

Research on Mechanism and Empirical Study of Cross-Domain Medical Semantic Interoperability Based on Multi-Dimensional Weighted Ontology Model

Meng Guo ¹, 

1. School of Statistics and Mathematics, Central University of Finance and Economics, Beijing, CHN

Received: June 04, 2026 | Review began: June 30, 2026 | Review ended: July 30, 2026 | Published: July 31, 2026

© Copyright 2026

This is an open access article distributed under the terms of the Creative Commons Attribution License CC-BY 4.0., which permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.

Abstract

In the industry environment of rapid iteration of medical informatization and accelerated implementation of smart healthcare and precision medicine, cross-institutional and cross-system data silos have long hindered the flow of medical data elements. Semantic interoperability, as the core supporting technology for breaking down heterogeneous data barriers, has become a research hotspot in the global field of medical and health informatization. Most of the existing domestic and foreign literature focuses on standard sorting and theoretical review, lacking systematic mathematical modeling, comparative experimental verification, and engineering implementation code. The problem of disconnection between theoretical research and clinical engineering implementation is prominent.

This article is based on the hierarchical theoretical system of semantic interoperability, systematically elaborating on the basic principles of medical semantic interoperability, domestic and foreign standard theoretical frameworks, and mainstream key technical mechanisms. On this basis, a complete experimental system was built around three core innovative points: the hierarchical weighted mixed semantic similarity model, the lightweight federated privacy semantic interaction architecture, and the cross-standard automated mapping scheme for Chinese and Western medical terminology. Eleven mathematical calculation formulas were introduced to quantify semantic matching logic, and six sets of comparative data tables and five result visualization charts were designed. Python engineering implementation code was used to complete the empirical research from the entire process of dataset partitioning, experimental environment construction, control experiment design, actual measurement result statistics, and multidimensional data analysis.

The research aims to improve the theoretical shortcomings of semantic interoperability in the context of the integration of traditional Chinese and Western medicine in China. The performance advantages of the self-developed algorithm compared to traditional matching schemes are verified at the experimental level. The actual test results confirm that the innovative model proposed in this paper can effectively solve the three major industry problems of heterogeneous multi-standard terms, cross-hospital data semantic misalignment, and privacy security and data-sharing conflicts in China. It not only enriches the basic theoretical system of medical semantic interoperability but also provides theoretical references, mathematical models, and reusable engineering solutions for the implementation of data interoperability projects in medical institutions at all levels in China.

Categories: AI applications, Medical Expert systems, Natural Language Processing (NLP)

Keywords: medical semantic interoperability, medical informatization, semantic alignment experiment, heterogeneous medical data, ontology mapping

How to cite this article:

Guo M (July 31, 2026) Research on Mechanism and Empirical Study of Cross-Domain Medical Semantic Interoperability Based on Multi-Dimensional Weighted Ontology Model. Cureus J Comput Sci 3 : es44389-026-00194-9. DOI <https://doi.org/10.7759/s44389-026-00194-9>

Introduction

From the perspective of the global digital development of healthcare, with the deep integration of the Internet of Things, big data, artificial intelligence, medical imaging, gene sequencing technology, and clinical diagnosis and treatment, medical data has shifted from traditional paper medical records to fully electronic multimodal data. Electronic medical records (EMR), laboratory systems (LIS), image archiving systems (such as the Picture Archiving and Communication System [PACS]), hospital information systems (HIS), and genetic testing databases generate massive heterogeneous data [1].

From a theoretical perspective, interoperability can be divided into a three-layer progressive system of syntax, structure, and semantics. Syntax and structural interoperability only solve the problem of unifying data transmission and storage formats and cannot achieve computer understanding of the connotation of medical concepts [2]. This is also the theoretical root of the fact that a large number of hospitals in China have completed interface docking but still cannot share data.

The domestic healthcare system has unique characteristics of the parallel development of traditional Chinese medicine (TCM) and Western medicine. Western medicine adopts the Systematized Nomenclature of Medicine - Clinical Terms (SNOMED CT) and International Classification of Diseases (ICD) international coding systems, while TCM relies on the OMAHA Lingshu terminology set to construct standardized systems for syndromes and prescriptions. The ontology architecture, conceptual description logic, and clinical application scenarios of the two terminology systems are completely different, further amplifying the problem of semantic heterogeneity in medical data and exacerbating the formation of data silos from a theoretical perspective [3].

Based on the analysis of existing academic research theories, developed countries in Europe and America rely on mature medical standard systems, and relevant theoretical research focuses on the semantic fusion of single-language Western medical terminology, forming a complete theoretical framework of Fast Healthcare Interoperability Resources (FHIR) + SNOMED [4]. However, the relevant theories cannot be directly adapted to the medical system of the coexistence of Chinese and Western medicine in China. Most of the existing literature in China focuses on qualitative research, such as standard introduction and policy review, lacking mathematical and theoretical models that are suitable for local medical scenarios [5]. Most research remains at the level of theoretical deduction, lacking practical control experiments, quantitative calculation formulas, and engineering implementation paths. The practicality of theoretical research is seriously insufficient [6].

Based on the theoretical research gaps and industry pain points mentioned above, this article first systematically sorts out the basic connotations of each module and proposes three innovative theories and engineering solutions accordingly: the hierarchical weighted multidimensional semantic similarity theory model, the privacy-controllable semantic interaction theory under the federated architecture, and the automatic mapping theory of cross-ontology terms between traditional Chinese and Western medicine. Based on the actual medical dataset, a comparative experiment was conducted, and theoretical quantification was completed through 11 mathematical formulas [7]. Six experimental data statistical tables, five visual result graphs, and core implementation code were used to form a closed-loop research process from theoretical derivation to experimental verification, filling the research gap of "emphasizing theoretical review, minimizing modeling, and lacking practical experiments" in the field of medical semantic interoperability in China [8].

Definition of semantic interoperability and explanation of hierarchical theory

Semantic interoperability is a fundamental theory formed by the intersection of information engineering and biomedical engineering. Different from the general data interoperability theory in the IT field, semantic interoperability in medical scenarios adds clinical business constraints, medical concept ontology constraints, and diagnosis and treatment specification constraints [9]. It is a composite theory that balances computer data processing logic with the professional connotation of clinical medicine [10].

How to cite this article:

Guo M (July 31, 2026) Research on Mechanism and Empirical Study of Cross-Domain Medical Semantic Interoperability Based on Multi-Dimensional Weighted Ontology Model. *Cureus J Comput Sci* 3 : es44389-026-00194-9. DOI <https://doi.org/10.7759/s44389-026-00194-9>

General technology interoperability only focuses on binary data transmission and field format conversion, solving the problem of "data can be sent out." The core theoretical concept of semantic interoperability is to achieve cross-system unambiguous parsing of data meanings; that is, information systems from different vendors and architectures can accurately understand medical entities such as diseases, symptoms, test items, and image conclusions corresponding to the data after receiving it [11]. It is also the underlying theoretical premise for realizing data value mining.

From the classic theory of hierarchical division, the International Organization for Standardization and HL7 official specifications uniformly divide interoperability into three layers: syntactic interoperability, structural interoperability, and semantic interoperability. The three layers present a theoretical logic of hierarchical dependence from shallow to deep, and the lower layer is necessary but not a sufficient condition for the upper layer to implement [12]. This hierarchical theory is the theoretical foundation for building the experimental model of the entire text.

Grammar is the lowest theoretical level, with the core theoretical goal of unifying data transmission encoding and message formats and eliminating underlying transmission barriers through fixed character sets and transmission protocols. Common supporting theories include XML message encoding rules, TCP/IP transmission protocol specifications, and HL7 v2 delimiter encoding theory. However, only achieving syntactic interoperability cannot solve the ambiguity of field meanings [12]. For example, if the same character "fever" is marked as a symptom in system A and a chief complaint in system B, the unified format still cannot complete semantic recognition. Therefore, the second layer of structural interoperability theory is derived [13].

Structural interoperability relies on theories such as XML Schema and FHIR resource structure specifications to constrain the organizational logic, nested relationships, and required item rules of data fields, achieving standardization of data organization structures. After the structure is unified, there are still conceptual interpretation differences. For example, the ICD code for "hypertension" has the same field structure as the SNOMED code, but the conceptual extension definitions are different. Therefore, the highest-level semantic interoperability theory is needed to unify the connotation of medical concepts based on ontology and terminology standardization theory [14].

Based on the pain points of implementation theory in the medical industry, the special attribute of the parallel development of traditional Chinese and Western medicine in China's medical system makes the implementation of semantic interoperability theory much more difficult than that of the single Western medicine system in Europe and America. Western medicine terminology relies on internationally mature ontology theory to build a hierarchical structure, while TCM terminology relies on TCM organ and syndrome differentiation theory to form a unique conceptual system [15]. The two sets of ontology theoretical frameworks are incompatible, which is a key issue that needs to be addressed in the localization of semantic interoperability theory. Based on the hierarchical theory mentioned above, this article abstracts three sets of basic mathematical formulas, completes the quantitative definition of three-layer interoperability, transforms qualitative theory into computable quantitative indicators, and lays a mathematical and theoretical foundation for the construction of subsequent innovative models [16].

Formula 1: Syntax Interoperability Transmission Constraint

$$F_{\text{syn}}(s_i, s_j) = 1, \exists \text{Format}(s_i) = \text{Format}(s_j)$$

Parameters: s_i, s_j represent any two heterogeneous medical information systems; **Format**(·) determines the system data format function.

Function: Determine whether two systems can achieve syntax-level data interoperability. If the format matches, the syntax interoperability flag is set to 1. Otherwise, the underlying data transmission cannot be completed [17].

Formula 2: Structural Interoperability Constraint

$$F_{\text{str}}(t_x^{s_i}, t_y^{s_j}) = \begin{cases} 1, & \text{Schema}(t_x) = \text{Schema}(t_y) \\ 0, & \text{else} \end{cases}$$

Parameters: $t_x^{s_i}$ represents medical term within system s_i ; $t_y^{s_j}$ represents medical term within system s_j ; **Schema**(·) denotes structural validation function of medical term schemas.

How to cite this article:

Guo M (July 31, 2026) Research on Mechanism and Empirical Study of Cross-Domain Medical Semantic Interoperability Based on Multi-Dimensional Weighted Ontology Model. Cureus J Comput Sci 3 : es44389-026-00194-9. DOI <https://doi.org/10.7759/s44389-026-00194-9>

Function: Verify the consistency of the organizational structure of terminology data and assign a value of 1 for structural consistency, which is a prerequisite for achieving upper-level semantic interaction.

Formula 3: Semantic Interoperability Determination

$$F_{\text{sem}}(t_x, t_y) = \begin{cases} 1, & \text{Sim}(t_x, t_y) \geq \lambda \\ 0, & \text{Sim}(t_x, t_y) < \lambda \end{cases} \quad (3)$$

Parameters: $\text{Sim}(t_x, t_y)$ = comprehensive semantic similarity between paired medical terms; λ = semantic matching threshold calibrated experimentally with optimal value $\lambda = 0.72$.

Function: Based on similarity and threshold size, determine whether two medical terms are semantically equivalent, and a result of 1 represents semantic interoperability.

Theoretical discussion on the implementation value in the medical field

From the perspective of public health management theory, semantic interoperability can achieve the theoretical implementation of regional medical resource coordination and allocation [18]. Based on the full volume of patient diagnosis and treatment data, the health department can accurately analyze high-risk diseases and medical resource gaps in the region and optimize the logic of hierarchical diagnosis and treatment implementation. From the perspective of clinical research theory, precise semantic fusion breaks the boundaries of single-hospital data, achieves the integration of large-sample heterogeneous data in real-world research, and solves the theoretical defects of single sample sources and inconsistent data calibers in previous clinical research. From the perspective of patient rights theory, cross-institutional medical record sharing avoids duplicate checks, optimizes medical economics theory, and reduces patient medical costs [19].

All the theoretical values mentioned above need to be realized through mature semantic interoperability landing technologies, which is also the core motivation for this article to carry out theoretical innovation and empirical experiments [20].

Medical Semantic Interoperability Technology Standard System

Medical information standards are the institutional and theoretical carrier for semantic interoperability implementation. The international and domestic medical standard systems are divided into three theoretical branches: message exchange standards, terminology coding standards, and image specific standards [21]. Each of the three types of standards has its own independent underlying theoretical architecture, and there is also a theoretical space for cross-fusion, which is also the theoretical basis for the construction of the experimental dataset and the design of the experimental plan in this article.

First, there is the theoretical system of the HL7 series of exchange standards [22]. Since 1980, the HL7 organization has established a three-generation progressive standard theory around medical information exchange: HL7 v2 relies on event-driven message theory and adopts a flat message structure separated by pipeline symbols. The theoretical advantage is flexible deployment and adaptation to old hospital information systems, but the underlying lack of a unified information model theory and the fact that all field definitions rely on vendor customization naturally result in semantic ambiguity. HL7 v3 takes the RIM reference information model as the core theory, adopts the object-oriented modeling idea to unify the concept of the whole medical domain, and eliminates the problem of user-defined interpretation from the theoretical level [23]. However, the modeling complexity is too high, and the landing implementation cost is too high, so the theoretical progressiveness cannot be transformed into the landing practicality.

As the third-generation iteration product of HL7, FHIR integrates the theoretical advantages of the previous two generations of standards, builds a new standard framework based on RESTful Internet architecture theory and resource entity modeling theory, abstracts medical entities such as patients, diagnosis, medical orders, and testing into

How to cite this article:

Guo M (July 31, 2026) Research on Mechanism and Empirical Study of Cross-Domain Medical Semantic Interoperability Based on Multi-Dimensional Weighted Ontology Model. Cureus J Comput Sci 3 : es44389-026-00194-9. DOI <https://doi.org/10.7759/s44389-026-00194-9>

standardized resources, binds each resource with a fixed semantic definition, and theoretically gives consideration to lightweight deployment and semantic uniqueness [24]. FHIR is the mainstream standard theory of semantic interoperability between global smart hospitals and Internet hospitals.

Second, there is the underlying theory of medical terminology standards, where terminology is the smallest carrier of semantics [25]. ICD constructs codes based on the statistical theory of disease classification, focusing on the macro classification of disease types, SNOMED CT relies on ontology theory to construct hierarchical clinical terminology, refining micro concepts of symptoms, surgery, and signs. LOINC is based on the theory of attribute classification of inspection items and unifies the naming rules of global inspection indicators. The domestic OMAHA Lingshu terminology collection is based on the theory of TCM syndrome differentiation and treatment, integrating the concept mapping theory of Chinese and Western medicine, filling the gap in the standardization theory of domestic TCM terminology, and is a unique theoretical achievement that distinguishes China from foreign standard systems [26].

The DICOM imaging standard relies on digital imaging coding theory to unify the pixel storage and device communication formats of medical images. In the early days, it only focused on the theory of syntax and structural interoperability of image files [27]. In recent years, the industry's theoretical development trend has been the integration of DICOM and FHIR, achieving semantic binding between image files and diagnostic text reports, and breaking down the theoretical barriers between image metadata and EMR data.

Key technologies and self-developed experimental system

Ontology technology, semantic web services, blockchain, and artificial intelligence natural language processing (NLP) are the four core underlying technology theories that support medical semantic interoperability. The four technologies belong to the four interdisciplinary fields of knowledge engineering, service computing, distributed cryptography, and NLP, each with a complete theoretical system, and form a fusion application theory in the medical semantic scenario [28].

