

Research on Identifying Psychological Health Risks of College Students Based on Knowledge Graph and Multimodal Text Analysis

Dongli Chen

School of Marxism, Zhengzhou Railway Vocational & Technical College, Zhengzhou 451460, Henan, China

Corresponding author: Dongli Chen (e-mail: zzuchendl@126.com)

ABSTRACT This research constructs a “five-in-one” framework—comprising knowledge graph construction, multimodal data preprocessing, feature fusion, risk identification, and empirical verification—to enable the early detection of mental health risks among college students in intelligent campus environments. By integrating psychological, academic, behavioral and environmental information with advanced modeling, the study supports more accurate risk monitoring. Firstly, the core influencing factors of college students’ mental health risks are systematically sorted out. Based on social ecology theory and stress coping theory, a knowledge graph of mental health risks covering five dimensions: individual psychology, academic development, social support, life behavior, and environmental adaptation is constructed; Secondly, collect multimodal text data from college students (social media texts, course reviews, psychological counseling records, voice logs), and construct a multimodal text dataset through preprocessing steps such as text cleaning, emotional annotation, and feature extraction; Once again, design a multidimensional feature system that integrates text semantic features, emotional features, behavioral features, and knowledge graph correlation features, and construct a mental health risk identification model based on deep learning models (BERT+BiLSTM).

INDEX TERMS knowledge graph, multimodal text analysis, mental health risk, bert-bilstm, intelligent campus

I. INTRODUCTION

A. RESEARCH BACKGROUND

As a vital talent pool for the modernization of the industry, the mental health of college students directly impacts the sustainable innovation of high-end manufacturing and the overall stability of industrial production systems. Ensuring their psychological well-being is essential for maintaining the high-quality human capital required to drive technical advancement and operational efficiency [1-3]. The “2024 Report on the Development of Mental Health of Chinese College Students” shows that the detection rate of mental health problems among college students in China has reached 24.7%, with depression, anxiety, stress overload, sleep disorders, and other high incidence problems. Some students have extreme behaviors due to the lack of timely intervention in psychological crises, which has attracted widespread social attention [4-6].

B. RESEARCH SIGNIFICANCE

(1) Constructing a theoretical framework for the knowledge graph of mental health risks among college students, clarifying the core dimensions and interrelationships of risk

factors, and enriching the theoretical system of mental health risk assessment;

(2) Propose a risk identification theory that integrates knowledge graph and multimodal text analysis, and improve the method system of multidimensional information fusion and risk association mining;

(3) Developing a correlation model that integrates multimodal text features with industrial labor intensity parameters (e.g., workshop noise decibels, shift-rotation frequency, and manufacturing excellence metrics). This clarifies the collaborative empowerment mechanism of work-environment stressor data, aligning campus psychological profiles with the specialized operational pressures found in advanced manufacturing sectors, thereby ensuring the model’s relevance to the future workforce. This framework provides methodological references for enhancing psychological monitoring and operational safety across advanced manufacturing sectors.

II. CORE CONCEPTS AND THEORETICAL FOUNDATIONS

A. DEFINITION OF CORE CONCEPTS

Knowledge graph is a structured knowledge representation method that uses the “entity relationship attribute” triplet as the basic unit to clearly present the relationship between entities through graph structure. The knowledge graph of mental health risks among college students in this article refers to a structured knowledge system that takes mental health risk factors as the core entity, covering dimensions such as individual psychology, academic development, social support, life behavior, and environmental adaptation, presenting hierarchical relationships, logical connections, and causal relationships between risk factors.

B. THEORETICAL BASIS

1) Social Ecological Theory

Social ecological theory was proposed by Bronfenbrenner, which holds that individual development is the result of the interaction between individuals and the environment. The environment is divided into microsystems (family, school, peers), mesosystems (connections between microsystems), external systems (parental workplace, community), and macrosystems (socio-cultural, policy systems). This theory provides theoretical support for constructing a multidimensional system of psychological health risk factors, namely that the psychological health risks of college students are influenced by multiple factors such as individual psychological traits, academic environment, social support, life behavior, and external environment [7].

2) Stress Coping Theory

The Stress Coping Theory was proposed by Lazarus, which suggests that stress is a product of the interaction between individuals and their environment. An individual's cognitive evaluation and coping strategies directly affect the degree to which stress affects their mental health. This theory emphasizes that in addition to the stressor itself, individual cognitive evaluation, coping strategies, social support, and other factors can also affect mental health risks, providing a theoretical basis for the screening and correlation analysis of risk factors.

