

KGML-xDTD: a knowledge graph-based machine learning framework for drug treatment prediction and mechanism description

Chunyu Ma^{1,*}, Zhihan Zhou², Han Liu² and David Koslicki^{1,3,4,*}

¹Huck Institutes of Life Sciences, Pennsylvania State University, State College, PA 16801, USA

²Department of Computer Science, Northwestern University, Evanston, IL 60208, USA

³Department of Computer Science and Engineering, Pennsylvania State University, State College, PA 16801, USA

⁴Department of Biology, Pennsylvania State University, State College, PA 16801, USA

*Correspondence address. Chunyu Ma, Westgate building, Pennsylvania State University, University Park, PA 16802 USA. E-mail: cqm5886@psu.edu; David Koslicki, Westgate building, Pennsylvania State University, University Park, PA 16802 USA. E-mail: dmk333@psu.edu

Abstract

Background: Computational drug repurposing is a cost- and time-efficient approach that aims to identify new therapeutic targets or diseases (indications) of existing drugs/compounds. It is especially critical for emerging and/or orphan diseases due to its cheaper investment and shorter research cycle compared with traditional wet-lab drug discovery approaches. However, the underlying mechanisms of action (MOAs) between repurposed drugs and their target diseases remain largely unknown, which is still a main obstacle for computational drug repurposing methods to be widely adopted in clinical settings.

Results: In this work, we propose KGML-xDTD: a Knowledge Graph-based Machine Learning framework for explainably predicting Drugs Treating Diseases. It is a 2-module framework that not only predicts the treatment probabilities between drugs/compounds and diseases but also biologically explains them via knowledge graph (KG) path-based, testable MOAs. We leverage knowledge-and-publication-based information to extract biologically meaningful “demonstration paths” as the intermediate guidance in the Graph-based Reinforcement Learning (GRL) path-finding process. Comprehensive experiments and case study analyses show that the proposed framework can achieve state-of-the-art performance in both predictions of drug repurposing and recapitulation of human-curated drug MOA paths.

Conclusions: KGML-xDTD is the first model framework that can offer KG path explanations for drug repurposing predictions by leveraging the combination of prediction outcomes and existing biological knowledge and publications. We believe it can effectively reduce “black-box” concerns and increase prediction confidence for drug repurposing based on predicted path-based explanations and further accelerate the process of drug discovery for emerging diseases.

Keywords: drug repurposing, reinforcement learning, biomedical knowledge graph

Introduction

Traditional drug development is a time-consuming process (from initial chemical identification to clinical trials and finally to Food and Drug Administration [FDA] approval) that takes around 10 to 15 years and also comes along with billions-of-dollars investments and high failure rates [1]. Considering the rapid pace of novel disease evolution, it is urgent to find a more efficient and economical drug discovery method. Fortunately, it has been observed that a single drug can often be effective in treating multiple diseases. For example, thalidomide was originally used as an antianxiety medication [2] and was later found to have the anticancer potential for the treatment of cancers [3, 4]. Hence, drug repurposing, also known as the identification of new uses for existing drugs/compounds, might bring us hope to address this urgent need with the advantage of a shorter research cycle, lower development cost, and more preexisting safety tests.

Existing drug repurposing approaches can roughly be categorized into experimental-based approaches (e.g., binding affinity assays [5], phenotypic screening [6]), clinical-based approaches (e.g., off-label drug use analysis [7]), and computational-based

approaches (e.g., chemical structure-based [8] and GWAS-based approaches [9]). Compared with the former 2 approaches, the computational approaches are more cost- and time-efficient, particularly when the goal is to prioritize a large number of target drugs/compounds for follow-up experimental investigation. Among all computational drug repurposing methods, the integration of multiple biomedical data sources into a so-called biomedical knowledge graph (BKG) for drug discovery has become popular in recent years [10] due to the increasing availability of curated biomedical databases such as DrugBank [11], ChEMBL [12], and HMDB [13] and the advancement of semantic web techniques [14]. There are 3 types of existing BKGs: database-based BKGs (e.g., Hetionet [15], BioKG [16], iBKH [17]) are constructed by integrating biomedical data and their relations stored in existing biological databases. The literature-based BKGs (e.g., GNBR [18]) are built by leveraging Natural Language Processing (NLP) techniques to extract semantic information from a large amount of available biomedical literature and electronic health record (EHR) data, which are mostly disease specific [19–21]. The mixed BKGs (e.g.,

Received: January 31, 2023. Revised: May 5, 2023. Accepted: July 4, 2023

© The Author(s) 2023. Published by Oxford University Press GigaScience. This is an Open Access article distributed under the terms of the Creative Commons Attribution License (<https://creativecommons.org/licenses/by/4.0/>), which permits unrestricted reuse, distribution, and reproduction in any medium, provided the original work is properly cited.

CKG [22], RTX-KG2 [23]) are generated by combining the knowledge sources from the above 2 methods.

Based on these BKGs, several machine learning methods have been proposed or implemented for drug repurposing prediction by treating it as a link prediction task in the BKGs. For example, Himmelstein et al. [15] used the so-called degree-weighted path count (DWPC) to assess the prevalence of 1,206 metapaths and then classified drug–disease treatment relations by fitting these DWPC features to a logistic regression model. Ioannidis et al. [24] proposed a novel graph neural network model I-RGCN to learn the node and relation embeddings for the Covid-19 drug repurposing task. Zhang et al. [19] recently predicted the possible drugs for Covid-19 with 5 existing popular knowledge graph completion methods (e.g., TransE [25], RotatE [26], DistMult [27], ComplEx [28], and STELP [29]). Although some of these models have shown good performance in drug repurposing prediction on the small-scale BKGs, none have been scaled to massive BKGs with more than millions of nodes and edges and make a comprehensive comparison. More importantly, most of them lack the biological explanatory ability for their predictions, which limits their applicability in clinical research.

Currently, there are few computational models designed for drug repurposing explanations. A common and intuitive explanation based on a biomedical knowledge graph for drug repurposing leverages the semantic BKG-based paths between given drug–disease pairs. Sosa et al. [30] applied a graph embedding model UKGE [31], which utilizes the weighted (the frequency of relation appeared in literature) relation edges in a literature-based KG GNBR to identify new indications of drugs for rare diseases and then explain the results via the highest-ranking paths based on confidence scores. However, this method is only applicable in the literature-based BKGs with the weighted edge information. Most BKGs using database-based knowledge do not contain such information. Sang et al. [32] proposed GrEDeL that combines the TransE embedding method with a long short-term memory (LSTM) recurrent neural network (RNN) model to predict drug–disease relation. By using the embeddings of BKG paths as model input for predictions, they can provide path-based explanations. However, they claimed that the effectiveness of the approach relies heavily on the NLP tool SemRep, which is reported to have high false positives in named entity recognition [33]. Also, they did not fully evaluate how biologically reasonable their predicted path-based mechanisms of action (MOAs) are.

Besides the existing methods above, we view reinforcement learning (RL) as a promising solution for drug repurposing explanation. RL models solve the decision-making problem, in which an agent learns how to take appropriate actions to maximize cumulative rewards through interactions with the environment. RL has achieved widespread success in various domains, including games, recommendation systems, health care, transportation, and so on [34]. Graph reinforcement learning (GRL), first proposed around in 2017, aims to solve graph mining tasks such as link prediction [35], adversarial attacks [36], and relational reasoning [37]. Unlike its applications in other domains, one of the biggest challenges in GRL is finding an appropriate reward to guide the path searching in specific domains. To address the issue of finding biologically reasonable BKG-based paths for drug repurposing, it is crucial to incorporate biomedical domain knowledge to guide the path-finding process. Liu et al. [38] developed an RL-based model “PoLo” that utilizes the biological meta-paths identified in Himmelstein et al. [15] via the “DWPC” method to supervise path searching for drug repurposing. However, the “PoLo” model does not scale to a massive and complex BKG (e.g., CKG and RTX-KG2)

due to its dependence on the “DWPC” method that is reported to be computationally inefficient [39].

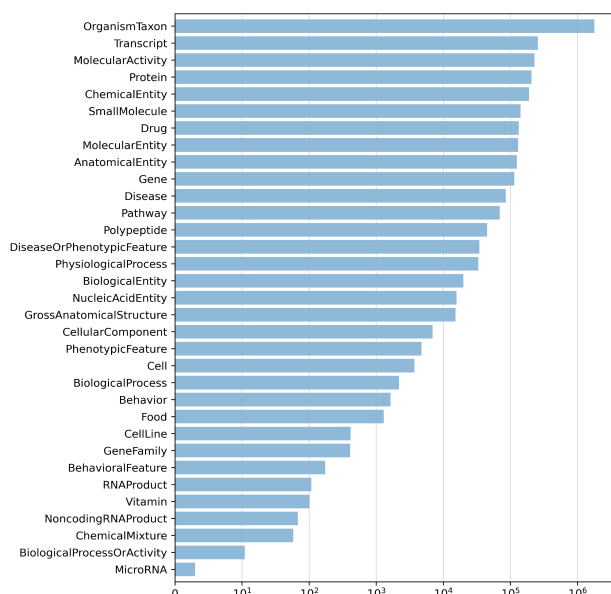
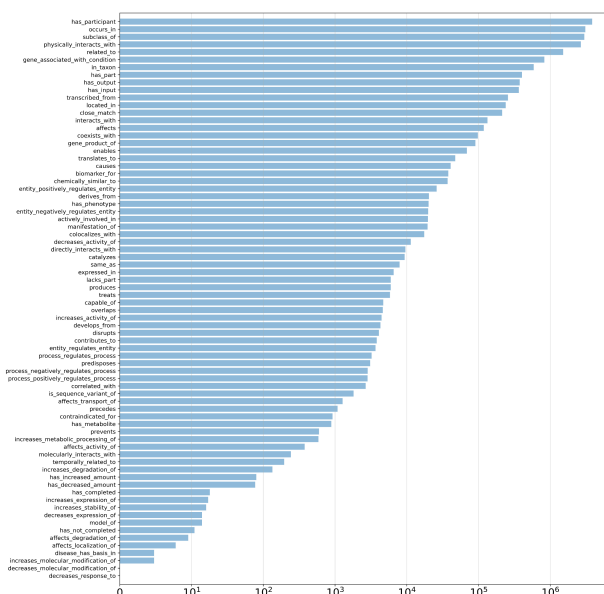
In this article, we describe KGML-xDTD: a **K**nowledge **G**raph-based **M**achine **L**earning framework for **e**xplainably predicting **D**rugs **T**reating **D**iseases, which contains 2 modules for both drug repurposing prediction and MOA explanation. We propose to amplify the ability of the RL model in biologically meaningful path searching by utilizing the biologically meaningful “demonstration paths” and pretrained drug repurposing model probability as rewards. We incorporate this idea into the appropriate models (e.g., GraphSAGE [40], random forest, and ADAC RL [41] models) and then make them applicable to the explainable drug repurposing problem at massive data scale and complexity. By comparing with the existing popular drug repurposing models and evaluating the predicted paths with an expert-curated path-based drug MOA database *DrugMechDB* [42, 43], we show that the proposed model framework can achieve state-of-the-art performance in both predictions of drug repurposing and recapitulation of human-curated drug MOA paths provided by *DrugMechDB*. In further case studies, by comparing the model predictions with the real regulatory networks, we show that the proposed framework effectively identifies biologically reasonable BKG-based MOA paths for real-world applications.

