

Schema Generation for Large Knowledge Graphs Using Large Language Models

Bohui Zhang¹ Yuan He² Lydia Pintscher³ Albert Meroño Peñuela¹ Elena Simperl^{1,4}

¹King’s College London ²University of Oxford ³Wikimedia Deutschland

⁴Technical University of Munich

{bohui.zhang, elena.simperl}@kcl.ac.uk

Abstract

Schemas play a vital role in ensuring data quality and supporting usability in the Semantic Web and natural language processing. Traditionally, their creation demands substantial involvement from knowledge engineers and domain experts. Leveraging the impressive capabilities of large language models (LLMs) in tasks like ontology engineering, we explore schema generation using LLMs. To bridge the resource gap, we introduce two datasets: YAGO Schema and Wikidata EntitySchema, along with novel evaluation metrics. The LLM-based pipelines utilize local and global information from knowledge graphs (KGs) to generate schemas in Shape Expressions (ShEx). Experiments demonstrate LLMs’ strong potential in producing high-quality ShEx schemas, paving the way for scalable, automated schema generation for large KGs. Furthermore, our benchmark introduces a new challenge for structured generation, pushing the limits of LLMs on syntactically rich formalisms.

1 Introduction

Graphs have emerged as a vital area of research in artificial intelligence and its foundational disciplines, significantly advancing progress across various domains, including knowledge representation and natural language processing (Sakr et al., 2021; Scherp et al., 2024; Hogan et al., 2025). This is especially evident with the rise of large language models (LLMs) (Brown et al., 2020), where graph-based methods enhance their reasoning capabilities for structured knowledge integration, and graphs serve as rich sources of structured and factual information (Zhang, 2023; Sun et al., 2023; Edge et al., 2024). Large knowledge graphs (KGs), such as Wikidata (Vrandečić and Krötzsch, 2014) and DBpedia (Auer et al., 2007), are compiled from heterogeneous sources, leading to significant quality issues like redundancy,

noise, and ambiguity (Shenoy et al., 2022). Beyond data noise, KGs frequently suffer from modeling issues. A survey by Wikimedia Deutschland on Wikidata’s ontology issues revealed conceptual ambiguity and inconsistent modeling in Wikidata, stemming from diverse contributor perspectives and inadequate guidelines—challenges common to large KGs (Ammalainen, 2023). This can manifest, for example, as entities for a company, its service and application being conflated, with predicates incorrectly shared between them. These quality deficits impede effective KG querying, sharing, and reuse. Critically, as KGs underpin tasks like pre-training (Chen et al., 2020; Pan et al., 2022; Yasunaga et al., 2022), retrieval-augmented generation (Xu et al., 2024; He et al., 2024a; Fang et al., 2024; Hu et al., 2025), and post-training (Agarwal et al., 2021; Li et al., 2023; Tang et al., 2024) for LLMs, ensuring KG quality is essential for maintaining the factual accuracy and reliability of these downstream AI systems (Pan et al., 2023).

As crucial resources for quality assurance, KG schemas are generally classified into semantic, validating, and emergent schemas (Hogan et al., 2021). In this work, we focus on **validating schemas**, which are particularly effective in detecting structural and semantic inconsistencies and thereby ensuring data integrity (Gayo et al., 2018; Scherp et al., 2024). These schemas are defined through *shapes*, where each shape identifies a set of focus nodes within a graph and specifies the constraints and patterns they must adhere to (Ahmetaj et al., 2025). To represent such schemas, several dedicated languages have been developed, including W3C standards like Shapes Constraint Language (SHACL) (Knublauch and Kontokostas, 2017), Shape Expressions (ShEx) (Prud’hommeaux et al., 2014), and PG-Schema (Angles et al., 2023). However, developing high-quality validation schemas aligned with user needs remains a substantial challenge (Rabbani et al., 2022). Current automatic

schema generation via pattern aggregation often inherits noise and errors from the source KGs (Rabbani et al., 2022). As a result, the primary approach involves manually writing schemas or refining those produced by basic pattern mining—a process that is both time-consuming and expensive, particularly for KGs with potentially millions of classes (Rabbani et al., 2022). The demonstrated success of LLMs in areas like ontology engineering (Lo et al., 2024; Zhang et al., 2025a) and graph-based reasoning (Wang et al., 2023) highlights their proficiency with structured data. Specifically, LLMs show aptitude for processing structured information, reasoning across graph and text modalities, and performing structured generation (He et al., 2024b; Jin et al., 2024). This aptitude strongly suggests their potential for automatically generating KG schemas.

Building on this, this paper investigates how LLMs can undertake this automatic schema generation for large KGs, especially where such schemas are missing or incomplete. Specifically, we address the research question: *How can LLMs generate high-quality ShEx shapes for a given KG?* We identify key challenges inherent to this task. First, due to the sheer scale of large KGs, pinpointing the essential information required for accurate ShEx generation by LLMs is difficult. In addition, even when relevant information is identified, investigating how LLMs can effectively leverage this structured data for schema generation is required. Moreover, the relative novelty and ongoing development of the ShEx standard means that publicly available ShEx corpora are scarce, potentially hindering LLM familiarity with the required syntax and conventions. To address these challenges, this paper makes the following main contributions:

- We formulate the novel task of using LLMs to generate ShEx schemas from KGs, relevant for both the Semantic Web and LLM research.
- We introduce a benchmark comprising two datasets derived from large KGs—YAGO and Wikidata—along with associated metrics designed for comprehensive evaluation from diverse perspectives¹.
- We propose and evaluate pipelines designed to efficiently generate high-quality ShEx

schemas by leveraging the structured generation capabilities of LLMs².

2 Related Work

Existing automatic ShEx schema extraction approaches operate on KGs provided as RDF data or via query endpoints, employing a two-stage process: collecting essential information from KGs, and then extracting and refining shapes based on this information. Mihindukulasooriya et al. (2018) approached shape generation by framing cardinality and range constraint prediction as machine learning (ML) classification problems on data profiling features. The sheXer system (Fernandez-Álvarez et al., 2022) extracts shapes by iteratively exploring triples associated with target nodes and mining KG structures. A key finding is that shapes derived from representative instance examples often converge with those from larger sets, suggesting sampling can be effective for accurate shape extraction. Influenced by graph pattern mining, QSE (Rabbani et al., 2023) extracts shapes from large KGs by computing support (number of entities satisfying a constraint) and confidence (proportion of entities satisfying a constraint among applicable entities). While defined differently, these metrics are conceptually similar to sheXer’s internal voting and trustworthiness scores.

Despite these advancements, studies (Rabbani et al., 2022; Fernandez-Álvarez et al., 2022; Rabbani et al., 2023) indicate that current methods often produce incomplete shapes, frequently missing essential elements like cardinality constraints, and lack standardized benchmarks for schema generation. Current evaluation practices primarily focus on generation efficiency, such as running time and memory usage. These limitations underscore the need for novel approaches to generate more comprehensive and accurate validating schemas for semantic quality and completeness.

3 Problem Formulation

The Resource Description Framework (RDF) is a standard model for representing data and is widely used in building KGs (Cyganiak et al., 2014). We define a KG as follows:

Definition 1 (Knowledge Graph) A KG in RDF is a directed, labeled graph $\mathcal{G}$, defined as a set of triples. Each triple $(u, p, o) \in \mathcal{G}$ consists of a

¹The data is available at: <https://doi.org/10.5281/zenodo.17128093>

²The code is available at <https://github.com/King-s-Knowledge-Graph-Lab/shapespresso>

subject u , a predicate p , and an object o . Given distinct sets of IRIs $\mathcal{I}$, blank nodes $\mathcal{B}$, and literals $\mathcal{L}$, the subject $u \in (\mathcal{I} \cup \mathcal{B})$ must be an IRI or a blank node, the predicate $p \in \mathcal{I}$ must be an IRI, and the object $o \in (\mathcal{I} \cup \mathcal{B} \cup \mathcal{L})$ may be an IRI, a literal, or a blank node.

In particular, we distinguish non-blank subject nodes based on their role in the taxonomy: classes and instances. Classes are primarily connected to their superclasses using predicates such as `subClassOf`. Instances are linked to their corresponding classes using predicates like `instanceOf`. To validate instances within a class, we can define a validating schema in ShEx as follows:

Definition 2 (Shape) A shape schema in ShEx consists of a set of shapes $\mathcal{S}$. Each shape is associated with a class c in a KG. A shape $s = (\alpha, \Psi) \in \mathcal{S}$ comprises a shape label α and a set of constraints $\psi \in \Psi$ in the form $\psi = (p, \tau, \kappa)$, where p is a predicate, τ is a node constraint, and κ specifies cardinality.

