

ARTIFICIAL INTELLIGENCE–DRIVEN MULTIMODAL FRAMEWORK FOR DRUG RECOMMENDATION AND DELIVERY PREDICTION USING CLINICAL AND DRUG INFORMATION

Shakila Shaikh¹(ORCID: 0000-0002-8719-8595), Mohini Reddy^{2*}(ORCID: 0009-0004-0962-7368), Dhanashree Kulkarni³(ORCID: 0009-0007-3643-6841)

Mukesh Patel School of Technology Management & Engineering, SVKM's NMIMS, Mumbai, India

¹Email: shakilashaikh48@gmail.com

²Email: mohini.reddy5@gmail.com (Corresponding author)

³Email: Dhanashree.huddedar@nmims.edu

Abstract—Drug recommendation and healthcare decision support require the integration of heterogeneous data sources including drug characteristics, clinical feedback, patient reviews, and therapeutic information. Conventional prediction approaches frequently rely on single-source information and may not adequately capture complex relationships influencing recommendation outcomes. This study proposes a multimodal machine learning framework for drug recommendation prediction through the integration of clinical reviews and drug-related information. Multiple publicly available datasets containing drug composition, therapeutic uses, side effects, manufacturer information, clinical reviews, and patient feedback were collected and fused into a unified multimodal dataset. The preprocessing pipeline included data cleaning, missing value handling, feature normalization, drug name standardization, multimodal fusion, and duplicate removal. Textual information was transformed using Term Frequency–Inverse Document Frequency (TF–IDF) and transformer-based semantic embeddings, while numerical clinical attributes were encoded and integrated using feature fusion techniques. Dimensionality reduction through Singular Value Decomposition (SVD) was additionally applied to improve representation efficiency. Several machine learning and deep learning approaches were evaluated, including Random Forest, XGBoost, SVD-enhanced XGBoost, weighted binary classification, and BERT-based multimodal classification. Experimental evaluation was performed using accuracy, precision, recall, and F1-score. Comparative analysis demonstrated that transformer-enhanced multimodal learning improved predictive performance, achieving a maximum classification accuracy of 68.14%, while the SVD-enhanced framework achieved the strongest balanced performance with an F1-score of 54.95%. Weighted learning further improved minority-class detection capability. The findings indicate that multimodal integration of clinical and drug information can support intelligent healthcare analytics and provide a foundation for future AI-assisted drug recommendation and decision-support systems.

Keywords—Multimodal Learning, Drug Recommendation, Machine Learning, Deep Learning, BERT, XGBoost, Healthcare Analytics, Feature Fusion, Clinical Reviews.

1. INTRODUCTION

The rapid growth of healthcare data and the increasing complexity of treatment selection have created a demand for intelligent systems capable of supporting clinical decision-making and improving patient outcomes. Drug recommendation and healthcare analytics are becoming increasingly important due to the large volume of drug-related information generated from clinical practice, patient experiences, pharmaceutical databases, and online medical platforms. Traditional drug selection approaches often depend on manual interpretation of clinical evidence and physician expertise, which may limit scalability and efficient utilization of available healthcare information.

Recent advances in Artificial Intelligence (AI), Machine Learning (ML), and Deep Learning (DL) have enabled automated analysis of large-scale healthcare datasets and contributed to the development of intelligent decision-support systems. Machine learning techniques have demonstrated promising performance in disease prediction, clinical risk assessment, drug discovery, and personalized healthcare applications. In particular, multimodal learning has emerged as an effective approach for integrating heterogeneous information sources such as textual clinical records, structured patient information, drug properties, and user feedback into a unified predictive framework [1–4].

Drug recommendation prediction remains a challenging task because healthcare information is highly heterogeneous and originates from multiple modalities. Drug characteristics such as composition, therapeutic usage, side effects, and manufacturer information contain valuable pharmaceutical knowledge, while clinical reviews and patient feedback provide real-world evidence regarding treatment effectiveness and user satisfaction. Existing studies frequently focus on single-source datasets and therefore may fail to capture comprehensive relationships among drug attributes and clinical outcomes [5].