3) Knowledge Graph Theory

The knowledge graph theory covers core technologies such as knowledge representation, knowledge extraction, knowledge fusion, and knowledge reasoning. By organizing knowledge in a triplet form of “entity relationship attribute”, it can effectively present the complex relationships between knowledge. This theory provides methodological support for the structured organization and association mining of mental health risk factors [8,9].

C. CONSTRUCTION OF A KNOWLEDGE GRAPH OF MENTAL HEALTH RISKS AMONG COLLEGE STUDENTS

PRINCIPLES FOR BUILDING KNOWLEDGE GRAPH

1) Principle of Comprehensiveness

Covering the main risk factors affecting the mental health of college students, including multiple dimensions such as individual psychology, academic performance, society, behavior, and environment, to avoid overlooking key risk points.

2) Principle of Scientificity

Based on mature theories such as social ecology theory and pressure response theory, combined with existing research results and practical experience, ensure the scientific and rational nature of risk factors and their interrelationships.

3) Principle of Operability

Risk factors should be measurable and able to obtain relevant information through multimodal text data or other channels, avoiding abstract and difficult to quantify factors.

4) Principle of Relevance

Highlight the correlation between risk factors, clarify the hierarchical, logical, and causal relationships between factors, and provide correlation support for risk identification [10-12].

5) Principle of Dynamism

Select risk factors that can reflect the dynamic changes in mental health status and adapt to the needs of dynamic risk identification.

D. KNOWLEDGE GRAPH ARCHITECTURE DESIGN

The knowledge graph of mental health risks among college students adopts a three-layer structure of “entity relationship attribute”, as follows:

The entity layer covers the core factors of psychological health risks for college students, which are divided into 5 categories and 89 specific entities. Although the ontology is streamlined, each entity represents a high-level condensed feature node derived from expert consensus, specifically designed to reduce noise in high-dimensional text data. This highly-distilled sparse structure provides a ‘schematic anchor’ for the BiLSTM layer, preventing model drift and explaining the 6.5% performance gain by focusing on critical risk nodes with the highest cross-modal correlation coefficients rather than redundant, low-impact variables:

Table 1 Entity Categories Level 2 Entity Categories
Specific entities Individual psychological entity psychological traits Neuroticism, extroversion, openness, agreeableness, conscientiousness, impulsivity, perfectionism Emotional state Depression, anxiety, stress, positive emotions, emotional stability Cognitive function Attention, memory, problem-solving ability, cognitive bias, self-awareness Academic development Entity academic pressure Course difficulty, exam pressure, GPA pressure, employment pressure, research pressure Academic performance GPA fluctuations, failing grades, absenteeism frequency, homework comple-

tion quality, learning efficiency Learning behavior duration, learning focus, self-learning ability, reasonable learning methods Social support Entity interpersonal interaction Social frequency, social scope, intimate relationship quality, interpersonal conflict frequency Support Source Family support Peer support, teacher support, psychological counseling support, social adaptation, social anxiety, loneliness, sense of belonging, group integration, life behavior, physical sleep schedule, sleep patterns, insomnia frequency, sleep quality, diet, exercise habits, exercise frequency, exercise duration, healthy eating habits, internet usage duration, internet purposes, internet addiction tendencies, social media dependence, consumption structure, consumption frequency, perceived economic pressure, environment adaptation, physical campus environment, campus life satisfaction, dormitory atmosphere, teaching facility satisfaction, life changes, major life events (heartbreak, family changes, academic setbacks), environmental changes, external pressure, economic pressure, competitive pressure, social expectation pressure [13].

Process and Implementation of Knowledge Graph Construction

1) Construction Process

Knowledge extraction: Using a hybrid method of “expert annotation+machine extraction”, risk entities, relationships, and attributes are extracted from mental health literature, clinical cases, and policy documents;

Knowledge Fusion: De duplication, conflict detection, and fusion of extracted knowledge to ensure consistency and accuracy;

Knowledge Storage: The use of a three-layer structure and Neo4j ensures extensibility for future crossinstitutional data integration while providing the structural rigor necessary for identifying hidden causal paths between sparse features. Neo4j has efficient graph query and association mining capabilities, which can support association analysis in risk identification;

Knowledge reasoning: Based on rule-based reasoning and machine learning reasoning, mining hidden risk associations and enriching the content of the knowledge graph;

Visual display: D3.js is used to visualize the knowledge graph, supporting functions such as risk factor query, relationship tracing, and path analysis [14-16].

Scale and Quality Assessment of Knowledge Graph

Scale: The constructed knowledge graph of college students' mental health risks includes 89 entities, 23 relationships, and 326 attributes, forming 1286 “entity relationship attribute” triplets, covering the core factors and associated relationships of college students' mental health risks [17-19].