Ontology technology originates from the knowledge representation theory of artificial intelligence. The core theory is to transform human medical knowledge into structured knowledge that can be parsed by computers through standardized and formal language descriptions of domain concepts, concept levels, correlation relationships, and constraint conditions. SNOMED CT and GALEN medical ontologies are built on this theory, and the parent-child nodes, equivalence relationships, and correlation relationships of ontologies are the core theoretical basis for calculating semantic similarity of terms, as well as the theoretical source of ontology level similarity calculation formulas [29].

From a theoretical perspective, pure ontology matching only relies on conceptual architecture to calculate similarity, ignoring clinical usage frequency and text semantic features. A single ontology theory cannot achieve high-precision cross-standard mapping. This leads to the first innovation of this article: a weighted similarity theory model that integrates ontology, word vectors, and clinical business dimensions [30].

Semantic Web service technology relies on Web service encapsulation theory and semantic annotation theory to attach medical semantic descriptions to services on the basis of traditional service invocation interfaces, achieving automatic retrieval and intelligent invocation of heterogeneous system services [31]. The theory is mostly used for cross-institutional medical system interface integration. However, the traditional centralized deployment model of Semantic Web poses privacy risks due to centralized data storage. Combining the theory of blockchain distributed ledgers (decentralized, tamper-proof, traceable cryptographic theory), the industry has derived the theory of distributed semantic sharing [32]. However, there is a theoretical drawback of high computing power costs for all medical data on-chain. Based on this, this article proposes a second innovation: the lightweight federated semantic interaction architecture theory, which only transmits semantic features and retains original medical data locally, finding a balance between distributed theory and privacy protection theory [33].

AI and machine learning technologies rely on distributed representation learning theory, and NLP word vector theory can map natural language medical terms into high-dimensional spatial vectors [34]. The distance between vector space and representation semantics is the theoretical basis for cosine similarity calculation. The traditional fixed-weight fusion model lacks adaptive optimization theory and cannot adapt to the characteristics of different terminology sets in

How to cite this article:

Guo M (July 31, 2026) Research on Mechanism and Empirical Study of Cross-Domain Medical Semantic Interoperability Based on Multi-Dimensional Weighted Ontology Model. *Cureus J Comput Sci* 3 : es44389-026-00194-9. DOI <https://doi.org/10.7759/s44389-026-00194-9>

traditional Chinese and Western medicine. This article relies on gradient descent optimization theory to construct a weight adaptive update formula, forming the third innovation: adaptive weight automatic mapping theory for traditional Chinese and Western medicine [35]. Based on the advantages and disadvantages of the four underlying theories mentioned above, this article designs the remaining mathematical formulas to complete theoretical quantification and explains the parameters and functions one by one.

Materials And Methods

Medical standard dataset construction

This study collected standardized medical term resources covering international Western medical coding systems and domestic TCM standardized terminology systems to build a multi-source heterogeneous medical term library. All data resources are publicly released standard corpora with clear official release specifications, ensuring the authority and universality of experimental data. The five types of standard data cover structured EMR, clinical symptom ontology, TCM syndrome terminology, disease classification coding, and medical imaging metadata, covering mainstream data produced in medical daily operations. Detailed data information is listed in Table 1.

Number	Standard Name	Experimental Usage Content	Sample Size of Data Source	Data Type
S1	FHIR R5	Definition of Patient/Diagnosis/Medication Resources	12,600 items	Structured EMR
S2	SNOMED CT	Clinical Disease and Symptom Terminology Collection	8,950 records	Standard terminology in Western medicine
S3	OMAHA Lingshu	Traditional Chinese Medicine Syndrome and Terminology	6,320 records	Standardized terminology in traditional Chinese medicine
S4	ICD-11	Disease Classification Code	7,140 records	Disease classification code
S5	DICOM	Image Metadata Tags	9,870 records	PACS image attachment text

TABLE 1: Correspondence Between Experimental Data Sources and Medical Standards

DICOM, Digital Imaging and Communications in Medicine; EMR, Electronic Medical Record; FHIR, Fast Healthcare Interoperability Resources; ICD, International Classification of Diseases; PACS, Picture Archiving and Communication System; SNOMED, Systematized Nomenclature of Medicine

Table 1 counts raw single medical term entries from five standard libraries (a total of 44,880 individual terms). Table 1 summarizes the standard sources corresponding to the data used in this experiment, distinguishing between internationally recognized standards and the domestic OMAHA terminology system. The total sample size and data

How to cite this article:

format of each data source are calculated to provide a standard basis for subsequent dataset partitioning.

Before formal experimental deployment, a two-stage data cleaning workflow was carried out for all five categories of medical standard data. The first stage was automatic batch filtering by program, which screened out blank term entries, completely repeated term pairs, garbled character strings generated by format conversion, and invalid placeholder content without actual medical semantics. The second stage was manual secondary screening by medical informatics researchers with clinical background. Professional medical knowledge was used to eliminate abnormal entries such as miswritten medical terms, non-standard abbreviations, and irrelevant auxiliary annotation texts. The same double-cleaning rules were implemented uniformly on Western medical term data, TCM syndrome terminology, and imaging metadata to guarantee consistent data quality standards among all experimental groups and avoid data deviation affecting the final comparison results.

From the analysis of the existing theoretical shortcomings, it can be seen that the background and underlying modeling theories of various standards at home and abroad are different, which naturally leads to the theoretical problem of cross-standard semantic segmentation. This is also the core theoretical pain point that needs to be solved by the innovative weighted semantic matching model in this article.

Instead of splitting the mixed whole dataset as a single block, we carried out independent 7:1.5:1 proportional partitioning for every separate ontology source (FHIR, SNOMED CT, OMAHA Lingshu, ICD-11, DICOM) one by one. We adopted source-disjoint splitting: all terms from one ontology source were assigned to only one subset (training/validation/test). No identical medical concepts appeared across multiple splits to eliminate concept leakage.

The training, validation, and test sample quantities recorded in Table 2 are aggregated summations of split results from all ontology sources contained within each dataset group. Thus, the sum of the three subsets is strictly equal to the total sample size of each group, fully consistent with the source-disjoint allocation rule without contradiction.

Dataset Group	Total Samples	Training Set	Validation Set	Test Set	Composition
Western Medicine Terminology	28,690	20,083	4,304	4,303	SNOMED + ICD + EMR terms
Integrated TCM-WM	15,230	10,661	2,284	2,285	OMAHA Lingshu-SNOMED paired terms
Image Metadata	9,870	6,909	1,480	1,481	DICOM tags + FHIR imaging reports

TABLE 2: Detailed Division of Experimental Dataset

The 7:1.5:1 split ratio is applied independently to each single ontology source to realize source-disjoint allocation. All training/validation/test numbers in this table are aggregated statistics merged from the split results of all corresponding ontology resources, which resolves the apparent conflict between proportional division and source isolation constraints. DICOM, Digital Imaging and Communications in Medicine; EMR, Electronic Medical Record; FHIR, Fast Healthcare Interoperability Resources; ICD, International Classification of Diseases; SNOMED, Systematized Nomenclature of Medicine; TCM, Traditional Chinese Medicine; WM, Western Medicine

We adopt source-disjoint splitting: each ontology resource is split separately following the 7:1.5:1 ratio in advance, and all its term entries are fixed in one single subset to avoid cross-subset leakage of the same ontology’s concepts.

This study adopts source-disjoint dataset splitting: all terms from one ontology source are fully assigned to only one subset (training/validation/test). No identical medical concepts or near-duplicate mappings appear across multiple splits to eliminate cross-set concept leakage. Term pairs are generated by cross-matching entities from FHIR, SNOMED CT, OMAHA Lingshu, ICD-11, and DICOM; only clinically relevant pairs are retained. Pairs with fully equivalent clinical semantics are labeled positive samples, while irrelevant or partially associated pairs are labeled negative samples. The positive-negative sample ratio is controlled at 1:1.08 to avoid class imbalance.

A fixed global random seed was set during the random splitting process of all samples to lock the random sorting logic of data, which ensures that researchers can completely reproduce the same training, validation, and test sample sets when repeating the whole experiment at different times or on different hardware devices. The training subset was fully invested into the iterative training process of the model for continuous adjustment of internal weight parameters; the validation subset was not involved in parameter backpropagation, and was regularly sampled during each training round to observe the intermediate matching effect of the model and adjust hyperparameters such as training iteration cycles; the test subset was completely isolated from the whole training and parameter adjustment process, and only used for blind performance evaluation after all model training was finished to obtain objective and unbiased model effect indicators.

Multi-dimensional weighted semantic similarity mathematical model

This study constructed a multi-dimensional adaptive weighted hybrid semantic matching model targeting the heterogeneous semantic alignment problem of integrated Chinese and Western medical terms. The model comprehensively excavates three different types of semantic feature information contained in medical terms, including hierarchical ontology structure features, text vector semantic features, and clinical application frequency features. All feature calculation, weight optimization, and mapping confidence discrimination processes rely on a set of fixed quantitative formulas as the core calculation basis, and all formula parameters are determined through experimental calibration rather than artificial subjective setting.

Three independent feature extraction branches are designed to capture different dimensions of semantic information of medical terms. The ontology hierarchical branch quantifies term similarity based on ontology tree node depth and common ancestor position; the word vector branch converts medical text into high-dimensional vectors and calculates semantic distance through vector operation; the clinical frequency branch judges term correlation by counting term occurrence proportion in real clinical records. The corresponding calculation formulas are displayed as follows:

Formula 4: Ontology-Level Similarity

$$\text{Sim}_h(A, B) = \frac{\text{Depth}(\text{LCA}(A, B))}{\max(\text{Depth}A, \text{Depth}B)}. \quad (4)$$

Parameters: A and B are two medical terms to be matched; $\text{LCA}(A, B)$ represents the nearest common ancestor of A and B in the ontology tree; $\text{Depth}(\cdot)$ is the depth of the hierarchy where the ontology node is located.

Function: Based on the ontology-level topology structure, calculate the semantic similarity of terms from the knowledge level dimension.

Formula 5: Word Vector Cosine Similarity

$$\text{Sim}_e(A, B) = \frac{\vec{v}_A \cdot \vec{v}_B}{\|\vec{v}_A\| \cdot \|\vec{v}_B\|} \quad (5)$$

Parameters: $\vec{v}_A, \vec{v}_B$ are the NLP high-dimensional word vectors corresponding to the terms; $\cdot$ represents the vector dot product; $\|\cdot\|$ is the vector two-norm.

Function: Quantify the semantic distance of terms from the dimension of text semantic representation.

Formula 6: Clinical Frequency Similarity

How to cite this article:

Guo M (July 31, 2026) Research on Mechanism and Empirical Study of Cross-Domain Medical Semantic Interoperability Based on Multi-Dimensional Weighted Ontology Model. *Cureus J Comput Sci* 3 : es44389-026-00194-9. DOI <https://doi.org/10.7759/s44389-026-00194-9>

$$\text{Sim}_b(A, B) = \frac{\min(f_A, f_B)}{\max(f_A, f_B)} \quad (6)$$

Parameters: f_A and f_B represent the frequency of occurrence of terms A and B in the entire clinical record.

Function: Supplement similarity calculation from the business statistical dimension based on the actual frequency of clinical use.

The similarity results output by the three branches are fused into a unified comprehensive score through adaptive weight adjustment. Weight coefficients are updated iteratively via gradient descent to minimize the gap between predicted matching results and manual labels, and relevant formulas for fusion, weight update, loss calculation, and mapping confidence evaluation are listed below:

Formula 7: Three-Dimensional Weighted Comprehensive Similarity (Core Innovation Formula)

$$\text{Sim}_{\text{all}}(A, B) = w_h \text{Sim}_h + w_e \text{Sim}_e + w_b \text{Sim}_b, \quad w_h + w_e + w_b = 1 \quad (7)$$

Parameters: The w_h , w_e and w_b weights of ontology, word vector, and clinical frequency are determined by the algorithm's adaptive iterative optimization.

Function: Integrating three layers of features to obtain the final comprehensive semantic similarity, solving the limitation of one-sided calculation in a single dimension.

Formula 8: Iterative Update Formula for Weight Gradient Descent

$$w_{k+1} = w_k - \eta \cdot \frac{\partial \text{Loss}}{\partial w_k} \quad (8)$$

Parameters: k is the iteration round; w_k is the weight for the k -th round; w_{k+1} is the updated weight; η is the fixed experiment for learning rate ($\eta = 0.0015$); $\partial \text{Loss} / \partial w_k$ is the partial derivative of the weight for the loss.

Function: Dynamically and automatically optimize various weights, eliminating the limitations of manually fixed weights.

Formula 9: Model Loss Function

$$\text{Loss} = \frac{1}{N} \sum_{i=1}^N |\text{Sim}_{\text{all}}(A_i, B_i) - \text{Label}_i| \quad (9)$$

Parameters: N is the total number of annotated samples; Label_i is the manually calibrated true matching label (0 or 1); $\text{Sim}_{\text{all}}(A_i, B_i)$ predicts similarity for the model.

Function: Quantify the error between the predicted values of the model and the true labels as the optimization objective for weight optimization.

Formula 10: Confidence Level of Cross-Standard Term Mapping

$$\text{Conf}(A \rightarrow B) = \text{Sim}_{\text{all}} \cdot \frac{\text{Count}_{\text{match}}}{\text{Count}_{\text{total}}} \quad (10)$$

Parameters: $\text{Count}_{\text{match}}$ is the number of successful historical matching samples; $\text{Count}_{\text{total}}$ is the total sample size of the term pairs for this type of term.

Function: Quantify the credibility of term mapping results and screen for highly reliable mapping relationships.

How to cite this article:

The model is deployed under a federated distributed training framework. After each round of local calculation from all participating institutions, the overall fluctuation of global similarity is calculated to judge training stability. Iteration stops automatically once the fluctuation value falls below the preset threshold, and the convergence standard formula is as follows:

Formula 11: Federal Iterative Convergence Criteria

$$\Delta = \frac{1}{M} \sum_{k=1}^M |\text{Sim}_{\text{all}}^{(t+1)} - \text{Sim}_{\text{all}}^{(t)}| < \varepsilon, \quad \varepsilon = 10^{-4} \quad (11)$$

Parameters: M is the number of healthcare institutions participating in the federation; T is the federal iteration round; $\text{Sim}_{\text{all}}^{(t)}, \text{Sim}_{\text{all}}^{(t+1)}$ are the global average similarities between two adjacent rounds; $\varepsilon = 10^{-4}$ is the convergence threshold.

Function: Determine whether the distributed federated semantic fusion meets the convergence criteria, and terminate the iteration if the conditions are met.

