

Title	知識グラフの誤り検出と訂正に向けた大規模言語モデルの活用に関する研究
Author(s)	董, 娜
Citation	
Issue Date	2026-03
Type	Thesis or Dissertation
Text version	ETD
URL	https://hdl.handle.net/10119/20586
Rights	
Description	Supervisor: 白井 清昭, 先端科学技術研究科, 博士

氏 名	DONG, Na
学 位 の 種 類	博士 (情報科学)
学 位 記 番 号	博情第 573 号
学 位 授 与 年 月 日	令和 8 年 3 月 25 日
論 文 題 目	Exploration of Large Language Models for Noise Detection and Refinement in Knowledge Graphs
論 文 審 査 委 員	白井 清昭 北陸先端科学技術大学院大学 教授 池田 心 北陸先端科学技術大学院大学 教授 井之上 直也 北陸先端科学技術大学院大学 准教授 KERTKEIDKACHORN, Natthawut 北陸先端科学技術大学院大学 准教授 SORNLERTLAMVANICH, Virach 武蔵野大学 教授

論文の内容の要旨

Knowledge graphs (KGs) have emerged as a powerful paradigm for representing structured knowledge about the real world. KGs encode information as a collection of triples, each consisting of a head entity, a relation, and a tail entity. This kind of structure enables not only the integration of heterogeneous data sources but also the support of downstream tasks such as question answering, recommender systems, and semantic search. However, the practical application of KGs is often hindered by several common challenges. KGs automatically constructed from large text corpora are prone to errors due to natural language ambiguity and extraction processes; while manually constructed KGs, while more reliable, are still susceptible to human error, have limited scalability, and are incomplete. As a result, KGs frequently contain misinformation, redundancy, and noisy triples, all of which can undermine their reliability and reduce their effectiveness in real-world applications.

This dissertation addresses the critical problem of improving the quality of noisy KGs through a two-step framework: (1) detecting noisy triples, and (2) refining the noisy triples by correcting the entities involved. Noisy KGs not only distort factual knowledge but also propagate errors to downstream reasoning and knowledge completion tasks. Unlike conventional approaches that rely heavily on hand-crafted rules, statistical heuristics, or embedding-based anomaly detection, our work leverages recent advances in large language models (LLMs), which possess strong semantic understanding and contextual reasoning capabilities, making them well-suited for identifying and repairing noisy knowledge structures at scale.

Within this framework, we propose two methods: *LLM_sim* and *LLM_rule*. The first method, *LLM_sim*, evaluates the plausibility of a candidate triple by comparing

it with existing triples in the KG, thereby capturing subtle semantic inconsistencies that traditional models often miss. In addition to detection, we designed *LLM_sim* to repair detected noisy triples by grouping and calculating similarity. The second method, *LLM_rule*, induces logical rules directly from the KG to capture semantic constraints between entities and relations. We generate relevant rules through *LLM_rule*, which can propose effective corrections, providing interpretability and transparency that are often missing in similarity-based or embedding-based methods.

To validate the effectiveness of the proposed methods, we conducted extensive experiments on both synthetic and real-world noisy KGs. Synthetic datasets allow controlled injection of noise. These datasets enable systematic evaluation of detection and refinement capabilities under varying noise ratios. Real-world datasets reflect the complex and heterogeneous nature of naturally occurring noise. Across both settings, our results demonstrate that *LLM_sim* and *LLM_rule* achieve strong performance in identifying and refining noisy triples.

Crucially, our comparative analysis reveals that **LLM_rule** generally emerges as the superior model, outperforming *LLM_sim* in both detection accuracy and refinement quality across fact-oriented datasets (e.g., FB15k-237, s-NELL). While *LLM_sim* remains competitive in linguistically rich domains like WN18RR, *LLM_rule*'s integration of explicit logical constraints offers greater robustness against hallucination and structural inconsistencies. We further investigated the subsequent impact of noise refinement on knowledge graph completion (KGC). KGC aims to infer missing triples in a KG by learning patterns from observed triples. However, noise in the training data severely impairs the generalization performance of KGC models. By incorporating refined KGC generated by *LLM_sim* and *LLM_rule*, we observed significant improvements in standard KGC metrics, including mean reciprocal rank (MRR) and hits@k. These improvements confirm that high-quality input data is a crucial factor in improving the performance of completion models and that noise refinement can serve as a key preprocessing step in the KGC pipeline.