Materials and Methods

Datasets

Customized biomedical knowledge graph

To accommodate biomedical-reasonable predictions of drugs' indications and their mechanisms of action, the ideal biomedical knowledge graph should integrate biomedical knowledge from comprehensive and diverse databases and publications, as well as accurately identify and merge different identifiers representing the same biological entity into one (e.g., “CHEBI:2367” and “ChEMBL455626” are 2 distinct identifiers separately presented in the ChEBI database [44] and ChEMBL database [12] but represent the same compound “abyssinone I”). Thus, we utilize the canonicalized version of the Reasoning Tool X Knowledge Graph 2 (RTX-KG2c) [23], one of the largest open-source BKGs that integrates knowledge from extensive human-curated and publication-based databases and has been widely used in the Biomedical Data Translator Project [45, 46]. Compared to other commonly used open-source BKGs mentioned above, RTX-KG2c is a biolink model-based standardized [47] and regularly updated BKG that efficiently merges biologically and semantically equivalent nodes and edges via multiple curation steps. The version 2.7.3 of RTX-KG2c that we use contains around 6.4 million nodes and 39.3 million edges with knowledge from 70 public biomedical sources, where all biological concepts (e.g., “ibuprofen”) are represented as vertices and all concepts–predicates–concept (e.g., “ibuprofen-increases activity of-GP1BA gene”) are presented as edges. For drug repurposing purposes, we customized RTX-KG2c with 4 principles (see more details in Supplementary Section S1): (i) excluding the nodes whose categories are irrelevant to drug repurposing explanation (e.g., “GeographicLocation” and “Device”), (ii) filtering out the low-quality edges based on our criteria, (iii) removing the hierarchically redundant edges, and (iv) excluding all drug–disease edges. After these processing steps, 3,659,165 nodes with 33 distinct categories (Fig. 1A) and 18,291,237 edges with 74 distinct types (Fig. 1B) are left in our customized biomedical knowledge graph, which is used for downstream model training.

A Number of Nodes by Category in Customized BKG**B** Number of Edges by Predicate in Customized BKG**Figure 1:** Number of nodes by category (A) and number of edges by predicate (B) in the customized BKG.

Data sources for model training

To train the KGML-xDTD framework for drug repurposing prediction and its MOA explanation, we utilize 4 high-quality and NLP-derived training datasets:

- **MyChem Data [48]** is provided by the BioThings API collection [49], which contains up-to-date annotations regarding indication and contraindication for chemicals collected from 11 reliable data resources (summarized in Supplementary Section S2). We use drug-disease pairs with the relation “indication” as true positives while those with “contraindication” as true negatives.
- **SemMedDB Data [50]** is provided by the Semantic MEDLINE Database (SemMedDB), which leverages NLP techniques to extract semantic triples with “treats” and “negatively treats” relations from PubMed abstracts. We use drug-disease pairs with the relation “treats” as true positives and those with “negatively treats” as true negatives.
- **NDF-RT Data [51]** is provided by National Drug File-Reference Terminology from the Veterans Health Administration (VHA), which contains FDA-approved information on drug interaction, indications, and contraindications. We use drug-disease with therapeutics label “indications” as true positives and those with “contraindications” as true negatives.
- **RepoDB Data [52]** is a standard set of successful and failed drug-disease pairs in clinical trials collected by the Blavatnik Institute at Harvard Medical School. We use drug-disease with the status “approved” as true positives and those with “terminated” as true negatives.

We further filter drug-disease pairs from SemMedDB Data due to publication bias and possible NLP mistakes by using both the co-occurrence frequency and the PubMed publication-based normalized Google distance (NGD) [53] defined below:

$$NGD(c1, c2) = \frac{\max\{\log N(c1), \log N(c2)\} - \log N(c1, c2)}{\log N - \min\{\log N(c1), \log N(c2)\}} \quad (1)$$

where $c1$ and $c2$ are 2 biological concepts used in the customized BKG, $N(c1)$ and $N(c2)$ respectively represent the total number of

Table 1: Pair count of true-positive (indications) and true-negative (contraindications or no effect) data from 4 data sources after data preprocessing

Source	True positive (treats)	True negative (not treat)
MyChem	3,663	26,795
SemMedDB	8,255	11
NDF-RT	3,421	5,119
RepoDB	2,127	738
Shared	3,971	526
Total	21,437	33,189

Note that “shared” means those pairs are from 2 or more data sources.

unique PubMed IDs associated with $c1$ and $c2$, $N(c1, c2)$ is the total number of unique PubMed IDs shared between $c1$ and $c2$, and N is the total number of pairs of Medical Subject Heading (MeSH) terms annotations in PubMed database. Only the SemMedDB drug-disease pairs with at least 10 supporting publications and an NGD score of 0.6 or lower are left for the downstream model training.

These datasets are pooled together and then processed by (i) mapping the raw identifiers of drugs and diseases to the identifiers used in the customized BKG and (ii) removing duplicate drug-disease pairs in both the true-positive set and the true-negative set. Table 1 shows the drug-disease pair count from each data source after data preprocessing.

DrugMechDB

DrugMechDB [42, 43], to our best knowledge, is the first human-curated path-based database for explaining the MOA from a drug to a disease in an indication, with 3,593 MOA paths for 3,327 unique drug-disease pairs. These paths are extracted from free-text descriptions from DrugBank, Wikipedia, and other literature sources and then have been curated by subject matter experts and also follow the schema of the Biolink model. Hence, we can match them to nodes used in the RTX-KG2 BKG via the *Node Synonymizer* function [23]. Since the maximum length of predicted MOA paths

generated by the KGML-xDTD framework is fixed to 3 in this study due to memory and training time constraints, we consider those 3-hop BKG-based paths as “correct” if all 4 of their nodes show up in the complete DrugMechDB-based MOA paths. Thus, we find 472 unique drug-disease pairs of which each has at least 1 such “correct” matched path in all possible 3-hop paths between drug and disease in the customized BKG. The large reduction in evaluation paths is likely due to incompleteness of the underlying knowledge graphs, imperfect bioentity matching, and possibility of disconnected drug and disease pairs in the customized BKG. However, these paths are used for additional, external validation data only. We use the matched paths as true-positive biologically meaningful paths for the evaluation of the model-predicted paths in the task of MOA prediction (introduced below).

Model framework

The model framework of KGML-xDTD consists of 2 modules: a drug repurposing prediction (DRP) module that combines the advantages of GraphSAGE [40] and a random forest model, and an MOA prediction module that utilizes an adversarial actor-critic RL model. We show the overview of the entire model framework in Fig. 2. The implementation details of each module in KGML-xDTD framework are presented in Supplementary Section S3.

Notations

Let $\mathcal{G} = \{\mathcal{V}, \mathcal{E}\}$ be a directed biomedical knowledge graph, where each node $v \in \mathcal{V}$ represents a biological entity (e.g., a specific drug, disease, gene, or pathway), and each edge $e \in \mathcal{E}$ represents a biomedical relationship (e.g., *interacts-with*; see more in Fig. 1B). We use $\mathcal{V}^{\text{drug}}$ to represent all the drug nodes (the nodes with the categories of “Drug” and “Small Molecule” in the customized BKG) and $\mathcal{V}^{\text{disease}}$ to represent all the disease nodes (the nodes with the categories of “Disease,” “PhenotypicFeature,” “BehavioralFeature,” and “DiseaseOrPhenotypicFeature” in the customized BKG). For each notation, we use bold formatting to represent its embedding (e.g., $\mathbf{v}$ represents the embedding of v).

DRP module

Drug repurposing aims to identify new indications of existing drugs/compounds. We solve it as a link prediction problem on the graph $\mathcal{G}$. Specifically, given any drug-disease pair (v_i, v_j) where $v_i \in \mathcal{V}^{\text{drug}}$ and $v_j \in \mathcal{V}^{\text{disease}}$, we predict the probability that drug i can be used to treat disease j . We first use GraphSAGE to calculate the embedding for each node. Ideally, the node embeddings should contain 2 kinds of information: node attributes and node neighborhoods.

To capture the neighborhood information, we optimize GraphSAGE to encourage neighbor nodes to have similar embeddings and nonneighbor nodes to have distinct embeddings. Specifically, we perform random walks for each node to collect its neighborhood information and train the model to maximize a node’s similarity with its neighbor nodes. For a node u , the loss is calculated as

$$L_G(\mathbf{z}_u) = -\log(\sigma(\mathbf{z}_u^T \mathbf{z}_v)) - k \cdot \mathbb{E}_{v_n \sim P_n(v)} \log(\sigma(\mathbf{z}_u^T \mathbf{z}_{v_n})) \quad (2)$$

where $\mathbf{z}_u, \mathbf{z}_v$ are respectively the embeddings of nodes u, v , σ is the sigmoid function; v is a node that co-occurs with u in fixed-length random walks; P_n represents negative sampling distribution; and k indicates the number of negative samples (nodes not in u ’s fixed-length neighborhood).

To capture the node attributes information, we utilize the PubMedBERT model [54], a pretrained language model designed for

biomedical texts, to generate a node attribute embedding for each node based on the concatenation of the node’s name and category. We further compress the embeddings to 100 dimensions with principal components analysis (PCA) to reduce memory usage and use them as the initial node feature for GraphSAGE. In this way, the final GraphSAGE embedding of each node should contain the information regarding both graph topology and node attributes. We concatenate the GraphSAGE embeddings of drug-disease pairs and use them as input of a random forest model to classify each drug-disease pair into one of the “not treat,” “treat,” and “unknown” classes. We obtain “treat” and “not treat” drug-disease pairs from 4 data sources (described in “Data sources for model training” section). We generate “unknown” drug-disease pairs through negative sampling [55], that is, replacing the drug or disease identifier in each “treat” drug-disease pair with a random drug or disease identifier to generate a new pair that does not appear in both the “treat” and “not treat” classes. Specifically, for each unique “treat” drug-disease pair, we respectively replace its drug identifier with 1 random drug identifier as well as replace its disease identifier with 1 random disease identifier to make the “unknown” drug-disease pairs.

MOA prediction module

When potential indications of a given drug are identified by the drug repurposing prediction module, a natural yet essential question is: can we biologically explain the predictions? We solve this by employing an RL model to predict the BKG-based MOA paths, which are the paths on the knowledge graph from drug nodes to disease nodes. These BKG-based MOA paths can semantically describe an abstract biological process of how a drug treats a disease.

Demonstration paths

To encourage the RL agent to terminate the path searching at the expected diseases through a biologically reasonable path, we leverage so-called “demonstration paths”, a set of biologically likely paths (e.g., drug1-gene1-protein3-disease1) that explains the underlying reasons for why a drug can treat a disease. We extract 396,705 demonstration paths from the customized BKG using the known drug-target interactions collected from 2 curated biomedical data sources: DrugBank (v5.1) and Molecular Data Provider (v1.2) (see “Data Availability” section), as well as the PubMed publication-based NGD (see Equation 1). We show more details regarding demonstration path extraction in Supplementary Section S4.

Adversarial actor-critic reinforcement learning

We formulate the MOA prediction as a path-finding problem and adapt the adversarial actor-critic reinforcement learning model [41] to solve it. Reinforcement learning is defined as a Markov decision process (MDP) that contains:

States: Each state s_t at time t is defined as $s_t = (v_{\text{drug}}, v_t, (v_{t-1}, e_t), \dots, (v_{t-k}, e_{t-(k-1)}))$, where $v_{\text{drug}} \in \mathcal{V}^{\text{drug}}$ is a given starting drug node; $v_t \in \mathcal{V}$ represents the node where the agent locates at time t ; and the tuple $(v_{t-k}, e_{t-(k-1)})$ represents the previous k th node and $(k-1)$ th predicate. For the initial state s_0 , the previous nodes and predicates are substituted by a special dummy node and predicate. We concatenate the embedding of all nodes and predicates of s_t to get the state embedding $\mathbf{s}_t$, where the node embeddings are node attribute embeddings generated with the PubMedBERT model (see “DRP module” section) and the predicate embeddings employ one-hot vectors.