III. MULTIMODAL TEXT DATA COLLECTION AND PREPROCESSING

A. MULTIMODAL TEXT DATA COLLECTION

1) Data Source

Collect multimodal text data from college students, covering both text and speech modalities. The specific sources are as follows:

Social media text: Text content posted by college students on campus social platforms, class group chats, forums, etc., including dynamic updates, comment replies, help seeking information, etc., collected a total of 1.2 million pieces, with a total word count of about 8 million words;

Course review text: A total of 350,000 comments were collected. Crucially, the analysis does not target the literal assessment of teaching quality, but rather the latent linguistic markers—such as cognitive distortion patterns, shifts in linguistic complexity, and social withdrawal signals—often embedded in academic feedback during periods of high psychological stress. By extracting emotional valence and self-referential frequency from these texts, the model identifies secondary indicators of academic burnout and environmental maladaptation, which serve as early-warning precursors to clinical depression;

Psychological counseling records: counseling text records from university mental health centers (desensitized), collected a total of 12000 pieces, with a total word count of approximately 600000 words;

Voice log: Daily voice clips voluntarily recorded by college students (such as emotional diaries and learning summaries), with a total of 85000 collected and a total duration of approximately 425 hours [20].

2) Data Collection Principles

Privacy protection principle: All data collection must obtain informed consent from students, and sensitive information must be anonymized (such as hiding names, student IDs, and contact information) to ensure compliance and legality of data collection;

Voluntary participation principle: Sensitive data such as voice logs are recorded voluntarily and students are not required to participate;

Principle of Proportional Representativeness: While participation is strictly voluntary, the study utilizes a weighted sampling strategy covering students of different grades, majors, and genders. We acknowledge the inherent participation bias where high-risk individuals may opt-out; therefore, the ‘comprehensiveness’ refers to the diversity of demographic backgrounds rather than total population saturation, and the dataset is balanced using synthetic oversampling techniques to mitigate the statistical silence of high-risk non-volunteers;

Dynamic collection principle: The data collection cycle is 12 months, updated once a month, to capture the dynamic changes in students' mental health status.

B. PREPROCESSING OF MULTIMODAL TEXT DATA

1) Preprocessing of Text Modal Data

Multi step preprocessing of textual data such as social media texts, course reviews, and psychological counseling records:

Text cleaning: Remove irrelevant information from the text (such as emoticons, special characters, advertising links, redundant content), and retain the core text;

Word segmentation processing: Using stuttering word segmentation tools to segment Chinese text, optimizing the segmentation results with professional dictionaries in the field of mental health to avoid splitting core risk vocabulary;

Stop word removal: Construct a stop word list in the field of mental health (such as meaningless words like “de”, “is”, “in”, etc.), remove stop words from the text, and reduce the computational complexity of the model;

Text standardization: standardize text capitalization, correct spelling errors, and normalize semantics (such as labeling “depressed” and “feeling down” as depression related expressions);

Emotional annotation: Using the method of “machine pre annotation+manual correction”, the text is annotated with emotions, including positive, neutral, negative, depressive tendency, anxiety tendency, and stress tendency. The annotation consistency Kappa coefficient reaches 0.87, ensuring the quality of annotation.

2) Preprocessing of Speech Modal Data

Preprocess the voice log data and extract voice features:

Voice cleaning: Remove background noise and silent segments from the voice, while preserving effective voice signals;

Feature extraction: The Librosa tool is used to extract acoustic features of speech, including fundamental frequency (F0), spectral centroid, spectral bandwidth, Mel frequency cepstral coefficients (MFCCs), speech rate, energy, and other 38 dimensional features;

Emotional feature annotation: Invite 3 speech emotion recognition experts to annotate the speech segments with emotional tags including calmness, joy, sadness, anxiety, and anger, with a consistent Kappa coefficient of 0.85;

Feature standardization: Z-score standardization is applied to the extracted speech features to unify the data scale and provide support for model training.

C. CONSTRUCTION OF MULTIMODAL TEXT DATASET

1) Dataset Partitioning

The preprocessed multimodal text data is divided into training set, validation set, and testing set in a ratio of 7:2:1. The dataset size is shown in the following table 1:

Table dataset text data volume (items) speech data volume (items) covering risk types training set 1.197 million 59500 Depression, anxiety, stress overload, sleep disorders, interpersonal sensitivity validation set 34217000 Same as test set 17108500 Same as total 1.7185 million Same as above

2) Dataset Quality Assessment

Evaluate the quality of the dataset from three dimensions: data integrity, annotation accuracy, and representativeness:

Data integrity: The missing rate of text data is less than 0.3%, and the missing rate of voice data is less than 0.5%, indicating good data integrity;

Annotation accuracy: Randomly select 10000 text data and 500 speech data for manual review, with a annotation accuracy rate of 92.3%, meeting the training requirements of the model;

Representativeness: The dataset covers college students of different grades, majors, genders, and origins, with a uniform distribution of risk types and good representativeness.