Multimodal learning offers an opportunity to address this limitation by combining complementary information from multiple data sources. Feature fusion methods, transformer-based embeddings, and advanced machine learning algorithms can improve representation quality and support more accurate healthcare predictions. Furthermore, dimensionality reduction techniques help reduce redundancy and improve computational efficiency when processing high-dimensional textual information [6].

Motivated by these challenges, this study proposes a multimodal machine learning framework for drug recommendation prediction through the integration of clinical reviews and drug information collected from multiple publicly available datasets. The proposed framework combines drug composition, therapeutic usage, side effects, manufacturer information, clinical reviews, and user feedback into a unified analytical pipeline. Data preprocessing, feature standardization, multimodal fusion, and dimensionality reduction techniques are incorporated to improve feature representation. Multiple predictive approaches including Random Forest, XGBoost, Singular Value Decomposition (SVD)-enhanced learning, weighted classification, and transformer-based embeddings are evaluated and compared [7].

The major contributions of this work are summarized as follows:

1. Development of a multimodal dataset by integrating heterogeneous drug and clinical information from multiple public sources.
2. Design of a complete preprocessing and feature fusion pipeline for healthcare prediction tasks.
3. Comparative evaluation of traditional machine learning and transformer-based approaches including Random Forest, XGBoost, SVD-enhanced models, and BERT embeddings.
4. Analysis of multimodal learning effectiveness using accuracy, precision, recall, and F1-score evaluation metrics.
5. Demonstration of the potential of multimodal artificial intelligence for intelligent drug recommendation and healthcare decision-support applications.

The remainder of this paper is organized as follows: Section 2 presents materials and methods, Section 3 describes the proposed methodology and experimental setup, Section 4 discusses the obtained results and ablation analysis, and Section 5 concludes the study with future research directions.

II. LITERATURE REVIEW

Artificial Intelligence (AI), Machine Learning (ML), and Deep Learning (DL) have increasingly transformed healthcare analytics and intelligent decision-support applications by enabling automated analysis of large-scale medical and pharmaceutical data. Several studies have demonstrated that ML algorithms such as Random Forest, Support Vector Machine, Decision Tree, and Gradient Boosting can effectively support prediction tasks using structured healthcare information and improve decision-making performance compared with conventional statistical approaches [1–3]. However, these approaches generally depend on manual feature engineering and often exhibit limitations when handling heterogeneous healthcare information collected from multiple sources.

With the advancement of deep learning, researchers introduced neural-network-based approaches to improve healthcare prediction and clinical analytics. Architectures including Convolutional Neural Networks (CNN), Long Short-Term Memory (LSTM), and transformer-based methods have demonstrated improved capability in extracting complex patterns from textual and clinical datasets [4–6]. In particular, transformer models such as Bidirectional Encoder Representations from Transformers (BERT) introduced contextual language understanding and achieved strong performance in medical text mining and healthcare prediction applications [7].

Recent studies have highlighted the importance of multimodal learning for integrating heterogeneous healthcare information. Multimodal frameworks combine structured attributes, textual descriptions, clinical feedback, and pharmaceutical information into unified predictive models and often outperform single-modal approaches [8–10]. Feature fusion and dimensionality reduction methods such as Singular Value Decomposition (SVD) have additionally been adopted to improve computational efficiency and reduce redundancy in high-dimensional representations [11].

Although existing research demonstrates promising outcomes, several limitations remain. Many published studies rely on isolated datasets or focus only on a single healthcare modality, limiting their ability to model real-world clinical complexity. In addition, relatively few studies integrate drug composition, therapeutic information, side effects, and patient reviews into a unified predictive framework while simultaneously comparing traditional machine learning and transformer-based methods [12–13].

A. Research Gap

From the existing literature, the following are the important gaps that can be identified:

- Lack of integration of satellite data, rainfall, and ground-based information at fine spatial scales
- Over-reliance on either physical-based models or data-based models
- Lack of physically consistent machine learning models for long-term groundwater forecasting
- Lack of focus on district-level risk assessment and decision-support systems

B. Contribution of this work

To address these gaps, the present study proposes a unified framework that:

- Existing drug recommendation studies primarily rely on single-modal datasets, limiting their ability to capture relationships among heterogeneous healthcare information.

