

A reinforcement learning model for the knowledge graph of imperial guangdong maritime customs archival translation

Global Journal of Social Sciences Studies

Vol. 12, No. 2, 1-17, 2026.

e-ISSN: 2518-0614



Chen Lilan

School of Foreign Languages, Guangdong Pharmaceutical University, Guangzhou, China.

Email: chenlilan@gdpu.edu.cn

ABSTRACT

The Guangdong Provincial Archives stores a large volume of valuable historical original Customs archives which are enormous in number and volume, formal in format, various in types and genres, rich in content, providing precious value for such studies as those of Imperial Chinese and foreign exchanges. It has always been the focus of research in the academic world at home and abroad. This paper, characterized by the inputting and outputting relationship among the entities (proper names of people, places, titles, weights and measures) that are most closely related to the translation of Imperial Guangdong Maritime Customs archives, constructs the knowledge graph of Imperial Guangdong Maritime Customs archival translation. In this paper, representative subgraphs being calculated based on the given node sets, a novel archival translation knowledge graph method is proposed, and the spatial components are integrated into the reinforcement learning framework to help understand the knowledge graph of archival translation. Through establishing the Imperial Guangdong Maritime Customs archival translation theory model, the internal structure and external information of archives are analyzed to obtain a more comprehensive and holistic view of the archival translation tasks. At last, the model proposed in this paper is proved to be reasonable by evaluating the data sets of Imperial Guangdong Maritime Customs archives.

Keywords: Archival translation, Imperial Guangdong maritime customs archives, Knowledge graph, Reinforcement learning, Spatial model.

DOI: 10.55284/gjss.v12i2.1893

Citation | Lilan, C. (2026). A reinforcement learning model for the knowledge graph of imperial guangdong maritime customs archival translation. *Global Journal of Social Sciences Studies*, 12(2), 1-17.

Copyright: © 2026 by the authors. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

Funding: This work was supported by Guangdong Provincial Foreign Language Programs Grant, Grant No. GD23WZXC02, and the Higher education project granted by Guangdong Provincial Education Science Planning Leading Group, and the Grant No. 2024GXJK544.

Institutional Review Board Statement: Not applicable.

Transparency: The author confirms that the manuscript is an honest, accurate, and transparent account of the study; that no vital features of the study have been omitted; and that any discrepancies from the study as planned have been explained. This study followed all ethical practices during writing.

Competing Interests: The author declares that there are no conflicts of interest regarding the publication of this paper.

History: Received: 6 April 2026/ Revised: 8 July 2026/ Accepted: 15 July 2026/ Published: 20 July 2026

Publisher: Online Science Publishing

Highlights of this paper

- This paper constructs a knowledge graph for Imperial Guangdong Maritime Customs archives by integrating archival entities such as personal names, place names, titles, weights, and measures.
- A spatially explicit reinforcement learning framework is proposed to generalize archive knowledge graphs by combining graph structure, textual semantics, and spatial relationships.
- Experimental results show that the proposed MDP-RL model improves archive knowledge representation and retrieval performance compared with non-spatial and benchmark strategies.

1. INTRODUCTION

This paper focuses on Imperial Chinese Maritime Customs archives, particularly represented by Imperial Guangdong Maritime Customs archives that are stored in Guangdong Provincial Archives, including the original Imperial Guangdong Customs archives and relevant Customs publications, publications about Chinese staff and physical archives. Based on the analysis of the language barriers encountered in the process of sorting out and developing these Maritime Customs archives, combined with other aspects, these Customs archives have the following characteristics: First, they are rich in content, all-encompassing and scattered. As historical archives, the Imperial Chinese Maritime Customs archives cover the original archives and archival publications formed in the course of handling Customs business and administrative affairs from 1854 to 1949, with the content containing Imperial Chinese Maritime Customs taxes, internal and external trade, preventative, Customs business, post, personnel, secretary, general affairs, appraisal, tidessurvey, medical treatment, exhibitions, the domestic and foreign affairs, social intelligence and so on. These archives are various in types and genres, rich in content, enormous in number and volume, scattered in storing, involving official archives, confidential archives of domestic and foreign affairs of China, and intelligence archives on important local affairs. Second, the languages used in these archives are diverse, with foreign languages, especially English, as the main language carrier, paper and microfilm as the main content carrier. Since China's Customs were dominated by foreigners in Imperial times, the foreign Inspectrate system was implemented, and most of the Customs staff were foreigners, especially in the late Qing Dynasty and early Republic of China. During that period, the Inspector General of China's Maritime Customs Robert Hart expressly stipulated that official business documents should be expressed in English. Under the unremitting efforts and resistance of the Chinese people, it was not until 1942 that China's Maritime Customs official business documents should be mainly in Chinese, English or other languages for supplement. In addition, in different places, due to different foreign forces, there are different foreign languages in the official business documents for different Customs. For example, the language in the official business documents issued by the Kiaochow Customs in Shandong Province was German. The language carrier of the Imperial Chinese Maritime Customs documents is mostly English, with other languages, such as German, French, Japanese and so on. Third, the Imperial Chinese Maritime Customs archives contain terms and proper names unique in the Imperial Chinese Maritime Customs and the modern Chinese Pinyin system. In the Imperial Chinese Maritime Customs archives have there are lots of proper names, such as the Customs agencies (preventive, secretary, general affairs, Chinese language, personnel, taxation, etc.), titles of positions in the Customs agencies (Inspector General, Assistant Inspector General, Commissioner, assistant, appraiser, pilot, etc.). These pinyin systems are different from modern Chinese Pinyin system, mainly spelt in the Wiles-Wade System and then the Postal pinyin spelling, also influenced by some different dialects, making it difficult for young Chinese archive users who are unfamiliar with these spelling rules to understand these archives, encountering such cases as the corresponding errors or non-corresponding for proper names in these archives. Fourthly, the utilization rate of these Imperial Chinese Maritime Customs archives is not high. Taking Imperial Guangdong Maritime Customs archives as an example, as these archives are mainly stored in

Guangdong Provincial Archives and mainly stored in microfilms, they can only be accessed on microfilm readers during working hours, which brings inconvenience to archive users in terms of access time and utilization. In addition, as mentioned above, these foreign language archives limit some archive users to search and understand these archives because of language barriers.

The Imperial Guangdong Maritime Customs archives (1861-1949) took foreign languages as the main language carrier and paper as the content carrier, and the time span of nearly one hundred years resulted in language changes and paper damage, brittleness and insect erosion. These archive content involves politics, economy, culture, society, personnel, the condition of the people, tax, preventative, foreign trade and so on. Because foreigners supervised the Imperial Chinese Maritime Customs, the system and institutions of the Customs are different from those before the Opium War and after the founding of People's Republic of China, under the influence of dialects on languages used in Imperial Chinese Maritime Customs archives, to better develop and utilize these archives and give full play of these archives' historical, economic, political, cultural value and value in other aspects, it's a must to digitalize and study these archives on the basis of these characteristics, to expand the scope of archives users. In the Internet context, archive users can be expanded from the original experts in archives and history majors to the Internet users, increasing the circulation and utilization channels of these archives, so that through the network, the archive users can have access to and use these archives at home, thus improving the utilization rate of these archives, making full use of the historical, economic and cultural value and other values of these archives.