IV. CONSTRUCTION OF FEATURE FUSION AND RISK IDENTIFICATION MODEL

A. MULTI DIMENSIONAL FEATURE SYSTEM CONSTRUCTION

Integrating knowledge graph related features with multimodal text features, a multidimensional feature system is constructed, including three categories: textual features, speech features, and knowledge graph related features, totaling 126 dimensional features.

Textual Features (68 Dimensions)

Semantic features: The BERT model is used to extract deep semantic features of the text, output 768 dimensional vectors, and reduce them to 30 dimensions through principal component analysis (PCA);

Emotional characteristics: including text emotional polarity (positive/neutral/negative), emotional intensity, probability of depression tendency, probability of anxiety tendency, and probability of stress tendency, totaling 5 dimensions;

KEYWORDS features: Extract the word frequency and TF-IDF weight of mental health risk related keywords (such as “depression”, “anxiety”, “insomnia”, “high stress”), totaling 20 dimensions;

Text structural features: including text length, number of sentences, average sentence length, punctuation frequency, etc., totaling 13 dimensions.

Speech Features (38 Dimensions)

The speech acoustic features extracted using the librosa tool include:

Fundamental frequency characteristics: fundamental frequency mean, fundamental frequency standard deviation, fundamental frequency maximum value, fundamental frequency minimum value, a total of 4 dimensions;

Spectral characteristics: Spectral centroid, spectral bandwidth, spectral roll off, and spectral flux, totaling 4 dimensions;

Mel frequency cepstral coefficients (MFCCs): 13 dimensional MFCC coefficients and their first-order and second-order differences, totaling 39 dimensions (reduced to 20 dimensions by PCA);

Energy characteristics: energy mean, energy standard deviation, energy peak, a total of 3 dimensions;

Speech speed characteristics: speech speed (words/second), pause frequency, average pause duration, a total of 3 dimensions;

TABLE 1. Performance Comparison of Different Models

Model	Accuracy (%)	Sensitivity (%)	Specificity (%)	False Positive Rate (%)	F1-Score (%)
Logistic Regression	72.5	68.3	75.2	24.8	70.8
SVM	76.8	73.5	79.1	20.9	75.2
Random Forest	80.3	77.6	82.1	17.9	78.9
BiLSTM (Single Text Modal)	83.6	81.2	85.3	14.7	82.5
BERT (Single Text Modal)	85.7	83.5	87.2	12.8	84.6
Single Speech Modal Model	81.5	79.3	83.1	16.9	80.4
Proposed Multimodal Fusion	90.2	88.7	91.5	8.3	89.5

Emotional features: Probability of speech emotion classification (calm/joy/sadness/anxiety/anger), a total of 5 dimensions.

Knowledge Graph Association Features (20 Dimensions)

Based on the knowledge graph of mental health risks, extract the following associated features:

Risk factor correlation strength: The correlation strength between risk factors appearing in the text and core risk entities (such as depression and anxiety), with a total of 8 dimensions;

Risk path length: The length of the associated path composed of risk factors in the text (the shorter the path, the higher the risk), totaling 3 dimensions;

Risk factor clustering degree: The clustering degree of risk factors in the text (the higher the clustering degree, the higher the risk), a total of 3 dimensions;

Coverage of protective factors: The degree to which protective factors (such as social support and positive emotions) appear in the text cover risks, with a total of 3 dimensions;

Risk Evolution Trend: Based on time-series text data, predict the evolution trend (upward/stable/downward) of risk factors, with a total of 3 dimensions.

B. CONSTRUCTION OF RISK IDENTIFICATION MODEL

1) Model Architecture Design

Based on BERT+BiLSTM, a multimodal fusion model for identifying mental health risks in college students is constructed. The model architecture includes four parts: input layer, feature fusion layer, encoding layer, and classification layer. The overall architecture is shown in Figure 1 (a standard architecture diagram needs to be inserted here in the EI paper, which is now described in text):

Input layer: Input a multidimensional feature system, including text features, speech features, and knowledge graph association features;

Feature fusion layer: using a Knowledge-Graph-Guided Attention (KGGA) mechanism to weight and fuse multimodal features. Unlike standard attention, the KGGA layer incorporates the directed causal paths defined in the KG as prior constraints, dynamically adjusting the weights of different modal features. This ensures that the deep learning output is not merely associative but is weighted by the hierarchical importance and causal influence of risk factors (e.g., environmental adaptation acting as a causal trigger for emotional traits), highlighting features with high structural contribution to risk identification;

Encoding layer: Using BERT model to encode text semantic features and capture deep semantic information of the text; Using BiLSTM model to encode temporal features and capture dynamic changes in mental health status;

Classification layer: Using a fully connected layer and Softmax activation function, output risk classification results (low risk, medium risk, high risk).