- Most published approaches focus on either clinical records or drug information independently, with limited integration of both data sources.
- Previous research provides limited utilization of multimodal feature fusion techniques for combining drug composition, therapeutic usage, side effects, and patient reviews.
- Several existing methods depend heavily on traditional machine learning models without extensive comparison with transformer-based approaches such as BERT.
- Current studies often lack effective handling of high-dimensional textual information through dimensionality reduction techniques.
- Comparative evaluation among Random Forest, XGBoost, SVD-enhanced learning, weighted classification, and transformer-based models remains limited in drug recommendation prediction tasks.
- There is insufficient investigation of integrated multimodal frameworks capable of supporting intelligent drug recommendation and healthcare decision support.
- Existing literature provides limited evidence regarding the contribution of individual modalities through ablation analysis.
- Many studies report predictive performance but provide limited discussion on feature interaction across heterogeneous healthcare data sources.
- Therefore, there remains a need for a multimodal machine learning framework that integrates clinical reviews and drug information while performing comprehensive comparative evaluation across multiple AI techniques for drug recommendation prediction.

III. DATA SOURCES AND PREPROCESSING

A. Data Sources

This study utilized four publicly available drug-related and clinical datasets collected from online repositories to construct a unified multimodal dataset for drug recommendation prediction. The selected datasets contain heterogeneous information including drug composition, therapeutic usage, side effects, clinical reviews, patient feedback, and pharmaceutical attributes.

The first dataset consisted of detailed drug information containing medicine names, composition, therapeutic uses, side effects, manufacturer details, and review statistics. This dataset contributed pharmaceutical and therapeutic knowledge required for drug representation. The second dataset contained drug side effects and medical condition information including generic names, drug classes, medical conditions, ratings, and related pharmaceutical attributes. This dataset provided supplementary clinical and drug-specific characteristics.

The third dataset included WebMD drug reviews and clinical feedback collected from patients and users. The dataset contained variables such as age, condition, effectiveness, satisfaction, sex, and textual reviews that contributed clinical and behavioural information. The fourth dataset consisted of drug review records containing drug names, medical conditions, patient reviews, ratings, and usefulness scores. This dataset supported textual analysis and recommendation prediction. The datasets were integrated to form a multimodal healthcare dataset that combines structured and unstructured information for predictive modelling [15-18].

B. Data Preprocessing

Data preprocessing was performed to improve data quality and prepare heterogeneous information for machine learning analysis. Initially, all datasets were loaded independently and standardized by converting column names into a uniform format. Duplicate records were removed to reduce redundancy and ensure consistency across integrated sources.

Missing values were handled using column-specific strategies. Textual attributes containing missing information were replaced with default values such as “unknown” or empty strings, while numerical variables were imputed using predefined values where appropriate. Drug names were standardized through normalization procedures including lowercasing, removal of dosage-related information, and elimination of unnecessary textual variations to improve matching across datasets [19]. After cleaning, multimodal data integration was performed through drug-based matching using standardized drug identifiers. Feature selection was applied to retain relevant attributes including composition, therapeutic usage, side effects, manufacturer information, medical condition, clinical reviews, effectiveness measures, ratings, and user feedback.

Textual attributes were transformed into machine-readable representations using Term Frequency–Inverse Document Frequency (TF–IDF) and transformer-based embeddings. Numerical features including effectiveness scores, ratings, satisfaction, and usefulness counts were encoded and normalized prior to model training. Dimensionality reduction using Singular Value Decomposition (SVD) was additionally applied to reduce feature sparsity and improve computational efficiency.

The final processed dataset was used for multimodal machine learning and deep learning experiments including Random Forest, XGBoost, weighted classification, and BERT-enhanced prediction models.

TABLE 1. SUMMARY OF DATA SOURCES

Dataset	Main Attributes	Purpose
Medicine Details	Composition, Uses, Side Effects, Manufacturer	Drug Information
Drug Side Effects	Medical Condition, Drug Class, Ratings	Clinical Context

WebMD Reviews	Age, Effectiveness, Reviews, Satisfaction	Clinical Feedback
Drug Reviews	Condition, Review, Rating	User Experience

C. Data Preprocessing Pipeline

The collected datasets originated from multiple heterogeneous sources and therefore required preprocessing before model development. A structured preprocessing pipeline was implemented to improve data quality, reduce inconsistencies, and prepare multimodal information for predictive analysis. Initially, all datasets were imported and inspected to identify differences in structure, attribute naming, and data completeness. Column names were standardized into a uniform format and duplicate records were removed to improve consistency.