In order to further study the Imperial Guangdong Maritime Customs archives, this paper will describe the problem of archival translation of the Imperial Guangdong Maritime Customs archives with knowledge graph, so that the subgraph retains the salient substructure and meaning, here with highlighting nodes and edges. However, this task faces many challenges, especially in the area of archival space. One of the challenges has to do with the inherently complex structure of graphical data. Different from other common structures such as one-dimensional sequence of natural language and two-dimensional grid of image, graphic structure has its own unique structure. For example, at the global level, two graphs can be isomorphic (isomorphic) (Sun et al., 2020) that is, have the same structure, but they have different representations (for example, markers and visual representations). At the local level, substructures (like homophily and structural equivalence (Li & Madden, 2019)) coexist as proxies in the graph to encode the underlying patterns. Moreover, since most knowledge graphs follow what is known as the Open World Assumption (OWA) – meaning that there are triples that can be lost in the knowledge graph, without assuming that these lost triples do not apply in reality – the original structure of the graph may not represent complete information. This adds another layer of complexity.

In view of this, we propose a method based on reinforcement learning to identify the context information of the Imperial Guangdong Maritime Customs archives. Our method integrates internal structure and external information to generalize the knowledge graph of the Imperial Guangdong Maritime Customs archives.

The research contributions of this paper are as follows:

- Using reinforcement learning, a knowledge graph process of the Imperial Guangdong Maritime Customs archives is proposed. Rather than rely mainly on internal information, such as the method based on the node groups in the grouping and aggregating methods, and digits needed in describing graphs in the compression method based on digits, our method is putting the task into a framework for a sequence of reinforcement learning optimal decision process, obtaining internal information from graph structure and complementary advantages of external knowledge in the use of proper names of people and place, titles, weights and measures.

- With the richness of spatial semantics in the Imperial Guangdong Maritime Customs archive knowledge graph, incorporate this information in the generalization process to better capture the relevance of the Imperial Guangdong Maritime Customs archive entities and provide better results. We do this by following an established scientific approach to the Imperial Guangdong Maritime Customs archives, which is to model distance decay, and from the perspective of information theory.
- We create an Imperial Guangdong Maritime Customs archive dataset, including proper names of people, places, titles, weights and measures, for the Imperial Guangdong Maritime Customs archive knowledge graph generalization task and make it publicly available. In short, the lack of standard data sets has been one of the obstacles to the generalization of Imperial Guangdong Maritime Customs archive knowledge graph and to some extent the development of the Imperial Guangdong Maritime Customs archive retrieval research. By proactively collecting this dataset, we hope it will spur further research in this area.

The rest of this paper is organized as follows: Section 2 reviews prior work on digital archives, multilingual semantic resources, knowledge graph representation and completion, reinforcement-learning-based graph reasoning, and spatial semantics. Section 3 describes the construction of the Imperial Guangdong Maritime Customs archival knowledge graph. Section 4 presents the proposed spatially explicit reinforcement learning method. Section 5 reports the experimental settings and results. Section 6 concludes the paper and outlines future work.

2. LITERATURE REVIEW

2.1. *Digital Archives and Multilingual Semantic Access*

The digitization of archival collections has moved archives from collection-level description toward content-based computational access. Artificial intelligence can support archival description, retrieval, and discovery, but historical records remain heterogeneous, incomplete, multilingual, and noisy (Colavizza, Blanke, Jeurgens, & Noordegraaf, 2021; Ehrmann, Hamdi, Pontes, Romanello, & Doucet, 2023; Muehlberger et al., 2019). Cultural-heritage Linked Data research has shown that RDF-based semantic integration can connect records across institutions and support exploratory access (Hyvönen, 2020; Isaac & Haslhofer, 2013). For cross-lingual entity normalization, DBpedia and Wikidata provide persistent identifiers and structured relations, BabelNet connects lexical and encyclopedic knowledge across languages, and DBpedia Spotlight links textual mentions to graph entities (Lehmann et al., 2015; Mendes, Jakob, García-Silva, & Bizer, 2011; Navigli & Ponzetto, 2012; Vrandečić & Krötzsch, 2014). However, these general resources do not fully model historical romanization, customs-specific titles, or period-specific weights and measures, so domain curation remains necessary for Imperial Guangdong Maritime Customs archives.

2.2. *Knowledge Graph Representation and Completion*

Knowledge graphs represent entities and relations in a form that supports integration, inference, and retrieval (Hogan et al., 2021; Ji, Pan, Cambria, Marttinen, & Yu, 2021). Because automatically constructed graphs are incomplete and may contain errors, refinement and completion are central concerns (Paulheim, 2016). Surveys of knowledge graph embedding methods show that latent representations can capture graph regularities while supporting scalable link prediction (Wang, Mao, Wang, & Guo, 2017). TransE models a relation as a translation between entity vectors (Bordes, Usunier, Garcia-Duran, Weston, & Yakhnenko, 2013) whereas ComplEx and RotatE improve the representation of asymmetric and compositional relation patterns (Sun, Deng, Nie, & Tang,

2019; Trouillon, Welbl, Riedel, Gaussier, & Bouchard, 2016). ConvE introduces convolutional feature interactions for link prediction (Dettmers, Minervini, Stenetorp, & Riedel, 2018) and R-GCN and CompGCN extend graph convolution to multi-relational structures (Schlichtkrull et al., 2018; Vashishth, Sanyal, Nitin, Agrawal, & Talukdar, 2020). These studies provide a technical basis for combining a lightweight TransE state representation with structural and semantic signals in the present archival setting.

2.3. Reinforcement Learning for Knowledge Graph Reasoning

Embedding models estimate plausibility, but they do not explicitly expose a sequence of decisions for expanding a subgraph. Reinforcement-learning approaches address this limitation by treating graph traversal as a sequential decision process. DeepPath learns multi-hop relation paths with rewards for accuracy, diversity, and efficiency (Xiong, Hoang, & Wang, 2017) while MINERVA trains an end-to-end agent to reach answer entities through graph walks (Das et al., 2018). Reward shaping can reduce false-negative supervision and discourage spurious paths (Lin, Socher, & Xiong, 2018) and M-Walk combines neural policies with Monte Carlo tree search to alleviate sparse rewards (Shen, Chen, Huang, Guo, & Gao, 2018). Unlike endpoint-oriented reasoning, the present study uses policy learning to construct a representative archival subgraph, and its reward jointly considers similarity, relation diversity, and connection to the focal place.

2.4. Spatial Semantics and Research Gap

Spatial context is especially important because customs archives encode places, administrative jurisdictions, routes, and part-whole relations. LinkedGeoData demonstrates how geographic entities can be interlinked in an RDF graph (Stadler, Lehmann, Höffner, & Auer, 2012) and GeoSPARQL supports formal spatial representation and querying (Battle & Kolas, 2012). More recent location-aware embeddings show that explicitly encoding coordinates and spatial extents improves geographic reasoning over knowledge graphs (Mai et al., 2020). Nevertheless, prior work has not integrated multilingual historical archival entities, spatial action selection, and reinforcement-learning-based subgraph generalization for archival translation. This gap motivates the framework developed in this study.