2) Model Parameter Optimization

Optimize key parameters of the model through grid search and cross validation:

BERT model: using pre trained BERT based Chinese model, with 768 hidden layer dimensions, 12 attention heads, and 12 layers;

BiLSTM model: 256 hidden layer neurons, 2 layers, dropout probability of 0.5;

Attention mechanism: 8 heads of attention, learning rate 0.001;

Batch size: 32;

Iteration times: 50;

Loss function: Cross entropy loss function;

Optimizer: Adam optimizer, with a learning rate decay coefficient of 0.99.

3) Model Training Process

Data preprocessing: Standardize multimodal features and divide them into training, validation, and testing sets;

Pre training: Using a pre trained BERT model to initialize the text feature encoding layer, improving the convergence speed of the model;

Fine tuning training: Train the model on the training set, monitor the model performance through the validation set, and use early stopping method to prevent overfitting;

Model evaluation: Evaluate model performance on the test set, adjust model parameters and fusion weights to ensure model generalization ability.

C. MODEL COMPARISON EXPERIMENT AND RESULT ANALYSIS

1) Experimental Environment and Evaluation Indicators

Experimental environment: The hardware environment consists of an Intel Core i9-13900K processor, 64GB of memory, and an RTX4090 graphics card; The software environment includes Python 3.9, TensorFlow 2.10, PyTorch 1.13, and HuggingFace Transformers;

Evaluation indicators: Precision (P), sensitivity (Recall, R), specificity (S), false positive rate (FPR), and F1 score (F1 Score) are used as performance evaluation indicators for the model. The calculation formula is as follows:

$$\text{Accuracy } P = \frac{TP + TN}{TP + TN + FP + FN} \times 100\%,$$

$$\text{Sensitivity } R = \frac{TP}{TP + FN} \times 100\%,$$

$$\text{Specificity } S = \frac{TN}{TN + FP} \times 100\%,$$

$$\text{False alarm rate } FPR = \frac{FP}{TN + FP} \times 100\%,$$

$$F1 = \frac{2 \times P \times R}{P + R} \times 100\%.$$

where TP is true positive (correctly identified as high-risk), TN is true negative (correctly identified as low-risk), FP is false positive (misjudged as high-risk), and FN is false negative (misjudged as low-risk).

2) Model Comparison Experimental Results

The multi-mode fusion model (BERT+BiLSTM+attention fusion) proposed in this paper is compared with the single mode model and the traditional machine learning model. The experimental results are shown in the following table2

Result analysis:

The multimodal fusion model proposed in this paper is significantly superior to the single modal model and the traditional machine learning model in all evaluation indicators, with an accuracy rate of 90.2% and an F1 value of 89.5%, which indicates that the introduction of multimodal feature fusion and knowledge map associated features effectively improves the accuracy of risk identification;

The performance of the BERT model with a single text modality is superior to that of the BiLSTM model, indicating that the BERT model can better capture deep semantic features of text;

The performance of a single speech modality model is slightly lower than that of a single text modality model, but it still has high recognition ability, indicating that speech features contain rich psychological health status information;

Traditional machine learning models (logistic regression, SVM, random forest) have relatively low performance, mainly due to the difficulty in capturing complex associations and deep semantic information between multimodal features.

3) Identification Results of Different Risk Types

The recognition performance of the model in this article for different types of mental health risks is shown in the following table 3.

Result analysis:

The model has the highest accuracy in identifying depression and anxiety (both exceeding 91%), and these risks exhibit significant features in multimodal text (such as frequent occurrence of depression related vocabulary, slow speech rate, and low energy), making them easy to identify;

The recognition accuracy of interpersonal sensitivity is relatively low (87.4%), mainly due to the hidden manifestations of such risks and the relatively weak correlation between multimodal features and risks; Overall, the model's performance in identifying various major mental health risks is at a high level, which can meet the core needs of identifying mental health risks in universities.

4) Contribution Analysis of Knowledge Graph Features

To verify the role of knowledge graph correlation features, ablation experiments were conducted to compare the performance changes of the model before and after adding knowledge graph features. The results are shown in the following table 4.