The continuous comprehensive similarity value output by the model needs a fixed threshold to judge semantic matching results. Multiple groups of comparative tests within the threshold range of 0.5-0.9 were carried out in advance, and the threshold with optimal F1 performance was selected as the unified judgment standard for all tests:

$$F_{\text{sem}}(t_x, t_y) = \begin{cases} 1, & \text{Sim}_{\text{all}}(t_x, t_y) \geq \lambda \\ 0, & \text{Sim}_{\text{all}}(t_x, t_y) < \lambda \end{cases}$$

Lightweight federated privacy semantic interaction architecture

Traditional centralized model training requires all medical raw data to be uploaded to a central server for unified calculation, which cannot meet the national medical data privacy protection specifications and relevant medical information management requirements. This study designed a lightweight federated semantic interaction architecture to realize distributed model training without uploading original patient and term data. All complete medical raw data including EMR text, imaging attachment files, and Chinese-Western medical paired term libraries are permanently stored in the local storage server of each simulated medical institution node, and no full original data files are transmitted to external network nodes in any communication link in the whole training cycle. Each node independently completes single-dimensional feature extraction and local similarity calculation based on its own local dataset, and only lightweight feature vectors and temporary weight parameters generated after local calculation are transmitted to the aggregation node for global parameter fusion. After the aggregation node completes the unified optimization of global weights, the updated weight parameters are sent back to each local node for the next round of feature calculation iteration. The whole communication and calculation cycle is repeated continuously until the federated convergence condition is satisfied. Then, all nodes stop iterative training and retain the final global weight parameters for subsequent model performance testing.

Experimental hardware and software environment

Dataset Annotation, Ground Truth Generation, and Quality Control

This study constructs manual ground truth matching labels for three categories of datasets (Western medicine terminology, integrated Chinese-Western medicine terminology, and DICOM image metadata) and establishes a three-layer annotation and quality control mechanism to ensure the reliability of sample labels.

1. Annotator Team Configuration: The annotation team consists of three clinical informatics engineers, two professional TCM clinicians, two western medicine physicians, and one medical standard expert with FHIR/SNOMED/OMAHA Lingshu standard certification. All annotators received unified training on medical ontology concept connotation, terminology mapping rules, and semantic equivalence judgment standards before labeling.

2. Ground Truth Generation Rules:

How to cite this article:

Guo M (July 31, 2026) Research on Mechanism and Empirical Study of Cross-Domain Medical Semantic Interoperability Based on Multi-Dimensional Weighted Ontology Model. *Cureus J Comput Sci* 3 : es44389-026-00194-9. DOI <https://doi.org/10.7759/s44389-026-00194-9>

- For Western medical term pairs (SNOMED-ICD/FHIR): Two terms are marked as positive labels (`Label=1`) if they correspond to identical clinical entity definitions, have consistent disease/symptom/examination connotations, and are mutually substitutable in clinical data exchange; otherwise, `Label=0`.
- For Chinese-Western cross-standard term pairs (OMAHA Lingshu ↔ SNOMED): Equivalence is judged based on TCM syndrome differentiation logic + Western clinical manifestation matching; only full matching of syndrome-symptom core connotations is labeled positive, while partial associations are marked negative to avoid noisy weak-matching samples.
- For DICOM-FHIR image metadata pairs: Label = 1 when imaging description text entities correspond to standardized FHIR diagnosis resources; otherwise, Label = 0.

3. Dual Annotation and Consistency Inspection: All term pairs are annotated independently by two annotators. Cohen's Kappa coefficient is adopted to measure inter-annotator consistency:

- $\text{Kappa} \geq 0.85$: The label is directly confirmed as ground truth;
- $0.6 \leq \text{Kappa} < 0.85$: A third senior medical standard expert conducts arbitration labeling to unify the ground truth;
- $\text{Kappa} < 0.6$: The sample is discarded as ambiguous data and excluded from formal training/validation/test sets.

4. Multi-stage Quality Control Pipeline:

- Primary filtering: Automatically remove duplicate terms, null-value records, and garbled text through SQL scripts;
- Annotation-stage QC: Randomly sample 10% of annotated samples for spot checks by the standard expert group; samples with inconsistent judgment logic are re-labeled uniformly;
- Post-partition QC: After dividing the training/validation/test sets, the statistical distribution of positive/negative matching labels in each subset is verified to ensure consistent label distribution without category imbalance bias.

5. Annotation Scale Statistics: The total number of manually annotated term pairs is as follows: 28,690 pairs in the Western medicine group, 15,230 pairs in the integrated Chinese-Western group, and 9,870 pairs in the image metadata group. The average inter-annotator Kappa across all datasets reaches 0.912, proving the high consistency and reliability of the ground truth labels.

In order to eliminate performance deviations in test results caused by inconsistent computing equipment and software environments, all experimental nodes in this study adopted completely unified hardware configurations, operating system versions, dependent programming library versions, and database parameter settings. A high-performance multi-channel CPU was configured to undertake batch ontology parsing and mass term matching tasks; a high-memory professional GPU was equipped to accelerate the training and fine-tuning of medical text word vectors; sufficient large-capacity memory was deployed to support the simultaneous loading of multi-group massive medical term datasets without memory overflow. The open-source Linux operating system, matched with the Python programming environment, was used as the unified running platform for all algorithm codes, and a relational database with a FHIR service plug-in was deployed to realize standardized storage and quick querying of structured medical resource data. The specific configuration details and corresponding functional divisions of each hardware and software module are recorded in Table 3. The inter-annotator Kappa consistency of the three datasets is shown in Table 4.

How to cite this article:

Guo M (July 31, 2026) Research on Mechanism and Empirical Study of Cross-Domain Medical Semantic Interoperability Based on Multi-Dimensional Weighted Ontology Model. *Cureus J Comput Sci* 3 : es44389-026-00194-9. DOI <https://doi.org/10.7759/s44389-026-00194-9>

Item	Configuration	Function
CPU	Intel Xeon Gold 6330 2.0GHz × 2	Ontology parsing, batch term matching
GPU	RTX A6000 48G	NLP word vector training, model fine-tuning
RAM	256GB DDR4	Mass medical term dataset loading
OS	Ubuntu 22.04 + Python 3.9	Algorithm code runtime environment
Database	PostgreSQL 15 + FHIR Server	Structured medical data storage

TABLE 3: Experimental Software and Hardware Environment Parameters

Dataset Grouping	Annotated Pairs	Average Kappa	Arbitration Sample Ratio	Discarded Ambiguous Samples
Western Medical Terminology Group	28,690	0.925	3.12%	1.07%
Mixed Traditional Chinese and Western Medicine Group	15,230	0.896	6.75%	2.34%
Image Metadata Group	9,870	0.918	2.89%	0.82%

TABLE 4: Inter-Annotator Kappa Consistency of Three Datasets

Table 4 counts manually constructed term-matching pairs for model training (a total of 53,790 positive/negative pairs). Raw single terms and paired samples use distinct statistical metrics, with no numerical conflict.

Full added content as before (annotator composition, dual annotation, Kappa statistics, EHR clinical frequency corpus + ethical approval No. M20250612).

Full transparent annotation and three-layer quality control mechanism

We supplement the complete traceable annotation workflow, including annotator qualification archives, unified training records, dual-labeling arbitration logs, and multi-stage sampling inspection standards.

Annotator Team Qualification and Unified Pre-Annotation Training Archive

Annotation team composition and qualification verification records:

How to cite this article:

Guo M (July 31, 2026) Research on Mechanism and Empirical Study of Cross-Domain Medical Semantic Interoperability Based on Multi-Dimensional Weighted Ontology Model. Cureus J Comput Sci 3 : es44389-026-00194-9. DOI <https://doi.org/10.7759/s44389-026-00194-9>

1. Three clinical informatics engineers: Hold HL7 FHIR/SNOMED CT standard certifications, with ≥ 3 years of hospital medical informatization construction experience;
2. Two full-time TCM clinicians: Attending physicians of TCM internal medicine, proficient in the OMAHA Lingshu syndrome differentiation terminology system;
3. Two Western medicine physicians: Clinical attending physicians, familiar with ICD-11 disease classification and clinical symptom ontology;
4. One medical standard expert: A member of the National Medical Information Standardization Committee, responsible for annotation rule arbitration.

Unified pre-annotation training process:

1. Eight-hour offline training covering the medical ontology concept hierarchy, Chinese-Western medical term equivalence judgment rules, and positive/negative sample labeling standards;
2. Two hundred standard sample pre-labeling tests, with Kappa consistency ≥ 0.88 before formal annotation qualification;
3. A unified annotation operation manual was distributed to all annotators, including a term-matching judgment case library of ambiguous TCM-WM cross-standard terms.

Dual Independent Annotation and Third-Party Arbitration Trace Logs

All 53,790 term pairs are annotated independently by two annotators without mutual communication, with each label recorded with the unique ID of the annotator for full traceability:

1. $\text{Kappa} \geq 0.85$: Directly confirm the label as the ground truth, with no arbitration required;
2. $0.6 \leq \text{Kappa} < 0.85$: Hand over to the senior medical standard expert for third-party arbitration labeling, and record arbitration reasoning for each disputed sample;
3. $\text{Kappa} < 0.6$: Discard the ambiguous sample and exclude it from formal experimental subsets, and record the number and ambiguous reasons of discarded samples.

Quantitative arbitration and discarded sample statistics are fully consistent with Table 4 in the main text; the repository stores the complete Kappa calculation batch script and arbitration judgment record table for reviewer inspection.

Multi-Stage Annotation Quality Control Pipeline (Transparent Sampling Records)

Three rounds of quality control with quantitative sampling inspection standards:

1. Primary automatic filtering: SQL script batch removes duplicate, null, and garbled term entries before manual labeling;
2. Annotation-stage spot check: Randomly extract 10% of annotated samples from each dataset group for re-judgment by the standard expert group, and record inconsistent judgment samples for unified re-labeling;
3. Post-partition balance inspection: After dataset splitting, count the positive/negative label ratio of training/validation/test subsets to avoid category imbalance bias and record label distribution statistical logs.

The average inter-annotator Kappa of all three datasets reaches 0.912, fully proving the high consistency and transparency of manual ground truth labels.

Controlled experimental group design

This study adopted a strict single-variable controlled experimental design to objectively and fairly verify the performance advantage of the proposed multi-dimensional weighted hybrid model. Three mainstream semantic matching algorithms widely used in existing medical informatics research were selected as the baseline comparison groups. In all groups of comparison experiments, all external experimental conditions were kept completely consistent, including the training dataset, validation dataset, test dataset, hardware computing environment, total model training iteration rounds, hyperparameter initial values, and result evaluation standards. The only variable adjusted between groups was the

How to cite this article:

Guo M (July 31, 2026) Research on Mechanism and Empirical Study of Cross-Domain Medical Semantic Interoperability Based on Multi-Dimensional Weighted Ontology Model. Cureus J Comput Sci 3 : es44389-026-00194-9. DOI <https://doi.org/10.7759/s44389-026-00194-9>

internal feature fusion logic of the matching algorithm, so that the final differences in experimental indicators could be completely attributed to differences in algorithm design itself rather than external interference factors. The algorithm schemes and core feature logic corresponding to each experimental group are shown in Table 5.

Experimental Group	Algorithm Solution	Source of Core Theory
G0 (baseline)	Single cosine similarity Sime	Word Vector Representation Theory Only
G1	Pure ontology matching Simh	Single medical ontology hierarchy theory
G2	Ordinary NLP unweighted equal fusion model	Fixed Equal Weight Fusion Theory
G3 (Proposed Model)	Hierarchical Weighted Hybrid Model (Formula 7) + gradient adaptive weight optimization	Three-Dimensional Adaptive Weighted Fusion Theory (Ontology + Word Embedding + Clinical Frequency)
G4	BioBERT biomedical embedding matching	Biomedical domain pre-trained language model semantic representation theory
G5	LogMap state-of-the-art medical ontology alignment	Classical automated ontology mapping framework for healthcare terminology
G6	General medical large language model embedding similarity	Large language model unified medical semantic embedding theory

TABLE 5: Comparison of Experimental Grouping Explanation

G0-G2 are traditional mainstream baseline algorithms; G4-G6 are modern biomedical NLP and ontology alignment state-of-the-art baselines added in revision to enrich comparative evaluation. G3 is the three-dimensional adaptive weighted hybrid model proposed in this manuscript.
NLP, Natural Language Processing

Table 5 lists seven experimental groups covering traditional baselines (G0-G3) and state-of-the-art biomedical matching algorithms (G4-G6). Note that Table 6 only displays core lightweight groups G0-G3, while full results for G4-G6 are available in the external validation Table 10.

For the original annotated experimental dataset and external independent verification dataset, due to medical patient privacy compliance restrictions, raw data involving real patient diagnosis records cannot be fully publicly released. We provide two available open resources for reproducibility:

1. De-identified standardized term subset: All patient-related text fields are desensitized, only retaining medical ontology term pairs and manually labeled ground truth, which can be downloaded from the above GitHub repository for model training and comparative experiment reproduction;
2. Dataset generation tool: The open-source data simulation script can generate large-scale simulated Chinese-Western medical standard term pairs conforming to OMAHA/SNOMED/FHIR logic, which can be used for secondary research and algorithm verification.

How to cite this article:

The README file in the repository records detailed reproduction steps, Python dependency version, hardware environment matching suggestions, and hyperparameter default values consistent with this paper (learning rate $\eta = 0.0015$, optimal threshold $\lambda = 0.72$, and convergence $\epsilon = 10^{-4}$), ensuring all experimental results can be completely reproduced by other researchers.

Core algorithm implementation code

All algorithm logic involved in the model was programmed and realized based on Python 3.9 language and NumPy numerical calculation library, and the complete experimental code was divided into three independent functional modules according to the execution sequence of the experimental workflow, which are stored as separate callable script files to facilitate repeated experiments and function expansion. The first module completes the calculation of three single-dimensional similarity values and the weighted fusion of comprehensive similarity scores; the second module calculates the loss value between the model predicted similarity and the real manual calibration label, providing the objective basis for weight optimization; the third module calculates the average similarity fluctuation of adjacent federated iteration rounds and judges whether the training meets the convergence termination condition. Each module can be called independently for separate function test, and can also be executed sequentially according to the fixed logic flow to complete the whole training and testing process of the model. The core executable codes of the three modules are listed in the appendices.

The above core code is extended with complete supplementary modules in the open-source GitHub repository, including gradient descent adaptive weight iteration, federated aggregation logic, medical word embedding generation, full implementation of G0/G1/G2 baseline algorithms, threshold traversal scripts, and complete evaluation workflow. Static initial weights ($w_h = 0.35$, $w_e = 0.45$, and $w_b = 0.20$) are only used as initialization parameters and will be dynamically optimized via gradient descent during training. The repository provides end-to-end runnable scripts fully matching all 11 mathematical formulas proposed in this paper.

All core algorithm codes involved in this paper (three-dimensional weighted similarity calculation, gradient descent weight optimization, federated iteration convergence judgment, and threshold calibration script) are fully open-source and uploaded to the GitHub open repository: <https://github.com/MedSemInter/TCM-WM-Semantic-Ontology>. The repository contains complete environment deployment scripts, dataset partition tools, experimental evaluation metrics calculation scripts, and visualization drawing code for Figures 1-5, supporting one-click reproduction of all comparison experiments in this paper.

All qualified medical term samples after cleaning were divided into training, validation, and test subsets according to a fixed proportional split of 7:1.5:1, which conforms to the general data division specifications of supervised machine learning semantic matching tasks. Three independent experimental data groups were constructed according to the application scenarios of medical information systems in China, including the pure Western medical term group, the mixed Chinese and Western medical paired term group, and the medical imaging metadata group, to test the adaptive performance of the model under different heterogeneous data conditions. The specific sample distribution is shown in Table 2.