Predicates used in shape constraints are drawn from the KG. Node constraints τ may belong to several categories, including (1) node kind constraints, (2) datatype constraints, (3) values constraints, (4) XML Schema string facet constraints and (5) XML Schema numeric facet constraints. In this work, we focus on the first three categories. Cardinality $\kappa = (n, m)$ specify the allowable number of occurrences of a predicate-object pair, where $n \in \mathbb{N}$ and $m \in \mathbb{N} \cup \{*\}$, with ‘*’ indicating an unbounded upper limit. ShEx examples are provided in Appendix B. The task of schema generation can now be defined as:

Definition 3 (Schema Generation) Given a KG $\mathcal{G}$, and a class $c \in \mathcal{G}$ representing a set of nodes, the objective is to generate a shape schema $\mathcal{S}$ describing the triples involving nodes in the KG.

Schema generation is a challenging task for both traditional rule-based methods (Rabbani et al., 2023) and ML approaches, due to several factors. Large KGs often suffer from quality issues, and pattern extraction may inadvertently propagate these issues into generated shapes. Moreover, the lack of benchmarks and high-quality ShEx ground truths makes training and evaluation of ML-based approaches particularly difficult.

Dataset		YAGOS	WES
Classes		36	50
Constraints	Sum	678	1,874
	Mean	18.83	37.48
	Median	18	35
Instances	Sum	1,227,509	2,127,696
	Mean	34,097.47	42,553.92
	Median	1,104	1,564

Table 1: Dataset statistics, including the number of classes, total number of constraints, average schema length, median constraint length, total number of instances across all classes, and the mean and median number of instances per class.

4 Schema Benchmark

4.1 Dataset

To address the need for benchmarks, we introduce two new dataset, YAGO Schema (YAGOS) and Wikidata EntitySchema (WES). Together, comprising 86 ShEx schema with a total of 2,552 constraints. Detailed statistics are presented in Table 1. Each schema in our datasets targets a specific class within the respective KG. While ShEx offers a rich syntax, the schemas in our datasets employ a community-prevalent subset to enhance simplicity and readability while remaining functional (detailed specification in Appendix B). In our setting, to simplify the syntax representation and ease evaluation and comparison, a shape schema $\mathcal{S}$ includes a *start shape* that targets the focus class. The remaining shapes in the schema serve as references, each describing one or more classes that specify the range of objects allowed for certain predicates in the start shape. The schemas predominantly use four fundamental node constraint types: node kind constraints (e.g., IRI), datatype constraints (e.g., `xsd:decimal`, `rdf:langString`), value set constraints (often specifying object values, e.g., `[schema:Organization]`), and shape references (e.g., `@<Person>`, which links to another defined shape like `<Person> { rdf:type [ schema:Person ] }`).

The YAGOS dataset was constructed using YAGO 4.5 (Suchanek et al., 2024) as the input KG. Ground truth ShEx schema construction was informed by the existing SHACL schema and the official YAGO 4.5 design document. While the documentation provided a starting point, it covered only a subset of properties and constraints, so the final schemas were manually refined by retaining

documented constraints and adding frequently used predicates identified through KG queries to match the intended scope.

The WES dataset was developed through a semi-automatic process. Target classes were selected from three sources: the community-curated Wikidata EntitySchema directory, classes mapped from YAGOS, and those mapped from Wikipedia categories. For each selected class, three knowledge engineers independently annotated its schema through an iterative process of automatic predicate compilation and prioritization, semantic refinement and deduplication, and cardinality and node constraint specification (datatype, value set, shape reference, or node kind). Constraints were retained if at least two annotators agreed, while disagreements were resolved through discussion or omission to ensure high-quality, consistent ShEx ground truth schemas. Further details of the dataset construction process are provided in Appendix A.

4.2 Evaluation Metrics

Novel metrics are essential for evaluating the quality of constructed schema. Simple text matching is unreliable, since the order of constraints and naming of shapes in ShEx affect textual similarity but not semantic correctness. In this work, we evaluate using the ShExJ (JSON format), converted from ShExC, which enables structured analysis at shape and constraint levels. Our evaluation uses two metric types: similarity, with graph edit distance at the shape level, and classification, with F1-score as the main metrics on the constraint level.

4.2.1 Similarity Metrics

Given that automatically generated schemas often require manual refinement (Rabbani et al., 2022), quantifying the structural similarity between a generated schema and its ground truth counterpart is crucial. To facilitate this comparison, we model each shape schema as a rooted, labeled tree graph. The shape label (or focus node) serves as the root, predicates form the first level of child nodes, node constraints linked to predicates occupy the next level, and cardinalities appear as leaf nodes, as shown in Figure 3b. Based on this graph representation, we employ the Graph Edit Distance (GED) (Sanfeliu and Fu, 1983), specifically the Tree Edit Distance (Zhang and Shasha, 1989), to measure the dissimilarity between a generated schema S' and the ground truth schema S . GED represents the minimum cost required to transform

S' into S using a sequence of edit operations:

$$\mathcal{D}(S', S) = \min_{e_1, \dots, e_{\mathcal{L}} \in \gamma(S', S)} \sum_{i=1}^{\mathcal{L}} c(e_i) \quad (1)$$

where $\gamma(S', S)$ is the set of all valid edit paths transforming S' to S , and $c(e_i)$ is the cost of an individual edit operation e_i . In this context, GED has a complexity of $O(|S'| \cdot |S|)$, where $|S|$ denotes the number of candidate constraints (i.e., distinct predicates), since schemas are represented as fixed-depth trees. Simply averaging raw GED scores across a dataset can be misleading. Smaller schemas naturally yield lower scores, possibly hiding significant relative errors, whereas larger schemas might skew the average. To mitigate this size bias, we normalize the GED by the maximum potential edit cost related to the ground truth schema's size, defined as $3 \cdot |S|$.

$$\tilde{\mathcal{D}}(S', S) = \frac{1}{3 \cdot |S|} \mathcal{D}(S', S) \quad (2)$$

The normalized GED (NGED) ranges from 0 (identical schemas) to potentially above 1 (if the transformation cost exceeds deleting the ground truth, e.g., due to many additions in S'). We report the average GED and average NGED over the dataset of N schemas:

$$\begin{aligned} \text{GED} &= \frac{1}{N} \sum_i^N \mathcal{D}(S'_i, S_i) \\ \text{NGED} &= \frac{1}{N} \sum_i^N \tilde{\mathcal{D}}(S'_i, S_i) \end{aligned} \quad (3)$$

The metrics offer an interpretable measure of the structural accuracy of generated schemas against ground truth, considering both the absolute number of edits and the relative error normalized by size.

4.2.2 Classification Metrics

Since each ShEx constraint comprises three elements (predicate, node constraint, and cardinality), we propose several levels of matching criteria, inspired by Fernandez-Álvarez et al. (2022), where constraints of different specificity can get positive votes. These criteria reflect practical usage scenarios and accommodate variations in large KGs modeling, ensuring meaningful evaluation without sacrificing flexibility.

Exact Matching We define an exact match between two constraints ψ and ψ' as $\mathcal{E}_{\text{exact}}(\psi, \psi')$, where all elements must match:

$$\mathcal{E}_{\text{exact}}(\psi, \psi') = \mathbb{I}[(p \equiv p') \wedge (\tau \equiv \tau') \wedge (\kappa \equiv \kappa')] \quad (4)$$

Approximate Class Matching Node constraints involving value shapes (e.g., class constraints) can be matched approximately due to the following reasons. First, upper ontologies in large KGs are often noisy and redundant. Substituting a class with a similar one may preserve semantics. Second, classes used in ground truth often represent broad coverage, not exhaustive correctness. Third, ShEx allows flexibility using the EXTRA keyword, which tolerates constraints beyond negative definitions. Thus, a generated constraint can be considered approximately equivalent to the ground truth if both involve value shapes and meet either of the following conditions: (1) the class specified in the ground truth constraint is a subclass of the one(s) in the generated constraint, or (2) in the WES dataset, the class used in the generated constraint corresponds to the value-type constraint defined for the predicate. We formalize the criteria as follows:

$$\mathcal{E}_{\text{subclass}}(\tau, \tau') = \begin{cases} 1, & \exists c \in \mathcal{C}(\tau), \exists c' \in \mathcal{C}(\tau'), c \sqsubseteq c' \\ 1, & \exists c \in \mathcal{C}(p), \exists c' \in \mathcal{C}(\tau'), c \sqsubseteq c' \\ 0, & \text{otherwise} \end{cases} \quad (5)$$

$$\mathcal{E}_{\text{subclass}}(\psi, \psi') = \mathbb{I}[(p \equiv p') \wedge \mathcal{E}_{\text{subclass}}(\tau, \tau') \wedge (\kappa \equiv \kappa')] \quad (6)$$

Here $\mathcal{C}(\tau)$ represents the set of classes defined in the node constraint τ , and $\mathcal{C}(p)$ refers to the list of value-type constraints for the predicate p from Wikidata.