3. KNOWLEDGE GRAPH DATABASE OF THE IMPERIAL GUANGDONG MARITIME CUSTOMS ARCHIVES

The 10 Customs archives stored in Guangdong Provincial Archives is the main component of the Imperial Guangdong Maritime Customs archives, with the time span from 1861 to 1949, huge in number, standard in format, variety in types and genres, rich in content, and has a total of 16115 volumes, with about 3.87 million pictures. The content of these archives involves Customs business, with a large part of the content being local Customs archives and information on local affairs in Guangdong, Hong Kong and Macao. These archives provide valuable historical original materials for the study of modern Sino-foreign exchange history, economic history, trade history, financial history, South China regional history, modern Chinese phonetic history, Guangdong dialect history, Guangdong local history, etc. It has always been a hot spot and focus of academic attention at home and abroad (Lilan & Yongsheng, 2021).

In view of the lack of knowledge graph research on existing archives and relevant benchmarks, we generalize Imperial Guangdong Maritime Customs archives dataset for this study. In this paper, by using unstructured human knowledge to guide the knowledge graph process of Imperial Guangdong Maritime Customs archives, the knowledge graph dataset of Imperial Guangdong Maritime Customs archives has two parallel parts: 1) the title of

each Imperial Guangdong Maritime Customs archive, and 2) the knowledge graph of each Imperial Guangdong Maritime Customs archive and its related entities. By looking up the corresponding Imperial Guangdong Maritime Customs archive, we extract core keywords in the Imperial Guangdong Maritime Customs archives (proper names of people and places, titles, weights and measures), including 182 keywords for people's names, 247 keywords for places, 51 titles and 63 keywords for weights and measures, 443 keywords in total. These keywords function as guidance of human generation for the Imperial Guangdong Maritime Customs archive knowledge graph.

For the knowledge graph of Imperial Guangdong Maritime Customs archives, we choose Canton Customs archives as the data source, because it has many entities of the Imperial Guangdong Maritime Customs archives. Through maintenance and update, each Imperial Guangdong Maritime Customs archive has a clear one-to-one correspondence, and provides diversified and comprehensive attribute coverage. In order to construct the knowledge graph of Imperial Guangdong Maritime Customs archives, we prepared the keywords in Imperial Guangdong Maritime Customs archives, selected part of these keywords and retrieved all the graph links that appeared in the 31504 archives. We generate maps to find the corresponding entities (keywords) and links to these places. After obtaining these seed entities, we generate SPARQL queries to iteratively retrieve neighbors of degree 1 and degree 2 to form subgraphs surrounding these seed nodes. In Imperial Guangdong Maritime Customs archives, all statements are organized as (Head, relation, tail) or (Subject, predicate, object) triples. Query 1 shows an example query that uses basic graph patterns to get entity (Inputting and outputting) neighbor nodes.

```

PREFIX dbr: Guangdong Maritime Customs archive

SELECT DISTINCT * WHERE {{
  Dbr: guangzhou? p1 ?o
  FIELTER ( CONTAINS (STR (?P1), ' Guangdong Maritime Customs archive' (?o) )
  UNION {
    ? S ? DBR:guangzhou
  }
  FIELTER ( CONTAINS (STR (?P1), ' Guangdong Maritime Customs archive' ( ? s) )

```

Figure 1. Sample SPARQL query to retrieve the 1-degree neighbor of dbr: Guangzhou as a title (Subject) and tail (Object) entity.

We only consider relationships with precursors because these map-based relationships are of higher quality. To implement our modeling strategy, we further removed duplicate triples and filtered entities that occurred less than 10 times. Finally, the archival dataset and subgraph containing 443 keywords were obtained. The subgraph connects these 443 entities with other spatial and non-spatial entities (such as people) through different relations, so as to form our knowledge graph of Imperial Guangdong Maritime Customs archives. Figure 2 illustrates the local knowledge graph generated by retrieving the one-degree and two-degree neighbors of the Guangzhou, China, and Huadu entities. It shows how these entities are connected to other archival entities through semantic relations, thereby providing contextual information for subsequent knowledge graph generalization.

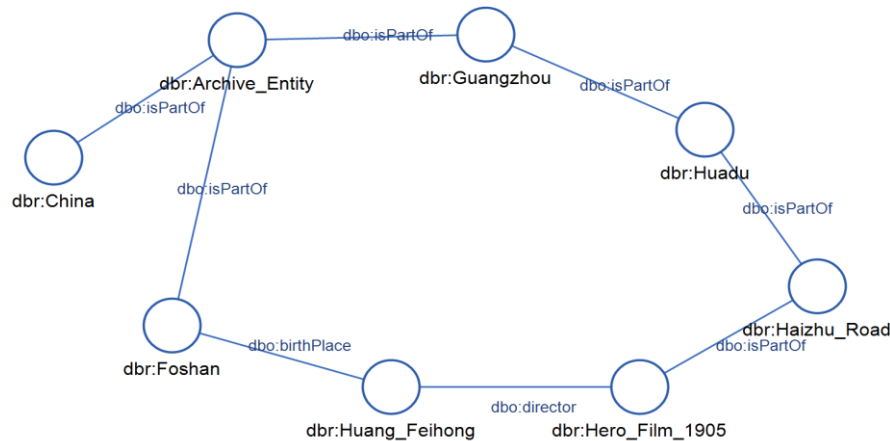


Figure 2. One- and two-degree neighborhood subgraph of the Guangzhou, China, and Huadu entities.

For these 443 keywords, the Imperial Guangdong Maritime Customs archive knowledge graph has a total of $443 \times 266 \times 18$ entities, 234 unique relations, and 443×234 triples. The data were divided into a training set with 402 keywords and a test set with 41 keywords.

4. METHODS

This paper will adopt the reinforcement learning method of Imperial Guangdong Maritime Customs archive knowledge graph. We start with the simplest graph, namely a single node (the Imperial Guangdong Maritime Customs archive entity in question), and through repeated testing, we successively add new relations (edges) and entities (nodes) until the representation of the graph is very consistent with the content of the Imperial Guangdong Maritime Customs archives. By explaining the basic components, such as actions, states and rewards, the reinforcement learning model of Imperial Guangdong Maritime Customs archive knowledge graph network is implemented (Gao, Chen, & Lu, 2004).

4.1. The Reinforcement Learning Frame

The generalization task of Imperial Guangdong Maritime Customs archive graph is formalized as a Markov decision process (S, A, P_a, R_a) , $S = \{s_1, s_2, \dots, s_n\}$ is a set of states containing useful information from the history of Imperial Guangdong Maritime Customs archives. $A = \{a_1, a_2, \dots, a_n\}$ is a series of action groups taken by the state provided by the Imperial Guangdong Maritime Customs archives. As there is no memory of MDP, the state transition probability matrix $P_a(s, s') = \Pr(s_{t+1} = s' | s_t = s, a_t = a)$ shows the probability of reaching state s' at the time $t + 1$ after taking action a at the time t and state s . $R_a(s, s')$ is the immediate reward after taking action a and transitioning from state s to state s' .