Actions: The action space A_t of each node v_t includes a self-loop action a_{self} and the actions to reach its outgoing neigh-

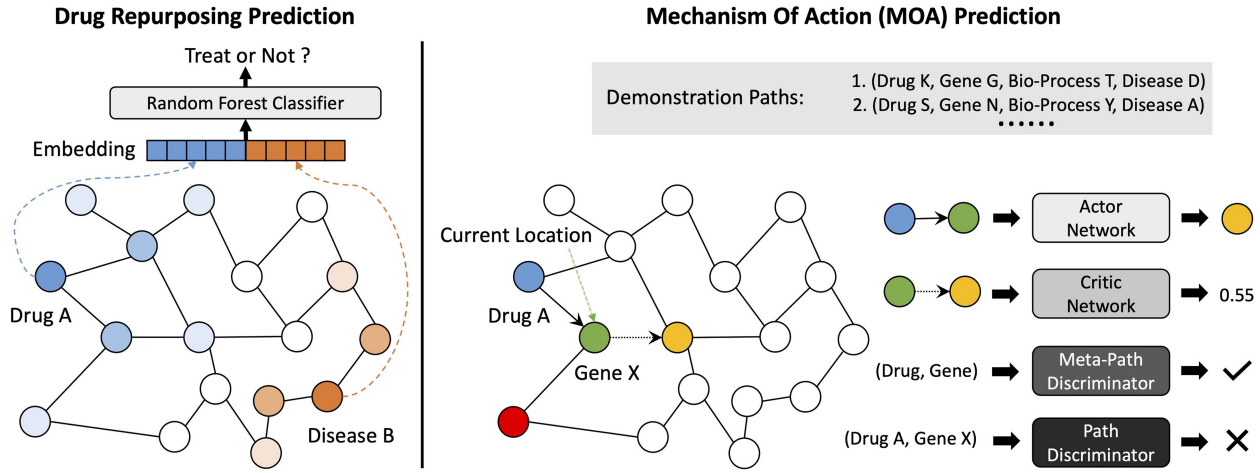


Figure 2: Illustration of the entire KGML-xDTD model framework: DRP module (left) and MOA prediction module (right).

bors in the graph $\mathcal{G}$. Due to memory limitation and extremely large out-degree of certain nodes in the knowledge graph, we prune the neighbor actions based on the PageRank scores if a node has more than 3,000 neighbors. Specifically, we let $A_t = (a_{\text{self}}, a_1, \dots, a_k, \dots, a_{n_{v_t}})$, where n_{v_t} is out-degree of node $v_t \in \mathcal{V}$. For each action $a_t = (e_t, v_{t+1}) \in A_t$ taken at time t , we concatenate its node and predicate embeddings to obtain action embedding $\mathbf{a}_t$. We learn 2 embedding matrices $\mathbf{E}^{N_n \times d}$ and $\mathbf{E}^{N_p \times d}$, respectively, for nodes and predicates (note that each subnetwork uses separate embedding matrices), where d represents the embedding dimension, N_n represents the number of nodes in the graph, and N_p represents the number of predicate categories in the graph.

Rewards: During the path-searching process, the agent only receives a terminal reward $R_{e,T}$ from the environment (i.e., there is no intermediate reward from environment: $R_{e,t} = 0, \forall t < T$). Let v_T be the last node of the path, and $\mathcal{N}_{\text{drug}}$ be the known diseases that drug v_{drug} can treat. The terminal reward $R_{e,T}$ from environment is calculated with the drug repurposing model via

$$R_{e,T} = \begin{cases} 1, & \text{if } v_T \in \mathcal{N}_{\text{drug}}. \\ p_{\text{treat}}, & \text{if } v_T \notin \mathcal{N}_{\text{drug}}; v_T \in \mathcal{V}^{\text{disease}} \text{ and } f(v_{\text{drug}}, v_T) \text{ is predicted as "treat"}. \\ 0, & \text{if } v_T \notin \mathcal{N}_{\text{drug}}; v_T \in \mathcal{V}^{\text{disease}} \text{ and } f(v_{\text{drug}}, v_T) \text{ is not predicted as "treat"}. \\ -1, & \text{if } v_T \notin \mathcal{V}^{\text{disease}}. \end{cases}$$

where p_{treat} is the “treat” class probability predicted by the drug repurposing model f .

The adversarial actor-critic RL model consists of 4 subnetworks that share the same model architecture MLP^i (note that i represents the id of each subnetwork described later, such as a for actor network, c for critic network, etc.) but with different parameters:

$$\text{MLP}^i(X) = \text{BA}(\text{BA}(XW_1^i + b_1^i)W_2^i + b_2^i)W_3^i + b_3^i \quad (3)$$

where $\{W_1^i, W_2^i, W_3^i, b_1^i, b_2^i, b_3^i\}$ are the parameters and biases of linear transformations, and BA represents a batch normalization layer followed by an ELU activation function.

Actor network: The actor network learns a path-finding policy π_θ (note that θ represents all parameters of the actor network) to guide the agent to choose an action a_t from the action space A_t based on the current state s_t :

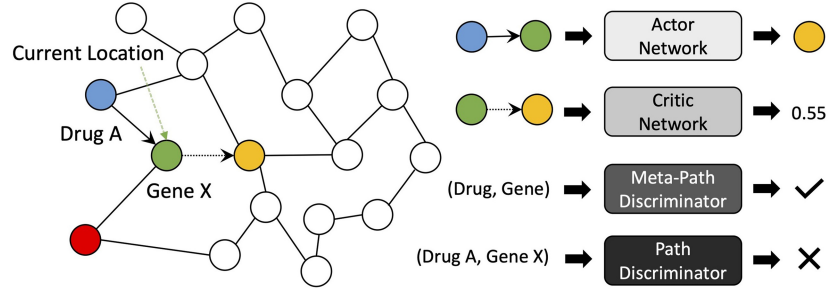
$$\pi_\theta(a_t | s_t, A_t) = \text{softmax}(\mathbf{A}_t \odot \text{MLP}^a(\mathbf{s}_t)) \quad (4)$$

where $\mathbf{A}_t$ is the embedding matrix of the action space A_t ; $\odot$ represents the dot product. Here, $\pi_\theta(a_t | s_t, A_t)$ represents the probability

Mechanism Of Action (MOA) Prediction

Demonstration Paths:

1. (Drug K, Gene G, Bio-Process T, Disease D)
2. (Drug S, Gene N, Bio-Process Y, Disease A)
-



of choosing action a_t at time t from the action space A_t given the state s_t .

Critic network: The critic network [56] estimates the expected reward $Q_\phi(s_t, a_t)$ (note that ϕ represents all parameters of the critic network) if the agent takes the action a_t at the state s_t by

$$Q_\phi(s_t, a_t) = \text{MLP}^c(\mathbf{s}_t) \odot \mathbf{a}_t \quad (5)$$

Path discriminator network: Since the RL agent only receives a terminal reward $R_{e,T}$ from the environment indicating whether it reaches an expected target, to encourage the agent to find biologically reasonable paths and provide intermediate rewards, we further guide it with demonstration paths. This network is essentially a binary classifier that distinguishes whether a path segment (s_t, a_t) is from demonstration paths or generated by the actor network. We treat all the known demonstration path segments (s_t^D, a_t^D) as positive samples and all actor-generated non-demonstration path segments (s_t^{ND}, a_t^{ND}) as negative samples. The path discriminator network $D_p(s, a) = \text{sigmoid}(\text{MLP}^p(\mathbf{s} \oplus \mathbf{a}))$, where $\mathbf{s}$ and $\mathbf{a}$ are respectively the embeddings of the state s and the action a ; $\oplus$ represents the concatenation operator and is optimized with

$$L_p = -\mathbb{E}_{(s,a) \sim P_D} [\log(D_p(s, a))] - \mathbb{E}_{(s,a) \sim P_A} [\log(1 - D_p(s, a))] \quad (6)$$

where P_D and P_A respectively represent the demonstration path segment distribution and the actor-generated non-demonstration path segment distribution. Based on the probability $D_p(s_t, a_t)$, the path discriminator-based intermediate reward $R_{p,t}$ is calculated as

$$R_{p,t} = \log(D_p(s_t, a_t)) - \log(1 - D_p(s_t, a_t)). \quad (7)$$

Meta-path discriminator network: Similar to the path discriminator, this network aims to judge whether the meta-path of the actor-generated paths is similar to that of demonstration paths. The meta-path is the path of node categories (e.g., [“Drug” → “Gene” → “BiologicalProcess” → “Disease”]). Similarly, the meta-path discriminator $D_m(M) = \text{sigmoid}(\text{MLP}^m(\mathbf{M}))$, where $\mathbf{M}$ is the embedding of the meta-path M defined as the concatenation of learned category embeddings of all nodes that appear in the path, is also a binary classifier where the meta-paths of demonstration paths are treated as positive samples while others are negative samples. We optimize it with the following loss:

$$L_m = -\mathbb{E}_{M \sim P_D^M} [\log(D_m(M))] - \mathbb{E}_{M \sim P_A^M} [\log(1 - D_m(M))] \quad (8)$$

where P_D^M and P_A^M respectively represent the demonstration meta-path distribution and the actor-generated non-demonstration meta-path distribution. The intermediate reward $R_{m,t}$ generated by the meta-path discriminator is calculated by

$$R_{m,t} = \log(D_m(M)) - \log(1 - D_m(M)). \quad (9)$$

The integrated intermediate reward R_t at time t is then calculated as

$$R_t = \alpha_p R_{p,t} + \alpha_m R_{m,t} + (1 - \alpha_p - \alpha_m) \gamma^{T-t} R_{e,T} \quad (10)$$

where $\alpha_p \in [0, 1]$ and $\alpha_m \in [0, 1 - \alpha_p]$ are hyperparameters, γ is the decay coefficient, and $R_{e,T}$ is defined in the “Rewards” section above.

To optimize the critic network, we minimize the temporal difference (TD) error [57] with loss:

$$L_c = TD^2 = [(R_t + Q_\phi(s_{t+1}, a_{t+1})) - Q_\phi(s_t, a_t)]^2. \quad (11)$$

Since the goal of the actor network is to achieve the largest expected reward by learning an optimal actor policy, we optimize the actor network by maximizing $J(\theta) = \mathbb{E}_{a \sim \pi_\theta} [Q_\phi(s_t, a)]$. We use the REINFORCE algorithm [58] to optimize the parameters. To encourage more diverse exploration in finding paths, we use the entropy of π_θ as a regularization term and optimize the actor network with the following stochastic gradient of the loss function L_a :

$$\nabla_\theta L_a = -\nabla_\theta J(\theta) = -\mathbb{E}_{\pi_\theta} [\nabla_\theta TD \log \pi_\theta(a_t | s_t)] - \alpha \nabla_\theta \text{entropy}(\pi_\theta) \quad (12)$$

where π_θ is the action probability distribution based on the actor policy, and α is the entropy weight.

We follow Zhao et al. [41] to train the adversarial actor-critic RL model in a multistage way. First, we initialized the actor network using the behavior cloning method [59] in which the training set of demonstration paths is used to guide the sampling of the agent with mean square error (MSE) loss. Then, in the first z epochs, we freeze the parameters of the actor network and the critic network and respectively train the path discriminator network and meta-path discriminator network by minimizing L_p and L_m . After z epochs, we unfreeze the actor network and the critic network and optimize them together by minimizing a joint loss $L_{\text{joint}} = L_a + L_c$.

Results

Evaluation settings

Data split

The post-processed drug-disease pairs (described in “Data sources for model training” section) are split into training, validation, and test sets where the drug-disease pairs of each unique drug are randomly split according to a ratio of 8/1/1. For example, let's say drugA has 10 known diseases that it treats (e.g., drugA-disease1, ..., drugA-disease10), 8 pairs are randomly split into the training set, 1 pair is to the validation set, and 1 pair to the test set. With this data split method, the model can be exposed to every drug in the training set, which complies with our goal of predicting new indications of known drugs and their potential MOAs based on the MOA of known target diseases.

Evaluation metrics

The proposed framework KGML-xDTD is evaluated on 2 types of tasks: **predicting drug-disease “treat” probability** (i.e., drug re-

purposing prediction) as well as **identifying biologically reasonable MOA paths from all BKG-based path candidates** (i.e., MOA prediction). These 2 tasks are evaluated based on classification accuracy-based metrics (e.g., accuracy, macro F1 score) and ranking-based metrics (e.g., mean percentile rank, mean reciprocal rank, and proportion of ranks smaller than K) defined as follows:

Accuracy (ACC) is the fraction of the model classification is correct, computed as

$$ACC = \frac{\text{Number of correct classifications}}{\text{Total number of drug-disease pair classifications}}. \quad (13)$$

Macro F1 score (Macro-F1) is the unweighted mean of all the per-class F1 scores:

$$F1^c = 2 * \frac{\text{precision}^c \times \text{recall}^c}{\text{precision}^c + \text{recall}^c}, \text{Macro-F1} = \frac{1}{|C|} \sum_{c \in C} F1^c \quad (14)$$

where C presents classification classes (e.g., “treat,” “not treat,” and “unknown”).