Datatype Matching For node constraints, we further define a relaxed criterion based on datatype compatibility. This criteria defines if the datatype of the node constraint matches, while requiring exact matches for the predicate and cardinality:

$$\mathcal{E}_{\text{datatype}}(\psi, \psi') = \mathbb{I}[(p \equiv p') \wedge (d(\tau) \equiv d(\tau')) \wedge (\kappa \equiv \kappa')] \quad (7)$$

where $d(\tau)$ extracts the datatype from the node constraint. Datatypes within the dataset's scope are converted into four general categories: xsd:dateTime, xsd:decimal, xsd:string, and

IRI. The list of datatypes is extensible and can be redefined depending on the characteristics of the input KGs.

Loosened Cardinality Besides node constraints, we relax the evaluation by allowing broader matches on cardinality. This is particularly useful when the exact cardinality range is less critical, and a looser bound—such as simply requiring optional presence—is sufficient for validation:

$$\mathcal{E}(\kappa, \kappa') = \begin{cases} 1, & 0 \leq n' \leq n \leq m \leq m' \\ 0, & \text{otherwise} \end{cases} \quad (8)$$

$$\mathcal{E}_{\text{cardinality}}(\psi, \psi') = \mathbb{I}[(p \equiv p') \wedge (\tau \equiv \tau') \wedge \mathcal{E}(\kappa, \kappa')] \quad (9)$$

Combined Relaxations The relaxation strategies described above can also be combined to accommodate a wider range of scenarios:

$$\mathcal{E}_{\text{subclass+cardinality}}(\psi, \psi') = \mathbb{I}[(p \equiv p') \wedge \mathcal{E}_{\text{subclass}}(\tau, \tau') \wedge \mathcal{E}(\kappa, \kappa')] \quad (10)$$

$$\mathcal{E}_{\text{datatype+cardinality}}(\psi, \psi') = \mathbb{I}[(p \equiv p') \wedge (d(\tau) \equiv d(\tau')) \wedge \mathcal{E}(\kappa, \kappa')] \quad (11)$$

Give a generated schema and a ground truth schema, we report macro-averaged precision, recall, and F1-score, defined by:

$$\begin{aligned} P &= \frac{|\{\psi \mid \mathcal{E}(\psi, \psi') = 1, \psi \in \Psi\}|}{|\Psi|} \\ R &= \frac{|\{\psi \mid \mathcal{E}(\psi, \psi') = 1, \psi \in \Psi\}|}{|\Psi'|} \\ F &= 2 \cdot (P \cdot R) / (P + R) \end{aligned} \quad (12)$$

5 Experimental Setup

5.1 Models

To evaluate the capabilities of LLMs in shape generation and the effectiveness of different information types, we compare against several baseline models across multiple input settings. Our primary non-ML baseline is sheXer (Fernandez-Álvarez et al., 2022), a well-established system for KG shape extraction that relies on graph structure mining and directly takes RDF sources as input. We also include RDFShapeInduction (Mihindukulasooriya et al., 2018), a feature learning and machine learning (ML)-based approach, specifically to benchmark performance on cardinality prediction. For this baseline, we re-implemented the model using

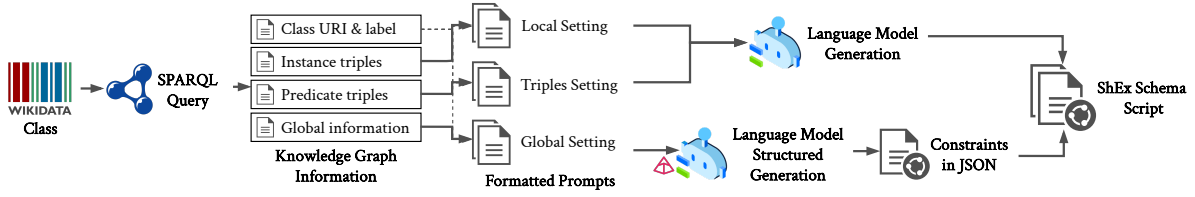


Figure 1: Experimental setup. In the local and triples settings, LLMs generate ShEx schema scripts end-to-end. In the global setting, LLMs first generate constraints in JSON format using their structured generation ability, which are then formulated into ShEx schema scripts.

Information Type	Local	Global	Triples
Input			
Class URI & label	✓	✓	✓
Class description		✓	
Predicate URI & label		✓	
Predicate description*		✓	
Instance triples	✓		
Predicate triples		✓	✓
Predicate frequencies		✓	
Datatype of objects		✓	
Cardinality distributions		✓	
Wikidata constraints*		✓	
Output			
Full ShEx schema script	✓		✓
Formatted constraints		✓	

Table 2: Prompt engineering settings. Information types marked with ‘*’ are primarily available for WES dataset.

the provided feature lists and replaced its classification components with more advanced models, including random forests (Breiman, 2001) and gradient boosting (Friedman, 2001).

For LLM-based comparisons, we selected GPT-4o mini (Hurst et al., 2024) and DeepSeek-V3 (DeepSeek-AI, 2024) as backbone models. Figure 1 illustrates the workflow of our LLM-based generation pipeline. Given a target class, the pipeline retrieves the required information according to the chosen setting—including triples and global context—via SPARQL queries over the endpoints associated with the KG. To ensure scalability to large KGs, the approach first samples the subgraph of the target class to optimize information extraction. Across all settings, the input prompt lengths remain relatively stable regardless of KG size, and the number of retrieved triples is constrained to a manageable limit (up to 5).

5.2 Prompt Engineering

To investigate how LLMs can effectively generate ShEx schemas and understand the impact of information type and generation mode, we designed experiments incorporating different types of infor-

mation extracted from a KG. Based on the nature of the information and existing KG-based prompt construction methods (Wen et al., 2024; Zhang et al., 2025b), we proposed three distinct few-shot prompt settings, summarized in Table 2.

Local Setting This setting provides LLMs with a number of representative instance examples of the target class, along with their associated triples (typically 1-hop neighbors) from KGs. Instances are selected based on specific criteria. Wikidata entities are sorted by ID length and numeric value of their ID. The popularity and importance of entities is loosely correlated with their IDs, as important or fundamental concepts were added early in Wikidata’s history. YAGO entities are sorted by the number of distinct predicates associated with them, prioritizing those with richer connections. The rationale is to allow LLMs to infer patterns directly from how entities of that class are described in the KG. In this setting, LLMs are asked to directly generate the complete ShEx schema, guided by few-shot examples of schemas.

Global Setting This setting focuses on aggregated, schema-level information about the target class and its predicates. As shown in Table 2, input comprises: (a) basic metadata (URIs, labels, and descriptions for the class and relevant predicates); (b) predicate usage statistics (predicate frequency and cardinality distributions); (c) datatype information (datatypes for predicate objects, and class distributions for objects that are URIs, which aids in identifying potential referenced shapes); (d) representative triple examples associated with the predicates; and (e) any available KG-specific constraints (for Wikidata, this includes predefined predicate constraints such as value-type³ (range) and subject-type⁴ (domain)). This setting aims to provide LLMs with a comprehensive, high-level

³<https://www.wikidata.org/wiki/Q21510865>

⁴<https://www.wikidata.org/wiki/Q21503250>

summary of the class’s structure and the characteristics of its predicates. LLMs are then tasked with generating these constraints in a structured JSON format, which is subsequently converted programmatically into ShEx. The structured generation process is detailed in Section 5.3.

Triples Setting Inspired by findings that representative samples can suffice for shape extraction (Fernandez-Álvarez et al., 2022), this setting provides LLMs with a set of triples focused on the usage of specific predicates relevant to the entities. Different from the local setting, this setting focuses on a set of predicates and provides example triples where those predicates appear, potentially sampled across many different instances. The goal is to highlight common patterns for individual predicates. Similar to the local setting, the LLM generates the ShEx schema with few-shot examples.

5.3 Structured Generation

Structured generation enables LLMs to produce outputs adhering to precise formats, vital for applications like code generation (Ugare et al., 2024) and tool calling (Zhang et al., 2024). This involves generating token sequences that satisfy specified constraints, often defined by regular expressions (regex) or context-free grammars (CFGs). A key challenge is efficiently applying these constraints over the LLM’s large vocabulary without degrading speed or performance (Dong et al., 2024; Koo et al., 2024). Even though LLMs excel at structured generation with well-defined JSON schemas, directly applying the full ShEx syntax to the LLM’s constraint decoding process is still challenging due to its complexity. To address this, our structured generation pipeline first simplifies the ShEx syntax within the context of the benchmark by converting it into a more manageable, JSON schema-like representation using ShExJ syntax. This process generates Pydantic models, each corresponding to a segment of ShEx syntax, to guide LLM constraint generation. We leverage the Instructor (Liu and Contributors, 2024) to enhance the LLM’s structured generation capabilities.