Result analysis:

After adding knowledge graph correlation features, the performance indicators of the model were significantly improved, with an accuracy increase of 6.5% and an F1 score increase of 6.1%. This indicates that the knowledge graph can effectively explore the correlation between risk factors and provide additional support for risk identification;

The contribution of knowledge graph is mainly reflected in the identification of hidden risks. Through association mining, potential risks that are difficult to capture with a single feature can be discovered.

V. EMPIRICAL RESEARCH DESIGN AND RESULT ANALYSIS

A. EMPIRICAL RESEARCH DESIGN

1) Research Object

Selecting undergraduate students from four different types of universities (comprehensive universities, science and engineering universities, liberal arts universities, and art universities) as research subjects, using stratified sampling method, a total of 3200 students were selected, including 1684 male students and 1516 female students; 823 freshmen, 956 sophomores, 812 juniors, and 609 seniors; There are 1728 students majoring in science and engineering, 1024 students majoring in humanities, and 448 students majoring in arts. All research subjects signed informed consent forms and voluntarily participated in this study.

2) Experimental Design

The experimental period is 12 months, and the research subjects will be randomly divided into an experimental group and a control group, with 1600 people in each group:

Experimental group: Adopting the multimodal fusion risk identification model proposed in this article for risk identification and early warning;

TABLE 2. Recognition Performance for Different Risk Types

Risk Type	Accuracy (%)	Sensitivity (%)	Specificity (%)	F1-Score (%)
Depression	92.6	91.3	93.5	92.4
Anxiety	91.5	89.8	92.7	90.7
Stress Overload	89.7	87.6	91.2	89.4
Sleep Disorder	88.3	86.2	89.8	87.5
Interpersonal Sensitivity	87.4	85.1	89.2	86.3

TABLE 3. Contribution of Knowledge Graph Features (Ablation Experiment)

Model Configuration	Accuracy (%)	Sensitivity (%)	Specificity (%)	F1-Score (%)	Performance Improvement
Without Knowledge Graph Features	84.7	82.3	86.5	83.4	-
With Knowledge Graph Features	90.2	88.7	91.5	89.5	6.5%

TABLE 4. Comparison of Recognition Efficacy Between Experimental and Control Groups

Evaluation Indicator	Experimental Group (%)	Control Group (%)	Improvement Rate
Recognition Accuracy	90.2	72.5	24.4%
Sensitivity	88.7	68.3	29.9%
Specificity	91.5	75.2	21.7%
False Positive Rate	8.3	24.8	-66.5%
False Negative Rate	11.3	31.7	-64.4%

Control group: Traditional mental health risk identification methods were used, specifically utilizing internationally recognized standardized scales including the Patient Health Questionnaire-9 (PHQ-9) for depression and the Generalized Anxiety Disorder-7 (GAD-7) for anxiety, supplemented by bi-annual manual screenings. These static, self-report-dependent tools serve as the baseline to evaluate the under-reporting rate in traditional campus counseling settings.

During the experimental process, variables such as teaching content, teaching duration, and intervention resources are controlled to ensure the fairness and effectiveness of the experiment.

3) Data Collection and Evaluation Indicators

Data collection: Relevant data will be collected before the start of the experiment, during the middle of the experiment (at the 6th month), and after the end of the experiment (at the 12th month), including students' mental health status assessment results, risk identification records, intervention effect feedback, etc; Evaluation indicators:

Recognition efficiency: including recognition accuracy, sensitivity, specificity, false positive rate, and false negative rate;

Timeliness of early warning: the time window for identifying risks in advance (such as 1 month or 2 months in advance);

Intervention effect: Improvement rate of psychological state and incidence of extreme events among high-risk students;

User satisfaction: The satisfaction of the experimental group students and teachers with the risk identification system.

1) Comparison of Recognition Efficiency

The comparison results of risk identification efficiency between the experimental group and the control group are shown in the following table5.

Table evaluation indicators: Experimental group, control group, improvement in recognition accuracy (%) 90.272.524.4% Sensitivity (%) 88.768.329.9% Specificity (%) 91.575.221.7% False alarm rate (%) 8.324.8-66.5% Missed alarm rate (%) 11.331.7-64.4% Result analysis:

The risk identification efficiency of the experimental group was significantly better than that of the control group, with a 24.4% increase in identification accuracy, a 29.9% increase in sensitivity, and a 66.5% and 64.4% decrease in false positive and false negative rates, respectively. This indicates that the model proposed in this paper can effectively improve the accuracy of risk identification and reduce misjudgments and missed detections;

The underreporting rate of the control group was as high as 31.7%, mainly due to the traditional method relying on students actively seeking help and regular surveys, which makes it difficult to capture hidden risks; The false alarm rate of the experimental group was only 8.3%, indicating that the model has good specificity and can effectively distinguish between real risks and temporary emotional fluctuations.