To visualize the process of the archive knowledge graph generalization, we assume that the MDP starts with a graph consisting of key words from Imperial Guangdong Maritime Customs archives. At each step, the current state of the process is analyzed (by considering information about the graph and the keywords in Imperial Guangdong Maritime Customs archives) and a decision is made to add a possible relationship to the graph, making it grow in order to more closely approximate the keywords in Imperial Guangdong Maritime Customs archives. The agent receives a certain amount of reward according to the degree of success in reaching the goal in this step. When the process is finished, that is, MDP has been carried out for some time, the graph is expected to be a good expression of the Imperial Guangdong Maritime Customs archive knowledge graph. The agent's goal is to maximize the reward it receives. In this process, the agent learns to find the optimal spot on the spectrum between

information scarcity (the Imperial Guangdong Maritime Customs archive entities) and information overload (the Imperial Guangdong Maritime Customs archive knowledge graph containing 443 nodes) by considering the keywords in the text, namely the title of the Imperial Guangdong Maritime Customs archives (Mackeprang, Müller-Birn, & Stauss, 2019). In order to maintain the balance between exploration and development, the agent's behavior is defined by the random strategy $\pi(a|s) = \Pr(a_t = a|s_t = s)$, which is a probability distribution, determining the probability of the agent taking action in state s , at the time step t . In our model, strategy network is used to study an approximate function capturing interactive dynamics and parameterize the strategy $\pi_\theta(a|s)$. It is a fully connected neural network with two hidden layers. The modified Linear Units (ReLU) are used as the activation function in the hidden layer (Xu, Wang, Chen, & Li, 2015) while the softmax function is used in the output layer to calculate the probability of each possible action. Before the training, we will further explain each concept in the context of the task.

4.2. States

These states are captured in the MDP. Since our model aims to capture both intrinsic and extrinsic information, we utilize the knowledge graph structure and semantic information of the Imperial Guangdong Maritime Customs archives.

Since there are 443 entities in our Imperial Guangdong Maritime Customs archive knowledge graph, it is not feasible to model them as discrete atomic symbols using a vector in the state. To provide a streamlined entity representation, we apply the Imperial Guangdong Maritime Customs archive Knowledge Graph Embedding method (TransE). The TransE (Lu, Liao, Zhang, & Song, 2020) model provides an scalable general method to embed nodes and edges in heterogeneous graphs into the same vector space. This model extracts local and global connectivity patterns between entities by treating the relationships in the graph as transformations embedded in the space. Thus, the intrinsic structure of the graph is embedded in the underlying representation of these entities and relationships. The states in the MDP should help the agent understand the current environment in order to make decisions about the next step. In this case, the entity embedding can help capture the development of the process related to the titles of the Imperial Guangdong Maritime Customs archives. We capture the internal structural information using the synthesis $z_t = \sum_{i \in E_t} e_i$ of entity embeddedness in the current Imperial Guangdong Maritime Customs archive graph at step t , where e_i is the embedding of entity i in a set of entities E_t .

Since these entities also appear as links in the titles of the Imperial Guangdong Maritime Customs archives, we denote the sum of the entity embedding from the target Imperial Guangdong Customs archive titles as $z_{target} = \sum_{i \in E_{target}} e_i$. Where E_{target} is a group of entities that appear in the target Imperial Guangdong Maritime Customs archive titles. The intrinsic components of the state representation are expressed as $s_t^{intrinsic} = (z_t, z_{target} - z_t)$, where the first component (left) encodes the structure of the title graph at step t , and the second component (right) encodes the gap between the current graph structure z_t and the expected structure z_{target} .

For the extrinsic components of the state representation $s_t^{extrinsic}$, we considered the entity and relationship labels in the graph and the Imperial Guangdong Maritime Customs archive titles. Neural word embedding has been proved to be an effective way to encode the meaning of words in natural languages. In this paper, the fastText word embedding model is adopted, after the word embedding is obtained by using the fastText model, the sum of entity and relation label embeddings $h_t = \sum_{l \in L_t} v_l$ is used to help capture the semantic information of the graph at step t . In order to obtain the potential representation of the Imperial Guangdong Maritime Customs archive titles, the SIF (Smooth Inverse Frequency) embedding method is used to generate paragraph embedding h_{target} using word embedding. A theoretical analysis of this method is provided by Cao, Zhong, Cao, and He (2019) with the idea of

multiplying each word vector $\mathbf{v}_w$ by the inverse of its occurrence probability $p(w)$. Here α is a smoothing constant with a default value of 0.001. Then these normalized, normalized and smooth word vectors are addition, divided by the number of words $|W|$ again, and we get the $\mathbf{h}'_{target}$ goal:

$$\mathbf{h}'_{target} = \frac{1}{|W|} \sum_{w \in W} \frac{\alpha}{\alpha + p(w)} \mathbf{v}_w \quad (1)$$

For all the 369 Imperial Guangdong Maritime Customs archive titles' matrix representation, delete the first major component from this matrix to generate the last embedding $\mathbf{h}_{target}$ for each Imperial Guangdong Maritime Customs archive title. Because the top singular vector tends to contain syntactic information, and removing it cleans up the embedding's ability to better express the semantic information.

Similar to the internal components, the external components of the state are represented as $s_t^{extrinsic} = (\mathbf{h}_t, \mathbf{h}_{target} - \mathbf{h}_t)$, and the state is concatenation of the two components:

$$s_t = (s_t^{intrinsic}, s_t^{extrinsic}) = (z_t, z_{target} - z_t, h_t, h_{target} - h_t) \quad (2)$$

After calculating the state representation, calculate the cosine distance between the current graph and the target Imperial Guangdong Maritime Customs archive title (entity embedding and label embedding), respectively, as $dist_{z_t} = 1 - \cos(z_t, z_{target})$ and $dist_{h_t} = 1 - \cos(h_t, h_{target})$. The termination of the process is determined by:

$$\mathcal{T} = \begin{cases} 1, & \text{if } dist_{z_t} \leq \frac{dist_{z_1}}{2} \text{ or } dist_{h_t} \leq \frac{dist_{h_1}}{2} \\ 0, & \text{otherwise} \end{cases} \quad (3)$$

Where $dist_{z_1}$ and $dist_{h_1}$ denote the initial cosine distance between the subgraph and the Imperial Guangdong Maritime Customs archive titles of entities and labels, respectively. When $\mathcal{T} = 1$, the process terminates. This means that if the cosine distance of the entity embedding or label embedding is at most half of the initial cosine distance, the process will terminate.

4.3. Actions

Given the Imperial Guangdong Maritime Customs archive entities and the Imperial Guangdong Maritime Customs archive titles, the agent's goal is to select actions that iteratively generalize the Imperial Guangdong Maritime Customs archive knowledge graph in discussion. Starting from the initial state s_0 , the policy network, as shown in Figure 3, outputs the probability of choosing each action a . Since there are 91 unique relations in the knowledge graph of Imperial Guangdong Maritime Customs archives, the size of normal action space is 692.