Mean percentile rank (MPR) is the average percentile rank of the 3-hop DrugMechDB-matched BKG-based path (described in “DrugMechDB” section) of true-positive drug-disease pairs:

$$MPR = \frac{1}{|PR|} \sum_{pr \in PR} pr \quad (15)$$

where PR is a list of percentile ranks of DrugMechDB-matched BKG-based paths of true-positive drug-disease pairs (“treat” category).

Mean reciprocal rank (MRR) is the average inverse rank of true-positive drug-disease pairs (“treat” category) or their 3-hop DrugMechDB-matched BKG-based paths:

$$MRR = \frac{1}{|R|} \sum_{r \in R} \frac{1}{r} \quad (16)$$

where R is a list of ranks of true-positive drug-disease pairs (for DRP task) or DrugMechDB-matched BKG-based paths (for MOA prediction task).

Hit@K is the proportion of ranks not larger than K for true-positive drug-disease pairs (“treat” category) or their 3-hop DrugMechDB-matched BKG-based paths:

$$\text{Hit@K} = \frac{1}{|R|} \sum_{r \in R} |r \leq K| \quad (17)$$

where R is a list of ranks of true-positive drug-disease pairs (for DRP task) or DrugMechDB-matched BKG-based paths (for MOA prediction task).

Drug repurposing prediction evaluation method

We utilize the metrics ACC and Macro-F1 to measure the accuracy of drug repurposing prediction of our KGML-xDTD framework while using ranking-based metrics MRR and Hit@K to show its capability in reducing false positives (i.e., the false drug-disease pairs ranking higher among possible drug-disease candidates). We use the following 3 methods to generate non-true-positive drug-disease candidates for each true-positive drug-disease pair to calculate the ranks that are employed in the MRR and Hit@K calculation:

- **Drug rank-based replacement:** For each true positive drug-disease pair, the drug rank-based replacement pairs are generated by replacing the drug entity with each of all 274,676

other drugs in the customized BKG while excluding all known true-positive drug-disease pairs.

- **Disease rank-based replacement:** For each true-positive drug-disease pair, the disease rank-based replacement pairs are generated by replacing the disease entity with each of all 124,638 other diseases in the BKG while excluding all known true-positive drug-disease pairs.
- **Combined replacement:** For each true-positive drug-disease pair, the combined replacement pairs are the combination of all replacement pairs of the above 2 methods. All known true-positive drug-disease pairs are excluded from these replacement pairs.

Due to the massive size of possible drug-disease candidates, some baseline models (e.g., GAT and GraphSAGE+SVM) are not applicable in this setting within a reasonable time (e.g., a week). Thus, we also employ a small subset of drug-disease replacements to calculate the MRR and Hit@K, allowing for comparison between KGML-xDTD with all baselines. Specifically, we utilize 1,000 random drug-disease pairs from the combined replacement set above: 500 with drug ID replacement and 500 with disease ID replacement. To enhance the robustness of results obtained through this random replacement method, we use this method to generate 10 sets of random drug-disease pairs (each with 1,000 pairs) independently and calculate the mean and standard deviation of the ranking-based metrics outcomes. In addition, since the drug repurposing prediction module of KGML-xDTD framework does 3-class classification while other baselines do 2-class classification, for a fair comparison, we recalculate ACC and Macro-F1 for KGML-xDTD by excluding the “unknown” class.

MOA prediction evaluation method

For the evaluation of MOA prediction, we use the DrugMechDB [42] to obtain the expert-verified MOA paths as ground-truth data, and match each biological concept in these verified MOA paths to the biological entities used in the customized BKG, and then generate the BKG-based matched paths of DrugMechDB drug-disease pairs (described in “DrugMechDB” section), which are considered biologically meaningful MOA paths. We first calculate the path scores for all 3-hop KG paths between drug and disease with the path-finding policy learned from adversarial actor-critic RL model using the following equation:

$$\text{path score} = \sum_{i=1}^k \delta^{i-1} \times \log(P_i \times N_i) \quad (18)$$

where k is the number of hops in this path, δ is a decay coefficient (we set it to 0.9 in this study), P_i represents the probability of choosing action a_i in the i^{th} hop following this path based on the trained RL model, and N_i is the number of possible actions in the i^{th} hop.

With the path scores, we obtain the ranks of the DrugMechDB-matched BKG-based paths and calculate their ranking-based metrics (e.g., MPR, MRR, and Hit@K). For those drug-disease pairs with multiple matched paths, we use the highest ranks of the matched paths as their ranks in the metrics calculation. We compare KGML-xDTD with the baseline models based on these metrics to show the capability of the MOA prediction module of KGML-xDTD in identifying biologically reasonable MOA paths from a massive and complex BKG with a comparably low false positive. In addition, we further perform 2 case studies to evaluate the effectiveness of KGML-xDTD in identifying the biologically reasonable MOA paths.

Drug repurposing prediction evaluation

For drug repurposing prediction evaluation, we compare the KGML-xDTD model framework against several state-of-the-art (SOTA) KG-based models and variants of KGML-xDTD for drug repurposing prediction based on the method described in “Drug repurposing prediction evaluation method” section.

We use 8 different SOTA KG-based models as baseline models that are commonly used for BKG-based drug repurposing [19, 60]. TransE [25], TransR [61], and RotatE [26] are the translation distance-based models that regard a relation (e.g., “treats”) as a “translation”/“rotation” (e.g., a kind of spatial transformation) from a head entity (e.g., a drug node) to a tail entity (e.g., a disease node). DistMult [27] is a bilinear model that measures the latent semantic similarity of a knowledge graph triple (head entity, relation/predicate, tail entity) with a trilinear dot product. ComplEx [28] and ANALOGY [62] are the extensions of DistMult that consider more complex relations (e.g., asymmetric relations). Simple [63] is a tensor factorization-based model to learn the semantic relation of a knowledge graph triple. GAT [64] is a popular graph neural model that leverages the important graph topology structure based on self-attention mechanism for graph-associated tasks (e.g., link prediction). Implementation details of these baselines are presented in Supplementary Section S5.

Besides these SOTA baseline models, we also compare the drug repurposing prediction module in KGML-xDTD with its several variants to show the effectiveness of model components. For example, to show efficacy of the combination of GraphSage and random forest (RF), we use a pure GraphSAGE model for link prediction (GraphSAGE-link), the combination of GraphSage and logistic model (GraphSAGE-logistic), and the combination of GraphSage and support vector machine (SVM) model (GraphSAGE-SVM). To demonstrate the effectiveness of node attribute embeddings (described in “DRP module” section) in improving repurposing prediction, we conduct an ablation experiment that replaces node attribute embeddings (NAEs) with random embeddings (initialized with the Xavier method [65]) as GraphSage initialized embeddings (KGML-xDTD w/o NAE); to support rationality of setting “unknown” class through negative sampling (described in “DRP module” section), we modify the drug repurposing prediction module for 2-class classification (i.e., true positive and true negative) (2-class KGML-xDTD) as a baseline comparison model.

Table 2 shows the performance of the KGML-xDTD model and all other baseline models in the task of drug repurposing prediction based on test set (described in “Data split” section). For the calculation of the MRR and Hit@K used in this table, we utilize the random subset replacement method described in the “Drug repurposing prediction evaluation method” section. As shown in the table, on the one hand, the KGML-xDTD outperforms most of the baseline models and achieves comparable performance as GAT in classification-based metrics (e.g., accuracy, macro F1 score), indicating its effectiveness in classifying known “treat” and “not treat” drug-disease pairs with both attribute and neighborhood information on the knowledge graph. On the other hand, KGML-xDTD’s exceptional performance in ranking-based metrics shows its superiority over baselines in identifying new indications of existing drugs out of a large number of possible drug-disease pairs with relatively low false positives, which is of great importance for guiding clinical research. Fig. 3 displays the comparison results where we calculate the MRR and Hit@K with 3 different “complete” replacement methods (described in “Drug re-

Table 2: The performance comparison of drug repurposing prediction (DRP) between KGML-xDTD and different baseline models based on test set (described in “Data split” section). The top panel shows the performance of state-of-the-art (SOTA) baseline models; the middle panel shows the performance of variants of the KGML-xDTD model framework; the bottom panel shows the performance of KGML-xDTD model framework

Model	Accuracy	Macro F1 score	MRR	Hit@1	Hit@3	Hit@5
TransE	0.708	0.708	0.301 (± 0.005)	0.134 (± 0.007)	0.327 (± 0.009)	0.482 (± 0.007)
TransR	0.858	0.855	0.329 (± 0.006)	0.150 (± 0.009)	0.378 (± 0.008)	0.542 (± 0.005)
RotatE	0.704	0.704	0.281 (± 0.007)	0.098 (± 0.008)	0.314 (± 0.007)	0.497 (± 0.009)
DistMult	0.555	0.495	0.182 (± 0.004)	0.042 (± 0.002)	0.157 (± 0.010)	0.292 (± 0.010)
ComplEx	0.624	0.460	0.138 (± 0.004)	0.026 (± 0.004)	0.106 (± 0.007)	0.205 (± 0.008)
ANALOGY	0.594	0.465	0.188 (± 0.004)	0.044 (± 0.004)	0.165 (± 0.009)	0.301 (± 0.008)
SimplE	0.599	0.472	0.167 (± 0.006)	0.036 (± 0.006)	0.140 (± 0.008)	0.259 (± 0.011)
GAT	0.936	0.934	0.002 (± 0.000)	0.000 (± 0.000)	0.000 (± 0.000)	0.000 (± 0.000)
GraphSAGE-link	0.919	0.915	0.002 (± 0.000)	0.000 (± 0.000)	0.000 (± 0.000)	0.000 (± 0.000)
GraphSAGE+logistic	0.791	0.784	0.002 (± 0.000)	0.000 (± 0.000)	0.000 (± 0.000)	0.000 (± 0.000)
GraphSAGE+SVM	0.807	0.793	0.002 (± 0.000)	0.000 (± 0.000)	0.000 (± 0.000)	0.000 (± 0.000)
KGML-xDTD w/o NAEs	0.909 (0.898*)	0.891 (0.892*)	0.159 (± 0.003)	0.035 (± 0.002)	0.143 (± 0.006)	0.262 (± 0.008)
2-class KGML-xDTD	0.929	0.925	0.278 (± 0.003)	0.183 (± 0.006)	0.321 (± 0.003)	0.389 (± 0.006)
KGML-xDTD (ours)	0.935 (0.930*)	0.923 (0.926*)	0.382 (± 0.004)	0.238 (± 0.007)	0.425 (± 0.006)	0.543 (± 0.006)

The values with * inside the parentheses are the adjusted results by excluding the “unknown” category for a fair comparison.

The ranking metrics (e.g., “MRR” and “Hit@K”) are calculated as the mean along with standard deviation based on 10 independent sets of non-true-positive drug-disease candidates generated by the random drug-disease replacement method (i.e., for each true-positive drug-disease pair in test set, we use 1,000 random drug-disease pairs as non-true-positive drug-disease candidates to calculate the rank). See more details in “Drug repurposing prediction evaluation method” section.

The abbreviation “w/o NAEs” in the name of model “KGML-xDTD w/o NAEs” represents without using node attribute embeddings.

