



AFRL-AFOSR-JP-TR-2018-0074

Autonomous Action by Learning Group Action Protocols and Case-Based Reasoning

Tu-Bao Ho
John von Neumann Institute
VNC IT Park, Linh Trung Ward, Thu Duc District
Ho Chi Minh City, 700000
VN

10/12/2018
Final Report

DISTRIBUTION A: Distribution approved for public release.

Air Force Research Laboratory
Air Force Office of Scientific Research
Asian Office of Aerospace Research and Development
Unit 45002, APO AP 96338-5002

REPORT DOCUMENTATION PAGE					Form Approved OMB No. 0704-0188	
The public reporting burden for this collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing the burden, to Department of Defense, Executive Services, Directorate (0704-0188). Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to any penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number. PLEASE DO NOT RETURN YOUR FORM TO THE ABOVE ORGANIZATION.						
1. REPORT DATE (DD-MM-YYYY) 12-10-2018		2. REPORT TYPE Final		3. DATES COVERED (From - To) 07 Sep 2017 to 06 Sep 2018		
4. TITLE AND SUBTITLE Autonomous Action by Learning Group Action Protocols and Case-Based Reasoning				5a. CONTRACT NUMBER		
				5b. GRANT NUMBER FA2386-17-1-4094		
				5c. PROGRAM ELEMENT NUMBER 61102F		
6. AUTHOR(S) Tu-Bao Ho				5d. PROJECT NUMBER		
				5e. TASK NUMBER		
				5f. WORK UNIT NUMBER		
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) John von Neumann Institute VNC IT Park, Linh Trung Ward, Thu Duc District Ho Chi Minh City, 700000 VN				8. PERFORMING ORGANIZATION REPORT NUMBER		
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES) AOARD UNIT 45002 APO AP 96338-5002				10. SPONSOR/MONITOR'S ACRONYM(S) AFRL/AFOSR IOA		
				11. SPONSOR/MONITOR'S REPORT NUMBER(S) AFRL-AFOSR-JP-TR-2018-0074		
12. DISTRIBUTION/AVAILABILITY STATEMENT A DISTRIBUTION UNLIMITED: PB Public Release						
13. SUPPLEMENTARY NOTES						
14. ABSTRACT The PI has been successful in their tasks for this research grant. They were able to develop and test a new learning theory and algorithms that can learn recommended sequence of actions for heterogeneous and temporal objects, and tested the new algorithm to create treatment actions for a patient. The method developed is generic enough to be applicable to many autonomous systems with appropriate adaptations. There was 1 conference paper and one submitted journal paper as a direct result of this grant.						
15. SUBJECT TERMS Action recommendation, heterogenous, temporal, similarity, measure, mixed data, clustering, complex object, nearest neighbor, nearest, neighbor						
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT	18. NUMBER OF PAGES	19a. NAME OF RESPONSIBLE PERSON	
a. REPORT	b. ABSTRACT	c. THIS PAGE			SINGLETON, BRIANA	
Unclassified	Unclassified	Unclassified	SAR		19b. TELEPHONE NUMBER (Include area code) 315-227-7007	

Final Report for AOARD Grant FA2386-17-1-4094
“Autonomous Action by Learning Group Action
Protocols and Case-Based Reasoning”

Name of Principal Investigator: Tu Bao Ho

E-mail: bao.ho@jvn.edu.vn

Institution: John von Neumann Institute

Vietnam National University Ho Chi Minh City

Mailing Address: IT Park, Linh Trung, Thu Duc, HCMC.

Phone and Fax: 84-97-7855-165

Period of Performance: 07/09/2017 - 06/09/2018

Abstract

Given a large set of heterogeneous and longitudinal objects each associated with longitudinal actions that follow a general action protocol, we consider the problem of recommending appropriate actions for a new object in its longitudinal process where current studies seem too rigid with fixed intervals of action periods corresponding to the varying-length longitudinal process objects. The proposal aims to develop machine learning methods for this recommendation task. The proposal is unique in its approach to deliver three new bodies of theory and techniques for: (a) Clustering heterogeneous and longitudinal objects into subgroups of similar objects. To this end, our learning framework overcomes the challenge of mixed data representation by adopting a mixed-variate restricted Boltzmann machine; (b) Learning an action protocol for each subgroup; (c) Recommend actions for a new object in its action process based on the learned protocol of the subgroup it belongs to and the actions of its nearest neighbors. To tackle the varying-length longitudinal records, we propose an algorithm taking into account selected actions and construct represented trees to leverage both frequent and infrequent features. Under the general framework, concrete techniques are developed to learn treatment protocols for patient subgroups of a disease and to recommend the treatment for a patient.

1. Introduction

Disease diagnosis and treatment are essential aspects of healthcare. Disease diagnosis is a process in which preexisting set of categories is agreed upon by the medical profession to designate a specific condition [1] while a comprehensive treatment is a normalized care plan where therapeutic interventions and medicines for a particular disease are organized on a timeline [2, 3]. While the diagnosis prediction problem has been studied extensively over the past decades [4, 5, 6, 7, 8, 9], the treatment learning and recommendation problem is still in the early development stage [10, 11, 12, 13, 14]. Recently, addressing the second problem becomes more urgent due to several reasons. First, there are constraints on budgets of medication resources allocated to hospitals [15, 16, 17]. As a result, in many cases, hospitals need to customize treatments to fit the available resources but still ensure the treatment quality. Second, even under similar diagnostic codes, patients often suffer various symptoms that should be treated flexibly. Thus, capturing the patterns among various treatments in practice turns out to be helpful to assist not only managers in managing their resources [18] but also the less working physicians in grasping treatment patterns often used in their organization.

Principally, one can learn treatment for a particular disease through medical domain knowledge. It can be a piece of written information available in the literature [19, 20, 21] or gained experiences. As there is a wide range of domain knowledge that should be taken into account, the knowledge-driven approach may require much time and effort to absorb such knowledge. In recent years, the fast development of electronic medical records (EMRs) has enabled addressing the problem via data-driven approaches where one can derive treatment patterns of patient cohort and recommend treatment for new patients automatically from a massive amount of patient medical records.

Under the light of the data-driven approach, it is obvious that patients who share many symptoms, laboratory indicators, and much demographic information are likely to share common treatment patterns. However, most current studies in the literature have exploited merely limited patient information. The unavailability of rich feature medical records could be attributed to the difficulty in data collection [22, 23]. Even if different kinds of patient data are made fully available, they usually exist in different data types that are difficult to feed current data mining methods. In addition to the data collection and representation issues, it is challenging to capture medical domain knowledge hidden in prescription records. For instance, it is well-known that

physicians usually deliver treatment through periods. Obviously, the derivation of treatment patterns over a particular period highly depends on the prescription drugs in that period. Unfortunately, in many cases, physicians did not point out explicitly where is the treatment period intervals in varying-length prescription records. This fact poses challenges in identifying suitable treatment periods and their associated treatment patterns.

The objective of this work is twofold. First, we aim to propose a treatment learning framework which addresses the above challenges. Our learning framework divides a patient cohort into clusters namely subcohort where treatment patterns over periods for each subcohort are discovered subsequently. To overcome the issue of representing mixed-type data, we employ a mixed-variate restricted Boltzmann machine (MV.RBM) [24]. The advantage of this model is at its robustness in transforming mixed-type objects to their homogeneous representation. To tackle the challenge of treatment period identification, we propose an algorithm which tries to capture significant changes in prescription indications. More interestingly, we construct prescription trees based on prescription drugs' frequencies to leverage treatment patterns for each subcohort. Second, we propose a treatment recommendation framework that suggests top M prescription drugs for new patients by taking into account the typical treatments extracted from learned prescription trees of similar patients. Our recommendation framework captures the intuition that the treatment of new patients can be learned from the prescription records of neighbor patients.

In short, we propose two frameworks, a treatment learning framework and a treatment recommendation framework. The main contributions of our work are summarized as follows.

1. First, by adopting an MV.RBM for learning a homogeneous representation of mixed-type patient records, we encourage exploiting more patient features for maximizing the capability of data utilization.
2. Second, we employ both the knowledge-driven approach and data-driven approach in the treatment learning framework. The exploited medical domain knowledge is prescription indications in treating a disease. The incorporated domain knowledge seems more interpretable to identify unseen treatment periods in varying-length prescription records.
3. Third, we propose a new way to represent treatment patterns flexibly by a prescription tree where each path summarizes the frequency of a sequence of

prescription drugs for each subcohort. The tree is not only meaningful for physicians to capture frequent and infrequent prescription drugs but also helpful to assist the treatment recommendation framework in finding typical drugs of neighbor patients to be recommended for new patients.

4. Fourth, we propose a flexible recommendation framework which allows taking into account prescription drugs from neighbor patients. Our experimental evaluation shows the appropriateness of the idea learning from the neighborhood patients for the treatment recommendation task.
5. Fifth, we propose a mechanism to label prescription drugs' indication from external medical domain resources. Our indication labeling framework is useful not only for identifying treatment periods of patients but also for interpreting typical treatment patterns of each subcohort.

2. Related Work

2.1 Clinical Pathway Mining

In the early development of healthcare mining, prescription records were not published widely for research purpose. The most related studies at that time focused on mining clinical pathway, a close concept of treatment where research objects are clinical procedures such as examination, treatment, prescription, nursing visit. Lin et al. [25] developed a graph mining technique to discover dependency patterns of clinical pathways for curing brain stroke. Haytham et al. [26] combined a b-color based framework with Markov model for clinical pathway clustering and prediction. Bouarfa et al. [27] developed tree-guide and global pair-wise multiple sequence alignment to detect consensus workflow and outliers from clinical activity logs. Chen et al. [28] proposed a model to learn and categorize workflows based on their duration for efficient workflow management.

We note that the treatment mining problem and the clinical pathway mining problem share some properties. Both address treatment data which could be treated as a sequence of events, i.e., clinical procedures or medications, at different granularity levels. However, while the research objects of the first problem are often a few clinical procedures, those of the second problem could be thousand of prescription drugs plus additional information on dosage, routes. These characteristics make the treatment mining problem generally more challenging to tackle compared to the clinical pathway mining problem. Another difference is the extent to which patient health status is

affected by solving each problem. It can be seen that compared to the clinical mining problem, the treatment mining problem whose research object are medications that are supposed to affect more directly to patient status.

2.2 Probabilistic-based Approach

In recent years, probabilistic models have emerged as a promising technique for solving many data mining problems. In the field of healthcare, a variety of probabilistic models has been proposed to learn patterns of clinical pathways or treatments. Huang et al. [29] extracted treatment pattern as latent variables generated from a probabilistic topic model which captured the link between patient features and treatments. Lu et al. [30] modeled the diagnosis, contextual information, medication by a multiple channel LDA approach. In that work, they assumed the co-occurrences among disease, medications were generated by latent health status group structure. Xu et al. [31] developed a topic model which exploited billing information and prescription records to discover the execution path of clinical pathways. Park et al. [32] proposed a disease-medicine topic model summarizing prescription records from insurance data. Recently, Yao et al. [33] et al. have developed a topic model which describes the generating process of prescription records from traditional Chinese medicine in books.

The main drawback of the studies following probabilistic approach is that they employed many hyperparameters which are assumed to follow some distributions without justification from medical domain knowledge. This limitation, therefore, weakens the interpretation of the developed probabilistic models considerably.

2.3 Deep learning-based Approach

Besides the probabilistic-based approach, deep learning recently has shown its promise in solving healthcare mining issue. Pham et al. [34] proposed Deepcare, a dynamic neuron network based on LSTM model to predict future medical outcomes. This framework is designed for multitasking including modeling disease progression, recommending necessary intervention, and predicting future risks. Le et al. [35] developed a memory-augmented neuron network to predict sequences of medications and procedure. The memory-augmented network is featured with dual controllers where one for encoding medical history data, and the other for decoding treatment sequences. More recently, Jin et al. [14] have addressed the treatment sequence prediction by a multifaceted LSTM framework designed for different types of EMRs.

Despite the outstanding performances obtained in many applications, the deep learning approach for healthcare still suffers limitations concerning model interpretation. Most of the current models have not paid attention to integrate medical domain knowledge that makes them convincing to physicians. Additionally, the developed models are usually black box models. The recommendation results are uneasy to be explained under the healthcare perspective.