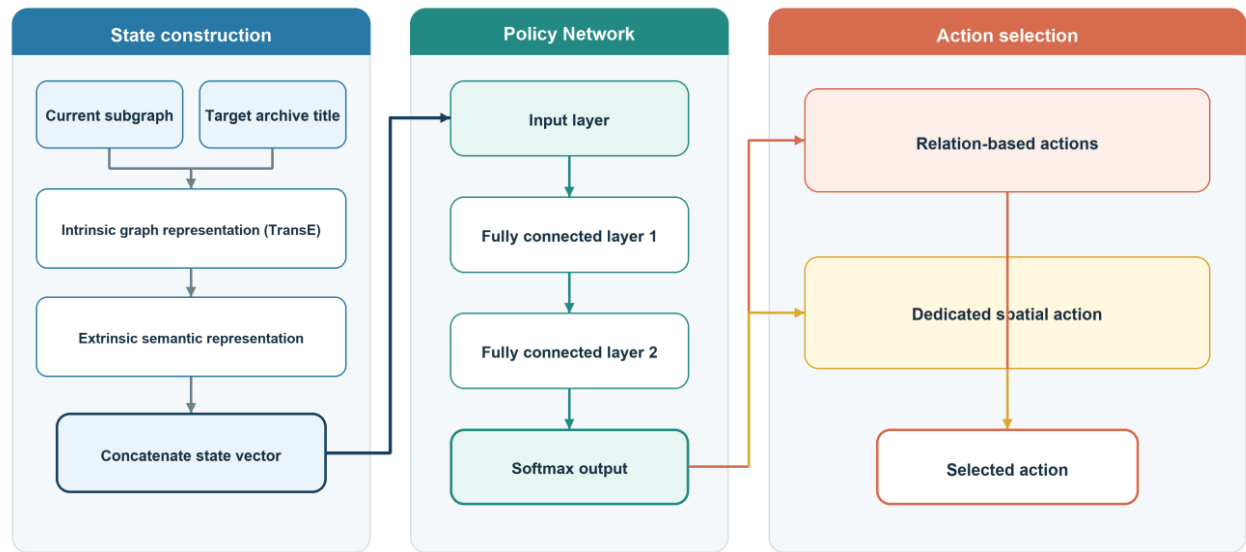


Figure 3. Overview of the state construction, policy network, and action-selection process in the proposed spatially explicit reinforcement learning framework.

After the agent performs an action and decides to add a relationship to the current subgraph, the environment examines possible ways to connect entities on the current subgraph with potential new entities through the selected relationship. Suppose (by examining the graph) that there are n potential triples to be added to the current subgraph, and each triplet contains an entity already in the graph, the selected relationship, and a new entity (head or tail entity) to add. We use index i to denote the new entity, $1 \leq i \leq n$ and $triple_i$ to denote the corresponding triple of entity i . Our model selects triples (and new entities) from all candidate triples of a distribution where the probability of each triplet ($p(triple_i)$) is proportional to the information content of the new entity:

$$p(triple_i) = \frac{-\log(p(i))}{\sum_{j=1}^n -\log(p(j))} \quad (4)$$

Where $p(i)$ is the probability of encountering entity i in the whole Imperial Guangdong Maritime Customs archives, $-\log(p(i))$ stands for the information. The rationale behind this approach is that entities with rich information content carry latent information that can more effectively enrich our knowledge about the entity which we desire.

In addition to the normal 692 actions, we propose a new step that includes a dedicated spatial action to make the model spatially explicit. The idea stems from the data-driven approach, which adopts the hidden patterns of the Imperial Guangdong Maritime Customs archive data, and the spatial actions themselves are modeled as additional actions that the agent can take at any step t . However, if the agent decides to take a spatial action, our model only collects candidates of the Imperial Guangdong Maritime Customs archive entities that are not directly connected to any entity in the current subgraph. We perform a spatial query to retrieve all archive entities within a k -meter radius of where we want to generalize. Our spatially explicit model selects an Imperial Guangdong Maritime Customs archive entity from these candidate Imperial Guangdong Maritime Customs archive entities, where the probability of each candidate Imperial Guangdong Maritime Customs archive entity $p(i)$ is inversely proportional to the distance between the candidate Imperial Guangdong Maritime Customs archive entity and the location q in question:

$$p(i) = \frac{d(j,q)^{-1}}{\sum_{j=1}^n d(j,q)^{-1}} \quad (5)$$

Where $d(j, q)$ represents the geodesic distance between candidate place i and place q . This inverse distance strategy is more beneficial to the neighboring Imperial Guangdong Maritime Customs archive entities. While the spatial radius buffer provides a local Imperial Guangdong Maritime Customs archive knowledge graph surrounding central place entities, this graph is modeled by the PageRank score of each entity in the entire Imperial Guangdong Maritime Customs archive knowledge graph. Intuitively, some places, such as landscape features, have generalized physical characteristics despite their distance due to their overall importance. By combining the global graph and the local Imperial Guangdong Maritime Customs archive graph, the weighted inverse distance is proposed in the probability calculation.

$$p(i) = \frac{pr_i d(j,q)^{-1}}{\sum_{j=1}^n pr_j d(j,q)^{-1}} \quad (6)$$

After deciding to add relationships and entities to a subgraph through spatial or non-spatial actions, a new state representation is generated, and then the new state is presented to the agent to help it decide on the next action.

4.4. Rewards

The reward function plays an important role in guiding the agent's knowledge graph of the Imperial Guangdong Maritime Customs archives. The goal of reinforcement learning model is to find the optimal behavior policy for the agent to obtain the optimal reward. In the model, the reward function has three components, namely, similarity score, diversity score and connection score.

To help the agent choose from such a large action space so that the subgraph represents actions (relations) similar to those to be represented by the Imperial Guangdong Maritime Customs archive titles, one direct approach is to incorporate this mechanism into the immediate reward. Besides calculating the cosine distance after the agent's action, the cosine similarity is also calculated. Then, the normal similarity score is specified as the sum of cosine similarity:

$$r_{similarity}^{normal} = \cos(z_t, z_{target}) + \cos(h_t, h_{target}) \quad (7)$$

The higher the value of cosine similarity, the higher the score of reward component $r_{similarity}^{normal}$. In addition, considering that the TransE models are sometimes not able to deal with the one-to-many and many-to-many relationships, the addition or division of the entity embedding may exacerbate such problems, because the connection information of a single node may be weakened by the coarse aggregation of other nodes, thus another measurement value is proposed to substitute the entity similarity measurement scoring $\cos(z_t; z_{target})$ to highlight the intrinsic structural differences between the subgraph and the Imperial Guangdong Maritime Customs archive knowledge graph. This measurement is inspired by the Hausdorff distance and is usually used to measure differences between two geometries. The cosine distance is used instead of Euclidean distance (as in Hausdorff distance), as it is not affected by magnitude changes of the embedded vectors. The measurement value is written as:

$$sim_{maxmin}(E_t, E_{target}) = 1 - \max_{i \in E_t} \min_{j \in E_{target}} (1 - \cos(e_i, e_j)) \quad (8)$$

Where E_t is the entity group on the subgraph at the time step t , E_{target} is the entity group in the Imperial Guangdong Maritime Customs archive titles, e_i and e_j are the entity embeddings of entities i and j respectively. Then the max-min similarity score is:

$$r_{similarity}^{maxmin} = sim_{maxmin}(E_t, E_{target}) + \cos(h_t, h_{target}) \quad (9)$$

Although there are 692 possible relationships (including spatial actions), these relationships follow a long-tail distribution, which can lead the agent to choose the most likely relationship to avoid penalties. This paper adds a diversity score to the reward function:

$$r_{diversity} = \begin{cases} +0.5 & \text{if relation is not already on the subgraph} \\ -0.5, & \text{otherwise} \end{cases} \quad (10)$$

Because the model may select relationships and entities that are not directly connected to the place entities in question, we want to discourage this behavior. To alleviate this potential problem, we suggest including the connection scores:

$$r_{connection} = \begin{cases} +0.5 & \text{if entity is directly connected to the place} \\ -0.5, & \text{otherwise} \end{cases} \quad (11)$$

The reward function is written as a fusion of three components:

$$R = r_{similarity} + r_{diversity} + r_{connection} \quad (12)$$

It is worth noting that simply reducing the relationship to be selected from a 1-degree query with respect to the entity to be summarized is not an appropriate solution. This restricts the profile subgraph to star-form.