Missing value handling was performed using attribute-specific strategies. Missing textual attributes such as reviews, side effects, and clinical descriptions were replaced using default placeholder values, while numerical variables including ratings and usefulness measures were imputed using predefined values to ensure complete records for analysis.

Drug information from different datasets showed inconsistencies due to dosage descriptions and naming variations. Therefore, drug normalization was performed through lowercasing, removal of dosage-related expressions, and generation of standardized drug identifiers. Multimodal dataset fusion was then applied using drug-based matching to combine pharmaceutical attributes, clinical reviews, patient feedback, and therapeutic information into a unified dataset.

Feature selection was performed to retain relevant variables for prediction. Textual features including composition, therapeutic usage, side effects, and reviews were transformed using TF-IDF and BERT embeddings, while numerical attributes were encoded and normalized. Singular Value Decomposition (SVD) was applied to reduce feature sparsity and improve efficiency.

Finally, the processed multimodal feature space was divided into training and testing subsets and evaluated using Random Forest, XGBoost, weighted classification, and BERT-enhanced prediction models.

IV. SYSTEM ARCHITECTURE AND PROPOSED METHODOLOGY

The proposed framework presents a multimodal machine learning architecture designed for drug recommendation prediction through the integration of heterogeneous drug-related and clinical information. The methodology combines pharmaceutical attributes, clinical reviews, patient feedback, and numerical healthcare indicators into a unified predictive pipeline to improve representation quality and model performance.

The proposed system consists of five major stages: data acquisition, preprocessing, multimodal feature generation, dimensionality reduction, and predictive modeling. Initially, drug-related and clinical datasets collected from multiple public sources are imported and integrated into a common analytical environment. The collected data include drug composition, therapeutic usage, side effects, manufacturer information, clinical reviews, ratings, and user feedback.

During preprocessing, data cleaning, duplicate removal, missing value handling, and drug name normalization are performed to standardize heterogeneous information across datasets. Drug identifiers are normalized to support dataset matching and improve integration quality. The processed datasets are then fused using multimodal integration techniques to generate a unified feature repository.

After integration, relevant features are selected and separated into textual and numerical modalities. Textual attributes including composition, uses, side effects, and review information are transformed using TF-IDF and BERT-based embedding methods to generate semantic representations. Numerical attributes such as effectiveness, satisfaction, ratings, and usefulness measures are encoded and normalized to support machine learning analysis.

To reduce feature sparsity and improve computational efficiency, Singular Value Decomposition (SVD) is applied before model training. The generated multimodal feature space is subsequently divided into training and testing subsets and evaluated using multiple predictive approaches including Random Forest, XGBoost, weighted classification, and BERT-enhanced XGBoost models.

Finally, prediction outputs are evaluated using accuracy, precision, recall, and F1-score metrics to determine the effectiveness of each learning strategy. Comparative analysis is conducted to identify the most suitable multimodal configuration for drug recommendation prediction.