We adopt a decomposed, two-step structured generation workflow specifically for the global setting (see Table 2), where the LLM processes global information from KGs. The first step is cardinality prediction. Given global information for a specific predicate relevant to the target class, the LLM is prompted to predict its cardinality, comprising the

minimum and maximum occurrence values. The second step is node constraint prediction, applying on the predicates accepted by the previous step. Based on the global information, the LLM predicts the node constraint for the predicate’s objects. If a datatype is evident from the input information, the LLM outputs this datatype. If the predicate’s objects are instances of another specific classes, the LLM outputs the URIs of these referenced classes, forming the basis for a ShEx shape reference. If the predicate’s objects are restricted to a fixed list of literal values, the LLM generates this complete list. Finally, if none of these conditions are strongly indicated, the LLM defaults to a general node kind constraint, which for our current datasets is typically fixed as IRI.

6 Results

Table 3 presents a comparative performance analysis of the baseline models on YAGOS and WES datasets. On YAGOS, GPT-4o mini achieved the highest F1 (0.591) in its triples setting, while DeepSeek-V3’s global setting showed the best structural similarity (lowest NGED of 0.295). YAGOS’s generally stronger results are attributed to its ontology being largely derived and refined from schema.org, along with its smaller scale—having only half the average number of constraints and 80% of the average number of instances per class compared to WES. For the more challenging WES dataset with a complex predicate vocabulary and twice the constraints per class, DeepSeek-V3 in the global setting yielded the highest F1 (0.318). Both GPT-4o mini and DeepSeek-V3 performed well on the NGED score in their global setting. Compared to the non-ML baseline sheXer, these results underscore the potent schema generation capabilities of LLM-based approaches.

We further analyze the performance of models under different matching criteria, using DeepSeek-V3 on the WES dataset as an example (Table 4). Full results for other models are in Appendix C. In general, loosening the matching criteria leads to improved evaluation scores, as expected. Notably, when combining datatype abstraction with loosened cardinality, DeepSeek-V3 in the global setting achieved impressive F1 scores (0.839). These results suggest that LLMs can generate schemas for certain practical validation scenarios, particularly where the primary goal is to ensure predicate completeness and general object constraints.

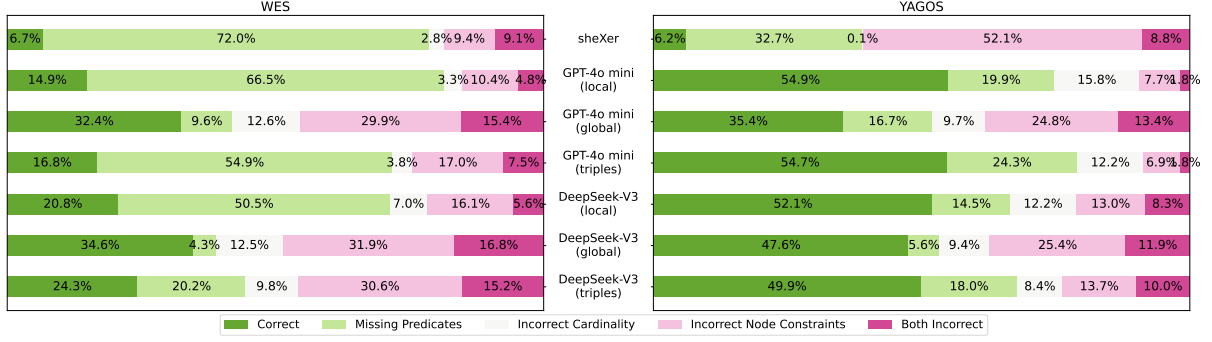


Figure 2: Error distribution across models and settings on WES (left) and YAGOS (right) datasets. The figure shows five categories: four error types and correctly generated constraints.

Models	Settings	YAGOS					WES				
		P	R	F1	GED	NGED	P	R	F1	GED	NGED
sheXer	/	0.111	0.081	0.092	34.08	0.581	0.106	0.099	0.096	90.30	0.833
GPT-4o mini	local	0.575	0.550	0.559	16.61	0.308	0.358	0.224	0.264	82.14	0.697
	global	0.421	0.362	0.388	21.03	0.366	0.306	0.328	0.312	54.44	0.484
	triples	0.631	0.564	0.591	18.36	0.344	0.328	0.196	0.237	72.44	0.577
DeepSeek-V3	local	0.536	0.524	0.526	18.56	0.348	0.322	0.240	0.263	66.16	0.585
	global	0.478	0.484	0.479	16.83	0.295	0.304	0.343	0.318	57.36	0.494
	triples	0.535	0.505	0.510	21.89	0.468	0.269	0.277	0.269	55.80	0.488

Table 3: Results of LLMs comparing with baseline models across different settings based on exact matching criteria. Five entities and their related triples are retrieved and feed into LLMs. Note that for YAGO Schema, entities of triple examples are sorted by predicate count, and for WES, entities are sorted by their IDs. The highest scores are set in **bold**.

Settings	Matching Criteria		WES		
	Node Constraint	Cardinality	P	R	F1
local	Exact	Exact	0.322	0.240	0.263
	Subclass	Exact	0.413	0.303	0.335
	Datatype	Exact	0.537	0.396	0.437
	Exact	Loosened	0.366	0.280	0.303
	Subclass	Loosened	0.466	0.352	0.384
	Datatype	Loosened	0.613	0.464	0.507
global	Exact	Exact	0.304	0.343	0.318
	Subclass	Exact	0.482	0.550	0.507
	Datatype	Exact	0.577	0.655	0.606
	Exact	Loosened	0.394	0.451	0.415
	Subclass	Loosened	0.638	0.738	0.676
	Datatype	Loosened	0.793	0.910	0.839
triples	Exact	Exact	0.269	0.277	0.269
	Subclass	Exact	0.371	0.377	0.367
	Datatype	Exact	0.536	0.549	0.534
	Exact	Loosened	0.320	0.334	0.322
	Subclass	Loosened	0.449	0.461	0.448
	Datatype	Loosened	0.675	0.699	0.677

Table 4: Results of DeepSeek-V3 on the WES dataset across different matching criteria.

As for individual criteria, allowing datatype abstraction provides the most significant performance gain, especially in the global setting. This indicates LLMs can effectively process given datatype information, finding it less challenging than inferring cardinality. Allowing approximate subclass matching improved classification scores more for

WES compared with YAGOS, suggesting predicting the precise class list for referenced shapes is harder for WES, while generating related object classes is not. These findings indicate that although LLM-generated schemas may benefit from refinement, the required adjustments are generally less intensive or domain-specific than those for existing automated approaches, while still supporting effective validation.

Figure 2 shows the distribution of results across five categories (correct and four error types) for models and settings on WES and YAGOS. Error types are: (1) missing predicates, (2) incorrect cardinality, (3) incorrect node constraint, and (4) both cardinality and node constraint are incorrect. Global settings effectively reduce missing predicates for models on both datasets. Missing predicates are more frequent on WES than YAGOS, likely due to WES’s larger set of candidate predicates. However, global settings show a higher rate of cases where both cardinality and node constraint are incorrect, especially on the WES dataset.

Specifically for cardinality prediction, we compare LLM results with RDFShapeInduction using different ML models based on the accuracy of the

Model	WES		
	Acc (min)	Acc (max)	Acc
sheXer	0.982	0.402	0.394
GPT-4o mini (local)	0.988	0.441	0.432
GPT-4o mini (global)	0.900	0.584	0.520
GPT-4o mini (triples)	0.988	0.467	0.457
DeepSeek-V3 (local)	0.975	0.490	0.474
DeepSeek-V3 (global)	0.978	0.499	0.485
DeepSeek-V3 (triples)	0.933	0.547	0.517
RDFShapeInduction (DT)	0.992	0.566	0.559
RDFShapeInduction (MLP)	0.988	0.606	0.596
RDFShapeInduction (RF)	0.994	0.645	0.640
RDFShapeInduction (GB)	0.993	0.662	0.656

Table 5: Average accuracy of cardinalities (minimum, maximum, and combined) for 10 sampled schemas, covering 612 candidate constraints from the WES dataset. RDFShapeInduction was evaluated with different classification models, including decision trees (DT), multi-layer perceptron (MLP), random forests (RF), and gradient boosting (GB).

lower limit, upper limit, and their combination (Table 5). ML models (e.g., gradient boosting) surpass LLMs, indicating that LLMs still have room for improvement on this subtask. To leverage this, we replace cardinality prediction components with ML models in LLM-based generation pipelines under the global setting. The results on the WES dataset are shown in Table 6, where LLMs augmented with ML models outperform LLMs alone. Similar improvements are observed on the YAGOS dataset (Tables 9 and 10, Appendix C). The strong performance of traditional ML models suggests that incorporating features learned from these models into LLM inputs is a promising strategy for improving end-to-end schema generation.