2) Comparison of Warning Timeliness

The comparison results of the warning timeliness between the experimental group and the control group are shown in the following table 6.

Table Warning Time Window Experimental Group Recognition Success Rate (%) Control Group Recognition Success Rate (%) Difference (%) Early 2 Months

B. EMPIRICAL RESULTS AND ANALYSIS

TABLE 5. Comparison of Early Warning Timeliness

Early Warning Time Window	Recognition Success Rate of Experimental Group (%)	Recognition Success Rate of Control Group (%)	Difference (%)
2 Months in Advance	78.5	23.6	54.9
1 Month in Advance	90.3	45.2	45.1
Real-Time (No Advance Warning)	95.7	72.5	23.2

TABLE 6. Comparison of Intervention Effects

Evaluation Indicator	Experimental Group (%)	Control Group (%)	Difference (%)
Improvement Rate of Mental State (High-Risk Students)	85.6	52.3	33.3
Extreme Event Incidence Rate	0.2	1.2	-83.3
Improvement Rate of Mental Health Literacy	21.5	10.7	10.8

78.523.654.9 Early 1 Month 90.345.245.1 No Early Warning (Real Time) 95.772.523.2 Result analysis:

The success rate of risk identification in the experimental group was significantly higher than that in the control group at 1 and 2 months in advance, with a success rate of 78.5% at 2 months in advance, an increase of 54.9% compared to the control group. This indicates that the model in this article has good prospective warning ability and can strive for sufficient time window for intervention work;

The control group performed poorly in early warning, with a recognition success rate of only 23.6% two months in advance. The main reason is that traditional methods lack dynamic monitoring capabilities and are difficult to capture risk trends in advance;

The real-time recognition success rate of the experimental group reached 95.7%, further verifying the high accuracy of the model.

C. EMPIRICAL CONCLUSION

The empirical research results indicate that:

The mental health risk identification model for college students based on knowledge graph and multimodal text analysis proposed in this article has high recognition accuracy (90.2%), sensitivity (88.7%), and specificity (91.5%), significantly better than traditional identification methods;

The model has good forward-looking warning ability, which can identify mental health risks 1-2 months in advance and buy a time window for intervention work;

The application of the model can significantly improve the intervention effect, with an improvement rate of 85.6% in the psychological state of high-risk students and a reduction of 83.3% in the incidence of extreme events;

Students and teachers have high satisfaction with the risk identification system, effective privacy protection measures, and widely recognized for its ease of operation and identification efficiency;

The model demonstrates robust generalization across diverse academic backgrounds and has broad promotion value for ensuring the mental well-being of the future workforce in high-precision manufacturing. This adaptability facilitates the implementation of standardized psychological monitoring systems across various sectors of the industry to enhance long-term industrial productivity and safety.

AUTHOR CONTRIBUTIONS

Chen Dongli designed, collected and analyzed the data, and drafted the manuscript. Chen Dongli conducted the study, critically revised the manuscript for important intellectual content, and gave final approval of the version to be published. Chen Dongli participated fully in the work, take public responsibility for appropriate portions of the content, and agreed to be accountable for all aspects of the work in ensuring that questions related to the accuracy or integrity of any part of the work are appropriately investigated and resolved.

CONFLICTS OF INTEREST

The author declares no conflict of interest.

FUNDING

This research received no external funding.

ACKNOWLEDGEMENTS

Not applicable.