2.4 Reinforcement-based Approach

Another approach to deal with the treatment recommendation problem is using the reinforcement learning methods. In this stream, most studies represented outcomes and treatments as sequences of states and actions. The reinforcement algorithms are then employed to find optimal treatment sequences. Zhao et al. [36] utilized a Q learning method to find optimal medications for non-small cell lung cancer from clinical trials. In that work, the authors employed a modified support vector regression to approximate the Q-function. Liu et al. [37] proposed a deep reinforcement learning framework to optimize dynamic treatment regimens from medical registry data. The authors predicted possible clinical procedures via a supervised learning step and estimated the long-term value function of dynamic treatment via a deep reinforcement learning step. Nemati et al. [38] investigated the optimal usage of heparin by deep reinforcement learning together with Hidden Markov Model. Recently, Wang et al. [12] have developed a framework which combines supervised learning and reinforcement learning to find optimal dynamic treatments using published real world electronic medical records.

Studies following this approach are typically conducted for clinical trials with available treatment outcomes. For prescription records, treatment outcomes could be reported in daily nursing notes or laboratory indicators. However, identifying the corresponding outcomes for every doctor’s order in the clinical context is challenging as it may require deep domain knowledge. Consequently, this drawback makes the implementation of reinforcement learning seems to be impractical for real-world EMRs.

2.5 Frequency-based Approach

The third approach derives treatment patterns based on prescription drugs’ frequencies. For clinical pathway mining, Lin et al. [25] constructed dependency graphs which took into account the frequency of clinical process. The dependency graphs

then were used to discover frequent clinical pathways. Hirano et al. [39] mined typical treatment processes based on occurrence and transition frequency maps. In that work, the author characterized each treatment sequence by a typicalness index and selected the highest typicalness values as candidate patterns. Sun et al. [40] proposed a similar measure among prescription records to divide them into clusters and find frequent drugs among core patients of each cluster as treatment patterns.

The limitation of the above works is that while they proposed many approaches to reveal frequent treatment patterns, they almost ignore infrequent patterns which may be useful for physicians to reduce the risk of prescribing drugs mistakenly.

3. Problem Formulation

We consider a set of patients $\{p_1, p_2, \dots, p_N\}$ who were diagnosed with similar disease codes. Each p_i is a heterogeneous object consisting of different data components such as basic demographic information $Info_{p_i}$, laboratory examination data Lab_{p_i} , nursing notes $Note_{p_i}$, and prescription data $Presc_{p_i}$. It is noted that the elements Lab_{p_i} , $Note_{p_i}$ and $Presc_{p_i}$ are longitudinal components over $n_{p_i}^l$, $n_{p_i}^n$, $n_{p_i}^p$ timestamps, respectively. Each component is a set of features that could be detailed further. For instance, the $Info$ component for patient p_i can be decomposed into:

$$Info_{p_i} = \{Info_{p_i}^{age}, Info_{p_i}^{gender}, Info_{p_i}^{marriage}, Info_{p_i}^{histIllness} \dots\}$$

describes in details the age, gender, marriage status, history of illness, to name a few, of patient p_i . The component $Presc_{p_i}$ describes information about prescription drugs over $n_{p_i}^p$ timestamps. It can be decomposed as follows.

$$Presc_{p_i} = \{Presc_{p_i}^{tp_1}, Presc_{p_i}^{tp_2}, \dots, Presc_{p_i}^{tp_{n_{p_i}^p}}\}$$

where each $Presc^j = \{dr_{j1}, dr_{j2}, \dots\}$ is a set of drugs prescribed at timestamp j . Each drug $dr = \{name, startdate, enddate, dosage, route\}$ is a compound object characterized by information about drug name, starting date, ending time of usage and the route delivered to patients. The component $Note_{p_i}$ contains nursing notes written in text format about p_i 's treatment progress over $n_{p_i}^n$ timestamps.

$$Note_{p_i} = \{Note_{p_i}^{tn_1}, Note_{p_i}^{tn_2} \dots, Note_{p_i}^{tn_{n_{p_i}^n}}\}$$

The component Lab_{p_i} describes different measurement values of patient condition in $n_{p_i}^l$ timestamps.

$$Lab_{p_i} = \{Lab_{p_i}^{tl_1}, Lab_{p_i}^{tl_2}, \dots, Lab_{p_i}^{tl_{n_{p_i}^l}}\}$$

where each $Lab^j = \{i_{j1}, i_{j2}, \dots\}$ is a set of indicator values at time stamp j . Each $i = \{name, value\}$ is an indicator characterized by its name and value.

This research aims to construct the following frameworks.

1. A treatment learning framework that utilizes all relevant features from the components $\{Info, Lab, Note, Presc\}$ to learn treatment patterns of each subcohort over n periods $\tau_1, \tau_2, \dots, \tau_n$. We note that the n treatment periods are not given in advance. Therefore, identifying n treatment periods for each patient is a subtask of this framework.
2. A treatment recommendation framework to recommend top M drugs that could be prescribed over n periods for new patients.

We denote the set of patients $\{p_1, p_2, \dots, p_N\}$ whose medical records are utilized to construct the treatment learning framework as training patients. The terms new patients and testing patients are used interchangeably in this report.

4. Treatment Learning Framework

In this section, we describe our framework for the treatment learning problem given medical records of training patients having the same diagnostic codes. The learning framework captures the intuition that patients who share as many latent features underlying their health condition and profiles as possible are likely to belong to the same subcohort. Patients in each subcohort, as a result, could share patterns in treatment. Figure 1 describes an overview of the proposed treatment learning framework. It consists of two major tasks: clustering patients into subcohorts and learning typical treatment patterns for each subcohort. We note that in this figure, the term regimens is equivalent to the term treatment patterns and the term regimen trees is equivalent to prescription trees. We present all relevant steps of the treatment learning framework in the following subsections.

4.1 Data Preprocessing

The framework collects two sets of training patients' medical records to accomplish the above tasks. One is the set of components $\{Lab, Info, Note\}$ named as non-treatment based records to serve for the subcohort construction and the other is the component $Presc$ named as treatment based records to serve for the derivation of treatment patterns of each subcohort. We note that non-treatment based data contains both static ($Info$) and longitudinal data ($\{Lab, Note\}$). As the ultimate goal of

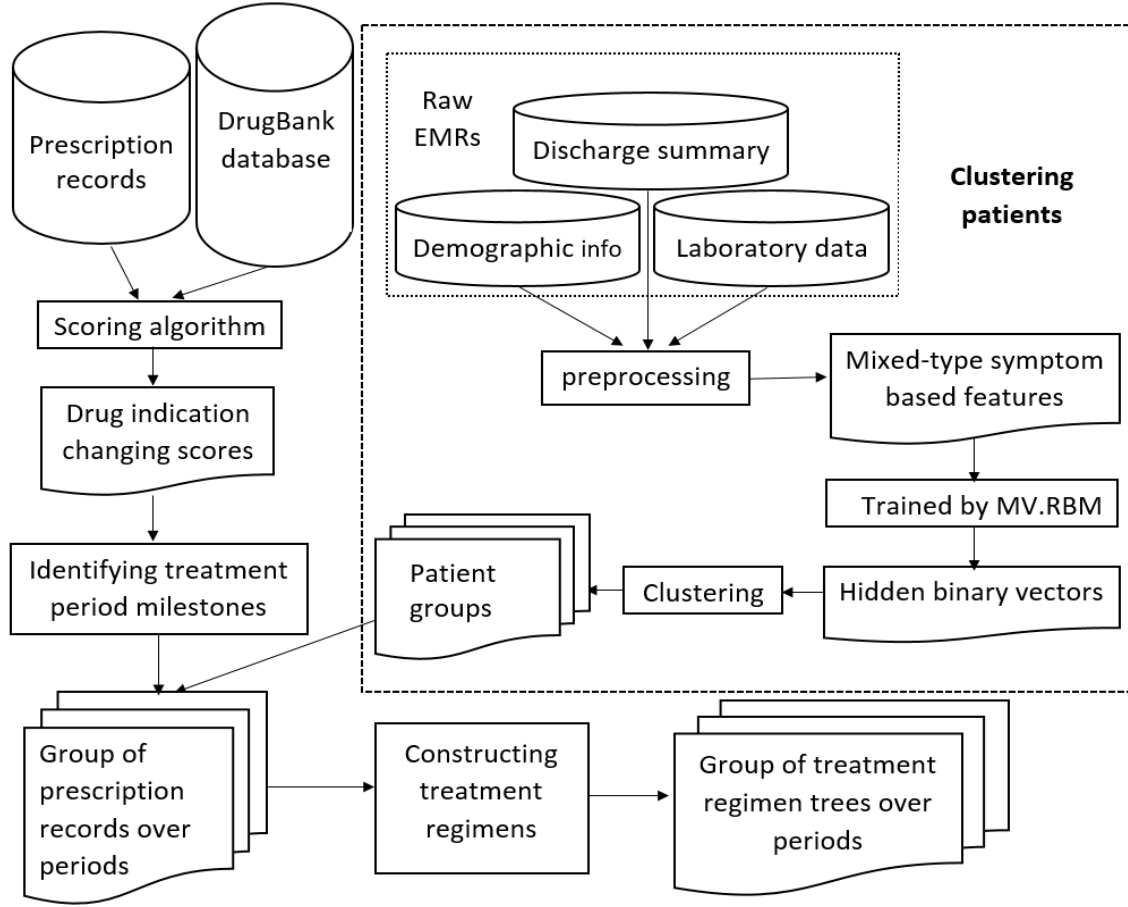


Figure 1: Overview of the treatment learning framework.

our work is to recommend prescription drugs for new patients from the early days of their admission, we only collect initial values of longitudinal non-treatment-base data of training patients, i.e. the data recorded at timestamp tl_1 , tn_1 , respectively. Our data collection’s strategy assumes that patients with similar initial signs, symptoms, or laboratory indicators are probably treated by similar care plans.

We encode categorical features as one-hot encoding vectors and normalize numerical features to zero-mean unit-variance. For text data, we extract initial relevant features by cTAKES [41], a well-known tool designed for clinical text processing. Its primary function is to identify clinical terms in a given text and link them to concepts in the Unified Medical Language System (UMLS) [42], a large ontology constructed for the biomedical domain. Not only do cTAKES normalize discovered clinical terms, but it also allows distinguishing semantic types of clinical terms. In our framework,

	$G_i(v_i)$	$H_{ik}(v_i)$	$P_i(v_i \mathbf{h})$
Binary	$a_i v_i$	$w_{ik} v_i$	$\frac{\exp(a_i v_i + \sum_k w_{ik} h_k v_i)}{1 + \exp(a_i + \sum_k w_{ik} h_k)}$
Gaussian	$-v_i^2/2\sigma^2 + a_i v_i$	$w_{ik} v_i$	$\mathcal{N}(\sigma_i^2(a_i + \sum_k w_{ik} h_k), \sigma_i)$
Categorical	$\sum_m a_{im} \delta_m[v_i]$	$\sum_{m,k} a_{imk} \delta_m[v_i]$	$\frac{\exp(\sum_m a_{im} \delta_m[v_i] + \sum_{m,k} w_{imk} \delta_m[v_i] h_k)}{\sum_l \exp(a_{il} + \sum_k w_{ilk} h_k)}$

Table 1: The type specific functions in the MV.RBM model

we are interested in the extracted terms with specific clinical meanings to represent training and testing patients. Thus, we select the terms linked to clinical concepts of which semantic types are signs/symptoms or diseases.

4.2 Data Representation and Patient Clustering

The encoded non-treatment based data is mixed of numerical, binary or categorical data types. Such heterogeneous input vectors are often difficult to feed to clustering methods. To this end, we employ an MV.RBM, a powerful unsupervised representation model, to transform encoded input vectors to their homogeneous representations.