7 Conclusion

The persistent challenge of ensuring data quality in large KGs necessitates effective and automated methods for generating validation schemas. This paper explored the application of LLMs to this task, introducing the first benchmark for ShEx schema generation from KGs. This benchmark, comprising two datasets and custom metrics, enabled a thorough assessment revealing LLMs’s ability to produce high-quality shape schemas. Beyond its direct contributions to Semantic Web practices, this work provides a new benchmark for evaluating the nuanced graph understanding and structured generation capabilities of LLMs. Future research could aim to broaden the benchmark’s scope by incorporating more ShEx features—such as additional

Model	WES				
	P	R	F1	GED	NGED
sheXer	0.083	0.071	0.069	116.40	0.849
GPT-4o mini (local)	0.329	0.139	0.188	104.60	0.704
GPT-4o mini (global)	0.239	0.290	0.259	80.50	0.609
GPT-4o mini (dt-global)	0.273	0.259	0.264	66.50	0.474
GPT-4o mini (gb-global)	0.283	0.327	0.301	67.30	0.500
GPT-4o mini (triples)	0.301	0.170	0.205	96.10	0.611
DeepSeek-V3 (local)	0.285	0.192	0.221	86.10	0.586
DeepSeek-V3 (global)	0.232	0.290	0.255	84.70	0.637
DeepSeek-V3 (dt-global)	0.284	0.272	0.275	61.70	0.441
DeepSeek-V3 (gb-global)	0.302	0.350	0.321	60.60	0.456
DeepSeek-V3 (triples)	0.254	0.250	0.244	81.50	0.601

Table 6: Results of hybrid approaches compared to baseline models across different settings, based on 10 sampled schemas covering 612 candidate constraints from the WES dataset using exact matching criteria. “dt_global” refers to incorporating decision trees-based cardinality prediction into the global setting pipeline, while “gb_global” does the same using gradient boosting.

constraint types (e.g., string and numeric facet constraints) and structural elements (e.g., imported schemas and annotations)—as well as by including diverse schema languages like PG-Schema. Moreover, advancing the models, particularly by enhancing their structured generation ability for complex schemas will be crucial.

Limitations

The size of the dataset is constrained by limited resources available for its construction. Creating high-quality shape schemas is a time-consuming process that requires multiple rounds of refinement by knowledge engineers. The limited size of the dataset makes it difficult to apply traditional ML approaches to the complete task.

Another limitation lies in the design of our evaluation metrics. While they capture the most essential features of schemas, they simplify schema structures and are therefore insufficient for handling more complex, deeply nested shapes. For example, our current formulation does not differentiate edit costs across specific constraint variations (e.g., predicates differing by the presence of the EXTRA keyword) and imported schemas.

Finally, our evaluation of LLMs is limited to those accessible via APIs. Models that can be run locally, such as those with 8B or 14B parameters, are not included. Preliminary experiments indicate that these locally runnable LLMs struggle to generate valid ShEx schemas and lack sufficient structured generation capabilities to reliably produce correct constraint syntax.

Ethics Statement

The YAGOS and WES datasets are developed with a strong commitment to ethical AI principles. They contain no personal, sensitive, or identifiable information and are free from harmful, offensive, or misleading content. Both datasets strictly comply with responsible AI guidelines.

Acknowledgments

This work was supported by the NMES Enterprise & Engagement Partnerships fund. Co-funded by the SIEMENS AG and the Technical University of Munich, Institute for Advanced Study, Germany.

References

- Oshin Agarwal, Heming Ge, Siamak Shakeri, and Rami Al-Rfou. 2021. [Knowledge Graph Based Synthetic Corpus Generation for Knowledge-Enhanced Language Model Pre-training](#). In *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 3554–3565, Online. Association for Computational Linguistics.
- Shqiponja Ahmetaj, Iovka Boneva, Jan Hidders, Katja Hose, Maxime Jakubowski, Jose Emilio Labra Gayo, Wim Martens, Fabio Mogavero, Filip Murlak, Cem Okulmus, Axel Polleres, Ognjen Savković, Mantas Šimkus, and Dominik Tomaszuk. 2025. [Common Foundations for SHACL, ShEx, and PG-Schema](#). In *Proceedings of the ACM on Web Conference 2025, WWW '25*, page 8–21, New York, NY, USA. Association for Computing Machinery.
- Daria Ammalainen. 2023. [Wikidata Ontology Issues — Suggestions for prioritisation 2023](#). Accessed: March 12, 2025.
- Renzo Angles, Angela Bonifati, Stefania Dumbrava, George Fletcher, Alastair Green, Jan Hidders, Bei Li, Leonid Libkin, Victor Marsault, Wim Martens, Filip Murlak, Stefan Plantikow, Ognjen Savkovic, Michael Schmidt, Juan Sequeda, Slawek Staworko, Dominik Tomaszuk, Hannes Voigt, Domagoj Vrgoc, and 2 others. 2023. [PG-Schema: Schemas for Property Graphs](#). *Proc. ACM Manag. Data*, 1(2).
- Sören Auer, Christian Bizer, Georgi Kobilarov, Jens Lehmann, Richard Cyganiak, and Zachary Ives. 2007. [DBpedia: a nucleus for a web of open data](#). In *Proceedings of the 6th International The Semantic Web and 2nd Asian Conference on Asian Semantic Web Conference, ISWC'07/ASWC'07*, page 722–735, Berlin, Heidelberg. Springer-Verlag.
- Leo Breiman. 2001. [Random forests](#). *Mach. Learn.*, 45(1):5–32.
- Tom B. Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, Sandhini Agarwal, Ariel Herbert-Voss, Gretchen Krueger, Tom Henighan, Rewon Child, Aditya Ramesh, Daniel M. Ziegler, Jeffrey Wu, Clemens Winter, and 12 others. 2020. [Language Models are Few-Shot Learners](#). *CoRR*, abs/2005.14165.
- Wenhu Chen, Yu Su, Xifeng Yan, and William Yang Wang. 2020. [KGPT: Knowledge-Grounded Pre-Training for Data-to-Text Generation](#). In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 8635–8648, Online. Association for Computational Linguistics.
- Richard Cyganiak, David Wood, and Markus Lanthaler. 2014. [RDF 1.1 Concepts and Abstract Syntax](#). W3c recommendation, W3C.
- DeepSeek-AI. 2024. [DeepSeek-V3 Technical Report](#). Preprint, arXiv:2412.19437.
- Yixin Dong, Charlie F. Ruan, Yaxing Cai, Ruihang Lai, Ziyi Xu, Yilong Zhao, and Tianqi Chen. 2024. [XGrammar: Flexible and Efficient Structured Generation Engine for Large Language Models](#). *CoRR*, abs/2411.15100.
- Darren Edge, Ha Trinh, Newman Cheng, Joshua Bradley, Alex Chao, Apurva Mody, Steven Truitt, Dasha Metropolitansky, Robert Osazuwa Ness, and Jonathan Larson. 2024. [From Local to Global: A Graph RAG Approach to Query-Focused Summarization](#).
- Jinyuan Fang, Zaiqiao Meng, and Craig Macdonald. 2024. [REANO: Optimising Retrieval-Augmented Reader Models through Knowledge Graph Generation](#). In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 2094–2112, Bangkok, Thailand. Association for Computational Linguistics.
- Daniel Fernandez-Álvarez, Jose Emilio Labra-Gayo, and Daniel Gayo-Avello. 2022. [Automatic extraction of shapes using sheXer](#). *Knowledge-Based Systems*, 238:107975.
- Jerome H Friedman. 2001. Greedy Function Approximation: A Gradient Boosting Machine. *Annals of statistics*, pages 1189–1232.
- Jose Emilio Labra Gayo. 2022. [WShEx: A language to describe and validate Wikibase entities](#). Preprint, arXiv:2208.02697.
- Jose Emilio Labra Gayo, Eric Prud'hommeaux, Iovka Boneva, and Dimitris Kontokostas. 2018. [Validating RDF Data](#). Springer International Publishing.