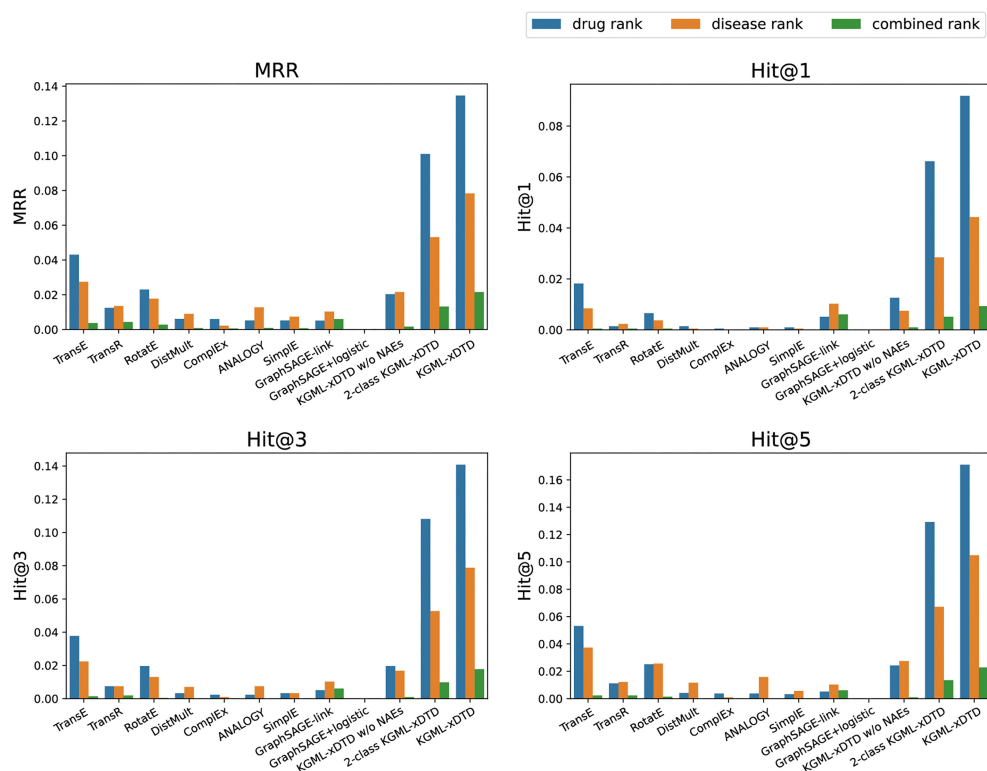


Figure 3: The performance comparison of DRP between KGML-xDTD and different baseline models (GAT and GraphSAGE+SVM are excluded due to computation time constraints) based on test set using 3 “complete” replacement methods (i.e., “drug rank-based replacement,” “disease rank-based replacement,” and “combined replacement” described in “Drug repurposing prediction evaluation method” section) to generate non-true-positive drug-disease candidates for each true-positive drug-disease pair for MRR and Hit@K calculation. “Drug rank,” “disease rank,” and “combined rank” respectively correspond to the methods of “drug rank-based replacement,” “disease rank-based replacement,” and “combined replacement.”

purposing prediction evaluation method” section). Although GAT and GraphSAGE+SVM are excluded in this comparison due to computation time constraints, we can see that the results presented in both Table 2 and Fig. 3 are consistent to demonstrate KGML-xDTD’s ability in reducing false positives. Therefore, excluding

the GAT and GraphSAGE+SVM from the comparison with “complete” replacement methods does not affect the conclusion. Besides, by comparing 2-class KGML-xDTD with the vanilla GraphSAGE model (e.g., GraphSAGE-link), we demonstrate the effectiveness of the random forest model over a neural network clas-

Table 3: The performance comparison of MOA prediction between KGML-xDTD and different baseline models (e.g., MultiHop and KGML-xDTD w/o DP) based on the test set (described in “Data split” section). The metrics in this table are calculated using path scores and all non-DrugMechDB-matched 3-hop paths between drug and disease as “negative” paths for each true-positive drug-disease pair (see more details in “MOA prediction evaluation method” section)

Model	MPR	MRR	Hit@1	Hit@10	Hit@50	Hit@100	Hit@500
MultiHop	61.400%	0.027	0.017	0.042	0.067	0.118	0.345
KGML-xDTD w/o DP	72.965%	0.015	0.008	0.017	0.067	0.160	0.403
KGML-xDTD (ours)	94.696%	0.109	0.059	0.193	0.496	0.613	0.849

The abbreviation “w/o DP” in the name of model “KGML-xDTD w/o DP” represents “without using demonstration paths.”

sifier in this task. The comparison between KGML-xDTD w/o NAE and KGML-xDTD shows that the KGML-xDTD benefits from the use of node attribute embeddings for drug repurposing prediction while the comparison with 2-class KGML-xDTD indicates the effectiveness of using negative sampling to generate “unknown” drug-disease pairs for model training. With the “unknown” drug-disease pairs, the KGML-xDTD model achieves significant improvement in ranking-based metrics, which is essential when applying to real-world drug repurposing because it can reduce the false positives.

MOA prediction evaluation

For MOA prediction, we evaluate how well the KGML-xDTD can identify the DrugMechDB-matched BKG-based MOA paths (described in “MOA prediction evaluation method” section) from a large number of possible paths in the customized BKG by utilizing ranking-based metrics (e.g., MPR, MRR, and Hit@K) and 2 specific case studies.

There are few machine learning models designed for the task of identifying biologically meaningful paths from biomedical knowledge graphs for explaining drug repurposing. Although the UKGE [30], GrEDeL [32], and Polo [38] models (all mentioned in “Introduction”) were proposed and can be used for this goal, they all have certain constraints and cannot be used as baseline models for comparison. The UKGE model cannot be applied to BKGs without weighted edge information (e.g., frequency of relation appeared in literature). The authors of the GrEDeL model do not provide the code to implement this model. The Polo model cannot be trained within a reasonable time (e.g., within 2 weeks) on a massive and complex BKG (e.g., RTX-KG2) due to its dependence on a computationally inefficient method “DWPC” [39]. Therefore, we choose the MultiHop reinforcement learning model [66] as a baseline model since it uses a similar LSTM model framework as the GrEDeL model and allows using a self-defined reward-shaping strategy in its reward function as what we do in the KGML-xDTD model (i.e., we can use the same reward strategy described in “Adversarial actor-critic reinforcement learning” section). Furthermore, we also compare with an ablated version of KGML-xDTD (i.e., KGML-xDTD w/o DP, which does not take advantage of the demonstration paths by setting α_p and α_m in Function 9 as 0) as another baseline model to show the importance of proposed demonstration paths.

We compare the MOA prediction performance between the KGML-xDTD model framework and different baseline models in Table 3. Although all the models use the same terminal reward function from the environment (i.e., the drug repurposing prediction module of KGML-xDTD), the MOA prediction module of KGML-xDTD achieves significantly better performance in identifying DrugMechDB-matched BKG-based MOA paths than the other 2 baselines across all ranking-based metrics. Comparison between

Table 4: Top 10 predicted drugs/treatments for hemophilia B (note that the red bolded drugs are used in the training set)

Drug/Treatment	Prob.	Publications
Eptacog Alfa (rFVIIa)	0.833	[67, 68]
Nonacog Alfa (rFIX)	0.803	[69]
Viral vector	0.780	[70]
Factor VIIa	0.748	[67, 71]
Recombinant FVIIa (rFVIIa)	0.724	[67, 71]
Thrombin	0.709	[72]
Factor IX	0.708	[73]
Epicriptine	0.702	
Hyperbaric oxygen	0.660	
Triamcinolone	0.649	

the KGML-xDTD with and without demonstration paths (i.e., KGML-xDTD w/o DP) further illustrates the great effectiveness of using proposed demonstration paths to guide the path-finding process. Due to the massive searching space and sparse rewards, the RL agent often fails to find biologically reasonable BKG-based MOA paths out of many possible choices, while our model KGML-xDTD, with the intermediate guidance provided by the demonstration path, is able to identify those biologically reasonable choices with a much higher probability. Moreover, comparing KGML-xDTD w/o DP and MultiHop reveals that the actor-critic model structure performs similarly to LSTM. However, incorporating the proposed demonstration paths can significantly enhance the effectiveness of the actor-critic model structure over LSTM for this task.

To further evaluate the performance of KGML-xDTD model framework in identifying biologically relevant MOA paths for drug repurposing, we present 2 different case studies to explore the potential repurposed drugs and their potential mechanism for 2 rare genetic diseases: hemophilia B and Huntington’s disease.

Case 1: Hemophilia B

Hemophilia B, also known as factor IX deficiency or Christmas disease, is a rare genetic disorder that results in prolonged bleeding in patients. It is caused by mutations in the factor IX (F9) gene, which is located on the X chromosome. Table 4 displays the top 10 drugs/treatments predicted by the KGML-xDTD model framework, including both those that are used in the training set (red bolded) and those that are not. Besides those known drugs/treatments used in the training set, the majority of the remaining 7 drugs/treatments on the list are supported by published research and have the potential to treat hemophilia B. For example, the activated human-derived coagulation factor VII (i.e., factor VIIa) or the recombinant activated factor VII (i.e., rFVIIa) is one of the proteins that can cause blood clots as an important part of the blood coagulation regulatory network (as shown in the

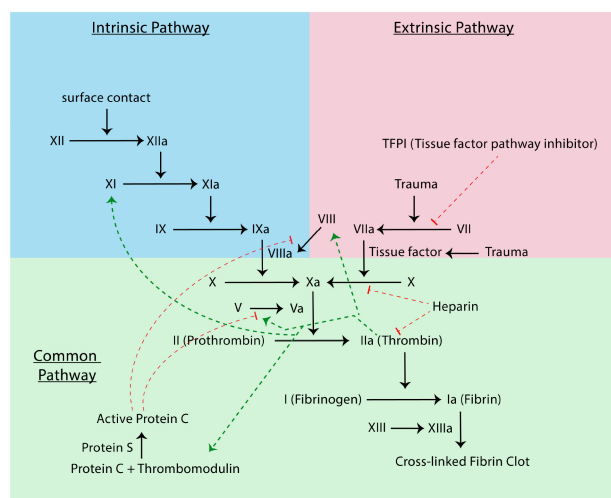


Figure 4: Blood coagulation regulatory network with arrows for molecular reactions (black), positive feedback (green), and negative feedback (red).

Fig. 4). This protein is used as an effective inhibitor in the treatment of patients with hemophilia B [67, 71]. Thrombin is a key enzyme in the maintenance of normal hemostatic function. It has been reported that using thrombin as a therapeutic strategy can help prevent bleeding in patients with hemophilia [72]. The use of recombinant factor IX therapy is a recommended treatment option for individuals with hemophilia B [73]. Some examples of recombinant factor IX products include BeneFIX, Rixubis, Ixinity, Alprolix Idelvion, and Rebinyn. These examples demonstrate the potential capability of KGML-xDTD for drug repurposing in real-world applications.

To further assess the biological explanations of the predicted 3-hop BKG-based MOA paths for the treatment of hemophilia B, we have used the curated DrugMechDB-based MOA paths, which are not used in the model training process. DrugMechDB contains relevant MOA paths of hemophilia B treatment only for Eptacog Alfa and Nonacog Alfa. We use the KGML-xDTD model to predict the top 10 potential 3-hop BKG-based MOA paths for these 2 drugs and compare them with the curated DrugMechDB-based MOA paths in Fig. 5 (for visualization purpose, we only display the top 5 predicted paths along with any available DrugMechDB-matched BKG-based paths in the top 10 predicted paths). The corresponding biological entities between the predicted paths and the curated DrugMechDB-based paths are highlighted in red. Although the predicted paths cannot exactly match the DrugMechDB-based MOA paths due to the limited path length and some missing semantic relationships in the customized biomedical knowledge graph, key biological entities (such as coagulation factor VII, coagulation factor X, and coagulation factor IX) that are important for the treatment of hemophilia B are present in the predicted paths. As shown in Fig. 4, the treatment of hemophilia B involves a complex molecular network of blood coagulation, and many of the coagulation factors (such as factor VII, factor III, factor II, factor VIII, factor IX, and factor X) present in the predicted paths are also part of this molecular network. In Supplementary Section S6, we also utilize the KGML-xDTD model framework to predict the top 10 three-hop BKG-based paths, which can serve as biological explanations of the predicted “treats” relationship between factor VIIa and hemophilia B (shown in Table 4). This particular drug/treatment-disease pair is not included in the training set and thus can be used to indicate how KGML-xDTD’s MOA path pre-

dictions can contribute to the explanation of the predicted drug repurposing results. The predicted paths show molecular details akin to those in Fig. 4 for treating hemophilia B. As a result, the predicted paths by KGML-xDTD model framework can help identify key molecules in the real drug action regulatory network, thereby aiding in explaining drug repurposing to some extent.