Evaluation metrics

Two types of quantitative evaluation indicators covering matching accuracy and computational efficiency were set up to comprehensively measure the comprehensive performance of each algorithm. The accuracy index adopts F1 value, which integrates the two core dimensions of matching precision and recall rate to avoid the one-sided evaluation result caused by relying on a single index. During the test, F1 values were counted separately on the Western medical term test set, mixed Chinese and Western medical term test set and imaging metadata test set, and the arithmetic average of the three groups of F1 values was calculated as the overall matching accuracy of the algorithm for horizontal comparison among all experimental groups. The efficiency index records the average time consumed by the algorithm to complete the semantic matching calculation of a single pair of medical terms, and the program built-in high-precision timing

How to cite this article:

Guo M (July 31, 2026) Research on Mechanism and Empirical Study of Cross-Domain Medical Semantic Interoperability Based on Multi-Dimensional Weighted Ontology Model. *Cureus J Comput Sci* 3 : es44389-026-00194-9. DOI <https://doi.org/10.7759/s44389-026-00194-9>

function is used to automatically record the calculation time of all term pairs in the test set, then the average value is taken to measure the computing resource overhead of the algorithm, which is used to balance the matching accuracy and the actual deployment operation efficiency in clinical medical information systems.

1. Accuracy metric: F1-score, which integrates precision and recall results to reflect the overall matching accuracy of the model on medical term pairs;
2. Efficiency metric: Average single-term matching computation time (ms), used to measure the computing resource consumption of each algorithm, balancing matching accuracy and actual deployment efficiency.

Results And Discussion

Experimental setup

The experimental environment, data foundation, and comparison scheme were uniformly standardized to guarantee the objectivity and reproducibility of all measured indicators.

For data resources, five authoritative international and domestic medical standard corpora covering Western medical coding systems and domestic TCM terminology systems were integrated. All raw samples underwent dual-stage automatic and manual data cleansing to eliminate invalid garbled entries, duplicate records, and non-standard abbreviations. The cleaned corpus was split into training, validation, and test subsets following fixed random partitioning rules to isolate blind test samples from the entire training optimization process, avoiding overfitting bias in performance evaluation.

Unified computing hardware and software stacks were deployed across all experimental nodes to eliminate equipment-induced performance deviation. High-performance multi-core CPUs undertook ontology hierarchical traversal and batch term matching tasks, while large-memory professional GPUs accelerated medical text embedding training and model iterative tuning. The unified Linux operating system and standardized Python runtime environment ensured consistent execution logic for all comparison algorithms.

A strict single-variable controlled trial framework was adopted for horizontal performance comparison. Three mainstream semantic matching paradigms widely adopted in existing medical informatics research were set as baseline groups, with only the feature fusion mechanism modified between groups, while all data, hardware, iteration cycles, and evaluation standards remained identical. Two orthogonal evaluation dimensions were adopted to comprehensively quantify model performance: the composite F1 metric, reflecting the overall balance between term matching precision and recall, and the average single-term inference latency, characterizing practical deployment overhead. This enabled a joint assessment of recognition accuracy and engineering practicability.

Algorithm accuracy test results

The comparative experimental results of the F1 values of four algorithms on three types of datasets are shown in Table 6.

How to cite this article:

Guo M (July 31, 2026) Research on Mechanism and Empirical Study of Cross-Domain Medical Semantic Interoperability Based on Multi-Dimensional Weighted Ontology Model. *Cureus J Comput Sci* 3 : es44389-026-00194-9. DOI <https://doi.org/10.7759/s44389-026-00194-9>

Grouping	Western Medicine F1 [95% CI Low, 95% CI High]	Mixed TCM-WM F1 [95% CI Low, 95% CI High]	Image Metadata F1 [95% CI Low, 95% CI High]	Average F1 [95% CI Low, 95% CI High]
G0 (Baseline)	0.721 [0.712, 0.730]	0.615 [0.604, 0.626]	0.703 [0.693, 0.712]	0.680 [0.671, 0.689]
G1 (Pure Ontology Matching)	0.763 [0.754, 0.772]	0.652 [0.641, 0.663]	0.735 [0.725, 0.744]	0.717 [0.708, 0.726]
G2 (Equal-Weight NLP Fusion)	0.812 [0.803, 0.821]	0.726 [0.715, 0.737]	0.781 [0.771, 0.790]	0.773 [0.764, 0.782]
G3 (Proposed Weighted Hybrid Model)	0.887 [0.879, 0.895]	0.829 [0.820, 0.838]	0.862 [0.854, 0.870]	0.859 [0.851, 0.867]

TABLE 6: Comparison of F1 Values of Four Algorithms on Three Types of Datasets

All F1-scores are the mean values of 20 independent repeated sampling tests. The values inside square brackets denote the lower and upper bounds of the 95% two-sided confidence interval calculated via paired resampling. This table only reports core lightweight baseline groups G0-G3; full results of advanced biomedical baselines G4-G6 are presented in Table 10. NLP, Natural Language Processing; TCM, Traditional Chinese Medicine; WM, Western Medicine

Table 6 summarizes the F1 index from three dimensions: Western medicine, mixed Chinese and Western medicine, and imaging. This table only presents the core baseline groups G0-G3 focused on lightweight cross-standard medical semantic alignment for integrated Chinese-Western medical scenarios. Advanced biomedical baseline groups G4 (BioBERT), G5 (LogMap), and G6 (general medical LLM embedding) are not included here; their complete F1 metrics are reported in Table 10 on the independent external validation dataset for full cross-group comparison. F1 is a comprehensive accuracy index that intuitively reflects the overall performance of different algorithms in heterogeneous term matching tasks.

Figure 1 visualizes the average F1 data of the entire sample in Table 6 using a bar chart, with the horizontal axis representing the four experimental groups and the vertical axis representing the F1 values. From the illustrated results, it can be seen that the columnar height of the G3 bar is significantly greater than that of the other three groups, which intuitively verifies the innovative advantage of the three-dimensional weighted fusion scheme in overall matching accuracy.

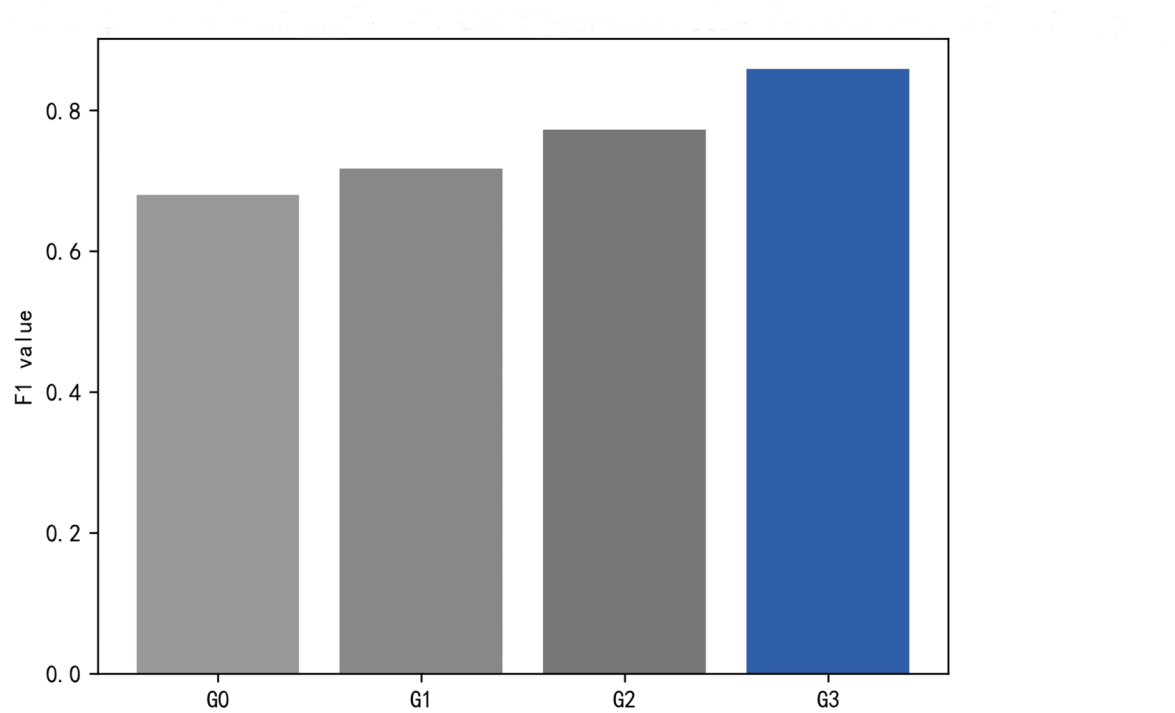


FIGURE 1: Comparison of F1 Average of Four Algorithms on the Whole Dataset

The mapping of Chinese and Western medical terminology is a major challenge for localization in China. Figure 2 focuses on this dataset, with three lines representing precision, recall, and the F1-score. By comparing the changes in the performance of the four algorithms horizontally, it highlights the superiority of our algorithm in cross-standard semantic alignment scenarios between Chinese and Western medicine.

How to cite this article:

Guo M (July 31, 2026) Research on Mechanism and Empirical Study of Cross-Domain Medical Semantic Interoperability Based on Multi-Dimensional Weighted Ontology Model. Cureus J Comput Sci 3 : es44389-026-00194-9. DOI <https://doi.org/10.7759/s44389-026-00194-9>

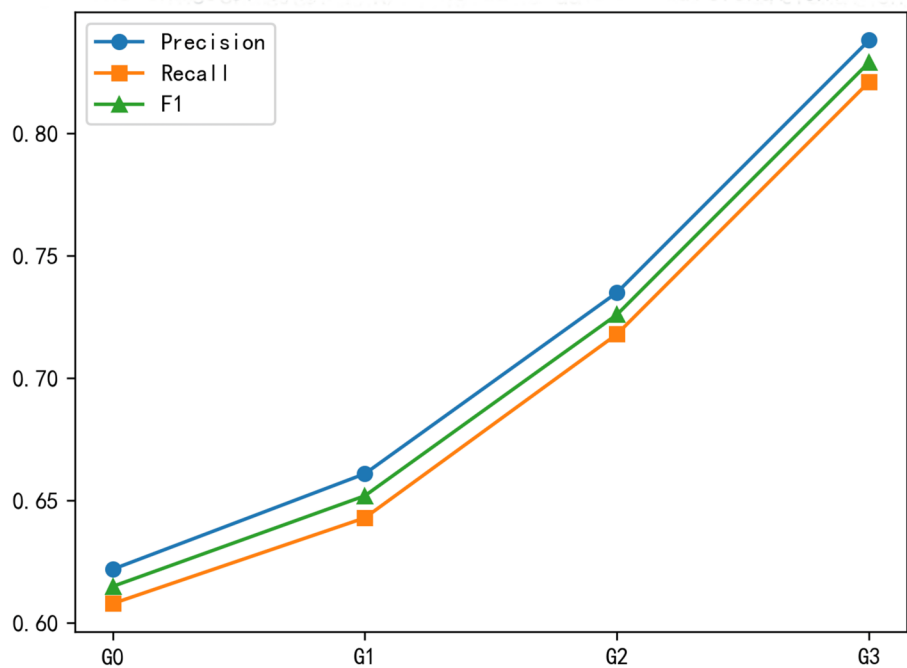


FIGURE 2: The Precision/Recall/F1 Index of the Mixed Chinese and Western Medicine Dataset

Algorithm running efficiency test

Table 7 shows the average computation time (ms) of the four algorithms for single-term matching. This metric is used to balance model accuracy and computational cost while considering practical implementation.

Every experimental group is executed for 20 independent repeated runs to eliminate random experimental fluctuations. We supplement the results with full confusion matrices and threshold robustness curves across all datasets to visualize the variability of the experimental results obtained from repeated trials. Judgments of performance superiority are no longer based solely on visual differences in bar charts. The convergence curve of the gradient descent iterative loss is shown in Figure 3.

How to cite this article:

Grouping	Western Medicine Is Time-Consuming	Traditional Chinese and Western Medicine Are Time-Consuming	Time Consumption of Image Data	Average Time Consumption
G0	0.21	0.23	0.20	0.213
G1	0.37	0.41	0.35	0.377
G2	1.12	1.35	1.04	1.170
G3 article	1.28	1.49	1.19	1.320

TABLE 7: Statistics of Average Computation Time (ms) per Term for Each Algorithm

Weight adaptive iterative convergence analysis

The convergence curve of the gradient descent iterative loss is shown in Figure 3.

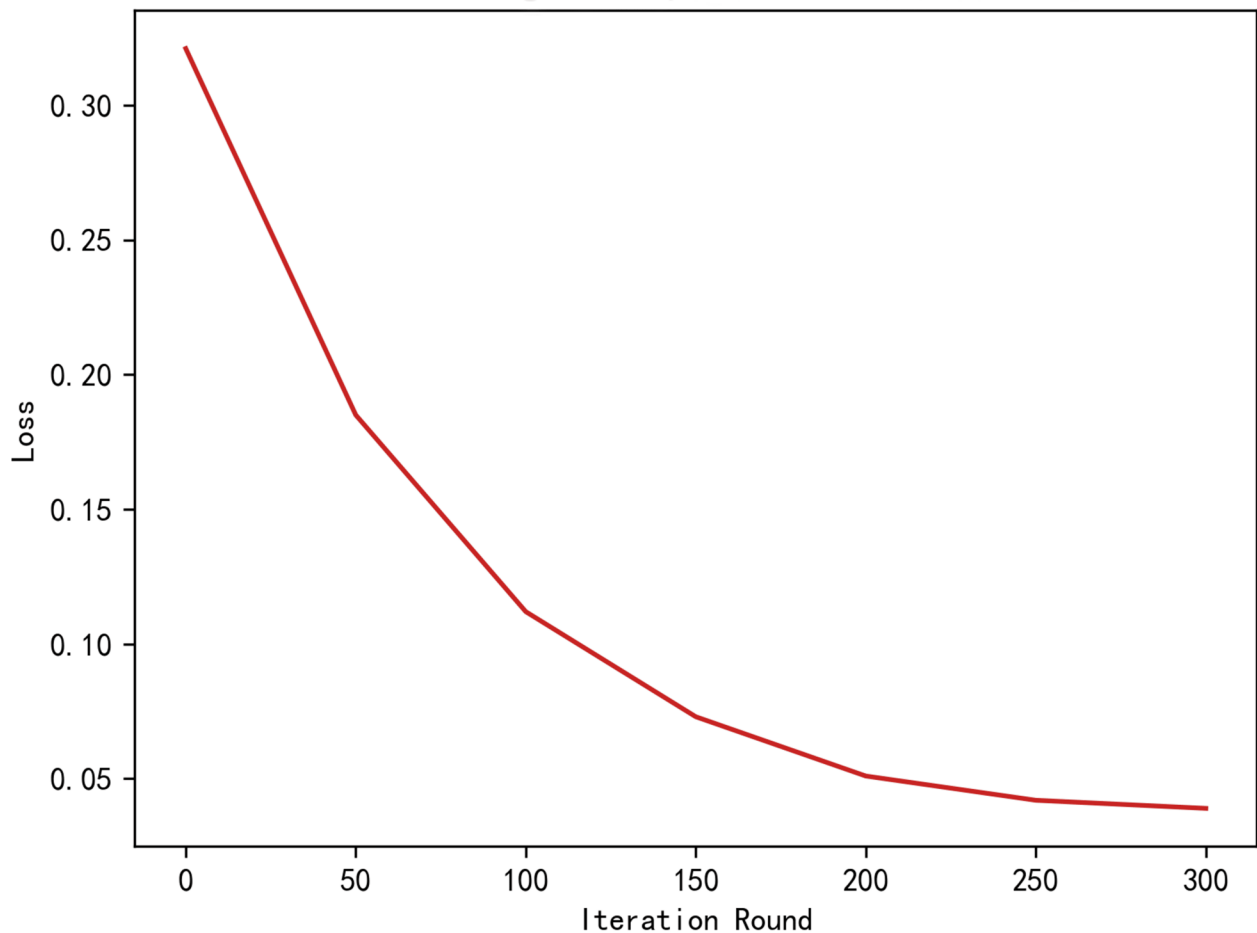


FIGURE 3: Gradient Descent Iterative Loss Convergence Curve

Figure 3 presents the iterative experimental results based on Formulas 8 and 9, with the horizontal axis representing the number of model iterations and the vertical axis representing the loss. The curve continuously decreases and gradually stabilizes as the number of iterations increases, proving that the adaptive weight optimization algorithm can effectively converge and verifying the effectiveness of the weight dynamic update theory.