MV.RBM is an extension of RBM model developed for heterogeneous input units. In the MV.RBM architecture, the data type of input layer is designed for not only binary units, but also numerical or categorical units. Denote $\mathbf{v} = (v_1, v_2, \dots, v_N)$ as the set of visible input units and $\mathbf{h} = (h_1, h_2, \dots, h_K)$ as the set of hidden units. MV.RBM defines a more deliberate energy function which covers the case of other data types in addition to binary data. The formula of the energy function is given as follows:

$$E(\mathbf{v}, \mathbf{h}) = -(\sum_i G_i(v_i) + \sum_k b_k h_k + \sum_{ik} H_{ik}(v_i) h_k)$$

where $\mathbf{b} = (b_1, b_2, \dots, b_N)$ are biases vectors for hidden layer, $G_i(v_i)$ and $H_{ik}(v_i)$ are specified-type functions. By exploiting the conditional independence property within nodes in a layer of bipartite structure, we can get the following factorization equations:

$$P(\mathbf{v}|\mathbf{h}) = \prod_{i=1}^N P_i(v_i|\mathbf{h}); \quad P(\mathbf{h}|\mathbf{v}) = \prod_{k=1}^K P(h_k|\mathbf{v})$$

$$P_i(v_i|\mathbf{h}) = \frac{1}{Z(\mathbf{h})} \exp(G_i(v_i) + \sum_k H_{ik}(v_i) h_k)$$

$$P(h_k|\mathbf{v}) = \frac{1}{1 + \exp(-w_k - \sum_i H_{ik}(v_i))}$$

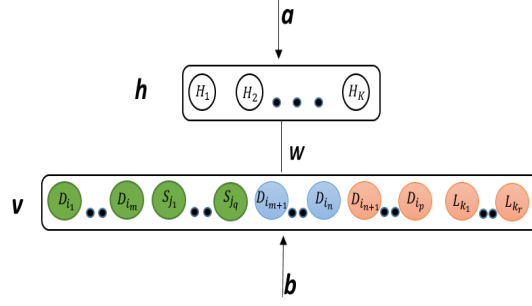


Figure 2: A MV.RBM for patient records. The green, blue and orange circles represent for binary, categorical and continuous input units. The circles with labels D , S , L indicate demographic, signs/symptoms and laboratory data, respectively.

where h_k^1 indicates the assignment $h_k = 1$, and $Z(\mathbf{h})$ is a normalization constant.

The functions $G_i(v_i)$, $H_{ik}(v_i)$ and corresponding $P_i(v_i|\mathbf{h})$ for each kind of data are given in Table 1 [24]. We note that in Table 1, a_i , a_{im} are input bias parameters, w_{ik} , w_{imk} are input-hidden weighting parameters. Those with extra subscript m are dedicated for categorical features. The inference step for model parameters $\theta = (b, a, w)$ can be found in [24].

Assuming input features are mutually independent given their latent factors, Figure 2 demonstrates how to feed non-treatment features by MV.RBM. Without losing generality, we suppose demographic features takes either numerical, binary or categorical values while indicator features take numerical values. For text data, we extract signs/symptom or disease UMLS concepts as described in the previous section and represent them by one-hot encoding vectors. After training the MV.RBM, we consider the hidden states as transformed features of encoded input vectors. The latent representation vectors then can be fed easily by clustering methods. In this particular work, we use Hierarchical clustering with Hamming distance to cluster representation vectors of training patients. Based on a survey about hierarchical clustering methods for binary vectors [43], we decide to use the complete linkage as it was reported to return low error rate when used in combination with symmetric measures.

4.3 Prescription Drugs' Normalization and Indication Assignment

This section describes our approach for preprocessing treatment based data. As physicians usually prescribe patients the same ingredient drugs under various names, it is crucial to perform drug normalization before doing further steps. For every prescription drug, we use cTAKES to identify its candidate UMLS clinical concepts and select the ones whose semantic is about medication. In case a few medication

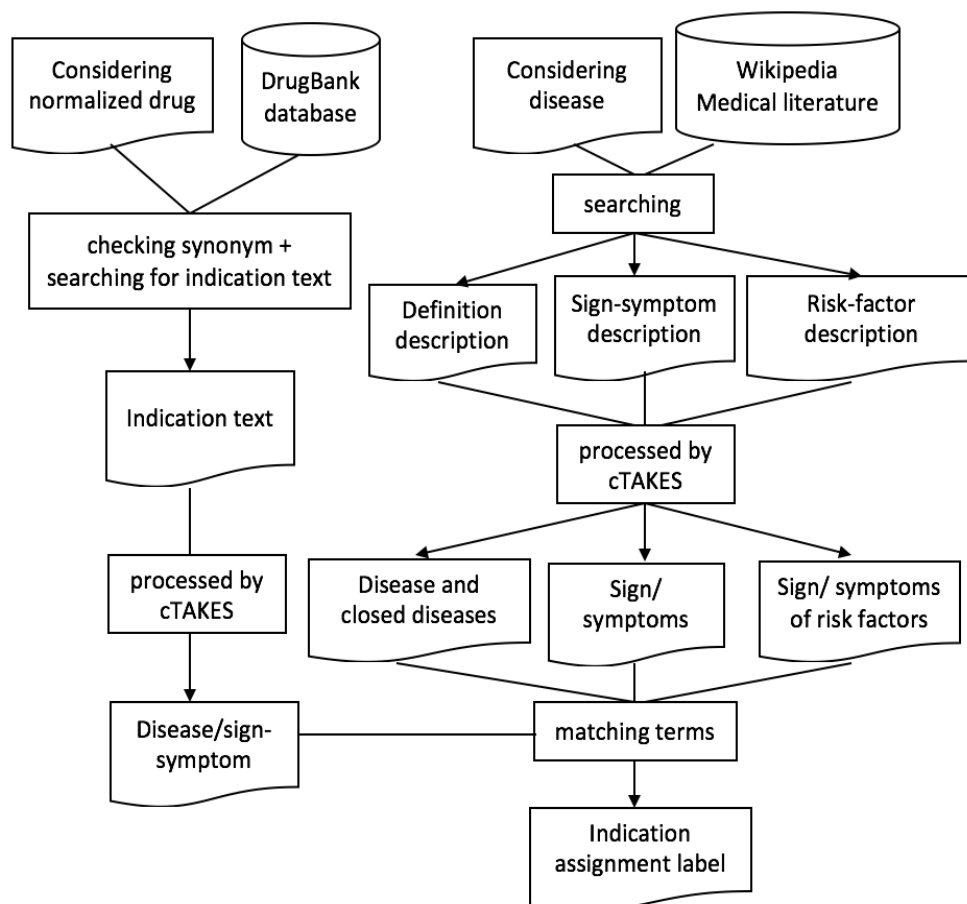


Figure 3: Overview of the framework for labeling prescription drug's indication

concepts are suggested for the same prescription drug, we revise the mapping by confirmation from domain experts and select the most appropriate one.

Besides normalizing drug name, labeling drug indication is one of the most important steps which assist the subsequent tasks in our work. Drugs with available indication labels are not only necessary for measuring the change of indication in prescription records but also meaningful in interpreting treatment patterns of patient subcohorts. For every normalized drug dr , we try to identify which diseases or symptoms are treated by dr and classify dr 's indication to one of the three following indication groups which are the primary group, symptom group or risk factor group. The primary group consists of drugs to treat at least one of the diagnostic diseases or their closed diseases. The symptom group is a set of drugs to treat typical symptoms of the considering diseases while the risk factor group is determined as a set of drugs to treat risk factors that may cause the considering diseases.

Figure 3 illustrates the idea of labeling drug indication. We collect from Wikipedia and biomedical literature text sections regarding the definitions, typical symptoms and risk factors of the considering diseases. These texts are then processed by cTAKES to extract UMLS concepts of which semantic types are about symptoms or diseases. For normalized drugs, we look for their indication description given in the DrugBank database. We take into account all synonym drugs to ensure that the indication of every normalized drug has been found in the database. The labeling mechanism works as follows. Given a normalized drug, we process its associated indication text by cTAKES and extract UMLS concepts which are diseases or symptoms which can be cured by the considering drug. We label all prescription drugs to the primary group, symptom healing group, risk factor healing group in that priority if their indication texts and the text sections describing the definition, the symptom, the risk factor of the considering diseases share common disease or symptom UMLS concepts. We note that our indication labeling mechanism also allows identifying in details the diseases or symptoms treated by every normalized drug. This feature is useful in interpreting the treatment patterns of each subcohort which will be presented in the later part of this report.

4.4 Treatment Period Identification

Patient prescription records are complex and varying-length objects. In the literature, researchers split prescription records by fixed intervals of treatment periods [40, 44]. As treatment period identification process may affect the derivation of treat-

ment patterns significantly, dividing treatment records by fixed intervals without a domain integration seems too rigid. In this work, we address this issue flexibly by a domain-based approach. The main idea is inspired by the observation that patients are possibly in a new treatment period whenever there is a significant change in prescription indication given to the patients. To characterize the strength of indication changes in prescription records, we construct an accumulated score for each timestamp which takes into account prescription drugs which are new, recent stoppage, redelivered with dosage changed. Based on the drugs' indication group, we count on the number of drugs among those drugs which belong to the primary indication, symptom group, or risk factor group. The accumulated scores aggregate these quantities weighted by their importance in treating the considering disease. We assign the weights in decreasing order of primary, symptom, or risk factor group.

Algorithm 1: Scoring prescription records

Data: Θ, T, MDB, SDB

Result: return *scores* as a list of accumulated scores

```

1 Initialize  $U$  as an empty set ; ▷ set of recently delivered drugs
2 Initialize scores as an empty list ;
3  $aScore := 0$  ; ▷ the accumulated score
4 for each  $d \in T$  do
5    $D := \{dr \mid \forall dr \in \Theta \wedge dr.startdate == d\}$  ; ▷ delivered drugs on date d
6    $N := \{dr \mid \forall dr \in D \wedge dr.name \notin U.name\}$  ; ▷ newly delivered drugs
7    $DC := \{dr \mid \forall dr \in D, \exists dr' \in U \text{ such that } dr.name == dr'.name \wedge dr.dosage <>$ 
    $dr'.dosage\}$  ; ▷ dosage changed drugs
8    $S := \{dr \mid \forall dr \in U \wedge dr.name \notin D.name \wedge dr.enddate < d\}$  ; ▷ recently
   stopped using drugs
9   for each  $d' \in U$  do
10    if  $\exists d'' \in D$  such that  $d'.name == d''.name$  then
11       $d' := d''$  ; ▷ update U with redelivered drugs
12    $U := (U \setminus S) \cup N$  ; ▷ update U with newly delivered drugs
13    $CD := N \cup DC \cup S$  ; ▷ considering drugs for calculating scores
14    $CMD := CD.name \cap MDB$  ; ▷ considering main drugs
15    $CSD := CD.name \cap SDB$  ; ▷ considering symptom-healing drugs
16    $UD := CD.name \setminus (CMD \cup CSD)$  ; ▷ unclassified drugs
17    $aScore = aScore + |CMD| \times w_{main} + |CSD| \times w_{symp} + |UD| \times w_{unk}$  ;
18   Add  $aScore$  to scores

```

Our notations for the treatment period identification algorithm are explained as follows. We denote $dr_j^{p,t} = \langle name, startdate, enddate, dosage \rangle$ to characterize every drug dr_j prescribed for patient p at specific timestamp t by its normalized drug

name, starting date, ending date of usage, and the dosage. Let $\Theta^p = \{dr_j^{p,t}.name\}$ be the set of drugs given to the patient, $T^p = \{dr_j^{p,t}.startdate\}$ be the ordered set of prescribed dates, and PD , SD , RD be sets of primary drug, symptom-healing drugs, risk factors healing drugs which have been identified in the previous section, respectively. The detailed algorithm for scoring changes in prescription indications for a patient p at the timestamp t is presented in Algorithm 3. For readability, we remove the superscript p, t, j and use Set notations in the pseudocode.

4.5 Prescription Tree Construction

We have demonstrated our domain-embedded algorithm for the treatment period identification. In this section, we describe how treatment patterns over a period for each patient subcohort are derived. It is worth noting that this step takes into account normalized drug names only. Information regarding dosage, route, chronological order among drugs in each period is supposed to be decided by physicians.

In the literature, treatment patterns are often discovered as a set of frequent prescription drugs from core patients of each subcohort [40]. This approach, however, often requires a minimum support threshold which is subjective and sensitive. Moreover, it is uneasy for physicians to figure out the set of infrequent prescription drugs under this approach.

Algorithm 2: Procedure for the construction of a prescription tree

```

Tree( $d, \nu, \Gamma, \Omega_{\delta_\nu}, \Lambda$ )
  1 if  $\Omega_{\delta_\nu}$  is  $\emptyset$  or  $d == \Upsilon$  then
  2   return
  3  $k := \arg \max_i \sum_{j=1}^n a_{ij}$  ;
  4  $\Gamma[\delta_\nu, k] := \text{“} \nwarrow \text{”}$ ;
  5  $\delta_k := \delta_\nu \cup k$  ;
  6  $\Omega_{\delta_k} := \Omega_{\delta_\nu}^k$  ;
  7  $\Omega_{\delta_{\nu+}} := \Omega_{\delta_\nu}^{-k}$  ;
  8  $\Lambda_{\delta_k} := \{j \text{ s.t. } \omega_{\delta_\nu}^{kj} = 1\}$  ;
  9 if  $|\Lambda_{\delta_k}| < \epsilon$  then
 10   Tree( $d, \Omega_{\delta_k}, \nu, \Lambda$ );
 11 else
 12   Tree( $d + 1, \Omega_{\delta_k}, k, \Lambda$ );
 13   Tree( $d, \Omega_{\delta_{\nu+}}, \nu, \Lambda$ );
 14 return ;

```

To this end, we suggest organizing treatment patterns in a tree form where each node represents a prescription drug. At a considering node, we extract prescription records of patients who were prescribed by the nodes from the root until the considering node. The leftmost child node is labeled with the most frequently prescribed drug apart from those linked with the parent nodes. Determining drug label of the next child node follows the same mechanism, but we exclude prescription records of patients who were treated by left-hand side nodes on the same level. We continue this procedure recursively until the number of patients treated with drugs from the root until that node is less than some threshold. Besides labeling drugs for prescription tree' nodes, we also save the IDs of the patients who were treated by the set of drugs from the root until each node. It is of interest to note that each patient is treated by nodes on a unique path in the tree. We named this path as treatment path of those patients linked with its node. This property is useful for the treatment recommendation task presented in the subsequent section. We denote the notations for the prescription tree construction algorithm as follows.