4.5. The Training

We use a policy-based approach to train our spatially explicit reinforcement learning model.

The goal of the policy-based approach is to maximize the total expected future reward J .

$$J(\theta) = \mathbb{E}_{s \sim P_{r(s)}, a \sim \pi_{\theta}(a|s)} R(s, a) \quad (13)$$

Following this reinforcement method, the policy network is updated using the following gradient:

$$\nabla_{\theta} J_{\theta} = \mathbb{E}_{s \sim P_{r(s)}, a \sim \pi_{\theta}(a|s)} Q(s, a) \nabla_{\theta} \log \pi_{\theta}(a|s) \quad (14)$$

$$\approx \frac{1}{N} \sum_{i=0}^N \sum_{a \in \text{eps}_i} Q(s, a) \nabla_{\theta} \log \pi_{\theta}(a|s)$$

Where, from sampling in the process of *eps*, $Q(s_t = s, a_t = a) = \mathbb{E}[G_t | s_t = s, a_t = a]$ is the expected gain from state s after taking action a . The gain $G_t = \sum_{k=0}^{\infty} \gamma^k R_{t+k+1}$ is the total discount gain of the discount factor $\gamma \in [0, 1]$ started from the time step t . A low value of γ implies that the agent is myopic in evaluating the situation and considers immediate rewards more important than delayed future rewards. In addition, similar to the idea of diversity rewards, we use the entropy of the strategy as a regularization term in optimization, where we encourage a more diverse set of actions. The entropy is:

$$H(\theta) = -\sum_{a \in A} \pi_{\theta}(a|s) \log \pi_{\theta}(a|s) \quad (15)$$

In order to maximize the total expected future rewards J and entropy H , the loss function is:

$$\mathcal{L}_{REINFORCE} = -(J + \alpha H) \quad (16)$$

Where α is the regular factor.

Because of the size of the action space, it is challenging for policy agents to learn to choose actions purely based on trial and error. To solve this problem, the model is first trained by supervised learning, and then the supervised strategy is retrained by the proposed reward function to learn and generalize the knowledge graph of the Imperial Guangdong Maritime Customs archives.

In the supervised learning phase, we used the links in the Imperial Guangdong Maritime Customs archive titles to help collect positive training samples. We search the entire graph to check whether links in the Imperial Guangdong Maritime Customs archive titles directly connect to the place entity itself, and track these connections to achieve the Imperial Guangdong Maritime Customs archive graph optimization.

5. PRACTICAL TESTING AND RESULT ANALYSIS

In this section, we use the proper names of people and place and terminology in the Imperial Chinese Maritime Customs archives to do experiments, especially the language knowledge sets, archive professional knowledge set and historical knowledge sets, to detect the performance of the MDP model for the real data sets. Totally there are 143 people participating in the experiment, including 35 sophomores and 108 freshmen. All the test questions are binary “true/false” tests. Some of the test questions from the total budget test questions in the three datasets are listed below, each of which contains different forms of language barrier questions. In the historical knowledge test, the network users complete a matching test between the vocabulary and historical events to match the people, places and time involved in Imperial Chinese historical events.

5.1. Implementation Rules

The 50-dimensional vectors are used in the entity and label embedding of the Imperial Guangdong Maritime Customs archives, so the obtained state representation is a 200-dimensional vector. For spatial actions, we can use a search radius of $k = 100,000$ meters in our archive query. The discount coefficient γ for the cumulative reward used in the model is 0.99. In the policy network, the first hidden layer has 512 units and the second hidden layer has 1024 units. The upper bound of the episode length is set as $max_eps_len = 20$.

Different alternative settings are proposed for actions and rewards respectively. Alternatives in the action component are: non-spatial vs. spatial, unweighted inverse distance. The alternative in the reward component is $r_{similarity}^{normal}$ vs. $r_{similarity}^{maxmin}$. The spatially explicit model of knowledge graph of the Imperial Guangdong Maritime Customs archives is superior in modeling. Different combinations of these alternatives are used to test the method, and five models are obtained, namely $RL_{nonspatial-normal}$ (model with spatial action using $r_{similarity}^{normal}$ score), $RL_{spatial-normal}$ (model with spatial action using $r_{similarity}^{normal}$ score), $RL_{nonspatial-maxmin}$ (model without spatial action using $r_{similarity}^{maxmin}$ score), $RL_{spatial-maxmin}$ (model with spatial action using $r_{similarity}^{maxmin}$ score), $RL_{spatial-maxmin-pr}$ (model with spatial action using $r_{similarity}^{maxmin}$ score and PageRank weighted inverse distance).

5.2. Results

To evaluate the model, consider the intrinsic and extrinsic components separately. Compared with the initial information, the calculation method of the results of the knowledge graph of the Imperial Chinese Maritime Customs archives is as follows:

$$diff_{entity} = \cos(z_T, z_{target}) \quad (17)$$

$$diff_{label} = \cos(h_T, h_{target}) - \cos(h_1, h_{target}) \quad (18)$$

Where $diff_{entity} \in [-2, 2]$, $diff_{label} \in [-2, 2]$ is the difference in cosine similarity between entity and label embeddings, z_T and h_T are the final entity and label embeddings of the generalization graph. The higher the diff difference, the better the generalization results.

Figure 1 and Table 1 show the performance of different policies in the three datasets, but only MDP-RL policy is run in this experiment because the user parameters of the real system are unknown. After using up all the budget test questions, MDP-RL generated 59.5% more cumulative rewards than the best policy for the historical expertise dataset (108.2 vs. 48.7) and 31% more cumulative rewards for the language knowledge dataset (121.5 vs. 90.5). In the archival professional dataset, all methods generated negative rewards, but MDP-RL generated negative rewards 10.8% lower than the best policy (-12.5 vs. -23.3). A two-tailed t-test for these rewards identified these differences for all datasets (the historical expertise dataset, the linguistic knowledge dataset, and the archival expertise dataset with t values of $t = 21.5$, 13.6 , and 46.7 , respectively; $p \ll 0.001$) are of significance.

Table 1. Comparison of rewards obtained by reinforcement learning system in three datasets.

Database	Policy	Rewards	Labels	Precision
Historical expertise	MDP-RL	*108.2	196	87.8
Historical expertise	The best policy	48.7	145	86.5
Historical expertise	The benchmark policy	-134.6	107	79.8
Linguistic knowledge	MDP-RL	*121.5	167	87.2
Linguistic knowledge	The best policy	90.5	129	86.1
Linguistic knowledge	The benchmark policy	-6	105	84.9
Archival expertise	MDP-RL	*-12.5	143	79.6
Archival expertise	The best policy	-23.3	122	78.4
Archival expertise	The benchmark policy	-38.2	99	76.3

Table 1 shows that the reinforcement learning policy achieves higher rewards than the benchmark policy in the three datasets of Imperial Guangdong Maritime Customs archives. Asterisks indicate significantly higher rewards.

The precision distribution of system users and the number of test questions answered by each system user can explain these experimental results. Some low-precision system users answered a large number of test questions in the historical expertise dataset. MDP-RL eliminated these system users through the test and produced a great gain.