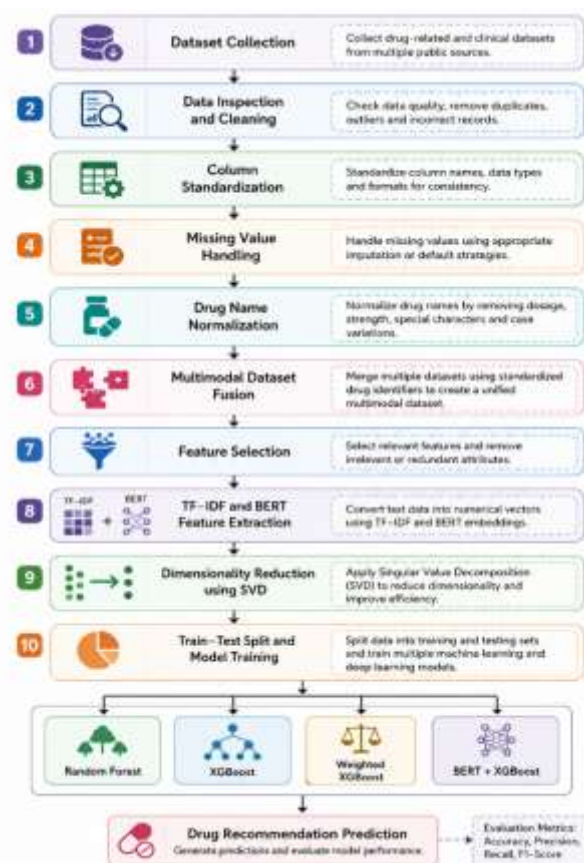


Fig. 1. **Proposed Multimodal Drug Recommendation Prediction Framework.**

V. EXPERIMENTAL SETUP

A. Experimental Environment

All experiments were conducted using the Kaggle Notebook environment to implement and evaluate the proposed multimodal machine learning framework. The implementation was developed using Python and executed with commonly used data analysis and machine learning libraries. Data preprocessing, feature extraction, dimensionality reduction, and model training were performed in an integrated computational workflow.

The experimental environment included libraries such as Pandas and NumPy for data manipulation, Scikit-learn for preprocessing and machine learning implementation, XGBoost for gradient boosting classification, and Sentence Transformers for generating transformer-based semantic embeddings. Sparse matrix operations and dimensionality reduction were performed using SciPy and Truncated Singular Value Decomposition (SVD).

The system was designed to process multimodal information consisting of textual and numerical attributes collected from multiple drug-related and clinical datasets.

B. Dataset Configuration

After preprocessing and multimodal fusion, the final integrated dataset contained 8,928 samples used for model development and evaluation. The dataset consisted of textual features including composition, therapeutic usage, side effects, and reviews, along with numerical features such as effectiveness, satisfaction, ratings, and usefulness measures. The processed feature space was divided into training and testing subsets using an 80:20 split ratio while maintaining class distribution through stratified sampling.

TABLE 2. DATASET CONFIGURATION

Parameter	Value
Final Samples	8,928
Text Features	TF-IDF + BERT
Numeric Features	5
Reduced Features	307
Training Ratio	80%
Testing Ratio	20%
Training Samples	7,142
Testing Samples	1,786

C. Hyperparameter Configuration

Model hyperparameters were selected experimentally to obtain stable performance.

TABLE 3. HYPERPARAMETER SETTINGS

Model	Parameters
Random Forest	n_estimators=200, max_depth=20
XGBoost	n_estimators=300, max_depth=8
Weighted XGBoost	scale_pos_weight applied
BERT	all-MiniLM-L6-v2
SVD	n_components=300

D. Evaluation Metrics

The proposed framework was evaluated using four performance metrics:

$$\text{Precision} = \frac{TP}{TP+FP} \text{-----}(1)$$

$$\text{Recall} = \frac{TP}{TP+FN} \text{-----}(2)$$

$$F1 = 2 \times \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \text{-----}(3)$$

where TP, TN, FP, and FN represent true positive, true negative, false positive, and false negative observations respectively.

VI.RESULTS AND DISCUSSION

A. Comparative Performance Analysis

The proposed multimodal framework was evaluated using multiple machine learning and deep learning approaches to analyze the effectiveness of integrating heterogeneous drug and clinical information. Performance evaluation was conducted using Accuracy, Precision, Recall, and F1-score metrics.

Experimental results indicate that different models exhibit varying predictive behavior depending on feature representation and class distribution. Traditional machine learning approaches achieved moderate performance, while transformer-based embeddings improved predictive accuracy through richer semantic understanding of textual information.

Random Forest produced baseline performance with an accuracy of 59.35% and an F1-score of 51.04%, demonstrating the capability of ensemble learning for multimodal healthcare prediction. XGBoost slightly improved the balanced performance and achieved an F1-score of 53.78%, indicating better handling of complex feature interactions.