- d : the current depth of the constructing prescription tree.
- ν : the constructing node.
- Γ : the constructing prescription tree.
- δ_ν : the treatment path from the root until ν .
- $\delta_{\nu+}$: the treatment path from the root until the next unlabeled child node of ν .
- Ω_{δ_ν} : the current patient-drug interaction matrix corresponding to treatment path δ_ν . Table 2 illustrates the initial interaction matrix Ω_{δ_ϕ} of the root node where

$$\begin{cases} \omega_{\delta_\phi}^{kj} = 1 \text{ indicates patient } p_j \text{ is treated with drug } dr_k \\ \omega_{\delta_\phi}^{ij} = 0 \text{ indicates patient } p_j \text{ is not treated with drug } dr_k \end{cases}$$

- $\Omega^k = (\omega^{ij})$: the interaction matrix of patients were treated with drug k , where

$$\begin{cases} i \text{ such that } i \neq k \\ j \text{ such that } a_{kj} = 1 \end{cases}$$

- $\Omega^{-k} = (\omega^{ij})$: the interaction matrix of patients were not treated with drug k , where

$$\begin{cases} i \text{ such that } i \neq k \\ j \text{ such that } a_{kj} = 0 \end{cases}$$

- Λ_{δ_ν} : the patient ids of patients who were treated by drugs on δ_ν .
- ϵ : the threshold to stop constructing prescription tree.
- Υ : the highest depth of the prescription tree.

Algorithm 2 demonstrates the detailed algorithm of the construction of the treatment tree for a patient subcohort over a specific period.

	p_1	p_2	$\dots$	p_j	$\dots$	p_n
dr_1	ω^{11}	ω^{12}	$\dots$	ω^{1j}	$\dots$	ω^{1n}
dr_2	ω^{21}	ω^{22}	$\dots$	ω^{2j}	$\dots$	ω^{2n}
$\dots$	$\dots$	$\dots$	$\dots$	$\dots$	$\dots$	1
dr_k	ω^{k1}	ω^{k2}	$\dots$	ω^{kj}	$\dots$	ω^{kn}
$\dots$	$\dots$	$\dots$	$\dots$	$\dots$	$\dots$	$\dots$
dr_m	ω^{m1}	ω^{m2}	$\dots$	ω^{mj}	$\dots$	ω^{mn}

Table 2: The initial drug-patient interaction matrix.

5. Treatment Recommendation Framework

This section presents a framework aiming to assist physicians for the treatment recommendation task. In the medical field, we assume that physicians have sufficient knowledge to filter out the necessary medications among those recommended. One of the most important things is to provide an explainable recommendation mechanism that is able to show which known patients' prescription records the recommendation is based on and how those records are integrated to derive the recommendation results. A treatment recommendation framework therefore, requires being good not only in terms of evaluation metrics but also in terms of explainability. Unfortunately, this issue is underestimated in most of the related studies.

Intuitively, one can suggest treatment for new patient based on prescription drugs of the most similar known patient, i.e. the nearest neighbor patient. However, it is

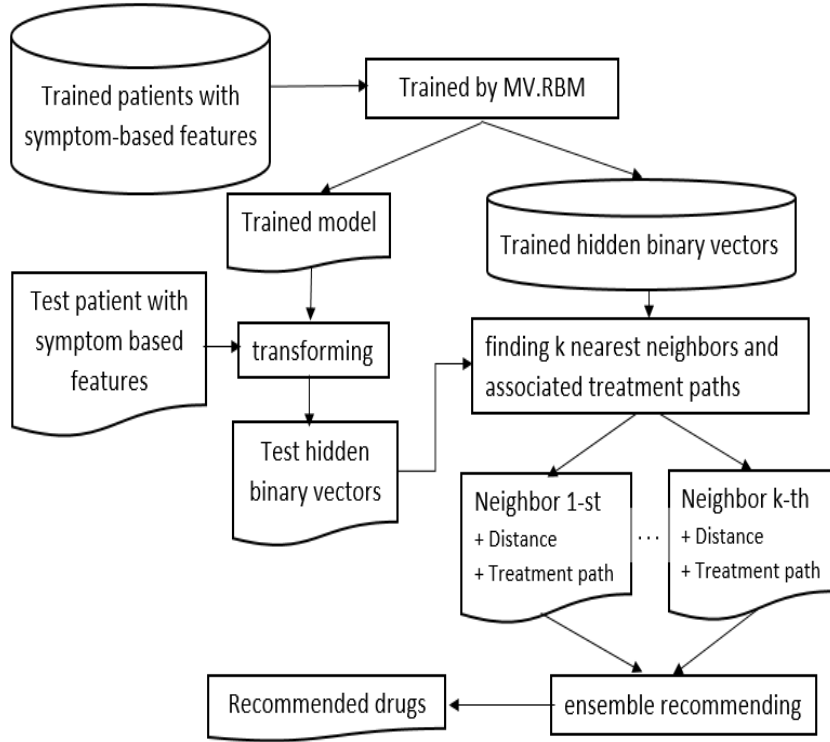


Figure 4: Overview of the recommendation framework.

uneasy to identify such patient since he may not have been recorded in the database, or he is just similar to the new patient partially in several aspects. As a result, the treatment of the new patient and the treatment of the nearest neighbor patient may not be identical in reality. To this end, we propose a neighbor-based treatment recommendation framework that mimics human intuition to suggest treatment for new patients based on their K neighbors' treatments. Figure 4 describes our idea for the treatment recommendation framework. Given a new patient p , we only consider his initial non-prescription based feature vector and transfer it to the binary form through the parameters $\hat{\theta} = (\hat{\mathbf{a}}, \hat{\mathbf{b}}, \hat{\mathbf{w}})$ learned by the trained MV.RBM used for representing mixed type vectors of training patients. Let h^p be the binary hidden vector of p , $h^{p'}$ be the binary hidden vector of training patient p' , the similarity between two patients is defined by the Hamming distance between their latent representations as below.

$$d_{p'}^p = \text{Hamming}(h^p, h^{p'}) = \frac{1}{|h|} \sum_{i=1}^{|h|} I(h_i^p \neq h_i^{p'})$$

In our proposed framework, we utilize the resulting prescription trees to find the K associated treatment paths of the neighbors. Each treatment path is considered

as a set of typical drugs treating one of the K neighbors and therefore contributes in the construction of treatment for patient p . It is worth noting that to capture the variant in treatments of similar patients, the K neighbor patients and their associated treatment paths can belong to different subcohorts. Let $\delta^{p'} = \{dr_1^{p'}, dr_2^{p'}, \dots, dr_M^{p'}\}$ denote the set of drugs linked with the treatment path of p' , we find the K similar patients $p_1, p_2, \dots, p_K$ of p and utilize the associated $\delta^{p_1}, \dots, \delta^{p_K}$ and the distances $d_{p_1}^p, d_{p_2}^p, \dots, d_{p_K}^p$ to recommend M drugs for p . Let $C = \{dr_1, dr_2, \dots, dr_j\}$ be the set of distinct drugs named as candidate drugs unioned from $\delta^{p_1}, \dots, \delta^{p_K}$. We propose two following recommendation approaches.

5.1 Non-weighting Approach

The intuition underlying this approach is that prescription drugs delivered to many neighbors are likely to be used for new patient. Therefore, for every candidate drug dr in C , we compute its path frequency $freq_{dr}^p$, i.e how many treatment paths contains dr as one of the criteria for recommendation. Drugs with higher path frequency indicates that it is prescribed for a greater number of neighbor patients and hence, have a higher probability to be recommended for new patients. The formula of $freq_{dr}^p$ is provided as follows.

$$freq_{dr}^p = \sum_{i=1}^K I(dr \in \delta^{p_i}) \quad (1)$$

To solve the case dr has the same path frequency with other drugs, we propose considering a distance priority metric d_{dr}^p , another measure which takes into account the distance from test patient to the neighbors whose treatment paths contain dr . The greater the sum of the inverse distance from those neighbors to p , the higher priority the drug is selected as candidate drugs. We provide the formula of d_{dr}^p as below.

$$dist_{dr}^p = \sum_{i=1}^K I(dr \in \delta^{p_i}) \times \frac{1}{d_{p_i}^p} \quad (2)$$

Algorithm 3 provides a pseudocode for recommending drugs for new patient p using non-weighting approach.

5.2 Weighting Approach

The non-weighting approach simply takes into account the frequency of candidate drugs among K treatment paths. This approach, however, seems to work effectively

Algorithm 3: Recommending prescription drugs for new patient p

Data: $\Gamma, \Lambda, \theta, v^p, H^{train}$

Result: return top M recommended drugs

- 1 Compute $h^p = P(h|v^p, \theta)$;
 - 2 Compute similarity between h^p and each training patient in H^{train} ;
 - 3 Select K most similar patients $p_1, p_2, \dots, p_K$;
 - 4 Trace associated treatment paths $\gamma^{p_1}, \gamma^{p_2}, \dots, \gamma^{p_K}$ through tracing variables Λ ;
 - 5 $C = \bigcup_{i=1}^K \gamma^{p_i}$;
 - 6 **for** each $dr \in C$ **do**
 - 7 Compute $freq_{dr}^p$ by equation (1);
 - 8 Compute $dist_{dr}^p$ by equation (2);
 - 9 Return top M drugs sorted by $(freq^p, dist^p)$;
-

only if the K neighbors are reliable neighbors, i.e their treatments are highly relevant to p 's treatment. In case the prescription drugs of K neighbors are quite different from p , we propose a more deliberate framework which assign a weight to each node on the K treatment paths to indicate how likely that node is selected when its associated treatment path is used as treatment for some neighbor patients. To do that, we split the training patients into several subsets where we consider each subset as sub-test patients and the rests are sub-training patients. For each patient in the sub-tests, we query his K_1 neighbors $p_1, p_2, \dots, p_{K_1}$ and their associated treatment paths $\delta^{p_1}, \dots, \delta^{p_{K_1}}$. For each patient p_j in the sub-training set, let S^{p_j} be the set of patients who have p as one of their K_1 neighbors. We assign a hitting-score $hit_{dr}^{\delta^{p_j}}$ to each drug dr on the treatment path δ^{p_j} of training patient p_j as follows.

$$hit_{dr}^{\delta^{p_j}} = \sum_{p_k \in S^{p_j}} d_{p_k}^{p_j} \times I(dr \in \delta^{p_k})$$

In the above formula, every time drug dr was used to treat a patient p_k in S^{p_j} , we add to the hitting score $hit_{dr}^{\delta^{p_j}}$ a reward equal to the distance $d_{p_k}^{p_j}$. The meaning is that when p_j and p_k are far neighbors and dr has been found in the treatment of p_k , it is added more weight than the closer neighbors as an compensation for the possibility of incorrectly identifying close neighbors. In case p_j and p_k are close neighbors and dr has been found in the treatment of p_k , as there is a high possibility dr can be found in the treatment of p_k , hence we add a relative small award equal to their distance to the hitting score.

After weighting all nodes in the prescription trees with sub train and sub test sets from training patients, we perform the procedure for selecting recommendation drugs for testing patient p . For each candidate drug dr in the set C , we compute an average hitting score $\overline{hit}_{dr}^p$ weighted by the distances from the test patient to neighbors whose treatment paths include dr . The formula of $\overline{hit}_{dr}^p$ is given as follows.

$$\overline{hit}_{dr}^p = \frac{1}{\sum_{i=1}^K I(dr \in \delta^{p_i})} \sum_{i=1}^K I(dr \in \delta^{p_i}) \times hit_{dr}^{\delta^{p_i}} \times d_{p_i}^p \quad (3)$$

Algorithm 4 summarizes the idea of the weighting approach in pseudocode.

