Table 20, the mean equivalent t-test significance p-value of the scores is less than 0.05, indicating that there is a significant difference between the scores of the control and experimental classes after the experiment.

Table 20. Independent sample T-test of English reading and writing post-test results

	Levin's test for equality of variances				T-test		
	Assuming	F	Sig.	t	Sig.(2-tailed)	95% confidence interval of the difference	
						Lower limit	Upper limit
Grade	Equal variances	49.205	0.000	-8.624	0.000	-0.983	-0.628
	Unequal variances			-8.976	0.000	-0.981	-0.639

5.3. Comparative Analysis Before and After the Experiment

The comparative analyses before and after the experiment focused on longitudinally comparing questionnaires, classroom observations, and English literacy scores.

5.3.1. Questionnaires:

1) Comparison of pre- and post-test data of the questionnaire in the experimental class. The pre- and post-test comparisons of English literacy, literacy interest and literacy attitude of the experimental class are shown in Table 21. It can be seen from Table 21 that the difference between the mean values of the three aspects of literacy interest, literacy ability and literacy attitude of the experimental class before and after the experiment is 1.04, 0.67 and 0.67 respectively. Meanwhile, the Sig.(2-tailed) of all the test items is less than 0.05, which indicates that there is a significant difference in the literacy ability, interest as well as attitude of the students in the experimental class after the experiment, which has been significantly enhanced or improved.

Table 21. Comparison of test data before and after questionnaire survey in experimental class

Dimensions	Group	Number of cases	Mean value	Standard deviation	t	Sig. (2-tailed)
Interests	Pre-test	52	3.38	0.921	-7.182	0.000
	Post-test	52	4.42	0.319		
Ability	Pre-test	52	3.70	0.702	-6.428	0.000
	Post-test	52	4.37	0.295		
Attitude	Pre-test	52	3.75	0.721	-5.941	0.000
	Post-test	52	4.42	0.297		

2) Comparison of pre- and post-test data of the questionnaire in the control class. The pre- and post-test comparisons of the control class students' English literacy ability, literacy interest and literacy attitude are shown in Table 22. It can be seen from the table that the difference between the mean values of the three aspects of reading and writing interest, ability and attitude of the control class students before and after the experiment is -0.23, -0.12 and -0.07 respectively, which indicates that after the experiment, the reading and writing interest and ability of the control class students not only did not improve, but also declined, and the attitude towards reading and writing was also not changed. After investigating the reasons, it can be found that there are three main reasons for the above phenomenon: firstly, in the traditional college English classroom, teachers mainly emphasize the structure and function of the language, and do not combine the meaning and form of the language, which can't stimulate the students' interest in learning and improve their enthusiasm to participate in the class. Secondly, classroom teaching is mainly dominated by the teacher, who is the center of the classroom, and classroom activities are basically completed under the control of the teacher, and students' initiative is not given full play. Finally, there is a lack of interactive communication among students, and students feel passive and bored in the classroom. From the T-test results, we know that the Sig.(2-tailed) of all the means is greater than 0.05, which indicates that there is no significant difference between the control class before and after the experiment in terms of literacy ability, literacy interest and literacy attitude.

Table 22. Comparison of test data before and after questionnaire survey in control class

Dimensions	Group	Number of cases	Mean value	Standard deviation	t	Sig. (2-tailed)
Interests	Pre-test	54	3.39	0.956	1.164	0.264
	Post-test	54	3.16	0.884		
Ability	Pre-test	54	3.76	0.675	1.089	0.297
	Post-test	54	3.64	0.668		
Attitude	Pre-test	54	3.82	0.753	0.857	0.401
	Post-test	54	3.75	0.729		

5.3.2. English Literacy Pre- and Post-test Scores:

1) Comparison of pre- and post-test scores in the experimental class. Comparison of pre- and post-test scores of the experimental class is shown in Table 23. From Table 23, it can be seen that the mean value of the experimental class in the pre-test is 3.58, and the mean value in the post-test is 4.42, and the mean value of the post-test scores is higher than that of the pre-test by 0.84. The standard deviations of the experimental class in the pre-test and post-test are 0.715 and 0.191 respectively, and the post-test scores are obviously more stable than that of the pre-test scores. And from the T-test results, we know that the Sig.(2-tailed) value of the mean of the scores is $0.000 < 0.05$, which indicates that there is a significant difference between the pre-test and post-test scores of the experimental class, and the English literacy scores of the students of the experimental class have been significantly improved after the experiment.

Table 23. Comparison of test results before and after experimental class

Dimensions	Group	Number of cases	Mean value	Standard deviation	t	Sig. (2-tailed)
Grade	Pre-test	52	3.58	0.715	-7.384	0.000
	Post-test	52	4.42	0.191		

2) Comparison of pre- and post-test scores in the control class. The comparison of the pre- and post-test scores of the experimental class is shown in Table 24. From Table 24, it can be seen that the mean of the control class in the pre-test and post-test is 3.67 and 3.53 respectively, and the mean of the post-test scores is lower than the pre-test by 0.14. The standard deviation of the experimental class in the pre-test is 0.648, and the standard deviation of the experimental class in the post-test is 0.644, which makes the pre- and post-test scores both unstable. And from the t-test results, the Sig.(2-tailed) of the mean of the scores is $0.235 > 0.05$, which means that there is no significant difference between the pre-test and post-test scores of the control class.