REFERENCES

- [1] F. Ciroku, J. Berardinis D, J. Kim, et al., "RevOnt: Reverse engineering of competency questions from knowledge graphs via language models," *Journal of Web Semantics*, vol. 82, pp. 100822–100822, 2024, doi: 10.1016/J.WEBSEM.2024.100822.
- [2] J. Liang, S. Zhang, Y. Zhang, et al., "A knowledge graph-based approach to modeling & representation for machining process design intent," *Advanced Engineering Informatics*, vol. 62, no. PA, pp. 102645–102645, 2024, doi: 10.1016/J.AEI.2024.102645.
- [3] P. Zhang, D. Chen, Y. Fang, et al., "A knowledge graphs representation method based on IsA relation modeling," *Expert Systems With Applications*, vol. 254, pp. 124468–124468, 2024, doi: 10.1016/J.ESWA.2024.124468.
- [4] J. Duan, G. Wang, X. Hu, et al., "Concept cognition for knowledge graphs: Mining multi-granularity decision rule," *Cognitive Systems Research*, vol. 87, pp. 101258–101258, 2024, doi: 10.1016/J.COGRYS.2024.101258.
- [5] W. Xie, F. Farazi, J. Atherton, et al., "Dynamic knowledge graph approach for modelling the decarbonisation of power systems," *Energy and AI*, vol. 17, pp. 100359–100359, 2024, doi: 10.1016/J.EGYAI.2024.100359.
- [6] W. Zhenhai, X. Yuhao, W. Zhiru, et al., "Knowledge Graph-Aware Deep Interest Extraction Network on Sequential Recommendation," *Neural Processing Letters*, pp. 56(4), 2024, doi: 10.1007/S11063-024-11665-2.
- [7] G. Lino, G. Gema, R. Maria A M, et al., "Guest editorial: Recent trends in semantic web and knowledge graphs," *The Electronic Library*, vol. 42, no. 3, pp. 365–367, 2024, doi: 10.1108/EL-07-2024-351.
- [8] L. Zhenghao, Q. Yuxing, L. Wenlong, et al., "Multi-feature fusion stock prediction based on knowledge graph," *The Electronic Library*, vol. 42, no. 3, pp. 455–482, 2024, doi: 10.1108/EL-02-2023-0053.

- [9] R. Enayat, G. Niya A, and K. Karishma, "The role of knowledge graphs in chatbots," *The Electronic Library*, vol. 42, no. 3, pp. 483–497, 2024, doi: 10.1108/EL-03-2023-0066.
- [10] M. J. Platow, G. C. Lee, C. Wang, et al., "The role of student and customer social identification on university students' learning approaches and psychological well-being," *Social Psychology of Education*, vol. 8, no. 1, pp. 1–30, 2025, doi: 10.1007/s11218-025-10060-6.
- [11] Z. Wang, F. Wang, B. Ma, and et al. The effect of physical activity and life events on mental health of college students: the mediating role of psychological vulnerability, "BMC Psychology," 2025; 13(1):1-14, doi: 10.1186/s40359-025-02539-w.
- [12] X. Su, B. Zhao, G. Li, et al., "Knowledge Graph Neural Network with Spatial-Aware Capsule for Drug-Drug Interaction Prediction," *IEEE journal of biomedical and health informatics*, to be published, 2024, doi: 10.1109/JBHI.2024.3419015.
- [13] D. Yiyang, S. Sicun, F. Junxuan, et al., "Paleontology Knowledge Graph for Data-Driven Discovery," *Journal of Earth Science*, vol. 35, no. 3, pp. 1024–1034, 2024, doi: 10.1007/S12583-023-1943-9.
- [14] J. Chen, X. Zhang, L. Xu, et al., "Trends of digitalization, intelligence and greening of global shipping industry based on CiteSpace Knowledge Graph," *Ocean and Coastal Management*, vol. 255, Art. no. 107206, 2024, doi: 10.1016/J.OCECOAMAN.2024.107206.
- [15] A. Ruschel, A. Gusmão C, and F. Cozman G, "Explaining answers generated by knowledge graph embeddings," *International Journal of Approximate Reasoning*, vol. 171, Art. no. 109183, 2024, doi: 10.1016/J.IJAR.2024.109183.
- [16] J. Yang, X. Ying, Y. Shi, et al., "Improving static and temporal knowledge graph embedding using affine transformations of entities," *Journal of Web Semantics*, vol. 82, Art. no. 100824, 2024, doi: 10.1016/J.WEBSEM.2024.100824.
- [17] Y. Peng, C. Yong A P, and N. Myeda E, "Knowledge graph of building information modelling (BIM) for facilities management (FM)," *Automation in Construction*, vol. 165, Art. no. 105492, 2024, doi: 10.1016/J.AUTCON.2024.105492.
- [18] W. Li, H. Zhong, J. Zhou, et al., "An attention mechanism and residual network based knowledge graph-enhanced recommender system," *Knowledge-Based Systems*, vol. 299, Art. no. 112042, 2024, doi: 10.1016/J.KNOSYS.2024.112042.
- [19] Z. Wei, K. Wang, F. Li, et al., "M3KGR: A momentum contrastive multi-modal knowledge graph learning framework for recommendation," *Information Sciences*, vol. 676, Art. no. 120812, 2024, doi: 10.1016/J.INS.2024.120812.
- [20] J. Zhu, X. Cai, E. Hussam, et al., "A novel cosine-derived probability distribution: Theory and data modeling with computer knowledge graph," *Alexandria Engineering Journal*, vol. 103, pp. 1–11, 2024, doi: 10.1016/J.AEJ.2024.05.114.