Algorithm 4: Recommending prescription drugs for new patient p with weighted nodes

Data: $\Gamma, \Lambda, \theta, v^p, H^{train}$

Result: return top M recommended drugs

- 1 Split training set into subtrain and subtest sets ;
 - 2 Initialize all nodes in the prescription trees with 0 hitting score;
 - 3 **for** each pair (subtrain , subtest) **do**
 - 4 **for** each p in the subtest **do**
 - 5 Select K_1 most similar patients $p_1, p_2, \dots, p_{K_1}$ among subtrain set ;
 - 6 Trace associated treatment paths $\gamma^{p_1}, \gamma^{p_2}, \dots, \gamma^{p_{K_1}}$ through tracing variables Λ ;
 - 7 **for** each γ^{p_i} **do**
 - 8 **for** each dr in γ^{p_i} **do**
 - 9 If p was treated with dr $hit_{dr}^{\delta^{p_i}} = hit_{dr}^{\delta^{p_j}} + d_{p_i}^p$
 - 10 Compute $h^p = P(h|v^p, \theta)$;
 - 11 Compute similarity between h^p and each training patient in H^{train} ;
 - 12 Select K most similar patients $p_1, p_2, \dots, p_K$;
 - 13 Trace associated treatment paths $\gamma^{p_1}, \gamma^{p_2}, \dots, \gamma^{p_K}$ through tracing variables Λ ;
 - 14 $C = \bigcup_{i=1}^K \gamma^{p_i}$;
 - 15 **for** each $dr \in C$ **do**
 - 16 Compute $\overline{hit}_{dr}^p$ by equation (3);
 - 17 Return top M drugs sorted by $\overline{hit}_{dr}^p$;
-

6. Experimental Evaluation

In this section, we focus on evaluating the efficacy of our proposed treatment recommendation framework in comparison to several explainable methods for medical

domain problem. We also present intermediate results expressing clinical meanings and the possible interpretation of resulting prescription trees.

6.1 Evaluation Metric

Let us denote the notations for this section are as follows:

- M : the number of recommended drugs.
- T : the test set, i.e set of new patients.
- n : the number of treatment periods.
- $\hat{D}_p^{\pi_j}$: the set of recommended drugs for the testing patient p in period π_j .
- $D_p^{\pi_j}$: the set of actual prescription drugs for p in period π_j .

We use precision, recall, and F1 score, the three well-known evaluation metrics, to evaluate the performance of our treatment recommendation framework. The formulas of these metrics are given as follows.

$$\begin{aligned}
 recall &= \frac{1}{|T| \times n} \sum_{p \in T} \sum_{j=1}^n \frac{|\hat{D}_p^{\pi_j} \cap D_p^{\pi_j}|}{|D_p^{\pi_j}|} \\
 precision &= \frac{1}{|T| \times n} \sum_{p \in T} \sum_{j=1}^n \frac{|\hat{D}_p^{\pi_j} \cap D_p^{\pi_j}|}{M} \\
 F1 &= \frac{2 \times precision \times recall}{precision + recall}
 \end{aligned}$$

6.2 Dataset

Our experiments were performed on MIMIC databases [45], a real world publicly available EMRs database. It consists of nearly 60000 patients who stayed in critical care units of the Beth Israel Deaconess Medical Center between 2001 and 2012.

Although our framework is designed for a set of patients who have the same diagnostic codes, it is uneasy to find such large datasets in reality since patients are usually diagnosed with non-identical series of ICD-9 codes. Therefore, we consider patients with the same first diagnostic code as a cohort. Table 3 reports top five single admission cohorts in the MIMIC III database.

We extracted patient records of the first, the second and the fifth cohort for our experimental evaluation as these cohorts seem to be less relevant to each other. For

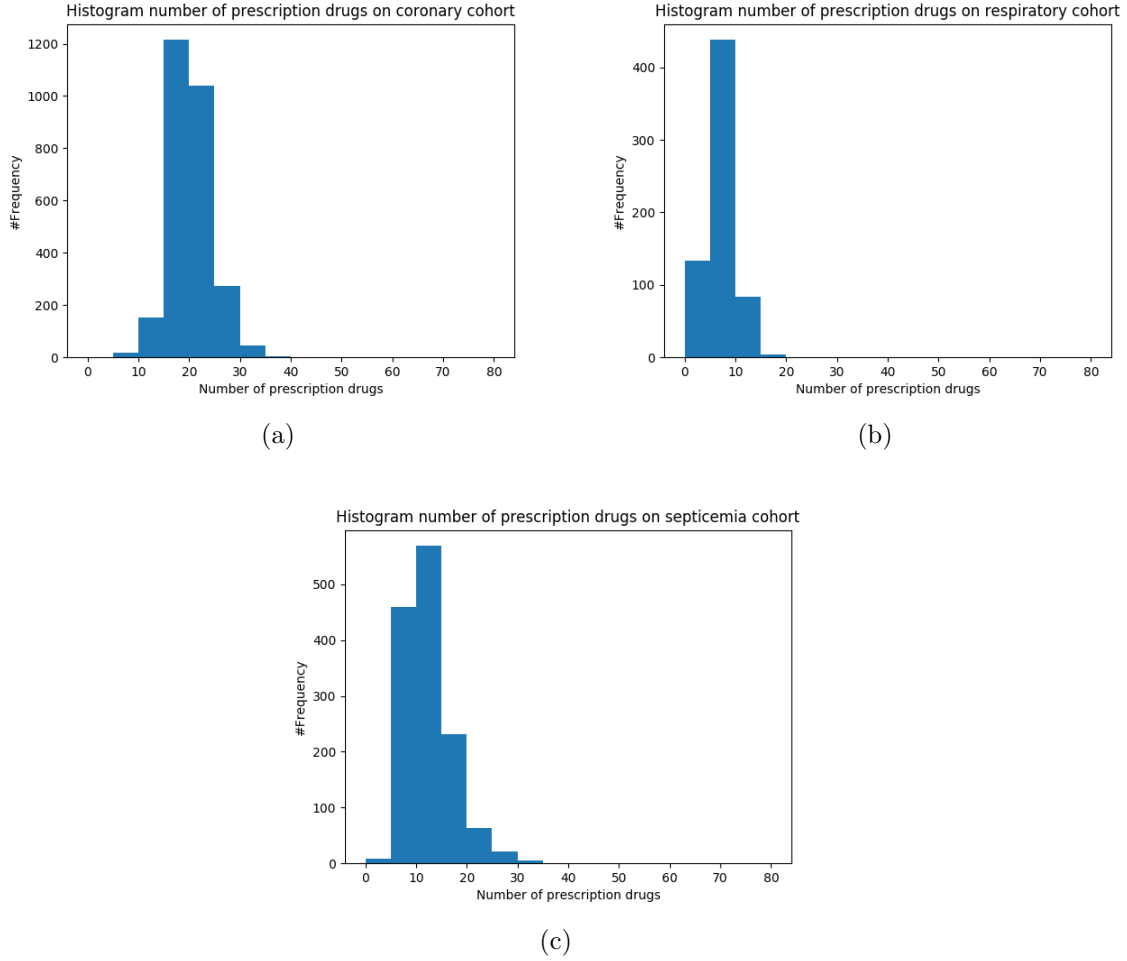


Figure 5: The histogram of number of prescription drugs per patient in three cohorts.

short, we name the three cohort as the coronary artery, septicemia, and respiratory cohort. We extract only patients who were prescribed more than three days for each subcohort. Those with no prescription data or nursing note data are not taken into account in our experiments. We use cTAKES to extract initial signs/symptoms. Sign/symptom features which appear less than 5% or greater than 95% in the cohort are excluded. As not all patients have laboratory tests on all indicators, we fill out the indicator features with unavailable values to 0. For demographic data, we only collect the features that probably affect patient health status for example gender, age, marriage status. For prescription data, we carefully normalize drug names by selecting the most appropriate UMLS term for each prescription drug. Some UMLS

Primary ICD 9	Name	Number of patients
41401	Coronary atherosclerosis of native coronary artery	3430
0389	Unspecified septicemia	1805
41071	Subendocardial infarction	1654
4241	Aortic valve disorders	1122
51881	Acute respiratory failure	945

Table 3: Top 5 primary codes with single admission

	Coronary artery	Septicemia	Respiratory
Number of patients	3430	1805	945
Number of processed patients	2751	1359	658
Number of prescription drugs	1251	1238	1047
Number of normalized drugs	599	630	537
Number of drugs with relevant indication labels	244	190	100

Table 4: Statistic about datasets used in our experimental evaluation

medication terms have equivalent terms in DrugBank database but it is hard to find their indication text if using the UMLS medication terms. For such cases, we use the DrugBank terms instead. It is noted that we only recommend the prescription drugs with indication labels as they are highly relevant drugs to treat the primary diagnosis code.

Table 4 presents an overview of the datasets before and after preprocessing. There are more than 500 prescription drugs delivered to each cohort. Figure 5 provides the histogram of the number of prescription drugs with indication labels delivered to each patient on the three cohorts and Table 5 provides a list of top 15 prescription drugs in each cohort. Most patients were treated with more than 10 labeled drugs.

6.3 Parameter setting

This section describes our parameter selection. We set the number of hidden units in the trained MV.RBM models to 100 units since the learning error rate does not decrease significantly with a larger size of hidden units. Regarding the number of subcohorts set in the clustering analysis step, Figure 6 shows the resulting dendrogram of coronary, respiratory and septicemia cohort, respectively. We find that in these cohorts, training patients are well separated at the distance above 0.6. However,

Coronary	potassium chloride; aspirin; metoprolol; acetaminophen; insulin lispro; docusate; sodium chloride; furosemide; oxycodone; magnesium sulfate
Septicemia	vancomycin; heparin; potassium chloride; sodium chloride; magnesium sulfate; insulin lispro; acetaminophen; furosemide; pantoprazole; docusate
Respiratory	heparin; potassium chloride; salbutamol; furosemide; insulin lispro; vancomycin; sodium chloride; ipratropium; docusate; metoprolol

Table 5: The top 10 prescription drugs in three cohorts

y

splitting the training patients at this distance will result in large size subcohorts where treatment in each subcohort may vary considerably. For this reason, we cut the dendrograms at the distance 0.4 to obtain more small-size clusters.

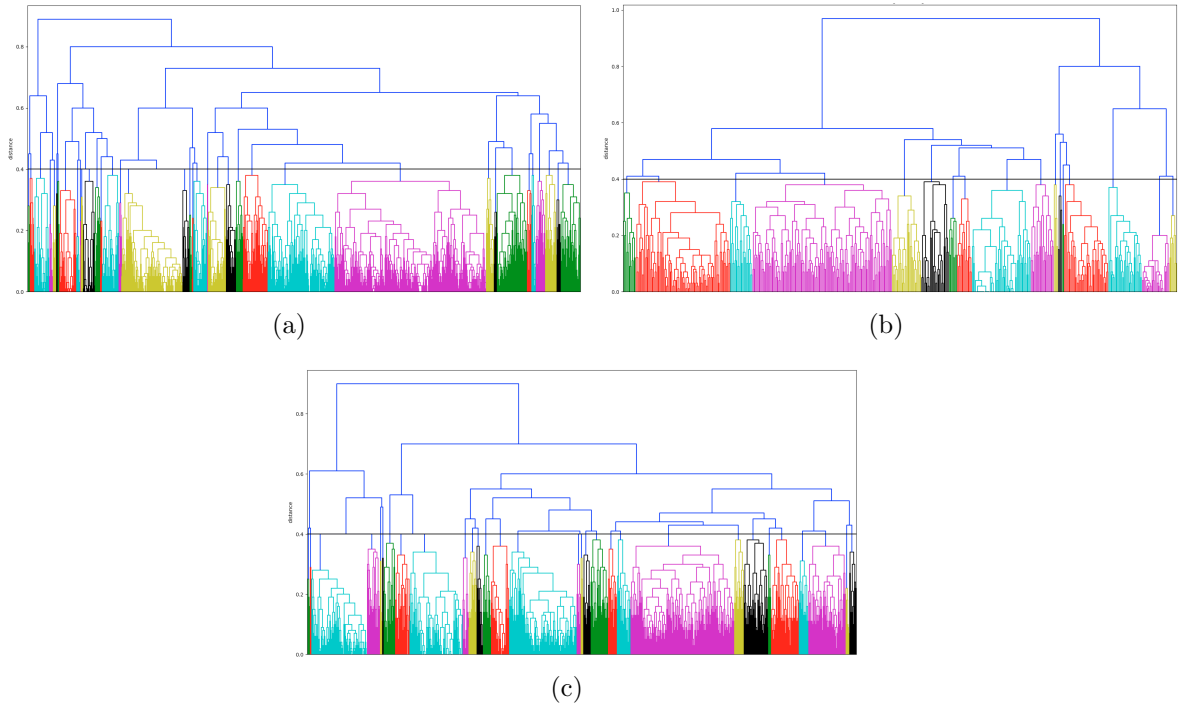


Figure 6: The clustering results of coronary cohort, respiratory cohort, and septicemia cohort, respectively.