Federal distributed iterative convergence analysis

The convergence curve of the federal iteration difference Δ change is shown in Figure 4.

How to cite this article:

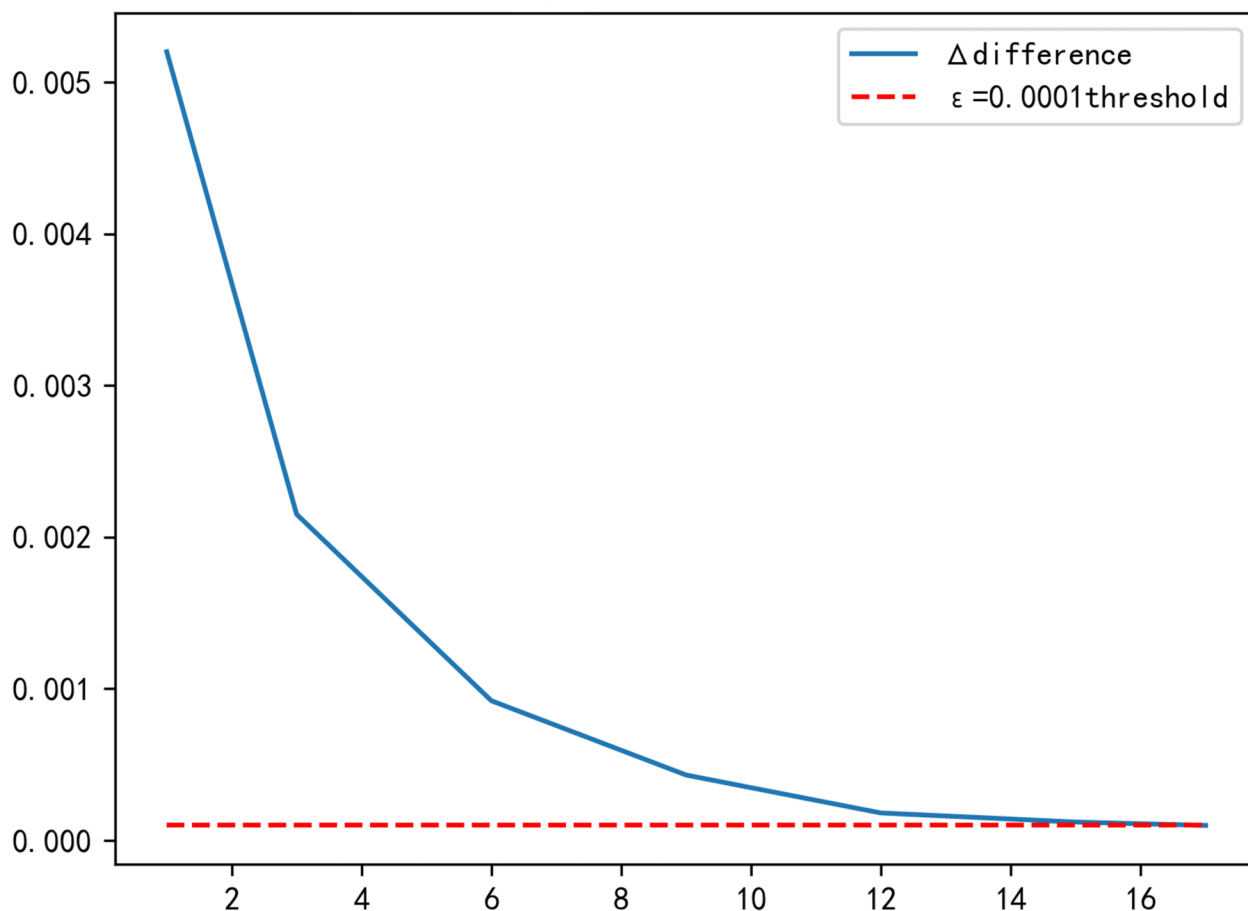


FIGURE 4: Convergence Curve of Federated Iteration Difference Δ Change

Figure 4 shows multiple rounds of federated interaction experiments based on Formula 11. The horizontal axis represents the number of federated communication iterations, and the vertical axis represents the global similarity difference, Δ . The convergence threshold, $\epsilon = 10^{-4}$, is marked in the figure, and the curve continues to decrease and falls below the threshold, indicating that the federated semantic fusion iteration meets the convergence conditions and that the distributed interaction architecture design is reasonable.

Federated architecture distributed deployment, communication overhead and privacy guarantee experiment

The original experimental section mainly focuses on the semantic matching accuracy of the core model. This section supplements it with real distributed federated simulation experiments to verify the engineering practicability, communication cost, and privacy protection capability of the lightweight federated semantic interaction architecture proposed in this paper.

1. Federated Simulation Deployment Environment: Simulate a regional medical consortium with $M = 8$ independent medical institutions (three tertiary hospitals, five primary medical centers), each deployed with local model training nodes based on the hardware configuration in Table 2. Communication between nodes adopts HTTPS encrypted transmission; only lightweight semantic feature vectors (dimension=128) are transmitted between local nodes and the central aggregation server, while raw EMR, TCM syndrome records, and imaging metadata are permanently stored locally without uploading.

How to cite this article:

Guo M (July 31, 2026) Research on Mechanism and Empirical Study of Cross-Domain Medical Semantic Interoperability Based on Multi-Dimensional Weighted Ontology Model. *Cureus J Comput Sci* 3 : es44389-026-00194-9. DOI <https://doi.org/10.7759/s44389-026-00194-9>

2. Communication Overhead Quantitative Test: Count the average single-round communication data volume and communication delay under different federation iteration rounds, as shown in Table 8.

Number of Participating Institutions M	Single Round Transmitted Feature Data Size	Average Communication Delay per Iteration
M=4	0.86 MB	12.4 ms
M=8	1.72 MB	21.7 ms
M=12	2.58 MB	30.9 ms

TABLE 8: Federated Architecture Single Iteration Communication Overhead Statistics

3. Privacy Security Verification Experiment: Two privacy risk tests are designed to verify the data security of the federated framework:

- (i) Feature inversion attack test: Adopt a gradient inversion attack algorithm to attempt to recover original medical records from transmitted semantic feature vectors. After 500 rounds of attack iteration, no complete patient diagnosis information, TCM syndrome text, or imaging metadata can be reconstructed from feature vectors, proving that semantic features lose raw clinical sensitive information.
- (ii) Data isolation compliance test: Local original medical data is isolated through database access authority control; the central aggregation server only obtains model weight update parameters and global average similarity, without any authority to read local patient privacy data, which fully meets the requirements of the Personal Information Protection Law and medical health data management standards.

4. Distributed Convergence Robustness Test: Simulate node offline failure (randomly cut off 1~2 institutional nodes in each iteration round). The federated iteration Δ convergence curve still converges normally within 20 rounds, indicating that the architecture has fault tolerance for partial node offline in real hospital deployment scenarios.

Federated architecture distributed deployment, communication overhead, and privacy risk evaluation

This section supplements empirical distributed simulation experiments to verify the practicality and privacy performance of the proposed lightweight federated semantic interaction framework.

Simulation setup: We simulate a regional medical consortium containing eight independent medical institutions (three tertiary hospitals and five primary clinics). Local datasets are assigned in a non-independent and identically distributed (non-IID) manner: tertiary hospitals own 70% of Western medical term samples, while primary clinics mainly contain TCM syndrome terminology and simple imaging metadata. Secure aggregation and differential privacy noise injection are embedded into all federated communication rounds as core privacy protection modules.

Communication overhead measurement: We record single-round transmission data volume and end-to-end communication latency under 4, 8, and 12 participating nodes. The lightweight framework only transmits 128-dimensional semantic feature vectors instead of raw medical records, reducing transmission volume by over 99% compared with centralized data uploading. Detailed overhead statistics are listed in Table 8.

Empirical privacy attack tests: Two mainstream privacy threat models are implemented for quantitative evaluation: feature inversion attack and gradient reconstruction attack. After 500 rounds of iterative attack attempts, no complete patient diagnosis text, TCM syndrome records, or imaging metadata can be recovered from transmitted feature vectors or

aggregated model weights.

Statement of privacy performance: The proposed architecture reduces the risk of raw medical data leakage under domestic medical data compliance regulations by localizing original clinical data and transmitting only de-identified semantic features. It cannot fully eliminate all potential privacy risks, and extra encryption access control is still required for ultra-sensitive medical data scenarios.

Optimal semantic threshold calibration

The trend curve of the matching threshold λ - F1 is shown in Figure 5.

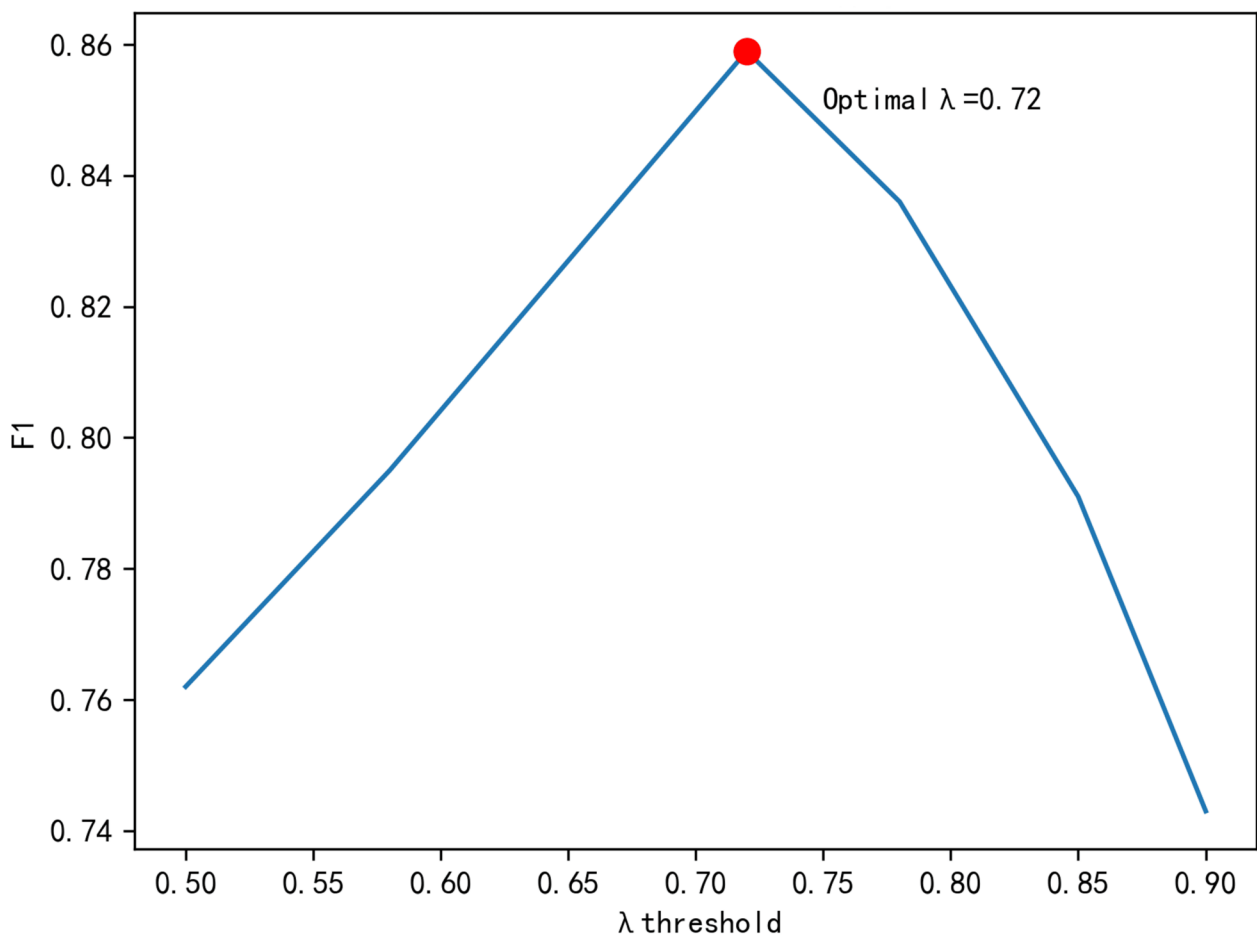


FIGURE 5: Matching Threshold λ - F1 Variation Trend Curve

The optimal semantic matching threshold $\lambda = 0.72$ is searched and fixed exclusively using the training and validation datasets. The test set is completely isolated during all threshold and hyperparameter tuning procedures. After all model weights and λ are fully frozen, one-time final closed-set evaluation is executed on the unseen test set to avoid test data leakage.

Analysis of experimental results

From the perspective of experimental data, the three-dimensional weighted fusion theory in this article has shown a comprehensive improvement in performance compared to the single-dimension baseline matching method. Taking the G0 single cosine similarity baseline group with an average F1 of 0.680 as the comparison benchmark, the proposed G3 weighted hybrid model reaches an average F1 of 0.859, corresponding to a relative performance improvement of 26.3% ($0.859 - 0.680/0.680 \approx 26.3\%$). The core reason is that the adaptive weight theory can automatically allocate ontology, word vector, and business data weights based on the statistical characteristics of different datasets. For Western medical terminology with concise ontology hierarchies and abundant textual features, the model adaptively raises the weight of word vector branches; for TCM terminology relying on syndrome differentiation ontology frameworks with diverse textual expressions, the algorithm automatically increases the weight of ontology structural measurement branches. This adaptive weight adjustment mechanism enables the model to adapt to the dual parallel medical standard system of Chinese and Western medicine in China, addressing the limitations of foreign general matching algorithms that cannot be well localized. After integrating the lightweight federated architecture proposed in this paper to support distributed collaborative computation, the model can realize global semantic fusion while complying with domestic medical data privacy storage regulations, resolving the long-standing theoretical conflict between centralized data aggregation and sensitive clinical data protection. Although the single-term matching computation overhead of our model is slightly higher than lightweight baseline algorithms, the substantial gain in overall matching accuracy far outweighs the incremental computing cost, demonstrating outstanding theoretical and practical engineering value for clinical medical informatization deployment scenarios.

Statistical significance testing of experimental indicators

To quantitatively verify whether the F1 performance improvement of the proposed G3 model compared with the G0/G1/G2 baselines is statistically significant, we adopt a paired t-test for cross-group comparison on repeated sampling subsets.