- Xiaoxin He, Yijun Tian, Yifei Sun, Nitesh V Chawla, Thomas Laurent, Yann LeCun, Xavier Bresson, and Bryan Hooi. 2024a. [G-Retriever: Retrieval-Augmented Generation for Textual Graph Understanding and Question Answering](#). In [The Thirty-eighth Annual Conference on Neural Information Processing Systems](#).
- Yuan He, Zhangdie Yuan, Jiaoyan Chen, and Ian Horrocks. 2024b. [Language Models as Hierarchy Encoders](#). In [Advances in Neural Information Processing Systems](#), volume 37, pages 14690–14711. Curran Associates, Inc.
- Aidan Hogan, Eva Blomqvist, Michael Cochez, Claudia D’amato, Gerard De Melo, Claudio Gutierrez, Sabrina Kirrane, José Emilio Labra Gayo, Roberto Navigli, Sebastian Neumaier, Axel-Cyrille Ngonga Ngomo, Axel Polleres, Sabbir M. Rashid, Anisa Rula, Lukas Schmelzeisen, Juan Sequeda, Steffen Staab, and Antoine Zimmermann. 2021. [Knowledge Graphs](#). [ACM Comput. Surv.](#), 54(4).
- Aidan Hogan, Xin Luna Dong, Denny Vrandečić, and Gerhard Weikum. 2025. [Large Language Models, Knowledge Graphs and Search Engines: A Crossroads for Answering Users’ Questions](#). [Preprint](#), arXiv:2501.06699.
- Yuntong Hu, Zhihan Lei, Zheng Zhang, Bo Pan, Chen Ling, and Liang Zhao. 2025. [GRAG: Graph Retrieval-Augmented Generation](#). In [Findings of the Association for Computational Linguistics: NAACL 2025](#), pages 4145–4157, Albuquerque, New Mexico. Association for Computational Linguistics.
- OpenAI: Aaron Hurst, Adam Lerer, Adam P. Goucher, Adam Perelman, Aditya Ramesh, Aidan Clark, AJ Ostrow, Akila Welihinda, Alan Hayes, Alec Radford, Aleksander Mądry, Alex Baker-Whitcomb, Alex Beutel, Alex Borzunov, Alex Carney, Alex Chow, Alex Kirillov, Alex Nichol, Alex Paino, and 399 others. 2024. [GPT-4o System Card](#). [Preprint](#), arXiv:2410.21276.
- Bowen Jin, Gang Liu, Chi Han, Meng Jiang, Heng Ji, and Jiawei Han. 2024. [Large Language Models on Graphs: A Comprehensive Survey](#). [IEEE Transactions on Knowledge and Data Engineering](#), 36(12):8622–8642.
- Holger Knublauch and Dimitris Kontokostas. 2017. [Shapes Constraint Language \(SHACL\)](#). W3c recommendation.
- Terry Koo, Frederick Liu, and Luheng He. 2024. [Automata-based constraints for language model decoding](#). In [First Conference on Language Modeling](#).
- Xin Li, Dongze Lian, Zhihe Lu, Jiawang Bai, Zhibo Chen, and Xinchao Wang. 2023. [GraphAdapter: Tuning Vision-Language Models With Dual Knowledge Graph](#). In [Advances in Neural Information Processing Systems](#), volume 36, pages 13448–13466. Curran Associates, Inc.
- Jason Liu and Contributors. 2024. [Instructor: A library for structured outputs from large language models](#).
- Andy Lo, Albert Q. Jiang, Wenda Li, and Mateja Jamnik. 2024. [End-to-End Ontology Learning with Large Language Models](#). In [The Thirty-eighth Annual Conference on Neural Information Processing Systems](#).
- Nandana Mihindukulasooriya, Mohammad Rifat Ahmad Rashid, Giuseppe Rizzo, Raúl García-Castro, Oscar Corcho, and Marco Torchiano. 2018. [RDF shape induction using knowledge base profiling](#). In [Proceedings of the 33rd Annual ACM Symposium on Applied Computing, SAC ’18](#), page 1952–1959, New York, NY, USA. Association for Computing Machinery.
- Jeff Z. Pan, Simon Razniewski, Jan-Christoph Kalo, Sneha Singhanian, Jiaoyan Chen, Stefan Dietze, Hajira Jabeen, Janna Omeliyanenko, Wen Zhang, Matteo Lissandrini, Russa Biswas, Gerard de Melo, Angela Bonifati, Edlira Vakaj, Mauro Dragoni, and Damien Graux. 2023. [Large Language Models and Knowledge Graphs: Opportunities and Challenges](#). [Transactions on Graph Data and Knowledge](#), 1(1):2:1–2:38.
- Xuran Pan, Tianzhu Ye, Dongchen Han, Shiji Song, and Gao Huang. 2022. [Contrastive Language-Image Pre-Training with Knowledge Graphs](#). In [Advances in Neural Information Processing Systems](#), volume 35, pages 22895–22910. Curran Associates, Inc.
- Eric Prud’hommeaux, Jose Emilio Labra Gayo, and Harold Solbrig. 2014. [Shape expressions: an RDF validation and transformation language](#). In [Proceedings of the 10th International Conference on Semantic Systems, SEM ’14](#), page 32–40, New York, NY, USA. Association for Computing Machinery.
- Kashif Rabbani, Matteo Lissandrini, and Katja Hose. 2022. [SHACL and ShEx in the Wild: A Community Survey on Validating Shapes Generation and Adoption](#). In [Companion Proceedings of the Web Conference 2022, WWW ’22](#), page 260–263, New York, NY, USA. Association for Computing Machinery.
- Kashif Rabbani, Matteo Lissandrini, and Katja Hose. 2023. [Extraction of Validating Shapes from Very Large Knowledge Graphs](#). [Proc. VLDB Endow.](#), 16(5):1023–1032.
- Sherif Sakr, Angela Bonifati, Hannes Voigt, Alexandru Iosup, Khaled Ammar, Renzo Angles, Walid Aref, Marcelo Arenas, Maciej Besta, Peter A. Boncz, Khuzaima Daudjee, Emanuele Della Valle, Stefania Dumbrava, Olaf Hartig, Bernhard Haslhofer, Tim Hegeman, Jan Hidders, Katja Hose, Adriana Iamnitchi, and 22 others. 2021. [The future is big graphs: a community view on graph processing systems](#). [Commun. ACM](#), 64(9):62–71.
- Alberto Sanfeliu and King-Sun Fu. 1983. [A Distance Measure Between Attributed Relational Graphs for](#)

- [Pattern Recognition](#). *IEEE Transactions on Systems, Man, and Cybernetics, SMC-13(3)*:353–362.
- Ansgar Scherpp, Gerd Groener, Petr Škoda, Katja Hose, and Maria-Esther Vidal. 2024. [Semantic Web: Past, Present, and Future](#). *Transactions on Graph Data and Knowledge*, 2(1):3:1–3:37.
- Kartik Shenoy, Filip Ilievski, Daniel Garijo, Daniel Schwabe, and Pedro Szekely. 2022. [A study of the quality of Wikidata](#). *Journal of Web Semantics*, 72:100679.
- Fabian M. Suchanek, Mehwish Alam, Thomas Bonald, Lihu Chen, Pierre-Henri Paris, and Jules Soria. 2024. [YAGO 4.5: A Large and Clean Knowledge Base with a Rich Taxonomy](#). In *Proceedings of the 47th International ACM SIGIR Conference on Research and Development in Information Retrieval, SIGIR '24*, page 131–140, New York, NY, USA. Association for Computing Machinery.
- Jiashuo Sun, Chengjin Xu, Luminyuan Tang, Saizhuo Wang, Chen Lin, Yeyun Gong, Heung-Yeung Shum, and Jian Guo. 2023. [Think-on-Graph: Deep and Responsible Reasoning of Large Language Model with Knowledge Graph](#). Preprint, arXiv:2307.07697.
- Jiabin Tang, Yuhao Yang, Wei Wei, Lei Shi, Lixin Su, Suqi Cheng, Dawei Yin, and Chao Huang. 2024. [GraphGPT: Graph Instruction Tuning for Large Language Models](#). In *Proceedings of the 47th International ACM SIGIR Conference on Research and Development in Information Retrieval, SIGIR '24*, page 491–500, New York, NY, USA. Association for Computing Machinery.
- Shubham Ugare, Tarun Suresh, Hangoo Kang, Sasa Misailovic, and Gagandeep Singh. 2024. [Improving LLM Code Generation with Grammar Augmentation](#). CoRR, abs/2403.01632.
- Denny Vrandečić and Markus Krötzsch. 2014. [Wikidata: a free collaborative knowledge base](#). *Commun. ACM*, 57(10):78–85.
- Heng Wang, Shangbin Feng, Tianxing He, Zhaoxuan Tan, Xiaochuang Han, and Yulia Tsvetkov. 2023. Can language models solve graph problems in natural language? In *Proceedings of the 37th International Conference on Neural Information Processing Systems, NIPS '23*, Red Hook, NY, USA. Curran Associates Inc.
- Yilin Wen, Zifeng Wang, and Jimeng Sun. 2024. [MindMap: Knowledge Graph Prompting Sparks Graph of Thoughts in Large Language Models](#). In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics*.
- Zhentao Xu, Mark Jerome Cruz, Matthew Guevara, Tie Wang, Manasi Deshpande, Xiaofeng Wang, and Zheng Li. 2024. [Retrieval-Augmented Generation with Knowledge Graphs for Customer Service Question Answering](#). In *Proceedings of the 47th International ACM SIGIR Conference on Research and Development in Information Retrieval, SIGIR '24*, page 2905–2909, New York, NY, USA. Association for Computing Machinery.
- Michihiro Yasunaga, Antoine Bosselut, Hongyu Ren, Xikun Zhang, Christopher D Manning, Percy S Liang, and Jure Leskovec. 2022. [Deep Bidirectional Language-Knowledge Graph Pretraining](#). In *Advances in Neural Information Processing Systems*, volume 35, pages 37309–37323. Curran Associates, Inc.
- Bohui Zhang, Valentina Anita Carriero, Katrin Schreiberhuber, Stefani Tsaneva, Lucía Sánchez González, Jongmo Kim, and Jacopo de Berardinis. 2025a. [OntoChat: A Framework for Conversational Ontology Engineering Using Language Models](#). In *The Semantic Web: ESWC 2024 Satellite Events: Hersonissos, Crete, Greece, May 26–30, 2024, Proceedings, Part I*, page 102–121, Berlin, Heidelberg. Springer-Verlag.
- Jiawei Zhang. 2023. [Graph-ToolFormer: To Empower LLMs with Graph Reasoning Ability via Prompt Augmented by ChatGPT](#). ArXiv, abs/2304.11116.
- K. Zhang and D. Shasha. 1989. [Simple fast algorithms for the editing distance between trees and related problems](#). *SIAM J. Comput.*, 18(6):1245–1262.
- Kexun Zhang, Hongqiao Chen, Lei Li, and William Wang. 2024. [Don't Fine-Tune, Decode: Syntax Error-Free Tool Use via Constrained Decoding](#). Preprint, arXiv:2310.07075.
- Qinggang Zhang, Junnan Dong, Hao Chen, Daochen Zha, Zailiang Yu, and Xiao Huang. 2025b. [KnowGPT: knowledge graph based prompting for large language models](#). In *Proceedings of the 38th International Conference on Neural Information Processing Systems, NIPS '24*, Red Hook, NY, USA. Curran Associates Inc.