Regarding the depth parameter of the prescription tree, we set it to the number of prescription drugs M to be recommended. For the number of treatment periods n , it is not easy to be confirmed by domain knowledge. Indeed, in our work, it plays a role as to which extent our recommendation framework concerns about the order of groups

of recommendation drugs. The higher value of n , the more important to force these groups to preserve the chronological order of delivered time. For this particular work, we split the prescription records of patients in all datasets into three periods where the splitting point are the timestamps of which the associated accumulated indication scores change significantly. For the weighting parameters of indication groups in the treatment period identification algorithm, we ask for an expert’s advice to assign the weight of main group, symptom group, risk factor group to 1, 0.7, 0.5, respectively.

In the treatment recommendation framework using both weighting and non-weighting approach, we vary the number of neighbors K with 5, 10, 15, 20, 30, 40, 50, 80, 100 and 200 neighbors. For the weighting approach, we split the training set into 5 subsets and then learn the hitting score of nodes in the treatment paths of training patients. The number of sub training neighbors K_1 is varied with 50, 100, 150, 200 neighbors.

6.4 Baseline

We consider our treatment recommendation problem as the top-M item recommendation problem where users are patients and items are prescription drugs. Although rich side information about patients such as patient demographic, indicator, progressive data is available, it is not straightforward to exploit such information to leverage user preferences, i.e how likely a prescription drug is given to a patient. For this reason, we compare our proposed recommendation framework to an API implemented for recommender system with implicit feedback in Graphlab library ¹. This API was implemented based on the idea presented in [46, 47, 48]. In the implicit case, there is no target value, the API uses the logistic loss to fit a model that attempts to predict all the given (user, item) pairs in the training data as 1 and all others as 0 ².

We report the results using three solvers: implicit alternative least square [47], stochastic gradient descent [49] and adaptive stochastic gradient descent. For short, we name these baselines as GraphLab + IALS, Graphlab + SGD and Graphlab + ASDG, respectively. We also test our framework when using $K = 1$, i.e recommending based on prescription drugs of the nearest neighbor patient only.

To estimate the hitting weight of the nodes in prescription tree more accurately, we prefer to set a relative large number of K_1 . We fix $K_1 = 100$ and report the results by using weighting and non-weighting approach in three scenario as follows.

1. <https://turi.com/products/create/docs/graphlab.toolkits.recommender.html>
2. https://turi.com/products/create/docs/generated/graphlab.recommender.ranking_factorization_recommender.RankingFactorizationRecommender.html

Sample text	Coronary artery disease (CAD), also known as ischemic heart disease (IHD), refers to a group of diseases which includes stable angina, unstable angina, myocardial infarction, and sudden cardiac death. It is within the group of cardiovascular diseases of which it is the most common type.
Extracted terms	coronary artery disease; coronary heart disease; cardiovascular diseases; coronary arteriosclerosis; myocardial ischemia; arteriopathic disease; myocardial infarction; stable angina; angina, unstable; coronary artery

Table 6: Sample extracted UMLS terms by using cTAKES.

- Using a relative small value of K ($K = 20$), named shortly as Weighting ($K=20$) and Non-Weighting ($K = 20$).
- Using a relative large value of K ($K = 50$), named shortly as Weighting ($K=50$) and Non-Weighting ($K = 50$).
- Using a large value of K ($K = 200$), named shortly as Weighting ($K=200$) and Non-Weighting ($K = 200$).

7. Result

We first present intermediate results followed by an example of a resulting prescription tree and its interpretation. Lastly, we show the performance of our proposed framework for the treatment recommendation task. Table 6 gives an example of UMLS construction extraction for the definition of coronary artery disease by using cTAKES. It can be seen that all extracted UMLS terms are relevant to coronary artery disease or its close diseases.

Table 7, 8, 9, 10, 11, 12 provide lists of extracted UMLS terms for the three text sections and medication terms belonging to the three indication groups of coronary artery cohort, respiratory cohort and septicemia cohort, respectively. We note that extraction of UMLS terms for the definition, the sign/symptom and the risk factor text sections of each disease cohort requires a little efforts to double check and eliminate possible irrelevant concepts that may be mixed in the results.

Figure 7 gives an example of resulting prescription tree over a specific period in a subcohort of the coronary cohort. It is noted that some nodes have a prefix “m”, “s”, “r” to indicate the associated prescription drugs are primary, sign/symptom or risk

Definition	coronary artery disease; coronary heart disease; cardiovascular diseases; coronary arteriosclerosis; myocardial ischemia; arteriopathic disease; myocardial infarction; stable angina; angina, unstable; coronary artery
Signs/symptoms	common cold; dyssomnias; coronary artery rupture; arteriopathic disease; myocardial infarction; stomach diseases; cardiac arrest; heart failure; stable angina; angina, unstable; cold sweat; feeling tired; vomiting; tired; fatigue; heartburn; chest pain; dyspepsia; pain; light-headedness; nausea; chest discomfort; dyspnea; actual discomfort; syncope; stress; does walk up stairs; palpitations; sweating; irregular heart beat; cardiac arrhythmia; angina pectoris; swelling; does climb; infarction; obstruction; occlusion of artery (disorder); thrombus
Risk factors	heart diseases; atherosclerosis; diabetes mellitus; obesity; rheumatoid arthritis; metabolic syndrome x; chronic kidney diseases; endometriosis; lupus erythematosus; hypercholesterolemia; cerebrovascular accident; hypertensive disease; arteriopathic disease; kidney diseases; arthritis; adult disease; overweight; malaise; stress; diet poor; hypokinesia; suicidal; depressive disorder; drug abuse; sexual abuse; insulin resistance

Table 7: Extracted UMLS terms for the text sections about the definition, the typical symptoms and the risk factors of coronary artery cohort.

Primary	nifedipine; eplerenone; aprotinin; timolol; isosorbide mononitrate; ibuprofen; ramipril; candesartan cilexetil; valsartan; atorvastatin; nicardipine; vitamin e; dopamine; cangrelor; abciximab; captopril
Signs/symptoms	sulindac; indapamide; sertraline; aluminum hydroxide; triamterene; tramadol; nitroglycerin; oxycodone; dipyridamole; dyclonine; nitroprusside; magnesium oxide; procainamide; cholestyramine; fentanyl; fosinopril
Risk factors	escitalopram; betaxolol; repaglinide; terazosin; dorzolamide; lanthanum carbonate; insulin glargine; pindolol; felodipine; brimonidine; prednisolone; desmopressin; vasopressin; sevelamer; glimepiride; ezetimibe

Table 8: Medication terms belonging to the primary, the sign/symptom, the risk factor indication groups of coronary artery cohort.

Definition	respiratory failure without hypercapnia; acute respiratory failure; acute-on-chronic respiratory failure; spastic ataxia, charlevoix-saguenay type; respiratory failure; chronic respiratory failure; respiratory depression; respiratory tract structure; respiratory tract infections; lower respiratory tract structure; lower respiratory tract infection
Signs/symptoms	increased sweating; restlessness; shallow breathing; cardiac arrhythmia; irregular heart beat; unconscious state; drowsiness; tachypnea; sweating; confusion; anxiety
Risk factors	chronic obstructive airway disease; lung diseases; communicable diseases; asthma; heart failure; airway obstruction; chronic lung disease; infectious disease of lung; respiratory failure; thrombophilia; pneumothorax; respiration disorders

Table 9: Extracted UMLS terms for the text sections about the definition, the typical symptoms and the risk factors of respiratory cohort.

Primary	clindamycin; ciprofloxacin; ceftazidime; amoxicillin; cephalixin; trimethoprim; tramadol; doxycycline; naloxone; guaifenesin; erythromycin; clarithromycin; aztreonam; benzocaine; dyclonine
Signs/symptoms	paroxetine; escitalopram; quinidine; chloral hydrate; oxazepam; modafinil; melatonin; diazepam; alprazolam; digoxin; doxepin; ropinirole; disopyramide; chlorpromazine; fluoxetine; lorazepam
Risk factors	theophylline; nesiritide; tobramycin; furosemide; cefepime; valsartan; hydrocortisone; nevirapine; vitamin e; irbesartan; captopril; ampicillin; torsemide; ethacrynate; nitroprusside; ribavirin

Table 10: Medication terms belonging to the primary, the sign/symptom, the risk factor indication groups of coronary septicemia cohort.

Definition	bacterial infections; bacteremia; sepsis; communicable diseases; bacterial sepsis
Signs/symptoms	oliguria; hyperglycemia; dehydration; common cold; alkalosis; diarrhea; lightheadedness; actual discomfort; chest pain; syncope; vomiting; dyspnea; nausea; pain; death anxiety; cold intolerance; exanthema; agitation; tremor; dizziness; weakness; chills; fever with chills; single organ dysfunction; myalgia; fever
Risk factors	infections, hospital; chronic disease; acquired immunodeficiency syndrome; diabetes mellitus; kidney diseases; sepsis due to fungus; infections of musculoskeletal system; pneumonia; soft tissue infections; hiv infections; urinary tract infection; candidiasis; senility

Table 11: Extracted UMLS terms for the text sections about the definition, the typical symptoms and the risk factors of septicemia cohort.

Primary	benzylpenicillin; cefepime; delavirdine; doxycycline; nevirapine; dopamine; raltegravir; tenofovir disoproxil; ampicillin; neomycin; ciprofloxacin; sulfamethoxazole; imipenem; cefalotin; lamivudine; oseltamivir
Signs/symptoms	clonidine; meperidine; escitalopram; nesiritide; paroxetine; nortriptyline; haloperidol; tramadol; cyclobenzaprine; loperamide; sodium bicarbonate; hydrocodone; oxycodone; ibuprofen; chlorpheniramine; hydromorphone
Risk factors	vasopressin; vancomycin; daptomycin; amikacin; lanthanum carbonate; bacitracin; furosemide; ramipril; repaglinide; levocarnitine; pentamidine; maraviroc; famciclovir; nitrofurantoin; amphotericin b; rosiglitazone

Table 12: Extracted septicemia's medication terms belonging to the primary, the sign/symptom, the risk factor indication groups.

factor drugs, respectively. The suffix number in each node named as node frequency indicates how many patients were prescribed by drugs on the path from the root node until that node excluding the drug with higher prescription frequency on the same level of that drug and those with higher prescription frequency on the same level of parent nodes. The sample prescription tree can be interpreted as follow. Starting from the root node, the primary drug metoprolol is the most prescription drug which was prescribed for 59 patients. Among the patients treated with metoprolol, the symptom drug furosemide is the most prescription drug treating 48 patients . Among the patients who were not treated with metoprolol, aspirin is the most prescription drug treating 3 patients. The remaining nodes in the tree can be explained in a similar way. It can be noticed that the treatment path (metoprolol, furosemide, potassium chloride, insulin lipspro, aspirin, acetaminophen, docusate, bisacodyl, magnesium hydroxide, oxycodone) was prescribed with high node frequency. This set can be considered as the primary treatment pattern of the considering subcohort.

In the indication assignment framework, we can recognize not only the indication group of a prescription drug but also the signs/symptoms treated by that drug. This property allows interpreting in depth the discovered treatment pattern sets. For example, besides the main drugs metoprolol , aspirin (cure myocardial infarction, the UMLS term highly relevant to the definition of coronary artery disease), the primary treatment pattern set also includes symptom drugs furosemide, acetaminophen, magnesium hydroxi (cure pain symptom, heart failure, heartburn symptom), risk factor drugs insulin lispro (cure diabetes mellitus disease). It can be inferred that most of the patients in the subcohort of sample prescription tree probably suffered symptoms of heart disease and diabetes disease.

Besides the set of frequent pattern drugs, the resulting prescription tree also allows recognizing a set of drugs which are not frequently described together. For instance, in the above example, the drug ibuprofen is not likely to be prescribed together with metoprolol. Another example is the case of potassium chloride. It can be seen that among patients who were prescribed with metoprolol, potassium chloride are rarely used without furosemide. The above example shows the usefulness of prescription trees constructed by our treatment learning framework. Such property is not easy to be recognized by learning frequent treatment patterns only.