Case 2: Huntington’s disease

Huntington’s disease (HD) is a rare neurogenetic disorder that typically occurs in midlife with symptoms of depression, uncontrolled movements, and cognitive decline. While there is currently no drug/treatment that can alter the course of HD, some drugs/treatments can be useful for the treatment of its symptoms in abnormal movements (e.g., chorea) and psychiatric phenotypes. We show 10 drugs/treatments with the highest predicted probability by the KGML-xDTD model framework after manual processing in Table 5. This processing involves excluding the chemotherapeutic drugs from the predicted drug candidate list due to their potential risk of cytotoxicity to normal cells (which could lead to false positives for drug repurposing of noncancer diseases [74, 75]), and only presenting the top 5 results in the training set and top 5 from the test or validation set. From this table, it can be observed that many of the top-ranked predicted drugs have been supported by publications as potential treatments for the symptoms of HD. Since there is currently no effective treatment for HD, DrugMechDB does not have a corresponding MOA path for comparison. To analyze the predicted paths by the KGML-xDTD model framework for the predicted nonchemotherapeutic drugs/treatments that are not included in the training set (shown in regular text in Table 5), we present their top 5 predicted paths in Fig. 6. From these predicted paths, we can see that most of them are biologically relevant. For example, Fig. 6A shows that risperidone is predicted to be useful for the treatment of HD by decreasing the activity of the genes associated with the 5-hydroxytryptamine receptor (e.g., HTR1A, HTR2A, HTR2C, HTR7) and dopamine receptor (e.g., DRD2), which have been proven to be involved in the pathogenesis of depressive disorders [76, 77]. The presence of depressive symptoms is a significant characteristic of HD [78]. Entinostat is predicted to have the potential to alleviate the symptoms of HD by inhibiting the functions of histone deacetylase genes such as HDAC1 and HDAC6 (see Fig. 6B), and one of the predicted 3-hop BKG-based MOA paths (“Entinostat” → “decreases activity of” → “HDAC1 gene” → “interacts with” → “Histone H4” → “gene associated with condition” → “Huntington’s disease”) is supported by the previous research [79, 80]. Primaquine is predicted to act on the IKBKG gene to potentially play a therapeutic role in neurodegenerative disease (see Fig. 6C), reported in [81]. According to the predicted BKG-based MOA paths (see Fig. 6D), isradipine may have a potential therapeutic effect for HD by mainly regulating the genes of the calcium voltage-gated channel, including CACNA1C and CACNB2. These genes may be associated with the symptoms of HD, such as depression and dementia [82]. Lastly, amifampridine is predicted to regulate the genes of the potassium voltage-gated channel (see Fig. 6E), which are potentially associated with HD [83]. All these examples indicate that the predicted BKG-based MOA paths can explain the mechanism of repurposed drugs to some extent.

Drug Class Analysis

The drug repurposing prediction of the KGML-xDTD model does not leverage any information regarding drug similarity such as drug classes, SMILES, drug side effects, and drug-related gene pro-

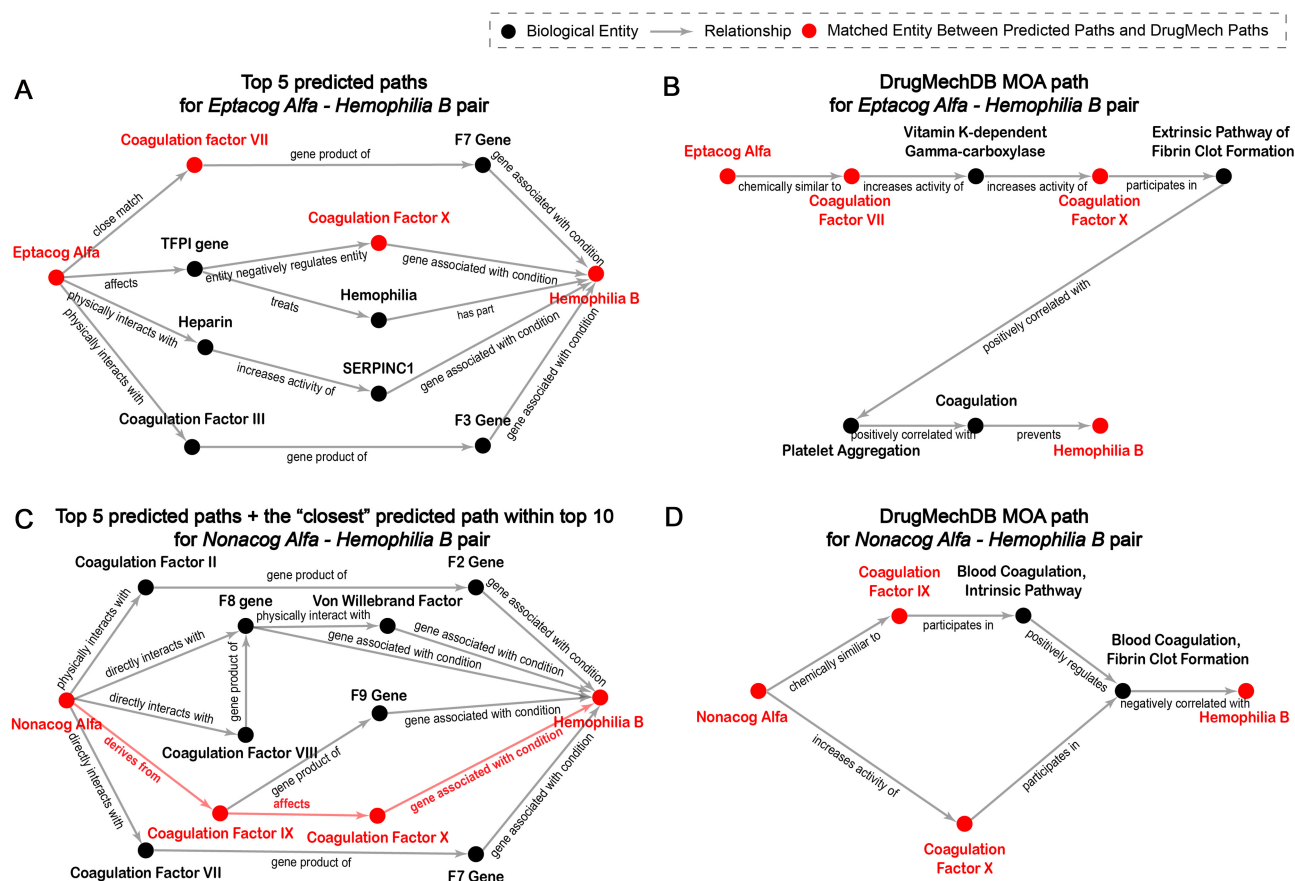


Figure 5: Comparison between the top 5 predicted 3-hop paths (including any available DrugMechDB-matched BKG-based paths in the top 10 predicted paths, highlighted in red) and the curated DrugMechDB-based MOA paths for Eptacog Alfa and Nonacog Alfa. Note that the RTX-KG2 paths and DrugMechDB paths might use different synonyms for the same biological concept. For better visualization and illustration, we utilize consistent entity synonyms between the predicted paths and the curated paths as well as present the predicted paths in a graph structure. (A, C) Graph representation of predicted paths generated by KGML-xDTD respectively for Eptacog Alfa and Nonacog Alfa. (B, D) Human-curated DrugMechDB MOA paths.

Table 5: Top 5 predicted drugs/treatments used in the training set (highlighted in red bold) and the top 5 nonchemotherapeutic predicted drugs/treatments that are not in the training set for Huntington's disease

Drug/Treatment	Prob.	Publications
Pimozide	0.939	[84, 85]
Therapeutic agent	0.939	
Olanzapine	0.938	[86, 87]
Riluzole	0.935	[88]
Antipsychotic agent	0.932	[89]
Risperidone	0.893	[78, 90]
Entinostat	0.888	[79]
Primaquine	0.887	
Isradipine	0.884	[91]
Amifampridine	0.882	

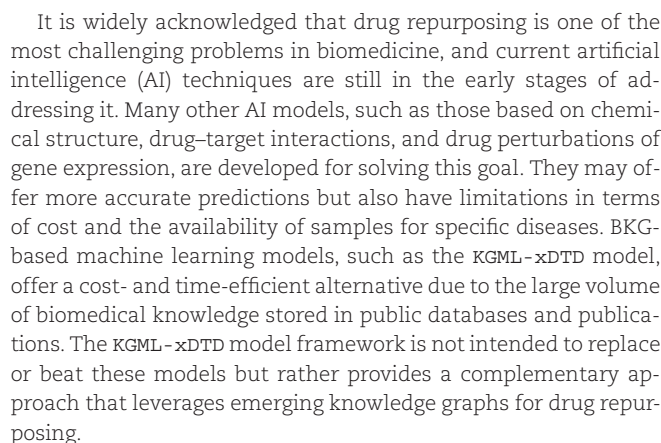
files/sequences, and we find that the distribution of drug classes in true-positive drug-disease pairs is similar between the training and test sets (see Fig. 7). In this section, we examine whether our model can only predict the drugs with the drug classes that it has seen in the training set.

To do this, we use the MyChem.info APIs [48] to retrieve the FDA's "Established Pharmacologic Class" (EPC) information for chemicals/drugs using their synonym identifiers. For the FDA-unapproved chemical/drug without such EPC information, we

consider it as a single class. We first utilize the KGML-xDTD model to predict the top 100 chemicals/drugs for each of the 1,140 diseases in the test set (described in "Data split" section) after excluding the drug-disease pairs presented in the training set. Then we count the number of drug classes among these 100 predicted drugs that are not seen in the training set for each disease. Fig. 8 shows the distribution of unseen drug classes in the top 100 predicted nontrain drugs across the 1,140 diseases in the test set. We can see that each disease has at least 70 different drug classes among the top 100 predicted drugs, indicating that the predictive power of the KGML-xDTD model is derived from the node attribute information and knowledge graph topology structure rather than any drug class information.

Discussion

In this work, we propose KGML-xDTD, a 2-module, knowledge graph-based machine learning framework that not only predicts the treatment probabilities between drugs/compounds and diseases but also provides biological explanations for these predictions through the predicted paths in a massive biomedical knowledge graph with comprehensive biomedical data sources as potential mechanisms of action. This framework can assist medical researchers in quickly identifying the potential drug/compound-disease pairs that might have a treatment relationship, which can accelerate the process of drug discovery for emerging diseases.



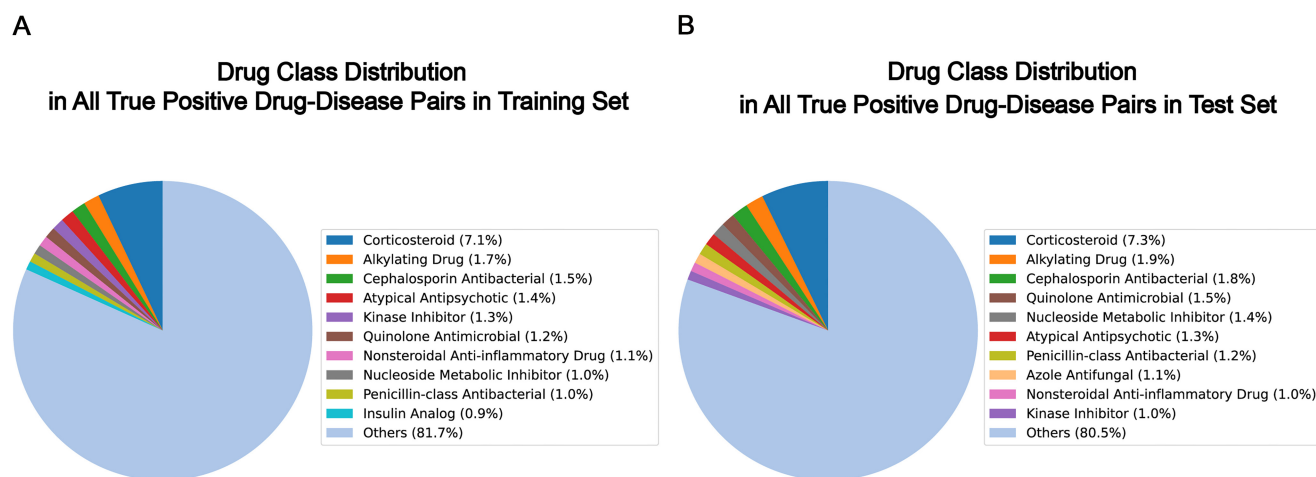


Figure 7: Comparison of the drug class distribution in true-positive drug-disease pairs between the training set and test set. The drug class of each drug/chemical in these pairs is determined based on the FDA “Established Pharmacologic Class” (EPC) accessed via MyChem.info APIs. There are 2,238 drug classes represented in the true-positive drug-disease pairs in the training set and 718 drug classes in the test set. For visualization purposes, we only show the top 10 drug classes in each set and the rest is classified into the “Others” class.