A Dataset Construction

The YAGOS dataset was constructed using YAGO 4.5 (Suchanek et al., 2024) as the input KG. Ground truth ShEx schema construction was informed by the existing SHACL schema and the official YAGO 4.5 design document⁵. While the provided documentation offered a starting point, it was not exhaustive, with only a subset of key properties and constraints defined. Therefore, the final ground truth ShEx scripts were manually crafted by retaining constraints explicitly mentioned in the design document and adding missing predicates relevant to each class. This was achieved by querying the KG to understand common property usage for entities belonging to the target classes and aligning these findings with the schema’s intended scope.

The WES dataset was developed through a semi-automatic process based on the Wikidata (Vrandečić and Krötzsch, 2014) truthy statements RDF dump (from qEndpoint’s Wikidata release 1.16.1⁶), involving knowledge engineers assisted by tools for pattern extraction and data querying. Target classes for the dataset were selected from three main sources:

1. Community-curated Wikidata EntitySchema directory: We drew upon schemas available in the Wikidata EntitySchema directory⁷. Due to the variability in quality and completeness of these community contributions, we manually selected a subset of those deemed relatively high-quality to serve as initial drafts. These selected drafts were then collaboratively reviewed, refined, and standardized by our knowledge engineers to meet consistent quality criteria.
2. Classes mapped from YAGOS: A set of classes was chosen by mapping them from the YAGOS dataset. Since YAGO facts are partially derived from Wikidata, creating schemas for these corresponding classes in WES allows for a comparison of modeling scope between the two KGs.
3. Classes mapped from Wikipedia Categories: Additional classes identified through map-

⁵<https://yago-knowledge.org/data/yago4.5/design-document.pdf>

⁶<https://github.com/the-qa-company/qEndpoint/releases/tag/v1.16.1>

⁷https://www.wikidata.org/wiki/Wikidata:Database_reports/EntitySchema_directory

Property Type	Count	Proportion (%)	Rank
WikibaseItem	1,607	14.30%	2
Quantity	646	5.75%	3
String	322	2.86%	4
Url	99	0.88%	5
Time	63	0.56%	7
Monolingualtext	60	0.53%	8

Table 7: Distribution of primary property datatypes in Wikidata.

pings from relevant Wikipedia Categories⁸, focusing on well-defined concepts suitable for schema modeling.

For each selected class, the following iterative process was undertaken by knowledge engineers:

1. Predicate identification and prioritization: A comprehensive list of predicates associated with entities of the target class was compiled, along with their occurrence frequencies. Based on usage frequency and modeling importance (as detailed in Table 7, which outlines key Wikidata property types considered), a working set of candidate properties was selected.
2. Predicate inclusion and refinement: Candidate predicates were evaluated for inclusion based on their semantic appropriateness (i.e., factual suitability for the class) and through property denoising and deduplication. The latter step involved identifying and consolidating functionally similar or overlapping properties common in Wikidata by selecting the most representative and predominantly used one. For instance, to model an item’s inception, one would choose the most suitable property from options like P580 (start time), P571 (inception), etc.
3. Cardinality determination: The cardinality for each included predicate was established. Predicates were generally assigned an optional cardinality (e.g., minimum 0 and maximum ‘*’ for zero or more, or minimum 0 and maximum 1 for at most one). This default was overridden if high frequency and semantic necessity strongly suggested a mandatory presence (a minimum of 1) or a more specific range.

⁸<https://en.wikipedia.org/wiki/Wikipedia:Contents/Categories>

4. Node constraint specification: The constraints on the object values of each predicate were defined.

- Datatype constraints: If a predicate consistently uses a specific datatype, a direct datatype constraint was applied.
- Value set: If objects were consistently drawn from a small, fixed list of literal values or specific URIs, a value set constraint was used.
- Shape reference: If objects were typically instances of or subclass of a few related classes, a shape reference was created. Knowledge engineers selected the most relevant object class(es) based on observed distributions and schema readability.
- Node kind: If neither a specific datatype, value set, nor a clear referenced shape was appropriate, a general node kind constraint (typically 'IRI') was applied.

To ensure the quality and consistency of datasets, three volunteers, all active researchers and engineers in ontology engineering with prior schema generation experience, were recruited from Wikidata community events. This process was supported by semi-automatic tools that incorporated a series of SPARQL queries to gather relevant statistics and patterns from Wikidata (representative examples of these queries are provided in Listings 1, 2, 3, and 4). Schema annotation involved three experts independently annotating each class's schema. Their annotations were collected, and a constraint was included in the final ground truth ShEx schema if at least two experts independently proposed an equivalent formulation. Discrepancies or cases with less than two-thirds agreement were resolved through discussion among the annotators or, if consensus could not be reached, the contentious constraint was omitted to maintain high confidence in the final dataset.

```
SELECT DISTINCT ?predicate (COUNT(
  DISTINCT ?subject) AS ?count)
WHERE {
  ?subject wdt:P31 wd:Q1248784 ;
           ?predicate ?object .
}
GROUP BY ?predicate
ORDER BY DESC(?count)
```

Listing 1: SPARQL query to retrieve predicates used with instances of Airport (Q1248784) and their usage frequency.

```
SELECT DISTINCT ?predicate ?datatype
WHERE {
  ?subject wdt:P31 wd:Q1248784 ;
           ?predicate ?object .
  BIND (datatype(?object) AS ?datatype)
}
```

Listing 2: SPARQL query to identify datatypes of objects for predicates associated with Airport (Q1248784).

```
SELECT (COUNT(DISTINCT ?subject) AS ?
  count)
WHERE {
  ?subject rdf:type schema:Book .
  FILTER NOT EXISTS {
    ?subject schema:illustrator ?object
  }
}
```

Listing 3: SPARQL query to count instances of schema:Book lacking a schema:illustrator predicate.

```
SELECT ?cardinality (COUNT(DISTINCT ?
  subject) AS ?count)
{
  SELECT DISTINCT ?subject (COUNT(?
    object) AS ?cardinality)
  WHERE {
    ?subject rdf:type schema:Book ;
             schema:illustrator ?object
          .
  }
  GROUP BY ?subject
}
GROUP BY ?cardinality
ORDER BY DESC(?count)
```

Listing 4: SPARQL query to determine the distribution of schema:illustrator predicate occurrences per schema:Book instance (i.e. cardinality distribution).

YAGO 4.5 is licensed under a CC BY-SA license, and Wikidata is licensed under the CC0 license⁹. Our datasets are under the same licenses as the KGs from which they were derived, respectively.

B ShEx Specification

The ShEx was initially proposed in 2014, with its current specification published in 2019. The language continues to evolve to incorporate new functionalities addressing the diverse requirements of KG validation, as evidenced by extensions like WShEx for Wikidata EntitySchemas (Gayo, 2022). ShEx is often preferred for KG validation over traditional ontologies for several reasons. First, while ontologies can perform some validation tasks, their primary design focus is typically on entailment and

⁹<https://www.wikidata.org/wiki/Wikidata:Licensing>

reasoning, which can lead to less expressive or less direct validation capabilities (Mihindukulasooriya et al., 2018). Furthermore, ontologies are not inherently designed to validate specific subsets of focus nodes extracted from a KG in the same granular way ShEx allows. Second, ShEx benefits from coexisting textual (ShExC) and JSON (ShExJ) syntaxes, with readily available tools like shex.js¹⁰ and PyShEx¹¹ for conversion between them. This ShExC-ShExJ interoperability is advantageous in the context of LLMs, which can effectively process and generate JSON-structured data. This makes ShEx a promising candidate for LLM-driven structured data generation and knowledge validation tasks, broadening its adoption and impact.