Table 13, 14, 15 report the performance of the proposed treatment recommendation framework and the baselines on coronary artery, septicemia and respiratory cohort, respectively. Note that in our experiments, we randomly selected 20% of pa-

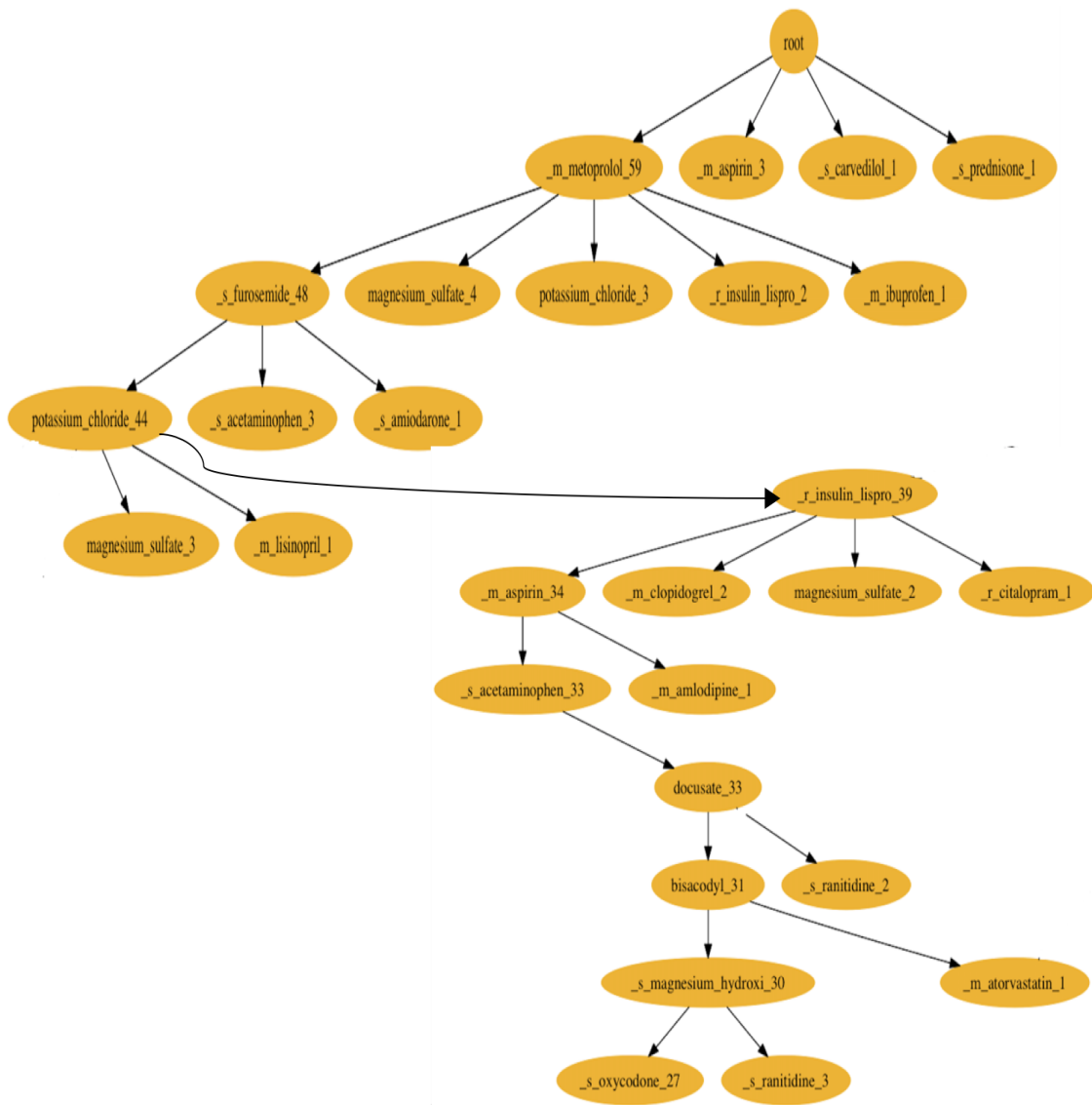


Figure 7: An example of resulting prescription tree.

tients in each cohort for testing. The obtained results show that in best cases, our framework achieves competitive results to the baselines over all evaluation metrics. Among three cohorts, the weighting approach seems to work well with the septicemia and respiratory cohort. This can be explained by the fact that treatments among similar patients in the coronary cohort may not be much different compared to those of the two other cohorts (higher obtained precision). In all cohorts, using a relative large enough number of neighbors ($K = 50$) yields almost as good as the best reported cases. This property is useful to show case-patients the recommendation results based on. Moreover, in the septicemia and respiratory cohorts, the higher obtained values when using weighting approach with a small number of neighbors ($K = 20, K = 50$) compared to the non-weighting approach indicate that it is better to employ the weighting approach to deal with the high variety of treatments in these cohorts. The poor performance when using the nearest neighbor approach confirms that combining treatment from neighbors are necessary for improving the efficacy of treatment recommendation frameworks, especially when identifying the nearest neighbor patient is almost uncertain.

8. Discussion

Our work obtains more interesting results in terms of domain exploitation and knowledge representation compared to related works in the literature. First, rather than defining treatment periods as fixed intervals, we track the change of drug indication in prescribed drugs as a hint to discover treatment periods. It can be seen that the idea fits our natural thinking on detecting patients' treatment periods given their prescription records. Second, by representing the learned treatment patterns in form of prescription trees, our treatment learning framework not only reflects the usage-frequency of drugs fully but also allows doctors to quickly recognize groups of frequent and infrequent prescription drugs in each patient subcohort. Therefore, in terms of knowledge representation, it could be said that our treatment learning framework is superior to most of the current studies which focused solely on frequent treatment patterns.

Another finding of our work is that the idea of learning from neighbor patients seems to work fine for the treatment recommendation task. Our experimental evaluation has shown that combining treatment patterns from many neighbors could be useful to improve the efficacy of the treatment recommendation framework. We have proposed the weighting and non-weighting recommendation approach with competi-

Metric	M Method	3	5	7	10
Precision	CF+IALS	0.444	0.445	0.481	0.462
	CF+ADA	0.768	0.702	0.652	0.58
	CF+SGD	0.778	0.697	0.648	0.577
	Nearest Neighbor	0.692	0.597	0.526	0.396
	Weighting (K=20)	0.763	0.685	0.635	0.567
	Non-Weighting (K=20)	0.757	0.7	0.661	0.59
	Weighting (K=50)	0.772	0.694	0.645	0.572
	Non-Weighting (K=50)	0.757	0.701	0.661	0.597
	Weighting (K=200)	0.768	0.688	0.641	0.569
	Non-Weighting (K=200)	0.752	0.696	0.656	0.594
Recall	CF+IALS	0.125	0.242	0.346	0.472
	CF+ADA	0.267	0.378	0.471	0.583
	CF+SGD	0.27	0.368	0.47	0.579
	Nearest Neighbor	0.237	0.317	0.377	0.4
	Weighting (K=20)	0.265	0.367	0.463	0.573
	Non-Weighting (K=20)	0.267	0.379	0.481	0.598
	Weighting (K=50)	0.27	0.372	0.469	0.577
	Non-Weighting (K=50)	0.267	0.38	0.48	0.603
	Weighting (K=200)	0.267	0.371	0.469	0.574
	Non-Weighting (K=200)	0.265	0.377	0.477	0.6
F1-score	CF+IALS	0.195	0.313	0.402	0.467
	CF+ADA	0.396	0.491	0.547	0.581
	CF+SGD	0.4	0.482	0.545	0.578
	Nearest Neighbor	0.353	0.414	0.439	0.398
	Weighting (K=20)	0.393	0.478	0.536	0.57
	Non-Weighting (K=20)	0.394	0.492	0.557	0.594
	Weighting (K=50)	0.4	0.485	0.543	0.574
	Non-Weighting (K=50)	0.395	0.493	0.557	0.6
	Weighting (K=200)	0.396	0.482	0.542	0.572
	Non-Weighting (K=200)	0.392	0.489	0.552	0.597

Table 13: Experimental results for treatment recommendation task performed on coronary artery cohort

Metric	M Method	3	5	7	10
Precision	CF+IALS	0.182	0.191	0.208	0.177
	CF+ADA	0.406	0.353	0.319	0.284
	CF+SGD	0.417	0.357	0.321	0.286
	Nearest Neighbor	0.261	0.188	0.139	0.097
	Weighting (K=20)	0.41	0.349	0.317	0.28
	Non-Weighting (K=20)	0.397	0.341	0.305	0.261
	Weighting (K=50)	0.413	0.356	0.322	0.285
	Non-Weighting (K=50)	0.398	0.351	0.311	0.271
	Weighting (K=200)	0.415	0.358	0.321	0.286
	Non-Weighting (K=200)	0.408	0.353	0.317	0.278
Recall	CF+IALS	0.119	0.203	0.287	0.346
	CF+ADA	0.242	0.352	0.439	0.553
	CF+SGD	0.25	0.353	0.44	0.558
	Nearest Neighbor	0.143	0.169	0.174	0.175
	Weighting (K=20)	0.244	0.342	0.434	0.543
	Non-Weighting (K=20)	0.235	0.338	0.42	0.504
	Weighting (K=50)	0.245	0.352	0.444	0.557
	Non-Weighting (K=50)	0.236	0.345	0.426	0.53
	Weighting (K=200)	0.248	0.356	0.443	0.557
	Non-Weighting (K=200)	0.242	0.347	0.437	0.546
F1-score	CF+IALS	0.144	0.196	0.241	0.235
	CF+ADA	0.303	0.353	0.37	0.376
	CF+SGD	0.313	0.355	0.371	0.379
	Nearest Neighbor	0.184	0.178	0.154	0.125
	Weighting (K=20)	0.306	0.346	0.366	0.369
	Non-Weighting (K=20)	0.296	0.34	0.353	0.344
	Weighting (K=50)	0.308	0.354	0.373	0.377
	Non-Weighting (K=50)	0.296	0.348	0.36	0.358
	Weighting (K=200)	0.31	0.357	0.372	0.378
	Non-Weighting (K=200)	0.304	0.35	0.368	0.368

Table 14: Experimental results for treatment recommendation task performed on septicemia cohort

Metric	M Method	3	5	7	10
Precision	CF+IALS	0.258	0.24	0.229	0.179
	CF+ADA	0.419	0.33	0.272	0.212
	CF+SGD	0.426	0.336	0.272	0.213
	Nearest Neighbor	0.233	0.153	0.112	0.078
	Weighting (K=20)	0.419	0.334	0.27	0.211
	Non-Weighting (K=20)	0.394	0.307	0.251	0.2
	Weighting (K=50)	0.425	0.332	0.272	0.211
	Non-Weighting (K=50)	0.405	0.322	0.26	0.207
	Weighting (K=200)	0.426	0.335	0.272	0.213
	Non-Weighting (K=200)	0.416	0.331	0.269	0.214
Recall	CF+IALS	0.242	0.397	0.532	0.589
	CF+ADA	0.421	0.543	0.612	0.672
	CF+SGD	0.426	0.554	0.616	0.672
	Nearest Neighbor	0.224	0.237	0.24	0.24
	Weighting (K=20)	0.42	0.548	0.609	0.664
	Non-Weighting (K=20)	0.381	0.499	0.566	0.643
	Weighting (K=50)	0.426	0.55	0.616	0.667
	Non-Weighting (K=50)	0.396	0.526	0.595	0.66
	Weighting (K=200)	0.426	0.549	0.613	0.676
	Non-Weighting (K=200)	0.414	0.545	0.613	0.682
F1-score	CF+IALS	0.249	0.299	0.321	0.275
	CF+ADA	0.42	0.41	0.377	0.323
	CF+SGD	0.426	0.418	0.377	0.323
	Nearest Neighbor	0.228	0.186	0.153	0.118
	Weighting (K=20)	0.42	0.415	0.375	0.32
	Non-Weighting (K=20)	0.388	0.38	0.348	0.305
	Weighting (K=50)	0.425	0.414	0.377	0.321
	Non-Weighting (K=50)	0.4	0.4	0.362	0.316
	Weighting (K=200)	0.426	0.416	0.377	0.324
	Non-Weighting (K=200)	0.415	0.412	0.374	0.326

Table 15: Experimental results for treatment recommendation task performed on respiratory cohort

tive results to the baseline. Each approach is appropriate for some different types of cohorts depending on the degree of treatment variants. More importantly, our case based framework seems to fit well for medical domain where the model's explainability is important. Based on these results we conclude that the proposed approach is generic enough as a case based recommendation framework. Some unsolved issues in this project will be addressed in another project titled "Autonomous action learning with new data-dependent similarity measure and dynamic action recommendation".