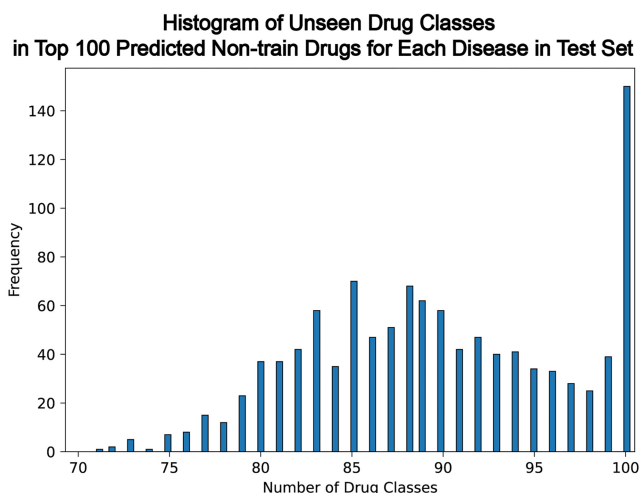


Figure 8: Distribution of unseen drug classes in the top 100 predicted nontrain drugs across the 1,140 diseases in the test set.

Future work to further enhance the KGML-xDTD model framework might include extending the predicted paths for more specific explanations and considering the negative drug-disease pairs so that the model can explain why certain drugs are harmful to diseases.

Availability of Source Code and Requirements

Project name: KGML-xDTD

- Project homepage: <https://github.com/chunyuma/KGML-xDTD>
- Operating system(s): Linux (Ubuntu)
- Resource usage in training step: A Linux (Ubuntu) system with at least 8 CPU cores, 800 GB of VRAM, and a 48 GB GPU card (48 GB Quadro RTX 8000 GPU card used in our training)
- Resource usage in inference step: A Linux (Ubuntu) system with at least 8 CPU cores and 50 GB of VRAM. The GPU card is not necessary, but if used, the GPU card needs at least 24

GB VRAM (48 GB Quadro RTX 8000 GPU card used in our inference)

- Time requirement: Based on our hardware performance and parameter settings (please see the scripts on GitHub), the training step takes approximately 2 weeks while the inference step takes approximately 25.42 seconds for 1 drug-disease pair with 3,320 potential paths. These time estimates may vary depending on the hardware performance, parameter settings, and the number of potential paths of a given drug-disease pair.
- Programming language: Shell Script (Bash) with Python 3.8.12
- Other requirements: Python 3.8.12 with GPU/CPU support (GraphSAGE training needs Python 2.7), neo4j-community 3.5.26, miniconda 4.8.2 (please see more requirements in the yaml files under “envs” folder on Github repository)
- Licenses: MIT license, DrugBank academic license, Apache 2.0 license, UMLS Metathesaurus license, CC-BY 4.0 license
- Research Resource Identifier (#RRID): SCR_023678

Data Availability

The data sets supporting the results of this article are publicly available in the Zenodo repository [92]. All supporting data and materials are available in the GigaScience GigaDB database [93].

Molecular data provider: A knowledge-centric data provider for systems chemical biology, as part of the NCATS Biomedical Data Translator (“Translator”). See more in <https://github.com/NCATSTranslator/Translator-All/wiki/Molecular-Data-Provider> [94].

Additional Files

Supplementary Section S1. Biomedical Knowledge Graph RTX-KG2c Preprocessing.

Supplementary Section S2. Summary of Data Resources Used by MyChem Data.

Supplementary Section S3. Implementation Details of KGML-xDTD Model Framework.

Supplementary Section S4. Implementation Details of Demonstration Path Extraction.

Supplementary Section S5. Implementation Details of Baseline Models.

Supplementary Section S6. Top 10 KGML-xDTD's Predicted Paths Serving as Biological Explanations for the Predicted "Treats" Relationship between Factor VIIa and Hemophilia B.

Supplementary Table S1. Eleven data resources used by MyChem data.

Supplementary Table S2. Hyperparameters used for baseline models.

Supplementary Fig. S1. Top 10 predicted 3-hop paths (integrated into a graph for better visualization) generated by the KGML-xDTD model framework serve as biological explanations of the predicted "treats" relationship between factor VIIa and hemophilia B. The path highlighted with red is the one in which all nodes can align with the key molecules in the real drug action regulatory network shown in Fig. 4 in the main text.

Abbreviations

ACC: accuracy; ADAC: adversarial actor-critic; BKG: biomedical knowledge graph; DWPC: degree-weighted path count; EHR: electronic health record; FDA: Food and Drug Administration; GRL: graph reinforcement learning; GWAS: Genome-wide association studies; KG: knowledge graph; LSTM: long short-term memory; Macro-F1: macro F1 score; MDP: Markov decision process; MeSH: Medical Subject Heading; MOA: mechanism of action; MPR: mean percentile rank; MSE: mean square error; NAE: node attribute embedding; NGD: normalized Google distance; NLP: Natural Language Processing; PCA: principal components analysis; RF: random forest; RL: reinforcement learning; RNN: recurrent neural network; RTX-KG2: Reasoning Tool X Knowledge Graph 2; RTX-KG2c: canonicalized version of the Reasoning Tool X Knowledge Graph 2; SemMedDB: Semantic MEDLINE Database; SMILES: Simplified Molecular Input Line Entry System; SVM: support vector machine; VHA: Veterans Health Administration.

Competing Interests

The authors declare no competing interests.

Funding

Research reported in this publication was supported by the National Center for Advancing Translational Sciences (NCATS) of the National Institutes of Health (NIH) under award numbers OT2-TR003428-01S3, OT2-TR003428-01S2, and OT2-TR003428-01; National Library of Medicine of the NIH under award number R01LM01372201; and the National Science Foundation under award number CAREER-1841569.

Authors' Contributions

D.K. and C.M. conceived the project and supervised the study. C.M. processed the raw data and designed model framework. C.M. and Z.Z. wrote code and trained models. C.M., Z.Z., and D.K. drafted the manuscript. All authors read and approved the final manuscript.

Acknowledgments

The authors thank Dr. Jared Roach from the Institute for Systems Biology for his assistance in manually evaluating the biological realism of the model-predicted paths in a double-blind manner, and

also the RTX-KG2 team (Stephen Ramsey, Amy Glen, E. C. Wood, Lili Acevedo) for their guidance in building RTX-KG2 and in resolving any issues that arose.

References

- Berdigaliyev N, Aljofan M. An overview of drug discovery and development. *Future Med Chem* 2020;12(10):939–47.
- Miller MT. Thalidomide embryopathy: a model for the study of congenital incomitant horizontal strabismus. *Trans Am Ophthalmol Soc* 1991;89:623–74.
- Verheul HMW, Panigrahy D, Yuan J, et al. Combination oral antiangiogenic therapy with thalidomide and sulindac inhibits tumour growth in rabbits. *Br J Cancer* 1999;79(1):114–8.
- Singhal S, Mehta J, Desikan R, et al. Antitumor activity of thalidomide in refractory multiple myeloma. *N Engl J Med* 1999;341(21):1565–71.
- Kairys V, Baranauskiene L, Kazlauskienė M, et al. Binding affinity in drug design: experimental and computational techniques. *Expert Opin Drug Discov* 2019;14(8):755–68.
- Aulner N, Danckaert A, Ihm J, et al. Next-generation phenotypic screening in early drug discovery for infectious diseases. *Trends Parasitol* 2019;35(7):559–70.
- Rusz CM, Ósz BE, Jitcă G, et al. Off-label medication: from a simple concept to complex practical aspects. *Int J Environ Res Public Health* 2021;18(19):10447.
- Swamidass SJ. Mining small-molecule screens to repurpose drugs. *Brief Bioinform* 2011;12(4):327–35.
- Sanseau P, Agarwal P, Barnes MR, et al. Use of genome-wide association studies for drug repositioning. *Nat Biotechnol* 2012;30(4):317–20.
- Bonner S, Barrett IP, Ye C, et al. A review of biomedical datasets relating to drug discovery: a knowledge graph perspective. *Brief Bioinform* 2022;23(6):bbac404.
- Wishart DS, Feunang YD, Guo AC, et al. DrugBank 5.0: a major update to the DrugBank database for 2018. *Nucleic Acids Res* 2017;46(D1):gkx1037.
- Gaulton A, Bellis LJ, Bento AP, et al. ChEMBL: a large-scale bioactivity database for drug discovery. *Nucleic Acids Res* 2012;40(D1):D1100–7.
- Wishart DS, Feunang YD, Marcu A, et al. HMDB 4.0: the human metabolome database for 2018. *Nucleic Acids Res* 2017;46(D1):gkx1089.
- Kanza S, Graham Frey J. Semantic technologies in drug discovery. In: O Wolkenhauer, ed. *Systems Medicine*. Oxford, UK: Academic Press; 2021:129–44.
- Himmelstein DS, Lizée A, Hessler C, et al. Systematic integration of biomedical knowledge prioritizes drugs for repurposing. *eLife* 2017;6:e26726.
- Walsh B, Mohamed SK, Nováček V. BioKG: a knowledge graph for relational learning on biological data. In: *Proceedings of the 29th ACM International Conference on Information & Knowledge Management, Virtual Event, Ireland*. p. 3173–3180. New York, NY, USA: ACM, 2020. <https://doi.org/10.1145/3340531.3412776>.
- Su C, Hou Y, Zhou M, et al. Biomedical discovery through the integrative biomedical knowledge hub (iBKH). *iScience* 2023;26(4):106460.
- Percha B, Altman RB. A global network of biomedical relationships derived from text. *Bioinformatics* 2018;34(15):2614–2624.
- Zhang R, Hristovski D, Schutte D, et al. Drug repurposing for COVID-19 via knowledge graph completion. *J Biomed Inform* 2021;115:103696.