To reduce feature sparsity and improve representation efficiency, Singular Value Decomposition (SVD) was integrated with XGBoost. Although overall accuracy slightly decreased, the SVD-enhanced model achieved the highest balanced performance with an F1-score of 54.95%, suggesting improved feature generalization.

A weighted binary classification strategy was further introduced to address class imbalance. This approach improved recall performance but reduced overall accuracy due to increased minority-class sensitivity. Finally, transformer-based semantic embeddings generated using BERT were integrated with XGBoost to capture contextual relationships among textual attributes. The BERT-enhanced multimodal framework achieved the highest predictive accuracy of 68.14%, demonstrating the effectiveness of transformer-assisted multimodal learning.

B. Model Performance Comparison

Table 4. Comparative Performance of Prediction Models

Model	Accuracy (%)	Precision (%)	Recall (%)	F1-score (%)
Random Forest	59.35	53.32	59.35	51.04
XGBoost	58.23	52.34	58.23	53.78
XGBoost + SVD	57.67	53.89	57.67	54.95
Weighted Binary XGBoost	58.79	34.26	44.64	38.77
BERT + XGBoost	68.14	40.33	18.77	25.62

C. Discussion

The obtained results demonstrate that multimodal integration improves predictive capability by combining complementary information from pharmaceutical and clinical sources. Traditional machine learning methods provided stable baseline performance; however, transformer-based representations improved semantic feature extraction and increased predictive accuracy. Although BERT-enhanced classification achieved the highest accuracy, lower recall indicates challenges in minority-class prediction and suggests opportunities for future optimization. Conversely, SVD-enhanced XGBoost demonstrated stronger balanced performance, indicating that dimensionality reduction can improve robustness in multimodal healthcare datasets.

Overall, the findings confirm that integrating drug information, clinical reviews, and feature fusion strategies contributes to more effective drug recommendation prediction and supports intelligent healthcare analytics.

D. Key Findings

1. Multimodal fusion improved predictive capability.
2. BERT achieved the highest accuracy (68.14%).
3. XGBoost + SVD achieved the highest balanced F1 (54.95%).
4. Weighted learning improved minority-class detection.
5. Transformer-based embeddings improved semantic understanding.

VII. CONCLUSION

This study proposed a multimodal machine learning framework for drug recommendation prediction through the integration of heterogeneous drug and clinical information collected from multiple publicly available datasets. The proposed framework combined pharmaceutical attributes including drug composition, therapeutic usage, side effects, and manufacturer information with clinical reviews, patient feedback, ratings, and effectiveness-related features to construct a unified predictive environment.

A structured preprocessing pipeline involving data cleaning, missing value handling, drug normalization, multimodal dataset fusion, feature selection, and representation learning was implemented to improve data quality and support predictive analysis. Textual information was transformed using TF-IDF and transformer-based embeddings, while dimensionality reduction through Singular Value Decomposition (SVD) was applied to reduce sparsity and improve computational efficiency.

Multiple machine learning and deep learning approaches including Random Forest, XGBoost, SVD-enhanced XGBoost, weighted classification, and BERT-based multimodal classification were experimentally evaluated. Comparative analysis demonstrated that transformer-assisted multimodal learning improved predictive performance, achieving a maximum classification accuracy of 68.14%, whereas the XGBoost + SVD model produced the strongest balanced performance with an F1-score of 54.95%. The results indicate that integrating heterogeneous healthcare modalities can improve drug-related prediction and support intelligent healthcare analytics.

Although the proposed framework demonstrated promising outcomes, several limitations remain. The study relied on publicly available datasets and simulated recommendation targets rather than real-world clinical deployment. In addition, class imbalance and missing clinical variables affected prediction robustness in certain experimental settings.

Future work may focus on incorporating real clinical records, advanced multimodal transformer architectures, explainable artificial intelligence techniques, and real-time healthcare decision-support systems to further improve recommendation quality and practical applicability in drug delivery and personalized medicine.