The action traces of system users and the test times of MDP-RL model, the archive translation actions and all the statistics of the best test show that the MDP-RL model can improve the static best test benchmark. When the MDP - RL transforms from basic model (static) into adaptive learning, the average number of tests each system user takes remains unchanged or slightly down, but the average times each system user performs archive translation decreases, the MDP - RL model removes system users more frequently, indicating that when performing archive translation action, the MDP-RL model performs more carefully and with higher accuracy.

The experiment results shows the knowledge graph generalization results of Imperial Guangdong Maritime Customs archives using *MDP - RL* model. The model learns to choose different relationships, such as dbo: person name, dbo: place name, dbo: weights and measures, dbo: title, dbo:is Part Of of, and spatial relationships. Our spatially explicit model shows advantages over the non-spatial model. In the dbr: place name example, when we include the connection reward $r_{connection}$ in the model, even if dbr: Guangzhou is included in the subgraph at some point, it does not generalize other entities.

5.3. The Analysis of the System

In this paper, the POMDP-RL model is first constructed and the appropriate model parameters are established. The experiment proves that the system can utilize the knowledge graph model of Imperial Chinese Maritime Customs archives. For the k^{th} system user, this paper randomly selects actions with probability $1/k$. MDP is constructed as Markov-like model to estimate parameters, and the parameters are used to initialize the unsupervised learning EM algorithm, with $P_{slip} < 0.5$, $P_{lose} < 0.5$, $P(\gamma_w) < 0.5$.

Figure 4 shows the learning performance knowing two knowledge points and a set of reasonable parameters ($P_{leave} = 0.01$, $P_{learn} = 0.4$, $P_{lose} = 0.05$, $P_{slip} = 0.1$, $P_{guess} = 0.5$, $P(\gamma_w) = 0.2$, $P(r) = 0.5$). When the parameters are true, the MDP-dataset model obtains 4.6 times the reward of the best policy in the benchmark space of this section (each knowledge point is trained twice).

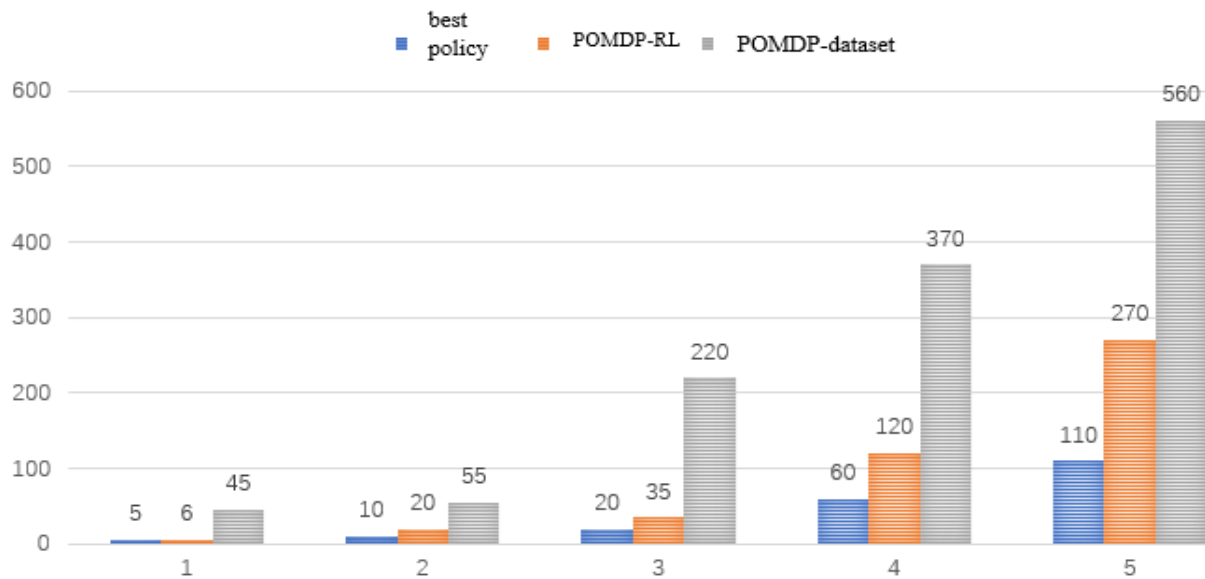


Figure 3. Cumulative rewards for simple (or two-point) problems.

Figure 4 shows the cumulative rewards for simple (or two-point) questions, averaging 1000 simulations with 95% confidence intervals. The POMDP-dataset uses real parameters, and MDP-RL reevaluates the parameters of each newly recruited system user. The best policy is also the benchmark policy, that is, each knowledge point is trained twice.

6. CONCLUSION AND FUTURE WORK

In this study, we introduce and motivate the need for knowledge graph of the Imperial Guangdong Maritime Customs archives, and propose a spatially explicit reinforcement learning framework to learn such graphs. Exploring the different possibilities of generalizing modeling, this paper proposes different actions and reward schemes in the model. By testing the model and using different combinations of alternative components, this paper proves that a spatially explicit reinforcement learning model produces better generalization results than the non-spatial model, so that space is indeed special in terms of knowledge graph generalization.

In the future, we would like to test whether reducing the variance of Monte Carlo policy gradient methods using a dominance function or an action-critic framework can help improve performance. Finally, our approach and other methods do not consider data type properties, which is an important goal for future research.

REFERENCES

- Battle, R., & Kolas, D. (2012). Enabling the geospatial semantic web with parliament and geosparql. *Semantic Web*, 3(4), 355-370. <https://doi.org/10.3233/SW-2012-0065>
- Bordes, A., Usunier, N., Garcia-Duran, A., Weston, J., & Yakhnenko, O. (2013). Translating embeddings for modeling multi-relational data. *Advances in Neural Information Processing Systems*, 26, 2787-2795.
- Cao, W., Zhong, J., Cao, G., & He, Z. (2019). Physiological function assessment based on kinect v2. *IEEE Access*, 7, 105638-105651. <https://doi.org/10.1109/ACCESS.2019.2932101>
- Colavizza, G., Blanke, T., Jeurgens, C., & Noordegraaf, J. (2021). Archives and AI: An overview of current debates and future perspectives. *ACM Journal on Computing and Cultural Heritage*, 15(1), 1-15. <https://doi.org/10.1145/3479010>
- Das, R., Dhuliawala, S., Zaheer, M., Vilnis, L., Durugkar, I., Krishnamurthy, A., . . . McCallum, A. (2018). *Go for a walk and arrive at the answer: Reasoning over paths in knowledge bases using reinforcement learning*. Paper presented at the The International Conference on Learning Representations.