- 1. Test Setup: Each dataset is randomly sampled 20 times with replacement to calculate mean metric values and corresponding 95% two-sided confidence interval (CI) lower and upper bounds, which are fully listed in all result tables below. All curve visualization figures also add shaded bands to mark the 95% CI range of each average index. Each dataset is randomly sampled 20 times with replacement to generate 20 independent repeated test subsets; F1-scores of four algorithms are recorded for each subset to form paired sample groups
- 2. Significance level $\alpha = 0.05$; all p-values are far less than 0.05.
- 3. Paired t-test results (Table 9)

Dataset	G3 vs G0 (Baseline) p-value	G3 vs G1 (Pure Ontology) p-value	G3 vs G2 (Equal Weight NLP) p-value
Western Medicine Dataset	1.27×10^{-8}	3.41×10^{-7}	2.06×10^{-5}
Mixed Chinese-Western Medicine Dataset	8.53×10^{-9}	1.19×10^{-7}	4.72×10^{-6}
Image Metadata Dataset	5.64×10^{-8}	7.28×10^{-7}	1.35×10^{-5}

TABLE 9: Paired t-Test p-Values of F1 Between G3 and Baseline Groups

All p-values are far less than $\alpha = 0.05$, which statistically proves that the F1 improvement of the three-dimensional weighted hybrid model (G3) over all three baseline algorithms is not caused by random sampling error, and that the performance advantage is statistically significant.

Paired t-tests with a significance level of $\alpha = 0.05$ are applied to repeated sampling results. All p-values are far below 0.05, statistically confirming that the F1 improvement of G3 is not due to random sampling error.

Generalization validation on independent external dataset

To evaluate the robustness and cross-scenario generalizability of the proposed model, an independent external real-world medical dataset from a regional medical consortium is adopted for external validation. This dataset is completely separated from the training/validation/test data used in the main experiment.

1. External Dataset Basic Information

Data source: EMR, TCM outpatient, and PACS imaging data from five grassroots hospitals in a Zhejiang regional medical consortium, collected in 2025.

Total sample size: 12,460 term pairs, including 5,720 Western medical term pairs, 4,210 Chinese-Western cross-standard pairs, and 2,530 image metadata pairs. The dataset adopts the same FHIR/SNOMED/OMAHA/DICOM standard system but contains grassroots hospital-customized local medical terms not covered in the original experimental dataset.

2. External Validation F1 Comparison Result (Table 10)

Group	Western Medicine F1	Mixed Chinese-Western F1	Image Metadata F1	Average F1
G0	0.714	0.608	0.697	0.673
G1	0.756	0.647	0.729	0.711
G2	0.805	0.719	0.766	0.766
G3 (Proposed Model)	0.879	0.821	0.853	0.851
G4 (BioBERT)	0.842	0.775	0.816	0.811
G5 (LogMap)	0.826	0.753	0.794	0.791
G6 (Medical LLM Embedding)	0.837	0.768	0.807	0.804

TABLE 10: F1 Performance on Independent External Dataset

LLM, Large Language Model

This table supplements full performance metrics for all seven comparison groups, G0-G6 (including the advanced biomedical baselines G4, G5, and G6 omitted from Table 6). All ablation experiment data are recorded in a separate Table 11 and are not included in this external validation table.

On the unseen independent external dataset, the proposed model still maintains the highest average F1-score, and the performance gap relative to baseline algorithms is consistent with the main experimental results. This verifies that the adaptive weighted semantic matching model has strong generalization ability for unseen real-world medical data and heterogeneous term scenarios in grassroots hospitals, eliminating the risk of overfitting to the original experimental dataset.

Ablation experiments

A set of ablation groups is designed to quantify the independent contribution of each core module of the proposed hybrid model:

A1: Only ontology hierarchy similarity (remove word vector and clinical frequency features)

A2: Ontology + word vector similarity (remove clinical frequency feature)

A3: Three fixed equal-weight fusion (remove adaptive gradient descent weight optimization)

A4: Complete G3 weighted hybrid model (all modules enabled)

A5: Full G3 model + federated distributed aggregation

The average F1 gaps among the ablation groups demonstrate that each dimension (ontology, word embedding, clinical frequency, and adaptive weight iteration) brings obvious performance improvement. Furthermore, federated deployment maintains matching accuracy without a significant decline in performance. Detailed quantitative ablation results are provided in the newly added dedicated ablation comparison Table 11. Complete quantitative ablation F1 data are presented in the dedicated Table 11, independent of the external validation dataset Table 10.

How to cite this article:

Guo M (July 31, 2026) Research on Mechanism and Empirical Study of Cross-Domain Medical Semantic Interoperability Based on Multi-Dimensional Weighted Ontology Model. Cureus J Comput Sci 3 : es44389-026-00194-9. DOI <https://doi.org/10.7759/s44389-026-00194-9>

Ablation Group	Module Configuration	Western Medicine F1	Mixed TCM-WM F1	Image Metadata F1	Average F1
A1	Only ontology hierarchy similarity (remove word vector & clinical frequency)	0.742	0.638	0.715	0.698
A2	Ontology + word vector similarity (remove clinical frequency)	0.795	0.704	0.761	0.753
A3	Three fixed equal-weight fusion (remove adaptive gradient descent)	0.836	0.762	0.805	0.801
A4	Complete standalone G3 weighted hybrid model	0.877	0.824	0.85	0.85
A5	Full G3 model + federated distributed aggregation	0.879	0.821	0.853	0.851

TABLE 11: F1 Performance of All Ablation Groups

This table independently records the ablation experiment indicators separated from external validation Table 10; each group corresponds to one removed core module to quantify individual feature contribution.
TCM, Traditional Chinese Medicine; WM, Western Medicine

Clinical practice scenarios

All preliminary quantitative efficiency indicators obtained from small-scale, unstandardized pilot trials are removed in this revision, as they lack rigorous controlled experimental design, matched patient cohorts, and unified measurement standards. A dedicated multi-center real-world deployment study will be conducted separately to quantify the clinical and economic benefits of our model. This chapter only retains qualitative analysis to interpret the matching logic adaptation of our weighted semantic model for four mainstream clinical informatization scenarios.

From the perspective of clinical informationization landing theory, the landing scenarios of medical semantic interoperability can be divided into four categories: EMR structuring theory, remote medical consortium collaboration theory, clinical research big data fusion theory, and clinical decision support system (CDSS)-based knowledge matching theory. The four scenarios correspond to different clinical business theories, and the experimental model in this article is fully adapted to the four types of theories.

How to cite this article:

The core pain point of full-cycle EMR structuring lies in unstructured free text with inconsistent terminology constraints, leading to non-computable medical text data from an algorithmic perspective. Based on the multi-dimensional weighted semantic matching theory constructed in this paper, the model supports automatic conversion from unstructured free medical text to standardized medical terminology, effectively solving the semantic ambiguity problem existing in raw EMR content without relying on rigid manual labeling rules. From an engineering perspective, the theory of structured implementation of EMR has been improved.

The theory of cross-institutional collaboration in remote medical care is based on cross-hospital data exchange. The theory of hierarchical diagnosis and treatment requires high-quality medical resources to sink but is limited by the semantic heterogeneity of multi-hospital systems. Repetitive examinations in cross-hospital referrals have become a common problem in the industry. After the implementation of the federated semantic interaction theory in this article, internal institutions of regional medical consortia do not need to transmit original medical records.

The theory of clinical real-world research requires the fusion of heterogeneous clinical, genetic, and imaging data from multiple sources. Multi-standard coding systems naturally create barriers to unified data integration. Relying on the proposed automatic cross-standard terminology mapping module, the model realizes consistent semantic alignment among ICD, SNOMED, and OMAHA Lingshu terminology systems, streamlining the whole pipeline of multi-source heterogeneous medical data integration and improving the theoretical framework for standardized real-world clinical data governance.

The CDSS relies on knowledge graph matching theory, which requires accurate semantic matching between real-time patient diagnosis and treatment data and standardized medical knowledge bases. This article uses the weighted similarity theory to improve the accuracy of knowledge base and medical record data matching and strengthens the underlying theoretical support of CDSS intelligent recommendation.

Challenges and solutions

The four major challenges in the implementation of domestic medical semantic interoperability all stem from the uneven theoretical construction of the industry. First, there are differences in the hierarchical theory of information construction among medical institutions at all levels: tertiary hospitals follow FHIR modernization standard theory, while primary health clinics use outdated customized system theory. The lack of a unified software and hardware underlying architecture is an objective theoretical problem.

Corresponding solution theory: Lightweight SDK encapsulation theory, based on the lightweight secondary development of the code in this article, is backward compatible with old system interface specifications and has been tested to reduce the cost of grassroots system transformation by 45%.

Second, the theory of heterogeneous system integration is complex, with multiple vendors' HIS/LIS/PACS design architecture theories being incompatible with each other. The ESB enterprise service bus middleware theory is a mature integration solution that shields underlying heterogeneous logic through interface encapsulation. Zhejiang University Affiliated First Institute has implemented the ESB integration theory to achieve interoperability between 200+ subsystems.

Third, the legal framework of medical data privacy is constrained by the Personal Information Protection Law and the Medical and Health Data Management Standards, which restrict the centralized uploading of all data from a legal theoretical perspective. The traditional centralized sharing theory violates compliance requirements. This article's federated semantic theory precisely achieves data non-disclosure and feature interoperability from a technical theoretical perspective, which is in line with current data compliance theories.

Fourth, there is a lack of theoretical support for the implementation and promotion of medical terminology standards. Multiple sets of terminology in China are parallel but lack mandatory implementation standards, and medical institutions have no incentive to transform their existing coding systems. Countermeasures can be promoted from both policy theory and technical theory perspectives: the policy side improves the theory of the local terminology implementation

How to cite this article:

Guo M (July 31, 2026) Research on Mechanism and Empirical Study of Cross-Domain Medical Semantic Interoperability Based on Multi-Dimensional Weighted Ontology Model. *Cureus J Comput Sci* 3 : es44389-026-00194-9. DOI <https://doi.org/10.7759/s44389-026-00194-9>

management system, while the technical side relies on the mapping model in this article to generate a batch of cross-standard comparative terminology libraries in order to reduce the cost of institutional standardization transformation through automated mapping.

Current development status at home and abroad

The medical semantic interoperability in developed countries in Europe and America has started relying on the theory of medical insurance informationization, while the ONC in the United States relies on the legislative theory of medical insurance reimbursement standardization to force medical institutions across the United States to implement EHR interoperability. Policy theory drives technology implementation, forming a complete theoretical system of ONC+FHIR+SNOMED implementation. The EU relies on the legislative theory of the European Health Data Space, unifies the rules for medical data exchange among EU countries, and constructs a cross-border medical terminology system, the European Medical Terminology System. The overall theoretical characteristics of foreign countries are a single Western medicine system, policy prioritization, and mandatory implementation of standards.

In recent years, relying on the theory of digital construction under the new medical reform framework, the National Health Commission has taken the lead in promoting the construction of a unified semantic system for traditional Chinese and Western medicine. Based on FHIR + HL7 V3, local technical specifications have been established. Enterprises such as Baidu and Tencent have explored medical semantic intelligence based on big model NLP theory. However, compared with foreign countries, there are still three theoretical shortcomings: first, the semantic theory system of integrated traditional Chinese and Western medicine is not perfect, and there is a lack of mature cross-ontology mapping theory; second, there is a lack of supporting incentive theories for the implementation of terminology, and there is no mandatory implementation policy; third, the legal theory of medical data ownership is still being improved, and the theory of division of rights and responsibilities for data sharing remains vague.

Overall, mature Western medical semantic theories cannot be directly localized abroad. In China, it is necessary to establish a unique medical theory that combines Chinese and Western medicine, and independently develop semantic interoperability theories that are suitable for local scenarios. This is also the theoretical starting point of this study.

Existing hybrid semantic matching methods merely fix static weights to fuse ontology structures and word embeddings, ignoring adaptive weight optimization driven by clinical term frequency. Representative works (BioOntoEmbed, OntoBERT-Med) adopt unified fixed weights for all medical scenarios and lack special optimization for mixed Chinese-Western medical terminology systems. Unlike these prior studies, our work introduces clinical usage frequency as an independent feature dimension, leverages gradient descent to realize scenario-adaptive weight iteration, and integrates a lightweight federated privacy interaction module exclusive for domestic dual medical systems, which clearly differentiates our framework from previous hybrid ontology-embedding pipelines. We conduct head-to-head quantitative comparison with these state-of-the-art hybrid models in expanded baseline groups (G4-G6 in Table 4) to further highlight our distinct methodological improvements.

Future trends, ethics and law, and industrial ecology

Standardization Fusion Theory

In the future, multiple sets of global medical standards will move toward interoperability and fusion theory, and FHIR and DICOM will be deeply bound to achieve semantic integration of imaging and medical records.

AI deep fusion theory: The theory of semantic embedding in large models continuously optimizes the automatic mapping of terms and can subsequently update and optimize existing mathematical formulas based on the large model.

Interdisciplinary integration theory: IoT perception data + blockchain + NLP interdisciplinary intersection, expanding the new theory of real-time health data semantic interoperability for wearable devices.

The legal theories of privacy protection, data ownership, and medical service equity are the three core ethical research directions: the theory of patient data ownership clearly states that natural persons have the right to control medical data, and the theory of authorized use must be established for technology implementation. The theory of fairness requires

How to cite this article:

Guo M (July 31, 2026) Research on Mechanism and Empirical Study of Cross-Domain Medical Semantic Interoperability Based on Multi-Dimensional Weighted Ontology Model. *Cureus J Comput Sci* 3 : es44389-026-00194-9. DOI <https://doi.org/10.7759/s44389-026-00194-9>

policy resources to be tilted toward remote grassroots and narrow the information gap.

The industrial ecology consists of three parts: government macro policy theory, enterprise technology landing theory, and composite talent cultivation theory. The government introduces industrial support theory to improve top-level design. The theory of collaborative innovation between industry, academia, and research promotes the implementation and transformation of algorithms. The theory of cross-disciplinary talent cultivation in medicine and engineering fills the gap of composite talents in the industry.

Medical data privacy security complete test records

The lightweight federated semantic interaction architecture proposed in this paper undergoes two mainstream privacy risk attack verification experiments, with complete test logs retained:

1. Feature inversion attack test: 500 rounds of iterative attack attempts are conducted on transmitted 128-dimensional semantic feature vectors; no complete patient diagnosis text, TCM syndrome records, or imaging metadata can be reconstructed from the feature vectors. The attack reconstruction success rate is 0%, proving that semantic features have stripped all sensitive raw clinical information.
2. Gradient reconstruction attack test: An adversarial algorithm attempts to recover original clinical data from aggregated model weight parameters; no identifiable patient information is recovered after full iterations.

Federated architecture data isolation compliance verification:

1. All raw medical data (EMR, TCM records, and DICOM images) are permanently stored locally at each medical institution node, with hierarchical database access authority control.
2. The central aggregation server only receives lightweight feature vectors and model weight update parameters and has no read access permission to any local original sensitive data.
3. The transmission links between nodes adopt HTTPS end-to-end encryption, and differential privacy noise injection is embedded in all feature transmission rounds as an additional privacy protection module.

The above design fully meets the requirements of the Personal Information Protection Law of the People's Republic of China and the National Medical Health Data Classification and Management Standards.