Table 24. Comparison of test results before and after control class

Dimensions	Group	Number of cases	Mean value	Standard deviation	t	Sig. (2-tailed)
Grade	Pre-test	54	3.67	0.648	1.305	0.235
	Post-test	54	3.53	0.644		

Summarizing the above analysis, it can be seen that the experimental group adopting the university English production-oriented method teaching mode based on artificial intelligence technology designed in this paper has effectively improved their interests, abilities, attitudes and their test scores in reading and writing teaching, which is obviously better than the traditional English teaching mode.

6. Conclusion

In this paper, combining the theory of taxonomy of educational objectives and the theory of production-oriented approach (POA), we designed a POA reading and writing teaching model for college English based on artificial intelligence technology, and experimentally verified its effectiveness.

In the 20 selected university English classroom cases, the classroom interaction distance is within the range of 0.73-1, the interaction density is 0.016-0.212, the network density is between 0.025-0.171, the minimum network diameter is 3, the maximum network diameter is 5, and the average degree is 3.631-6.553. It can be seen that most of the courses are inclined to the lecture mode, which is currently the most common university English classroom mode, which has less interaction between students and is dominated by teacher lecturing and teacher-student interaction in the classroom. However, at the same time, many teachers have begun to try the discussion mode, which increases teacher-student interaction and student-student interaction in the classroom, and the indoctrination mode is less and less used.

In this paper, we designed a teaching control experiment, the experimental class used the POA teaching mode designed in this paper, while the control class used the traditional lecture teaching mode, and the two classes were homogeneous before the experiment ($P > 0.05$). After the experiment, the experimental class showed significant improvement in the dimensions of interest, ability, attitude, and test scores in English reading and writing instruction ($P < 0.05$), while the control class showed less significant improvement ($P > 0.05$). Meanwhile, after the experiment, the experimental class was significantly better than the control group in all dimension scores, which verified the effectiveness of the POA reading and writing teaching model supported by human-computer collaboration designed in this paper.

References

- [1] Natasha, P. (2023). An production-oriented approach to the impact of online written languaging on form-

- focused writing tasks. *Journal of Language and Education*, 9(1 (33)), 112-127.
- [2] Liu, G., & Zhang, G. (2023). Teaching reform of moral education based on the concept of production-oriented education. *International Journal of Electrical Engineering & Education*, 60(2_suppl), 455-464.
 - [3] Larsen - Freeman, D. (2018). Looking ahead: Future directions in, and future research into, second language acquisition. *Foreign language annals*, 51(1), 55-72.
 - [4] Zhang, W., Tian, H., & Li, Y. (2022). Study on College English Writing Based on Production-Oriented Approach. *Journal of Education, Humanities and Social Sciences*, 4, 268-271.
 - [5] Wang, X. (2022). THE ROLE AND APPLICATION OF PRODUCTION ORIENTED METHOD IN COLLEGE ENGLISH TEACHING UNDER THE BACKGROUND OF THINKING DISORDER. *Psychiatria Danubina*, 34(suppl 1), 154-155.
 - [6] Li, Z. H. A. N. G. (2020). Exploring the Production-oriented Approach in the Teaching of Academic Writing and Presentation. *Contemporary Foreign Languages Studies*, 20(2), 61.
 - [7] Xu, T. (2023). Continuous Improvement and Practice of Production-Oriented Talent Training Model. *International Journal of Mathematics and Systems Science*, 6(4).
 - [8] Wang, Z. (2021). A study on the application of production-oriented approach in college English writing in private colleges. *Journal of Frontiers in Educational Research*, 1(2), 69-71.
 - [9] Yang, P. (2016). An integrated mode study on college English teaching of listening and speaking: Based on production-motivated, input-enabled hypothesis. *Theory and Practice in Language Studies*, 6(6), 1236.
 - [10] Wang, L., & Chen, L. (2022). FROM INPUT-BASED TO PRODUCTION-MOTIVATED: THE TEACHER'S ROLES IN TEACHING COMPREHENSIVE ENGLISH COURSE THROUGH PRODUCTION-ORIENTED APPROACH. *IETI Transactions on Social Sciences and Humanities*, 16, 8-16.
 - [11] Fan, L., Ye, S., & Wang, Y. (2021, June). Construction of an Production-oriented Teacher Evaluation System based on Normal Professional Certification. In *2021 IEEE 3rd International Conference on Computer Science and Educational Informatization (CSEI)* (pp. 41-45). IEEE.
 - [12] Wang, Z. (2023). Instructional Design of Multimodal Input Based on Production-oriented Approach in Teaching English Writing in High School. *Journal of Education and Educational Research*, 2(1), 13-16.
 - [13] Jiang, T. (2019, July). Application of "Production-motivated, Input-enabled Hypothesis" in College English Oral Teaching Design. In *4th International Conference on Contemporary Education, Social Sciences and Humanities (ICCESSH 2019)* (pp. 519-522). Atlantis Press.
 - [14] Li, H. (2024). Construction of College English teaching effect evaluation model based on artificial intelligence and production-oriented approach. *Informatica*, 47(10).
 - [15] Chiaki Konishi. (2010). Do School Bullying and Student—Teacher Relationships Matter for Academic Achievement? A Multilevel Analysis. *Canadian Journal of School Psychology*(1),19-39.
 - [16] Alhassan Yakubu Alhassan. (2025). Rethinking participation in urban planning: analytical and practical contributions of social network analysis. *Frontiers of Urban and Rural Planning*(1),1-1.
 - [17] Ömer Eltas. (2021). Comparison of the Mann-Whitney U test and Independent Samples t-test (Independent Student's t-test) in terms of Power in Small Samples Used in Bioistatistic Studies. *Veterinary Sciences and Practices*(1),88-94.