- Dettmers, T., Minervini, P., Stenetorp, P., & Riedel, S. (2018). Convolutional 2D knowledge graph embeddings. *Proceedings of the AAAI Conference on Artificial Intelligence*, 32(1), 1811-1818. <https://doi.org/10.1609/aaai.v32i1.11573>
- Ehrmann, M., Hamdi, A., Pontes, E. L., Romanello, M., & Doucet, A. (2023). Named entity recognition and classification in historical documents: A survey. *ACM Computing Surveys*, 56(2), 1-47. <https://doi.org/10.1145/3604931>
- Gao, Y., Chen, S.-F., & Lu, X. (2004). Research on reinforcement learning technology: A review. *Acta Automatica Sinica*, 30(1), 86-100.
- Hogan, A., Blomqvist, E., Cochez, M., d'Amato, C., Melo, G. D., Gutierrez, C., . . . Neumaier, S. (2021). Knowledge graphs. *ACM Computing Surveys*, 54(4), 1-37. <https://doi.org/10.1145/3447772>
- Hyvönen, E. (2020). Using the semantic web in digital humanities: Shift from data publishing to data-analysis and serendipitous knowledge discovery. *Semantic Web*, 11(1), 187-193. <https://doi.org/10.3233/SW-190386>
- Isaac, A., & Haslhofer, B. (2013). Europeana linked open data—data. europeana. eu. *Semantic Web*, 4(3), 291-297. <https://doi.org/10.3233/SW-120092>
- Ji, S., Pan, S., Cambria, E., Marttinen, P., & Yu, P. S. (2021). A survey on knowledge graphs: Representation, acquisition, and applications. *IEEE Transactions on Neural Networks and Learning Systems*, 33(2), 494-514. <https://doi.org/10.1109/TNNLS.2021.3070843>
- Lehmann, J., Isele, R., Jakob, M., Jentzsch, A., Kontokostas, D., Mendes, P. N., . . . Auer, S. (2015). Dbpedia—a large-scale, multilingual knowledge base extracted from wikipedia. *Semantic Web*, 6(2), 167-195. <https://doi.org/10.3233/SW-140134>
- Li, D., & Madden, A. (2019). Cascade embedding model for knowledge graph inference and retrieval. *Information Processing & Management*, 56(6), 102093. <https://doi.org/10.1016/j.ipm.2019.102093>
- Lilan, C., & Yongsheng, C. (2021). Spatial-temporal adaptive optimal allocation of archival tasks. *IEEE Access*, 9, 25809-25817. <https://doi.org/10.1109/ACCESS.2021.3057362>
- Lin, X. V., Socher, R., & Xiong, C. (2018). *Multi-hop knowledge graph reasoning with reward shaping*. Paper presented at the Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing (pp. 3243-3253).
- Lu, M., Liao, L., Zhang, F., & Song, D. (2020). *Exploration of transE in a distributed environment*. Paper presented at the IEEE 40th International Conference on Distributed Computing Systems (ICDCS) 2020 Nov 1 (pp. 1173-1174). IEEE.
- Mackeprang, M., Müller-Birn, C., & Stauss, M. T. (2019). Discovering the sweet spot of human-computer configurations: A case study in information extraction. *Proceedings of the ACM on Human-Computer Interaction*, 3(CSCW), 1-30.
- Mai, G., Janowicz, K., Cai, L., Zhu, R., Regalia, B., Yan, B., . . . Lao, N. (2020). SE-KGE: A location-aware knowledge graph embedding model for geographic question answering and spatial semantic lifting. *Transactions in GIS*, 24(3), 623-655. <https://doi.org/10.1111/tgis.12629>
- Mendes, P. N., Jakob, M., García-Silva, A., & Bizer, C. (2011). *DBpedia spotlight: shedding light on the web of documents*. Paper presented at the Proceedings of the 7th International Conference on Semantic Systems (pp. 1-8).
- Muehlberger, G., Seaward, L., Terras, M., Ares Oliveira, S., Bosch, V., Bryan, M., . . . Fiel, S. (2019). Transforming scholarship in the archives through handwritten text recognition: Transkribus as a case study. *Journal of Documentation*, 75(5), 954-976. <https://doi.org/10.1108/JD-07-2018-0114>
- Navigli, R., & Ponzetto, S. P. (2012). BabelNet: The automatic construction, evaluation and application of a wide-coverage multilingual semantic network. *Artificial Intelligence*, 193, 217-250. <https://doi.org/10.1016/j.artint.2012.07.001>
- Paulheim, H. (2016). Knowledge graph refinement: A survey of approaches and evaluation methods. *Semantic Web*, 8(3), 489-508. <https://doi.org/10.3233/SW-160218>

- Schlichtkrull, M., Kipf, T. N., Bloem, P., van den Berg, R., Titov, I., & Welling, M. (2018). *Modeling relational data with graph convolutional networks*. In A. Gangemi, R. Navigli, M.-E. Vidal, P. Hitzler, R. Troncy, L. Hollink, A. Tordai, & M. Alam (Eds.), *The Semantic Web: ESWC 2018 (Lecture Notes in Computer Science, Vol. 10843, pp. 593–607)*. Cham: Springer.
- Shen, Y., Chen, J., Huang, P.-S., Guo, Y., & Gao, J. (2018). M-walk: Learning to walk over graphs using monte carlo tree search. *Advances in Neural Information Processing Systems*, 31, 6786-6797.
- Stadler, C., Lehmann, J., Höffner, K., & Auer, S. (2012). Linkedgeodata: A core for a web of spatial open data. *Semantic Web*, 3(4), 333-354. <https://doi.org/10.3233/SW-2011-0052>
- Sun, Z., Deng, Z.-H., Nie, J.-Y., & Tang, J. (2019). *RotatE: Knowledge graph embedding by relational rotation in complex space*. Paper presented at the The International Conference on Learning Representations.
- Sun, Z., Wang, C., Hu, W., Chen, M., Dai, J., Zhang, W., & Qu, Y. (2020). *Knowledge graph alignment network with gated multi-hop neighborhood aggregation*. Paper presented at the Proceedings of the AAAI Conference on Artificial Intelligence (Vol. 34, No. 01, pp. 222-229).
- Trouillon, T., Welbl, J., Riedel, S., Gaussier, E., & Bouchard, G. (2016). *Complex embeddings for simple link prediction*. Paper presented at the 33rd International Conference on Machine Learning (pp. 2071-2080). Proceedings of Machine Learning Research, 48.
- Vashishth, S., Sanyal, S., Nitin, V., Agrawal, N., & Talukdar, P. (2020). *Interact: Improving convolution-based knowledge graph embeddings by increasing feature interactions*. Paper presented at the Proceedings of the AAAI Conference on Artificial Intelligence (Vol. 34, No. 03, pp. 3009-3016).
- Vrandečić, D., & Krötzsch, M. (2014). Wikidata: A free collaborative knowledgebase. *Communications of the ACM*, 57(10), 78-85. <https://doi.org/10.1145/2629489>
- Wang, Q., Mao, Z., Wang, B., & Guo, L. (2017). Knowledge graph embedding: A survey of approaches and applications. *IEEE Transactions on Knowledge and Data Engineering*, 29(12), 2724-2743. <https://doi.org/10.1109/TKDE.2017.2754499>
- Xiong, W., Hoang, T., & Wang, W. Y. (2017). *DeepPath: A reinforcement learning method for knowledge graph reasoning*. Paper presented at the Conference on Empirical Methods in Natural Language Processing (pp. 564-573).
- Xu, B., Wang, N., Chen, T., & Li, M. (2015). Empirical evaluation of rectified activations in convolutional network. *arXiv preprint arXiv:1505.00853*. <https://doi.org/10.48550/arXiv.1505.00853>

Online Science Publishing is not responsible or answerable for any loss, damage or liability, etc. caused in relation to/arising out of the use of the content. Any queries should be directed to the corresponding author of the article.