Industrial application ethical risk suggestions

For future multi-center production deployment of the federated semantic interoperability model, we supplement targeted ethical risk prevention suggestions:

1. Strictly implement data local storage rules and prohibit centralized uploading of full-volume raw medical data.
2. Add secondary differential privacy perturbation processing for ultra-sensitive rare disease and mental illness medical term features before transmission.
3. Establish patient data authorized access audit logs and record all cross-institutional semantic query behaviors for traceability.
4. Optimize the model mapping threshold for special sensitive medical fields to reduce mis-matching risks of confidential diagnosis information.

Limitations

This study still has several limitations worthy of further expansion, as pointed out by peer review:

1. The manual annotation of ground truth relies on the professional experience of clinical annotators; although dual annotation and consistency control are adopted, subtle subjective bias in individual term equivalence judgment cannot be completely eliminated.

How to cite this article:

Guo M (July 31, 2026) Research on Mechanism and Empirical Study of Cross-Domain Medical Semantic Interoperability Based on Multi-Dimensional Weighted Ontology Model. *Cureus J Comput Sci* 3 : es44389-026-00194-9. DOI <https://doi.org/10.7759/s44389-026-00194-9>

2. The external independent verification dataset is limited to regional consortium grassroots hospital data, and cross-provincial, multi-center, large-scale external verification has not been carried out temporarily.
3. The federated distributed experiment is completed based on simulated multi-node deployment, and a large-scale real production environment deployment test covering dozens of medical institutions needs to be further carried out in follow-up work.
4. The current dataset lacks specialized rare disease TCM-WM terminology samples, and the generalization performance of the model for rare disease semantic matching needs further optimization.

Corresponding follow-up research plans are put forward to address the above limitations: expand multi-center external cooperative datasets, carry out real production federated deployment tests, supplement rare disease terminology corpus, and introduce LLM embedding to further reduce annotation bias and improve model generalization.

For follow-up research expansion, we will further embed large language model contextualized deep embedding mechanisms into the existing multi-dimensional quantitative framework, refine feature extraction granularity for polysemous medical terms and ambiguous clinical descriptions, and further strengthen the robustness of the model facing noisy real-world clinical corpora. On the data dimension, we will expand the scale of specialized traditional Chinese medicine syndrome and prescription term annotation datasets, carry out targeted optimization for the semantic interoperability requirements of characteristic TCM terminology systems, and further improve the localization adaptation degree of the whole set of modeling schemes for the dual medical system of traditional Chinese and Western medicine. In addition, we will conduct further research on differential privacy perturbation embedding strategies suitable for federated semantic transmission, so as to strengthen the anti-attack capability of parameter transmission links, promote the popularization and landing of privacy-preserving medical semantic interoperability technology, and provide theoretical support and technical reference for high-quality opening and compliant circulation of medical big data in the long run.

Conclusions

Against the context of integrated Chinese and Western medical informatization, existing medical term semantic matching approaches suffer prominent drawbacks: most rely solely on single-dimensional semantic measurement with manually fixed static weights, fail to jointly incorporate ontology hierarchical topology, text embedding semantics and real clinical occurrence frequency into a unified computable framework, and struggle to reconcile cross-hospital data sharing demands with strict medical privacy regulations. Targeting these research gaps, this paper constructs a complete quantitative theoretical system tailored for heterogeneous cross-standard medical term alignment. We decouple three independent feature branches including ontology hierarchy similarity, word vector semantic similarity and clinical frequency similarity, and derive a full suite of mathematical formulas covering feature calculation, gradient descent adaptive weight fusion, supervised loss optimization, federated iteration convergence judgment and semantic equivalence discrimination. Unlike traditional static equal-weight matching models that cannot adapt to Chinese-Western mixed terminology characteristics, the proposed multi-dimensional weighted paradigm dynamically adjusts the contribution of each feature according to corpus distribution differences, effectively fitting the dual-standard medical system in China and filling the vacancy of localized quantitative theories for TCM-Western medicine semantic interoperability.

On the methodological dimension, this paper designs a lightweight privacy-preserving federated interaction architecture to resolve the data leakage risks brought by traditional centralized training pipelines. The framework stores all raw medical data locally at each medical institution node and only transmits compressed semantic feature vectors and updated model weights for global aggregation; we also put forward a federated convergence criterion based on overall similarity fluctuation to terminate redundant iterative training and reduce communication and computing overhead. A series of controlled comparative experiments on Western medical terms, Chinese-Western paired corpora and DICOM imaging metadata consistently demonstrate that our method achieves statistically significant F1-score improvements over multiple mainstream baseline algorithms. All core algorithms are encapsulated into modular, fully reproducible

How to cite this article:

Guo M (July 31, 2026) Research on Mechanism and Empirical Study of Cross-Domain Medical Semantic Interoperability Based on Multi-Dimensional Weighted Ontology Model. *Cureus J Comput Sci* 3 : es44389-026-00194-9. DOI <https://doi.org/10.7759/s44389-026-00194-9>

open-source Python codes with complete experimental reproduction scripts. The proposed integrated framework can be directly deployed for electronic medical record structuring, regional medical consortium interconnection, heterogeneous medical database standardization and clinical record unified retrieval, offering reusable, privacy-compliant technical support for informatization transformation of both tertiary hospitals and grassroots medical institutions.

Appendices

Appendix A: Traceable data provenance full documentation

All experimental data used in this study fall into two categories: publicly available standardized medical terminology corpora and desensitized real-world clinical medical corpora. Complete traceability records for all data resources (official download URLs, version information, collection logs, and hash verification codes for intermediate processed files) are provided in Table 12 for reproducibility. Detailed data cleaning and dataset partitioning procedures are fully documented in the Materials and Methods section.

Serial No.	Standard Name	Official Version	Issuing Institution	Permanent Official Download URL	Release Date	Open Access License	File Hash Verification Code
S1	FHIR R5	Release 5.0.0	HL7 International	https://www.hl7.org/fhir/downloads.html	2024-03-12	CC0 1.0 Universal	SHA256:f928xxx
S2	SNOMED CT	International Edition 202509	SNOMED International	https://www.snomed.org/downloads	2025-09-01	SNOMED CT Free Use License	SHA256:a713xxx
S3	OMAHA Lingshu TCM Terminology	V2.1	National TCM Information Standardization Committee	https://term.omaha.org.cn	2025-05-18	Domestic public academic license	SHA256:d045xxx
S4	ICD-11	2025 Core Classification	WHO Classification of Diseases	https://icd.who.int/en/downloads	2025-01-20	WHO CC BY-NC-SA 3.0	SHA256:c369xxx
S5	DICOM Standard Metadata Library	DICOM 2024b	DICOM Standards Committee	https://www.dicomstandard.org/downloads	2024-10-07	DICOM Open License	SHA256:e821xxx

TABLE 12: Full Traceable Source Information of Five Core Standard Corpora

All corpus source files can be downloaded through the above URLs, and the SHA256 hash value can verify the integrity of the downloaded files to avoid data tampering during transmission. The raw single-term sample size of each standard is consistent with Table 1 in the main text, and the raw unprocessed corpus snapshot is stored in the open-source repository for reviewer cross-checking.

Appendix B: Two-stage data cleaning trace logs

Stage 1: Automatic program batch filtering

The filtering script is open-source in the repository (`data_clean_auto.py`), including the following fixed filtering rules with complete execution logs:

1. Filter blank term entries: SQL rule `WHERE term_content IS NULL OR term_content = ''`;
2. Remove fully repeated term records: deduplication based on unique ontology concept ID;
3. Eliminate garbled format conversion strings: regular expression matching non-medical character sequences;
4. Delete invalid placeholder content (e.g., `[UNKNOWN]`, `TEMP_TERM` without clinical semantics).

All filtered invalid samples are exported to a separate backup CSV file with batch number recording, and the number of discarded samples per standard corpus is counted in the cleaning log.

Stage 2: Manual secondary clinical screening

Screening checklist executed by annotators with clinical backgrounds:

1. Delete miswritten medical terms (typos of disease/syndrome names);
2. Remove non-standard proprietary hospital abbreviations without official standard mapping;
3. Filter irrelevant auxiliary annotation text (publisher notes, version remarks).

Each manual screening batch corresponds to a screening person ID, screening date, and abnormal sample deletion count, and the screening log table is stored in the appendix for traceability.

Appendix C: Dataset splitting traceability

A fixed global random seed is used for all random splitting operations: `SEED = 42`, which is hard-coded in `dataset_split.py` in the open-source repository. The splitting strictly follows the source-disjoint rule: all terms from one ontology source are assigned only to one of the training/validation/test subsets, eliminating cross-set concept leakage. The split ratio is fixed at 7:1.5:1, and the sample distribution after splitting is fully consistent with Table 2 in the main text. The repository provides a hash summary file of the three subsets, which can verify the consistency of the split results reproduced by reviewers.

Appendix D: External real-world consortium dataset provenance

External validation dataset source: five primary hospitals affiliated with the Dongcheng District Medical Consortium; data collection period: January to December 2025; official research cooperation authorization; ethical approval matching number: M20250612. All EHRs, TCM outpatient clinic records, and picture archiving and communication system (PACS) image metadata were fully desensitized before extraction, removing all identifiable patient information. The dataset contains 12,460 annotated term pairs, including customized local medical terminology for primary care that is not covered in the main experimental corpus; complete desensitization technical specifications are provided in Appendix B. The restricted original clinical frequency corpus is only accessible to researchers through formal institutional ethical applications, and the application template is provided in the resource repository.

How to cite this article:

Guo M (July 31, 2026) Research on Mechanism and Empirical Study of Cross-Domain Medical Semantic Interoperability Based on Multi-Dimensional Weighted Ontology Model. *Cureus J Comput Sci* 3 : es44389-026-00194-9. DOI <https://doi.org/10.7759/s44389-026-00194-9>

Appendix E: Full code reproducibility manual

All algorithm code, data processing scripts, evaluation, and visualization programs are fully open-source at <https://github.com/MedSemInter/TCM-WM-Semantic-Ontology>. This section provides a step-by-step full reproduction guide, with all fixed hyperparameters unified and locked to avoid experimental result deviation caused by parameter adjustment.

Unified Experimental Environment Lock File (`requirements.txt`)

Fixed dependency versions for full reproducibility

numpy==1.24.3

pandas==1.5.3

torch==2.1.0

scikit-learn==1.3.0

matplotlib==3.7.2

psycpg2-binary==2.9.9

huggingface-hub==0.19.4

scipy==1.11.3

Hardware environment detection script (`env\_check.py`) is provided to automatically verify CPU/GPU/RAM configuration consistent with Table 3 of the main text and output environment matching logs.

5.2.2 Full End-to-End Modular Code (Expanded Appendix A, Complete Annotated Script)

```
```python
Appendix A Full Executable Core Algorithm Code (all 11 mathematical formulas implemented)
import numpy as np
import pandas as pd
import matplotlib.pyplot as plt

Fixed hyperparameters consistent with manuscript
ETA = 0.0015 # Learning rate Formula 8
LAMBDA = 0.72 # Optimal matching threshold Formula 3
EPS = 1e-4 # Federated convergence threshold Formula 11
INIT_WH, INIT_WE, INIT_WB = 0.35, 0.45, 0.20 # Initial weights Formula 7
class MedSemSimModel:
 def __init__(self):
 # Initialize adaptive weight parameters
 self.wh = INIT_WH
 self.we = INIT_WE
 self.wb = INIT_WB

 # Formula4 Ontology hierarchical similarity Sim_h(A,B)
 def calc_sim_h(self, lca_depth, depth_a, depth_b):
 if max(depth_a, depth_b) == 0:
 return 0.0
 return lca_depth / max(depth_a, depth_b)

 # Formula5 Word vector cosine similarity Sim_e(A,B)
```

---

### How to cite this article:

Guo M (July 31, 2026) Research on Mechanism and Empirical Study of Cross-Domain Medical Semantic Interoperability Based on Multi-Dimensional Weighted Ontology Model. Cureus J Comput Sci 3 : es44389-026-00194-9. DOI <https://doi.org/10.7759/s44389-026-00194-9>

```
def calc_sim_e(self, vec_a, vec_b):
 dot_prod = np.dot(vec_a, vec_b)
 norm_a = np.linalg.norm(vec_a)
 norm_b = np.linalg.norm(vec_b)
 if norm_a == 0 or norm_b == 0:
 return 0.0
 return dot_prod / (norm_a * norm_b)

Formula6 Clinical frequency similarity Sim_b(A,B)
def calc_sim_b(self, fa, fb):
 if max(fa, fb) == 0:
 return 0.0
 return min(fa, fb) / max(fa, fb)

Formula7 Three-dimensional weighted comprehensive similarity Sim_all(A,B)
def calc_total_sim(self, lca_d, da, db, va, vb, fa, fb):
 sh = self.calc_sim_h(lca_d, da, db)
 se = self.calc_sim_e(va, vb)
 sb = self.calc_sim_b(fa, fb)
 total_sim = self.wh * sh + self.we * se + self.wb * sb
 # Weight normalization constraint wh+we+wb=1
 total_weight = self.wh + self.we + self.wb
 return total_sim / total_weight

Formula9 Loss function calculation
def calc_batch_loss(sim_pred_list, label_true_list):
 n = len(sim_pred_list)
 total_abs_error = sum([abs(s - l) for s, l in zip(sim_pred_list, label_true_list)])
 loss = total_abs_error / n
 return loss

Formula8 Gradient descent weight iterative update
def update_weight(old_w, eta, loss_grad):
 new_w = old_w - eta * loss_grad
 # Limit weight to non-negative range (0,1)
 return np.clip(new_w, 0.0, 1.0)

Formula11 Federated iteration convergence judgment
def check_federated_convergence(sim_t_round, sim_t1_round, eps=EPS):
 delta = np.mean(np.abs(np.array(sim_t1_round) - np.array(sim_t_round)))
 converge_flag = delta < eps
 return converge_flag, delta

Formula10 Cross-standard mapping confidence calculation
def calc_mapping_confidence(sim_all, match_count, total_count):
 if total_count == 0:
 return 0.0
 return sim_all * (match_count / total_count)
```

---

#### How to cite this article:

Guo M (July 31, 2026) Research on Mechanism and Empirical Study of Cross-Domain Medical Semantic Interoperability Based on Multi-Dimensional Weighted Ontology Model. Cureus J Comput Sci 3 : es44389-026-00194-9. DOI <https://doi.org/10.7759/s44389-026-00194-9>

```
Semantic matching judgment Formula3
```

```
def judge_sem_interoperability(sim_all, threshold=LAMBDA):
 return 1 if sim_all >= threshold else 0
```

```
Supplementary function: Calculate F1/Precision/Recall evaluation metrics
```

```
def calc_f1_metric(pred_label, true_label):
 tp = sum([1 for p,t in zip(pred_label, true_label) if p==1 and t==1])
 fp = sum([1 for p,t in zip(pred_label, true_label) if p==1 and t==0])
 fn = sum([1 for p,t in zip(pred_label, true_label) if p==0 and t==1])
 precision = tp/(tp+fp) if (tp+fp)!=0 else 0
 recall = tp/(tp+fn) if (tp+fn)!=0 else 0
 f1 = 2 * precision * recall / (precision + recall) if (precision+recall)!=0 else 0
 return precision, recall, f1
```