9. List of Publications and Significant Collaborations that resulted from our AOARD supported project

9.1 List of peer-reviewed journal publications

9.2 List of peer-reviewed conference publications

- [1] Hoang Hung, Tu Bao Ho: Learning Treatment Regimens from Electronic Medical Records, PAKDD 2018, June 3rd - 6th, 2018, Melbourne.

9.3 Papers published in non-peer-reviewed journals and conference proceedings

9.4 Conference presentations without papers

9.5 Manuscripts submitted but not yet published

Hoang Hung, Tu Bao Ho: Learning and Recommending Treatments from Electronic Medical Records. Submitted to journal Knowledge-Based System.

9.6 Provide a list any interactions with industry or with Air Force Research Laboratory scientists or significant collaborations that resulted from this work

10. Attachments

Publications in sections 9.1, 9.2 and 9.3 listed above if possible.

References

- [1] Annemarie Jutel. Sociology of diagnosis: a preliminary review. *Sociology of health & illness*, 31(2):278–299, 2009.

- [2] William L Healy, Michael E Ayers, Richard Iorio, Douglas A Patch, David Appleby, and Bernard A Pfeifer. Impact of a clinical pathway and implant standardization on total hip arthroplasty: a clinical and economic study of short-term patient outcome. *The Journal of arthroplasty*, 13(3):266–276, 1998.
- [3] Carol L Ireson. Critical pathways: effectiveness in achieving patient outcomes. *Journal of Nursing Administration*, 27(6):16–23, 1997.
- [4] Fenglong Ma, Radha Chitta, Jing Zhou, Quanzeng You, Tong Sun, and Jing Gao. Dipole: Diagnosis prediction in healthcare via attention-based bidirectional recurrent neural networks. In *Proceedings of the 23rd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pages 1903–1911. ACM, 2017.
- [5] Sellappan Palaniappan and Rafiah Awang. Intelligent heart disease prediction system using data mining techniques. In *Computer Systems and Applications, 2008. AICCSA 2008. IEEE/ACS International Conference on*, pages 108–115. IEEE, 2008.
- [6] Shadab Adam Pattekari and Asma Parveen. Prediction system for heart disease using naïve bayes. *International Journal of Advanced Computer and Mathematical Sciences*, 3(3):290–294, 2012.
- [7] Carlos Ordonez. Association rule discovery with the train and test approach for heart disease prediction. *IEEE Transactions on Information Technology in Biomedicine*, 10(2):334–343, 2006.
- [8] Peter WF Wilson and Jane C Evans. Coronary artery disease prediction. *American journal of hypertension*, 6(11_Pt_2):309S–313S, 1993.
- [9] Jyoti Soni, Ujma Ansari, Dipesh Sharma, and Sunita Soni. Predictive data mining for medical diagnosis: An overview of heart disease prediction. *International Journal of Computer Applications*, 17(8):43–48, 2011.
- [10] Felix Gräßer, Stefanie Beckert, Denise Küster, Jochen Schmitt, Susanne Abraham, Hagen Malberg, and Sebastian Zaunseder. Therapy decision support based on recommender system methods. *Journal of healthcare engineering*, 2017, 2017.
- [11] J Mei, H Liu, X Li, G Xie, and Y Yu. A decision fusion framework for treatment recommendation systems. *Studies in health technology and informatics*, 216:300, 2015.

- [12] Lu Wang, Wei Zhang, Xiaofeng He, and Hongyuan Zha. Supervised reinforcement learning with recurrent neural network for dynamic treatment recommendation. In *Proceedings of the 24th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, pages 2447–2456. ACM, 2018.
- [13] Tu Bao Ho Khanh Hung Hoang. Learning treatment regimens using electronic medical records. In *The Pacific-Asia Conference on Data Mining and Knowledge Discovery*, page to appear. Springer, 2018.
- [14] Bo Jin, Haoyu Yang, Leilei Sun, Chuanren Liu, Yue Qu, and Jianing Tong. A treatment engine by predicting next-period prescriptions. In *Proceedings of the 24th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, pages 1608–1616. ACM, 2018.
- [15] Peter A Ubel, Michael L DeKay, Jonathan Baron, and David A Asch. Cost-effectiveness analysis in a setting of budget constraints is it equitable? *New England Journal of Medicine*, 334(18):1174–1177, 1996.
- [16] William A Rushing. Public policy, community constraints, and the distribution of medical resources. *Soc. Probs.*, 19:21, 1971.
- [17] James F Blumstein. Rationing medical resources: A constitutional, legal, and policy analysis. *Tex. L. Rev.*, 59:1345, 1980.
- [18] Christopher Newdick et al. Who should we treat?: rights, rationing, and resources in the nhs. *OUN Catalogue*, 2005.
- [19] Kidney Disease: Improving Global Outcomes (KDIGO) CKD-MBD Work Group et al. Kdigo clinical practice guideline for the diagnosis, evaluation, prevention, and treatment of chronic kidney disease-mineral and bone disorder (ckd-mbd). *Kidney international. Supplement*, (113):S1, 2009.
- [20] British Medical Association et al. *Withholding and withdrawing life-prolonging medical treatment: guidance for decision making*. John Wiley & Sons, 2008.
- [21] H Wada, J Thachil, M Di Nisio, P Mathew, S Kurosawa, S Gando, HK Kim, JD Nielsen, C-E Dempfle, M Levi, et al. Guidance for diagnosis and treatment of disseminated intravascular coagulation from harmonization of the recommendations from three guidelines. *Journal of thrombosis and haemostasis*, 11(4):761–767, 2013.

- [22] Danton S Char, Nigam H Shah, and David Magnus. Implementing machine learning in health care addressing ethical challenges. *The New England journal of medicine*, 378(11):981, 2018.
- [23] Peter B Jensen, Lars J Jensen, and Søren Brunak. Mining electronic health records: towards better research applications and clinical care. *Nature Reviews Genetics*, 13(6):395–405, 2012.
- [24] Truyen Tran, Dinh Phung, and Svetha Venkatesh. Mixed-variate restricted boltzmann machines. *arXiv preprint arXiv:1408.1160*, 2014.
- [25] Fu-ren Lin, Shien-chao Chou, Shung-mei Pan, and Yao-mei Chen. Mining time dependency patterns in clinical pathways. *International Journal of Medical Informatics*, 62(1):11–25, 2001.
- [26] Haytham Elghazel, Veronique Deslandres, Kassem Kallel, and Alain Dussauchoy. Clinical pathway analysis using graph-based approach and markov models. In *Digital Information Management, 2007. ICDIM'07. 2nd International Conference on*, volume 1, pages 279–284. IEEE, 2007.
- [27] Loubna Bouarfa and Jenny Dankelman. Workflow mining and outlier detection from clinical activity logs. *Journal of biomedical informatics*, 45(6):1185–1190, 2012.
- [28] You Chen, Wei Xie, Carl A Gunter, David Liebovitz, Sanjay Mehrotra, He Zhang, and Bradley Malin. Inferring clinical workflow efficiency via electronic medical record utilization. In *AMIA annual symposium proceedings*, volume 2015, page 416. American Medical Informatics Association, 2015.
- [29] Zhengxing Huang, Wei Dong, Peter Bath, Lei Ji, and Huilong Duan. On mining latent treatment patterns from electronic medical records. *Data Mining and Knowledge Discovery*, 29(4):914–949, 2015.
- [30] Hsin-Min Lu, Chih-Ping Wei, and Fei-Yuan Hsiao. Modeling health-care data using multiple-channel latent dirichlet allocation. *Journal of biomedical informatics*, 60:210–223, 2016.
- [31] Xiao Xu, Tao Jin, Zhijie Wei, Cheng Lv, and Jianmin Wang. Tcpcm: topic-based clinical pathway mining. In *Connected Health: Applications, Systems and Engineering Technologies (CHASE), 2016 IEEE First International Conference on*, pages 292–301. IEEE, 2016.
- [32] Sungrae Park, Doosup Choi, Minki Kim, Wonchul Cha, Chuhyun Kim, and Il-Chul Moon. Identifying prescription patterns with a topic model

- of diseases and medications. *Journal of biomedical informatics*, 75:35–47, 2017.
- [33] Liang Yao, Yin Zhang, Baogang Wei, Wenjin Zhang, and Zhe Jin. A topic modeling approach for traditional chinese medicine prescriptions. *IEEE Transactions on Knowledge and Data Engineering*, 30(6):1007–1021, 2018.
 - [34] Trang Pham, Truyen Tran, Dinh Phung, and Svetha Venkatesh. Deep-care: A deep dynamic memory model for predictive medicine. In *Pacific-Asia Conference on Knowledge Discovery and Data Mining*, pages 30–41. Springer, 2016.
 - [35] Hung Le, Truyen Tran, and Svetha Venkatesh. Dual control memory augmented neural networks for treatment recommendations. In *Pacific-Asia Conference on Knowledge Discovery and Data Mining*, pages 273–284. Springer, 2018.
 - [36] Yufan Zhao, Donglin Zeng, Mark A Socinski, and Michael R Kosorok. Reinforcement learning strategies for clinical trials in nonsmall cell lung cancer. *Biometrics*, 67(4):1422–1433, 2011.
 - [37] Ying Liu, Brent Logan, Ning Liu, Zhiyuan Xu, Jian Tang, and Yangzhi Wang. Deep reinforcement learning for dynamic treatment regimes on medical registry data. In *Healthcare Informatics (ICHI), 2017 IEEE International Conference on*, pages 380–385. IEEE, 2017.
 - [38] Shamim Nemati, Mohammad M Ghassemi, and Gari D Clifford. Optimal medication dosing from suboptimal clinical examples: A deep reinforcement learning approach. In *Engineering in Medicine and Biology Society (EMBC), 2016 IEEE 38th Annual International Conference of the*, pages 2978–2981. IEEE, 2016.
 - [39] Shoji Hirano and Shusaku Tsumoto. Mining typical order sequences from ehr for building clinical pathways. In *Pacific-Asia Conference on Knowledge Discovery and Data Mining*, pages 39–49. Springer, 2014.
 - [40] Leilei Sun, Chuanren Liu, Chonghui Guo, Hui Xiong, and Yanming Xie. Data-driven automatic treatment regimen development and recommendation. In *KDD*, pages 1865–1874, 2016.
 - [41] Guergana K Savova, James J Masanz, Philip V Ogren, Jiaping Zheng, Sunghwan Sohn, Karin C Kipper-Schuler, and Christopher G Chute. Mayo clinical text analysis and knowledge extraction system (ctakes): architecture, component evaluation and applications. *Journal of the American Medical Informatics Association*, 17(5):507–513, 2010.

- [42] Olivier Bodenreider. The unified medical language system (umls): integrating biomedical terminology. *Nucleic acids research*, 32(suppl_1):D267–D270, 2004.
- [43] Darius Tamasauskas, Virgilijus Sakalauskas, and Dalia Kriksciuniene. Evaluation framework of hierarchical clustering methods for binary data. In *Hybrid Intelligent Systems (HIS), 2012 12th International Conference on*, pages 421–426. IEEE, 2012.
- [44] Jingfeng Chen, Leilei Sun, Chonghui Guo, Wei Wei, and Yanming Xie. A data-driven framework of typical treatment process extraction and evaluation. *Journal of biomedical informatics*, 83:178–195, 2018.
- [45] Alistair EW Johnson, Tom J Pollard, Lu Shen, Li-wei H Lehman, Mengling Feng, Mohammad Ghassemi, Benjamin Moody, Peter Szolovits, Leo Anthony Celi, and Roger G Mark. Mimic-iii, a freely accessible critical care database. *Scientific data*, 3, 2016.
- [46] Yehuda Koren, Robert Bell, and Chris Volinsky. Matrix factorization techniques for recommender systems. *Computer*, (8):30–37, 2009.
- [47] Yifan Hu, Yehuda Koren, and Chris Volinsky. Collaborative filtering for implicit feedback datasets. In *Data Mining, 2008. ICDM’08. Eighth IEEE International Conference on*, pages 263–272. Ieee, 2008.
- [48] Steffen Rendle. Factorization machines. In *Data Mining (ICDM), 2010 IEEE 10th International Conference on*, pages 995–1000. IEEE, 2010.
- [49] Léon Bottou. Stochastic gradient descent tricks. In *Neural networks: Tricks of the trade*, pages 421–436. Springer, 2012.