The contributions of this work can be summarized as follows. First, we formulate a general framework for addressing noise in KGs that combines detection and refinement in a unified process. Second, we propose two novel methods, *LLM_sim* and *LLM_rule*, that leverage the semantic reasoning and rule induction capabilities of LLMs to achieve robust noise handling. Third, we provide extensive empirical evidence from both synthetic and real-world datasets, demonstrating that our methods consistently outperform baseline approaches in both detection accuracy and refinement quality. Finally, we show that our noise refinement framework directly enhances the performance of downstream KGC models, thereby underscoring its practical significance for KG applications. This study advances the state of the art in KG refinement by integrating the

semantic power of LLMs with interpretable rule-based reasoning. The proposed methods offer a scalable, effective, and interpretable solution to the longstanding problem of noise in KGs. By improving KG quality, our work not only enhances their reliability but also paves the way for more accurate and trustworthy knowledge-driven AI systems. The findings of this dissertation hold broader implications for the development of robust intelligent systems that depend on high-quality structured knowledge.

Keywords: Knowledge Graph, Knowledge Acquisition, Error Correction, Large Language Model, Rule Induction, Knowledge Graph Completion

論文審査の結果の要旨

知識グラフはエンティティをノード、エンティティ間の関係をリンクとするグラフ構造を持つ知識データベースであり、 (h, r, t) という 3 つ組 (h, t はエンティティ、 r は関係) の集合として定式化される。特に自動構築された知識グラフは誤った (h, r, t) を含むことも多く、これを自動検出し修正することは重要な研究課題である。本論文は大規模言語モデル (LLM) を活用して知識グラフの誤りを検出しこれを訂正する新しい 2 つの手法を提案している。

第 1 の手法は `LLM_sim` と呼ばれる。まず誤り検出を行う。判定対象の (h, r, t) が与えられたとき、それを自然言語文に変換し、さらに埋め込み表現に変換する。次に、知識グラフにおける他の 3 つ組から判定対象の 3 つ組と似ているものを検索する。そして、LLM にプロンプトを与えて判定対象の 3 つ組が正しいか否かを判定させる。このプロンプトに取得した類似の 3 つ組を含めることで LLM による判定の精度を高める。次に誤り訂正を行う。誤りと判定した 3 つ組に対して、LLM によってそれを訂正した 3 つ組を複数出力させる。これらの候補に対して既存の知識グラフにおける 3 つ組との類似度を計算し、その類似度が高い候補を選択する。

第 2 の手法は `LLM_rule` と呼ばれる。まず、関係 r に対するルールを LLM に生成させる。ここでのルールとは関係 r を持つ h と t が満たすべき制約を表す文である。誤り検出では、生成した 1 つのルールを参照して 3 つ組が正しいかを指示するプロンプトを LLM に与える。これをルール毎に繰り返し、判定結果 (誤りか否か) の `voting` によって最終的な判定を決定する。誤り訂正では、生成した 5 つのルールを全て与え、これを参照しつつ訂正後の (h, r, t) を生成するプロンプトを LLM に与える。

評価実験では、既存の 3 つの知識グラフに対し、10%、20%、30% の 3 つ組における h または t を別のエンティティに置換することで人工的に誤りを生成し、異なる誤り率の知識グラフに対する提案手法の有効性を検証した。これに加え、人手で誤りタグを付与し、またそれを訂正したデータセットも評価に用いる。実験の結果、いずれの実験設定でも、誤り検出タスク、誤り訂正タスクの両方において、`LLM_sim` と `LLM_rule` は LLM を用いない従来手法ならびに LLM を用いたゼロショットプロンプトのベースラインよりも高い性能を示した。

以上、本論文は、知識グラムの誤り検出ならびに訂正について LLM を効率的に活用する 2 つの有望なアプローチを示したものであり、学術的に貢献するところが大きい。よって博士 (情報科学) の学位論文として十分価値あるものと認めた。