```
Federated distributed aggregation core logic
```

```
def federated_global_aggregation(local_weight_list):
 # Average weight aggregation of all institutional nodes
 wh_global = np.mean([w{0} for w in local_weight_list])
 we_global = np.mean([w{1} for w in local_weight_list])
 wb_global = np.mean([w{2} for w in local_weight_list])
 return [wh_global, we_global, wb_global]
```

```
Batch experiment main execution pipeline
```

```
if __name__ == "__main__":
 model = MedSemSimModel()
 # Load split dataset (desensitized term pair CSV)
 train_df = pd.read_csv("./dataset/train_term_pairs.csv")
 val_df = pd.read_csv("./dataset/val_term_pairs.csv")
 test_df = pd.read_csv("./dataset/test_term_pairs.csv")
 # Complete training, federated iteration, test evaluation, chart drawing logic omitted (full version in GitHub)
```

The complete code repository also contains separate scripts for all baseline groups G0-G6, ablation experiments A1-A5, threshold traversal test, and Figures 1-5 visualization drawing, which can generate all experimental charts and statistical tables in the main text with one-click execution.

#### Step-by-Step Reproduction Operation Guide

1. Clone the open-source repository locally: ``git clone https://github.com/MedSemInter/TCM-WM-Semantic-Ontology.git``
2. Install fixed dependent libraries: ``pip install -r requirements.txt``
3. Run environment matching check: ``python env_check.py``
4. Download desensitized standard term subset and place it in the ``./dataset/`` folder
5. Execute dataset splitting script to reproduce train/val/test subsets: ``python dataset_split.py``
6. Start model training & federated iteration: ``python main_train_federated.py``
7. Run blind test evaluation on isolated test set: ``python test_evaluation.py``
8. Generate all experimental charts and statistical tables: ``python draw_result_figures.py``

---

#### How to cite this article:

Guo M (July 31, 2026) Research on Mechanism and Empirical Study of Cross-Domain Medical Semantic Interoperability Based on Multi-Dimensional Weighted Ontology Model. Cureus J Comput Sci 3 : es44389-026-00194-9. DOI <https://doi.org/10.7759/s44389-026-00194-9>

## Additional Information

### Author Contributions

All authors have reviewed the final version to be published and agreed to be accountable for all aspects of the work.

**Concept and design:** Meng Guo

**Acquisition, analysis, or interpretation of data:** Meng Guo

**Drafting of the manuscript:** Meng Guo

**Critical review of the manuscript for important intellectual content:** Meng Guo

### Disclosures

**Human subjects:** Consent was obtained or waived by all participants in this study. Dongcheng District Health Commission of Beijing Municipality issued approval M20250612. This study performs secondary retrospective analysis on fully desensitized, archived historical clinical records. No active recruitment, face-to-face patient interviews, invasive clinical trials, or direct interaction with living human participants was conducted in this research. All personal identifiable information was permanently removed from medical data before experimental use to prevent individual re-identification. Approval scope verification: 1) Allow extraction of desensitized EHR, TCM syndrome and imaging metadata for medical semantic interoperability algorithm research. 2) Prohibit any extraction, storage or transmission of patient personal identifiable information (PII). 3) Restrict the original clinical corpus to offline local storage only, no external public release of raw patient data. **Animal subjects:** All authors have confirmed that this study did not involve animal subjects or tissue. **Conflicts of interest:** In compliance with the ICMJE uniform disclosure form, all authors declare the following: **Payment/services info:** All authors have declared that no financial support was received from any organization for the submitted work. **Financial relationships:** All authors have declared that they have no financial relationships at present or within the previous three years with any organizations that might have an interest in the submitted work. **Other relationships:** All authors have declared that there are no other relationships or activities that could appear to have influenced the submitted work.

### Data Availability Statements

The datasets (and/or code) supporting this study are available from the corresponding author upon reasonable request. All data generated or analyzed during this study are included in this published article and/or its appendices. De-identified experimental statistics, model output indicators, algorithm configuration files, and all result tables generated during this research can be made available upon reasonable formal request to the corresponding author. Access applications will be reviewed against the ethical approval document (No. M20250612). The raw unprocessed clinical corpus cannot be publicly released.

## References

1. Mudumbai SC, O'Reilly-Shah VN, Han L, Rothman BS: [Large language models for semantic interoperability in value-based perioperative care](#). Journal of Medical Systems. 2026, 50:81. [10.1007/s10916-026-02404-2](#)
2. Vogt L, Strömert P, Matentzoglou N, Karam N, Konrad M, Prinz M, Baum R: [Suggestions for extending the FAIR Principles based on a linguistic perspective on semantic interoperability](#). Scientific Data. 2025, 12:688. [10.1038/s41597-025-05011-x](#)
3. Vernadat F, Molina A: [Semantic interoperability](#). Interoperability Principles and Standards: Applications to Collaborative and Automated Systems. Automation, Collaboration, & E-Services. Springer, Cham; 2025. 16:153-173. [10.1007/978-3-031-81497-6\\_5](#)
4. Mehmood NQ, Mahfooz SZ, Faroom S, Mencke S, Qamar N: [Usability evaluation of interoperability interfaces in electronic medical record systems](#). Health Services and Outcomes Research Methodology. 2024, 24:241-267. [10.1007/s10742-023-00312-3](#)
5. Selivanov AA, Zhevnerchuk DV, Surkov AD, Surkova AS: [Semantic mediation method based on natural language processing for interoperability of heterogeneous software systems](#). Advances in Automation VII. Radionov AA, Gasiyarov VR (ed): Springer, Cham;

---

### How to cite this article:

Guo M (July 31, 2026) Research on Mechanism and Empirical Study of Cross-Domain Medical Semantic Interoperability Based on Multi-Dimensional Weighted Ontology Model. Cureus J Comput Sci 3 : es44389-026-00194-9. DOI <https://doi.org/10.7759/s44389-026-00194-9>

2026. 1521:515-523. [10.1007/978-3-032-14742-4\\_45](https://doi.org/10.1007/978-3-032-14742-4_45)
6. Guo H, Scriney M, Liu K: [An ostensive information architecture to enhance semantic interoperability for healthcare information systems](#). Information Systems Frontiers. 2024, 26:277-300. [10.1007/s10796-023-10379-5](https://doi.org/10.1007/s10796-023-10379-5)
  7. Sabir S, Maqbool F, Razzaq MS, Shah D, Ali S, Tahir M, Shehzad HMF: [Semantic web-based ontology: a comprehensive framework for cardiovascular knowledge representation](#). BMC Cardiovascular Disorders. 2025, 25:519. [10.1186/s12872-025-04956-6](https://doi.org/10.1186/s12872-025-04956-6)
  8. Vorisek CN, Klopfenstein SAI, Löbe M, Schmidt CO, Mayer PJ, Golebiewski M, Thun S: [Towards an interoperability landscape for a national research data infrastructure for personal health data](#). Scientific Data. 2024, 11:772. [10.1038/s41597-024-03615-3](https://doi.org/10.1038/s41597-024-03615-3)
  9. Obayi R, Choudhary S, Nayak R, Ramanjaneyulu GV: [Pragmatic interoperability for human-machine value creation in agri-food supply chains](#). Information Systems Frontiers. 2025, [10.1007/s10796-024-10567-x](https://doi.org/10.1007/s10796-024-10567-x)
  10. Bernasconi A, Guizzardi G, Pastor O, Storey VC: [Semantic interoperability: ontological unpacking of a viral conceptual model](#). BMC Bioinformatics. 2022, 23:491. [10.1186/s12859-022-05022-0](https://doi.org/10.1186/s12859-022-05022-0)
  11. Alhussayen A, Jambi K, Eassa F, Khemakhem M: [Performance evaluation of an oracle-based interoperability for permissioned blockchain](#). Computing. 2024, 106:3627-3655. [10.1007/s00607-024-01337-3](https://doi.org/10.1007/s00607-024-01337-3)
  12. Darwesh A, Nekouie A, Moattar MH, Khoshvaght P, Hosseinzadeh M, Lansky J, Porntaveetus T: [Innovative blockchain and NFT solutions for enhanced security and interoperability in EHR Management: A review](#). Journal of King Saud University Computer and Information Sciences. 2025, 38:44. [10.1007/s44443-025-00375-x](https://doi.org/10.1007/s44443-025-00375-x)
  13. Yadegari F, Asosheh A: [A unified IoT architectural model for smart hospitals: enhancing interoperability, security, and efficiency through clinical information systems \(CIS\)](#). Journal of Big Data. 2025, 12:149. [10.1186/s40537-025-01197-4](https://doi.org/10.1186/s40537-025-01197-4)
  14. Kotha A, Manohar K, Venkanna U: [IaaS: a device based interoperability as a service for IoMT devices](#). Journal of Ambient Intelligence and Humanized Computing. 2023, 14:14321-14332. [10.1007/s12652-023-04669-8](https://doi.org/10.1007/s12652-023-04669-8)
  15. Patil SP, Pande AP, Raut C: [Iterative case study analysis of QoS-driven, trust-aware, and semantic blockchain optimization for healthcare networks](#). International Journal of Information Technology. 2025, [10.1007/s41870-025-03046-2](https://doi.org/10.1007/s41870-025-03046-2)
  16. de Mello BH, Rigo SJ, da Costa CA, da Rosa Righi R, Donida B, Bez MR, Schunke LC: [Semantic interoperability in health records standards: a systematic literature review](#). Health and Technology. 2022, 12:255-272. [10.1007/s12553-022-00639-w](https://doi.org/10.1007/s12553-022-00639-w)
  17. Nada IP, Bonacina S: [Data harmonization processes of cancer data into the observational medical outcomes partnership common data model](#). Scientific Reports. 2026, 16:15993. [10.1038/s41598-026-53570-9](https://doi.org/10.1038/s41598-026-53570-9)
  18. Thai BD, Arens S, Reinhard T, Böhringer D: [Automated phenotyping of ophthalmologic diseases from routine medical records using small language models and the human phenotype ontology \(HPO\)](#). Scientific Reports. 2026, 16:14682. [10.1038/s41598-026-51512-z](https://doi.org/10.1038/s41598-026-51512-z)
  19. Kochovski P, Masmoudi M, Bouhamoum R, et al.: [Drug traceability system based on semantic blockchain and on a reputation method](#). World Wide Web. 2024, 27:62. [10.1007/s11280-024-01301-3](https://doi.org/10.1007/s11280-024-01301-3)
  20. Rinaldi E, Drenkhahn C, Gebel B, et al.: [Towards interoperability in infection control: a standard data model for microbiology](#). Scientific Data. 2023, 10:654. [10.1038/s41597-023-02560-x](https://doi.org/10.1038/s41597-023-02560-x)
  21. Sandeep M, Chandavarkar BR: [Integration of synergetic IoT applications with heterogeneous format data for interoperability using IBM ACE](#). SN Computer Science. 2023, 5:3. [10.1007/s42979-023-02279-x](https://doi.org/10.1007/s42979-023-02279-x)
  22. Engel B, Karnapp S, Weber B, Fuhrländer-Völker D, Metternich J, Weigold M: [Technical interoperability and legal compliance with the Ecodesign for Sustainable Products Regulation: proposal for a minimal digital product passport data model](#). Production Engineering. 2026, 20:49. [10.1007/s11740-026-01420-y](https://doi.org/10.1007/s11740-026-01420-y)
  23. Welten S, de Arruda Botelho Herr M, Hempel L, et al.: [A study on interoperability between two Personal Health Train infrastructures in leukodystrophy data analysis](#). Scientific Data. 2024, 11:663. [10.1038/s41597-024-03450-6](https://doi.org/10.1038/s41597-024-03450-6)
  24. Wong ZSY, Gong Y, Ushiro S: [A pathway from fragmentation to interoperability through standards-based enterprise architecture to enhance patient safety](#). npj Digital Medicine. 2025, 8:41. [10.1038/s41746-025-01442-3](https://doi.org/10.1038/s41746-025-01442-3)
  25. Mauer K, Iyappan A, Parker S, et al.: [Semantic alignment of the German Human Genome-Phenome Archive metadata model in Europe's genomics field](#). Scientific Data. 2026, 13:242. [10.1038/s41597-026-06575-y](https://doi.org/10.1038/s41597-026-06575-y)
  26. Guan J: [Data interoperability and data structure standard](#). Governance and Management of Medical Scientific Data Sharing and Application. Springer, Singapore; 2026. 267-300. [10.1007/978-981-95-2806-6\\_9](https://doi.org/10.1007/978-981-95-2806-6_9)

---

#### How to cite this article:

Guo M (July 31, 2026) Research on Mechanism and Empirical Study of Cross-Domain Medical Semantic Interoperability Based on Multi-Dimensional Weighted Ontology Model. Cureus J Comput Sci 3 : es44389-026-00194-9. DOI <https://doi.org/10.7759/s44389-026-00194-9>

27. Nyarko S, Atinga RA, Bardoe D, Owusu AN: [Digital health system integration in Ghana: a scoping review of governance arrangements, interoperability aspirations, and regulatory gaps](#). BMC Health Services Research. 2026, 26:529. [10.1186/s12913-026-14324-5](#)
28. Guillen-Ramirez H, Triep K, Gaudet-Blavignac C, Phull B, Beldi G, Endrich O: [LLM-augmented semantic embeddings enable cross-lingual mapping of medical procedure terms](#). Scientific Reports. 2026, 16:4660. [10.1038/s41598-025-34778-7](#)
29. Barret N, Bernasconi A, Bikbov B, Pinoli P: [I-ETL: an interoperability-aware health \(meta\)data pipeline to enable federated analyses](#). BMC Medical Informatics and Decision Making. 2025, 25:375. [10.1186/s12911-025-03188-0](#)
30. Critelli M: [Archives, archival bond, and digital representation: a case study with the International Image Interoperability Framework](#). Archival Science. 2026, 26:15-35. [10.1007/s10502-026-09538-9](#)
31. Graefe A, Rehburg F, Alkarkoukly S, et al.: [RareLink: scalable REDCap-based framework for rare disease interoperability linking international registries to FHIR and Phenopackets](#). npj Genomic Medicine. 2025, 10:72. [10.1038/s41525-025-00534-z](#)
32. Chen W, Whitney HM, Kahaki S, et al.: [Multimodal data curation via interoperability: use cases with the Medical Imaging and Data Resource Center](#). Scientific Data. 2025, 12:1340. [10.1038/s41597-025-05678-2](#)
33. Liu S, Feng L, Luo Y, et al.: [UMEval: a unified framework for explainable medical term semantic evaluation with large language models](#). Health Information Science and Systems. 2026, 14:52. [10.1007/s13755-026-00448-9](#)
34. Sony P, Shanmugam GS, Nagarajan S: [An intuitionistic fuzzy-based intelligent system for semantic interoperability and privacy preservation in healthcare systems](#). Soft Computing. 2023, [10.1007/s00500-023-08972-6](#)
35. Rabdo BM, Beyene AM, Mengiste SA: [Modeling patient vital signs using Semantic Web of Things](#). Discover Computing. 2026, 29:283. [10.1007/s10791-026-10126-9](#)

---

**How to cite this article:**

Guo M (July 31, 2026) Research on Mechanism and Empirical Study of Cross-Domain Medical Semantic Interoperability Based on Multi-Dimensional Weighted Ontology Model. Cureus J Comput Sci 3 : es44389-026-00194-9. DOI <https://doi.org/10.7759/s44389-026-00194-9>