20. Wang Q, Li M, Wang X, et al. COVID-19 literature knowledge graph construction and drug repurposing report generation. In: Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies: Demonstrations. p. 66–77. Association for Computational Linguistics, 2021.
21. Li N, Yang Z, Luo L, et al. KGHC: a knowledge graph for hepatocellular carcinoma. *BMC Med Inform Decis* 2020;20(suppl 3):135.
22. Santos A, Colaço AR, Nielsen AB, et al. A knowledge graph to interpret clinical proteomics data. *Nat Biotechnol* 2022;40(5):692–702.
23. Wood EC, Glen AK, Kvarfordt LG, et al. RTX-KG2: a system for building a semantically standardized knowledge graph for translational biomedicine. *BMC Bioinform* 2022;23(1):400.
24. Ioannidis VN, Zheng D, Karypis G. Few-shot link prediction via graph neural networks for Covid-19 drug-repurposing. 2020. arXiv:2007.10261. <https://doi.org/10.48550/arXiv.2007.10261>.
25. Bordes A, Usunier N, Garcia-Duran A et al., Translating embeddings for modeling multi-relational data. In: CJC Burges, L Bottou, M Welling, eds. Proceedings of the 26th International Conference on Neural Information Processing Systems. p. 2787–2795. Curran Associates, Inc, 2013.
26. Sun Z, Deng ZH, Nie JY, et al. RotatE: knowledge graph embedding by relational rotation in complex space. 2019. arXiv:1902.10197. <https://doi.org/10.48550/arXiv.1902.10197>.
27. Yang B, Yih Wt, He X, et al. Embedding entities and relations for learning and inference in knowledge bases. 2014. arXiv:1412.6575. <https://doi.org/10.48550/arXiv.1412.6575>.
28. Trouillon T, Welbl J, Riedel S, et al. Complex embeddings for simple link prediction. 2016. arXiv:1606.06357. <https://doi.org/10.48550/arXiv.1606.06357>.
29. Wang B, Shen T, Long G, et al. Structure-augmented text representation learning for efficient knowledge graph completion. In: Proceedings of the Web Conference 2021. p. 1737–1748. New York, NY, USA: Association for Computing Machinery, 2021.
30. Sosa DN, Derry A, Guo M, et al. A literature-based knowledge graph embedding method for identifying drug repurposing opportunities in rare diseases. *Pacific Symp Biocomputing*. 2020;25:463–74.
31. Chen X, Chen M, Shi W, et al. Embedding uncertain knowledge graphs. *Proc AAAI Conf Artif Intell* 2019;33:3363–70.
32. Sang S, Yang Z, Liu X, et al. GrDeL: a knowledge graph embedding based method for drug discovery from biomedical literatures. *IEEE Access* 2019;7:8404–15.
33. Kilicoglu H, Roseblat G, Fiszman M, et al. Broad-coverage biomedical relation extraction with SemRep. *BMC Bioinform* 2020;21(1):188.
34. Li Y. Reinforcement learning applications. 2019. arXiv:1908.06973. <https://doi.org/10.48550/arXiv.1908.06973>.
35. Chen L, Cui J, Tang X, et al. RLPath: a knowledge graph link prediction method using reinforcement learning based attentive relation path searching and representation learning. *Appl Intell* 2022;52(4):4715–26.
36. Sun Y, Wang S, Tang X, et al. Adversarial attacks on graph neural networks via node injections: a hierarchical reinforcement learning approach. In: Proceedings of The Web Conference 2020. p. 673–683. New York, NY, USA: Association for Computing Machinery, 2020.
37. Zhou X, Wang P, Luo Q, et al. Multi-hop knowledge graph reasoning based on hyperbolic knowledge graph embedding and reinforcement learning. In: The 10th International Joint Conference on Knowledge Graphs. p. 1–9. New York, NY, USA: Association for Computing Machinery, 2021.
38. Liu Y, Hildebrandt M, Joblin M, et al. Neural multi-hop reasoning with logical rules on biomedical knowledge graphs. In: R Verborgh, K Hose, H, Paulheimeds. The Semantic Web. Cham, Switzerland: Springer International Publishing; 2021:375–91.
39. Womack F, McClelland J, Koslicki D. Leveraging distributed biomedical knowledge sources to discover novel uses for known drugs. *bioRxiv*. 2019:765305. <https://doi.org/10.1101/765305>.
40. Hamilton WL, Ying R, Leskovec J. Inductive representation learning on large graphs. 2017. arXiv:1706.02216. <https://doi.org/10.48550/arXiv.1706.02216>.
41. Zhao K, Wang X, Zhang Y, et al. Leveraging demonstrations for reinforcement recommendation reasoning over knowledge graphs. In: Proceedings of the 43rd International ACM SIGIR Conference on Research and Development in Information Retrieval. p. 239–248. New York, NY, USA: Association for Computing Machinery, 2020.
42. Mayers M, Steinecke D, Su AI. Database of mechanism of action paths for selected drug-disease indications. Zenodo 2020. <https://doi.org/10.5281/zenodo.3708278>.
43. Mayers M, Tu R, Steinecke D, et al. Design and application of a knowledge network for automatic prioritization of drug mechanisms. *Bioinformatics* 2022;38(10):btac205.
44. Degtyarenko K, Matos Pd, Ennis M, et al. ChEBI: a database and ontology for chemical entities of biological interest. *Nucleic Acids Res* 2008;36(Database issue):D344–50.
45. Consortium TBDT. Toward a universal biomedical data translator. *Clin Transl Sci* 2019;12(2):86–90.
46. Translator Consortium. The Biomedical Data Translator Program: conception, culture, and community. *Clin Transl Sci* 2019;12(2):91–4.
47. Unni DR, Moxon SA, Bada M, et al. Biolink model: a universal schema for knowledge graphs in clinical, biomedical, and translational science. *Clin Transl Sci* 2022;15(8):1848–1855.
48. Xin J, Afrasiabi C, Lelong S, et al. Cross-linking BioThings APIs through JSON-LD to facilitate knowledge exploration. *BMC Bioinform* 2018;19(1):30.
49. Xin J, Mark A, Afrasiabi C, et al. High-performance web services for querying gene and variant annotation. *Genome Biol* 2016;17(1):91.
50. Kilicoglu H, Shin D, Fiszman M, et al. SemMedDB: a PubMed-scale repository of biomedical semantic predications. *Bioinformatics* 2012;28(23):3158–60.
51. Brown SH, Elkin PL, Rosenbloom ST, et al. VA National Drug File Reference Terminology: a cross-institutional content coverage study. *Stud Heal Tech Inform* 2004;107(Pt 1):477–81.
52. Brown AS, Patel CJ. A standard database for drug repositioning. *Sci Data* 2017;4(1):170029.
53. Cilibiasi RL, Vitanyi PMB. The Google similarity distance. *IEEE Trans Knowl Data Eng* 2007;19(3):370–83.
54. Gu Y, Tinn R, Cheng H, et al. Domain-specific language model pretraining for biomedical natural language processing. *ACM Trans Comput Healthc* 2022;3(1):1–23.
55. Mikolov T, Sutskever I, Chen K et al., Distributed representations of words and phrases and their compositionality. In: Proceedings of the 26th International Conference on Neural Information Processing System, vol. 2, p. 3111–3119. Red Hook, NY, United States: Curran Associates Inc., 2013.
56. Lillcrap TP, Hunt JJ, Pritzel A, et al. Continuous control with deep reinforcement learning. 2015. arXiv:1509.02971. <https://doi.org/10.48550/arXiv.1509.02971>.
57. Sutton RS. Learning to predict by the methods of temporal differences. *Mach Learn* 1988;3(1):9–44.

58. Williams RJ. Simple statistical gradient-following algorithms for connectionist reinforcement learning. *Mach Learn* 1992;8(3–4):229–56.
59. Pomerleau DA. Efficient training of artificial neural networks for autonomous navigation. *Neural Comput* 1991;3(1):88–97.
60. Hsieh K, Wang Y, Chen L, et al. Drug repurposing for COVID-19 using graph neural network and harmonizing multiple evidence. *Sci Rep* 2021;11(1):23179.
61. Lin Y, Liu Z, Sun M, et al. Learning entity and relation embeddings for knowledge graph completion. In: *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 29(1). 2015. <https://doi.org/10.1609/aaai.v29i1.9491>.
62. Liu H, Wu Y, Yang Y. Analogical inference for multi-relational embeddings. In: *Proceedings of the 34th International Conference on Machine Learning*, vol. 70. p. 2168–2178. Sydney, Australia: PMLR, 2017.
63. Kazemi SM, Poole D. Simple embedding for link prediction in knowledge graphs. *Adv Neural Inform Process Syst* 2018;31:4284–95.
64. Veličković P, Cucurull G, Casanova A, et al. Graph attention networks. 2017. arXiv:1710.10903. <https://doi.org/10.48550/arXiv.1710.10903>.
65. Glorot X, Bengio Y. Understanding the difficulty of training deep feedforward neural networks. In: *YW Teh, DM. Tittertoned. AISTATS*, vol. 9 of JMLR Proceedings. 2010:249–256.
66. Lin XV, Socher R, Xiong C. Multi-hop knowledge graph reasoning with reward shaping. In: *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing*. p. 3243–3253. Brussels, Belgium: Association for Computational Linguistics, 2018.
67. Croom KF, McCormack PL. Recombinant factor VIIa (eptacog alfa): a review of its use in congenital hemophilia with inhibitors, acquired hemophilia, and other congenital bleeding disorders. *BioDrugs* 2008;22(2):121–36.
68. Minno GD. Eptacog alfa activated: a recombinant product to treat rare congenital bleeding disorders. *Blood Rev* 2015;29:S26–S33.
69. Rendo P, Smith L, Lee HY, et al. Nonacog alfa: an analysis of safety data from six prospective clinical studies in different patient populations with haemophilia B treated with different therapeutic modalities. *Blood Coagul Fibrinolysis* 2015;26(8):912–8.
70. Driessche T, Collen D, Chuah M. Viral vector-mediated gene therapy for hemophilia. *Curr Gene Ther* 2001;1(3):301–15.
71. Roberts HR, Monroe DM, White GC. The use of recombinant factor VIIa in the treatment of bleeding disorders. *Blood* 2004;104(13):3858–64.
72. Negrier C, Shima M, Hoffman M. The central role of thrombin in bleeding disorders. *Blood Rev* 2019;38:100582.
73. Goodeve AC. Hemophilia B: molecular pathogenesis and mutation analysis. *J Thromb Haemost* 2015;13(7):1184–95.
74. Sourimant J, Aggarwal M, Plemper RK. Progress and pitfalls of a year of drug repurposing screens against COVID-19. *Curr Opin Virol* 2021;49:183–93.
75. Gysi DM, do Valle Í, Zitnik M, et al. Network medicine framework for identifying drug-repurposing opportunities for COVID-19. *Proc Natl Acad Sci U S A* 2021;118(19):e2025581118.
76. Yohn CN, Gergues MM, Samuels BA. The role of 5-HT receptors in depression. *Mol Brain* 2017;10(1):28.
77. Delva NC, Stanwood GD. Dysregulation of brain dopamine systems in major depressive disorder. *Exp Biol Med* 2021;246(9):1084–93.
78. Coppen EM, Roos RAC. Current pharmacological approaches to reduce chorea in huntington's disease. *Drugs* 2017;77(1):29–46.
79. Shukla S, Tekwani BL. Histone deacetylases inhibitors in neurodegenerative diseases, neuroprotection and neuronal differentiation. *Front Pharmacol* 2020;11:537.
80. Yu IT, Park JY, Kim SH, et al. Valproic acid promotes neuronal differentiation by induction of proneural factors in association with H4 acetylation. *Neuropharmacology* 2009;56(2):473–80.
81. Singh S, Singh TG. Role of nuclear factor kappa B (NF-KB) signalling in neurodegenerative diseases: an mechanistic approach. *Curr Neuropharmacol* 2020;18(10):918–35.
82. Yagami T, Kohma H, Yamamoto Y. L-type voltage-dependent calcium channels as therapeutic targets for neurodegenerative diseases. *Curr Med Chem* 2012;19(28):4816–27.
83. Noh W, Pak S, Choi G, et al. Transient potassium channels: therapeutic targets for brain disorders. *Front Cell Neurosci* 2019;13:265.
84. Arena R, Iudice A, Virgili P, et al. Huntington's disease: clinical effects of a short-term treatment with pimozide. *Adv Biochem Psychopharmacol* 1980;24:573–5.
85. Videnovic A. Treatment of Huntington disease. *Curr Treat Options Neurol* 2013;15(4):424–38.
86. Paleacu D, Anca M, Giladi N. Olanzapine in Huntington's disease: olanzapine in Huntington's disease. *Acta Neurol Scand* 2002;105(6):441–4.
87. Squitieri F, Cannella M, Porcellini A, et al. Short-term effects of olanzapine in Huntington disease. *Neuropsychiatry Neuropsychol Behav Neurol* 2001;14(1):69–72.
88. Group HS. Dosage effects of riluzole in Huntington's disease: a multicenter placebo-controlled study. *Neurology* 2003;61(11):1551–6.
89. Unti E, Mazzucchi S, Palermo G, et al. Antipsychotic drugs in Huntington's disease. *Exp Rev Neurother* 2017;17(3):227–37.
90. Duff K, Beglinger LJ, O'Rourke ME, et al. Risperidone and the treatment of psychiatric, motor, and cognitive symptoms in Huntington's disease. *Ann Clin Psychiatry* 2008;20(1):1–3.
91. Miranda AS, Cardozo PL, Silva FR, et al. Alterations of calcium channels in a mouse model of Huntington's disease and neuroprotection by blockage of CaV1 channels. *ASN NEURO* 2019;11:1759091419856811.
92. Ma C, Zhou Z, Liu H, et al. Relevant datasets and software used for paper “KGML-xDTD: a knowledge graph-based machine learning framework for drug treatment prediction and mechanism description (1.0.0) [Data set]. Zenodo 2023. <https://zenodo.org/record/7582233>.
93. Ma C, Zhou Z, Liu H, et al. Supporting data for “KGML-xDTD: A Knowledge Graph-Based Machine Learning Framework for Drug Treatment Prediction and Mechanism Description.” GigaScience Database. 2023. <http://dx.doi.org/10.5524/102404>.
94. Molecular Data Provider Team. <https://github.com/NCATSTranslator/Translator-All/wiki/Molecular-Data-Provider>. Accessed 9 November 2021.