A. Future Scope

- Integration of real-world clinical and pharmaceutical datasets
- Adoption of advanced multimodal transformer architectures
- Development of explainable AI-based healthcare models
- Real-time drug recommendation and decision-support deployment
- Extension toward personalized and precision healthcare applications

REFERENCES

1. Breiman L. Random forests. *Machine Learning*. 2001;45(1):5–32.
2. DOI: 10.1023/A:1010933404324
3. Chen T, Guestrin C. XGBoost: A scalable tree boosting system. In: *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. New York: ACM; 2016. p. 785–794. DOI: 10.1145/2939672.2939785
4. Devlin J, Chang MW, Lee K, Toutanova K. BERT: Pre-training of deep bidirectional transformers for language understanding. In: *Proceedings of NAACL-HLT*. 2019. p. 4171–4186.
5. DOI: 10.48550/arXiv.1810.04805
6. Reimers N, Gurevych I. Sentence-BERT: Sentence embeddings using siamese BERT-networks. In: *Proceedings of EMNLP-IJCNLP*. 2019. p. 3982–3992. DOI: 10.48550/arXiv.1908.10084
7. Deerwester S, Dumais ST, Furnas GW, Landauer TK, Harshman R. Indexing by latent semantic analysis. *Journal of the American Society for Information Science*. 1990;41(6):391–407.
8. DOI:10.1002/(SICI)1097-4571(199009)41:6<391::AID-I1>3.0.CO;2-9
9. Ramos J. Using TF-IDF to determine word relevance in document queries. In: *Proceedings of the First Instructional Conference on Machine Learning*. 2003.
10. Pedregosa F, Varoquaux G, Gramfort A, Michel V, Thirion B, Grisel O, et al. Scikit-learn: Machine learning in Python. *Journal of Machine Learning Research*. 2011;12:2825–2830.
11. Vaswani A, Shazeer N, Parmar N, Uszkoreit J, Jones L, Gomez AN, et al. Attention is all you need. *Advances in Neural Information Processing Systems*. 2017;30.
12. Miotto R, Wang F, Wang S, Jiang X, Dudley JT. Deep learning for healthcare: Review, opportunities and challenges. *Briefings in Bioinformatics*. 2018;19(6):1236–1246. DOI: 10.1093/bib/bbx044
13. Rajkomar A, Dean J, Kohane I. Machine learning in medicine. *New England Journal of Medicine*. 2019;380(14):1347–1358. DOI: 10.1056/NEJMra1814259

14. Topol EJ. High-performance medicine: The convergence of human and artificial intelligence. *Nature Medicine*. 2019;25(1):44–56.
15. DOI: 10.1038/s41591-018-0300-7
16. Beam AL, Kohane IS. Big data and machine learning in healthcare. *JAMA*. 2018;319(13):1317–1318. DOI: 10.1001/jama.2017.18391
17. Ching T, Himmelstein DS, Beaulieu-Jones BK, Kalinin AA, Do BT, Way GP, et al. Opportunities and obstacles for deep learning in biology and medicine. *Journal of The Royal Society Interface*. 2018;15(141):20170387. DOI: 10.1098/rsif.2017.0387
18. Esteva A, Robicquet A, Ramsundar B, Kuleshov V, DePristo M, Chou K, et al. A guide to deep learning in healthcare. *Nature Medicine*. 2019;25(1):24–29. DOI: 10.1038/s41591-018-0316-z
19. Jordan MI, Mitchell TM. Machine learning: Trends, perspectives, and prospects. *Science*. 2015;349(6245):255–260. DOI: 10.1126/science.aaa8415
20. Ravi D, Wong C, Deligianni F, Berthelot M, Andreu-Perez J, Lo B, et al. Deep learning for health informatics. *IEEE Journal of Biomedical and Health Informatics*. 2017;21(1):4–21. DOI: 10.1109/JBHI.2016.2636665
21. Kourou K, Exarchos TP, Exarchos KP, Karamouzis MV, Fotiadis DI. Machine learning applications in cancer prognosis and prediction. *Computational and Structural Biotechnology Journal*. 2015;13:8–17.
22. DOI: 10.1016/j.csbj.2014.11.005
23. Litjens G, Kooi T, Bejnordi BE, Setio AAA, Ciompi F, Ghafoorian M, et al. A survey on deep learning in medical image analysis. *Medical Image Analysis*. 2017;42:60–88. DOI: 10.1016/j.media.2017.07.005