In our datasets, the ground truth ShEx schemas utilize a simplified yet functional subset of the ShEx language. This approach prioritizes core validation requirements while ensuring the schemas remain human-understandable. An illustrative example, the schema for ‘Museum (Q33506)’ from the WES dataset, is shown in Figure 3.

Node Constraints A key feature adopted in our dataset, and emphasized in the Wikidata EntitySchema initiative, is the use of shape references to create modular and interconnected schemas, often facilitated by IMPORT declarations in more complex scenarios. As shown in the example, the constraint on the predicate ‘country (P17)’ specifies that its object values must belong to the class representing countries. This is achieved using a shape reference, @<Country>, which points to the definition of the ‘Country’ shape (Figure 3a, lines 17-20), ensuring that objects of the ‘country (P17)’ predicate are instances of ‘country (Q6256)’.

Cardinality Cardinalities in ShExC are represented by the strings ‘?’, ‘+’, ‘*’ (following notation similar to the XML specification) and ranges such as {*m*,} to indicate that at least *m* elements are required. In ShExJ, they are represented by ‘min’ and ‘max’ values indicating their lower and upper bounds.

C Experimental Details

LLM-based experiments were conducted by accessing the LLM APIs. The similarity metrics leverage the Zhang-Shasha algorithm, implemented using

¹⁰<https://github.com/shexjs/shex.js>

¹¹<https://github.com/hsolbrig/PyShEx>

Settings	Matching Criteria		YAGOS		
	Node Constraint	Cardinality	P	R	F1
local	Exact	Exact	0.536	0.524	0.526
	Subclass	Exact	0.551	0.540	0.542
	Datatype	Exact	0.634	0.625	0.626
	Exact	Loosened	0.589	0.582	0.580
	Subclass	Loosened	0.604	0.598	0.596
	Datatype	Loosened	0.753	0.751	0.746
global	Exact	Exact	0.434	0.439	0.435
	Subclass	Exact	0.521	0.528	0.523
	Datatype	Exact	0.698	0.709	0.701
	Exact	Loosened	0.441	0.446	0.442
	Subclass	Loosened	0.532	0.541	0.535
	Datatype	Loosened	0.751	0.764	0.755
triples	Exact	Exact	0.535	0.505	0.510
	Subclass	Exact	0.554	0.523	0.528
	Datatype	Exact	0.630	0.591	0.600
	Exact	Loosened	0.581	0.550	0.554
	Subclass	Loosened	0.604	0.570	0.575
	Datatype	Loosened	0.759	0.720	0.725

Table 8: Results of DeepSeek-V3 on the YAGOS dataset across different matching criteria.

Model	YAGOS		
	Acc (min)	Acc (max)	Acc
sheXer	0.961	0.658	0.618
GPT-4o mini (local)	0.934	0.779	0.722
GPT-4o mini (global)	0.892	0.744	0.644
GPT-4o mini (triples)	0.942	0.684	0.658
DeepSeek-V3 (local)	0.880	0.786	0.729
DeepSeek-V3 (global)	0.969	0.768	0.737
DeepSeek-V3 (triples)	0.918	0.713	0.688
RDFShapeInduction (DT)	0.980	0.808	0.788
RDFShapeInduction (MLP)	0.827	0.689	0.579
RDFShapeInduction (RF)	0.989	0.809	0.798
RDFShapeInduction (GB)	0.976	0.797	0.773

Table 9: Average accuracy of cardinalities (minimum, maximum, and combined) for 9 sampled schemas, covering 221 candidate constraints from the YAGOS dataset.

an open-source package¹².

Examples of input prompts for the local, triples, and global settings are provided in Listings 5, 6, 7 for the class ‘film award (Q4220917)’.

Table 8 presents the results of DeepSeek-V3 on the YAGOS dataset across different matching criteria. Table 9 and 10 report the evaluation results for cardinality prediction and the performance of hybrid approaches on the YAGOS dataset. Table 11 and 12 show the performance of sheXer and GPT-4o mini models across six different compositions of matching criteria, evaluating their precision, recall, and F1-score on both datasets. A clear trend observable for GPT-4o mini is the consistent improvement in F1-scores as the matching criteria are relaxed. For instance, in the triples setting on the YAGO dataset, the F1-score climbs from 0.591

¹²<https://pypi.org/project/zss/>

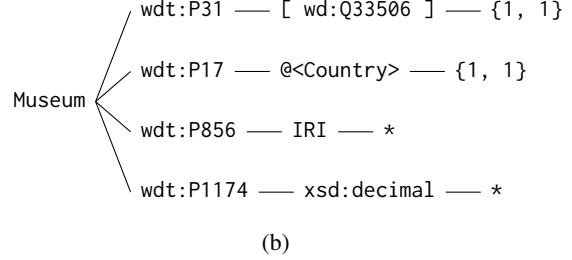
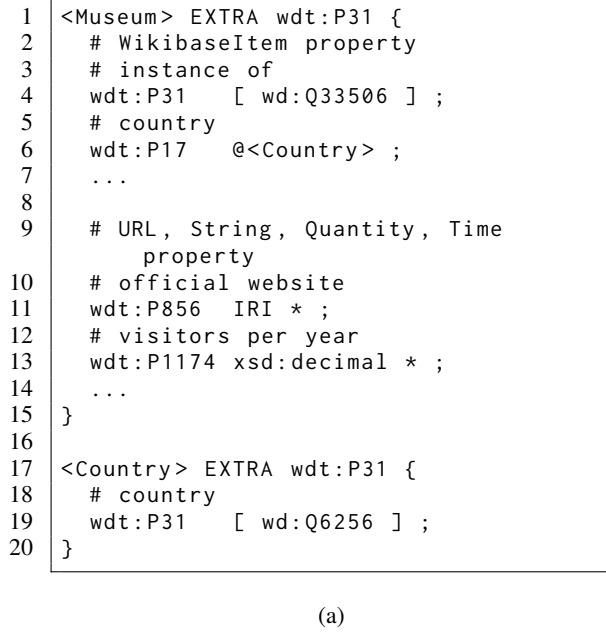


Figure 3: Example ShEx schema fragment for ‘Museum (Q33506)’ from the WES dataset: (a) the ShExC textual representation, where comments above each constraint indicate the label of its predicate, and (b) its corresponding tree structure representation used for similarity metrics.

Model	YAGOS				
	P	R	F1	GED	NGED
sheXer	0.085	0.056	0.066	43.67	0.616
GPT-4o mini (local)	0.669	0.631	0.649	14.89	0.218
GPT-4o mini (global)	0.342	0.294	0.314	27.56	0.398
GPT-4o mini (dt-global)	0.373	0.360	0.366	21.89	0.306
GPT-4o mini (gb-global)	0.356	0.353	0.354	22.33	0.315
GPT-4o mini (triples)	0.667	0.573	0.612	20.89	0.307
DeepSeek-V3 (local)	0.635	0.588	0.608	19.56	0.286
DeepSeek-V3 (global)	0.427	0.432	0.429	19.89	0.281
DeepSeek-V3 (dt-global)	0.466	0.448	0.456	19.22	0.269
DeepSeek-V3 (gb-global)	0.447	0.442	0.444	19.44	0.276
DeepSeek-V3 (triples)	0.603	0.488	0.535	24.11	0.348

Table 10: Results of hybrid approaches compared to baseline models across different settings, based on 9 sampled schemas covering 221 candidate constraints from the YAGOS dataset.

under ‘Exact’ node constraint and ‘Exact’ cardinality matching to 0.695 when ‘Datatype’ abstraction is allowed for node constraints and cardinality is ‘Loosened’. Similarly, on the more challenging Wikidata EntitySchema dataset, the global setting sees its F1-score rise from 0.312 (Exact/Exact) to 0.651 (Datatype/Loosened). This demonstrates that while GPT-4o mini produces a solid number of perfectly accurate constraints, its output contains an even larger proportion of constraints that are valid under more flexible, practical interpretations, particularly when considering datatype abstractions.

Compared to DeepSeek-V3, GPT-4o mini performs better under strict exact-matching condi-

tions. However, as the matching criteria are relaxed, DeepSeek-V3 surpasses GPT-4o mini. This indicates that while GPT-4o mini excels at generating precisely accurate constraints, DeepSeek-V3 is more capable of producing constraints that fulfill the core functional requirements of ShEx, making it more practical for real-world applications.

Matching Criteria		YAGO Schema			Wikidata EntitySchema		
Node Constraint	Cardinality	P	R	F1	P	R	F1
Exact	Exact	0.111	0.081	0.092	0.106	0.099	0.096
Subclass	Exact	0.111	0.081	0.092	0.106	0.099	0.096
Datatype	Exact	0.536	0.393	0.445	0.211	0.189	0.188
Exact	Loosened	0.111	0.081	0.092	0.149	0.129	0.129
Subclass	Loosened	0.111	0.081	0.092	0.151	0.131	0.131
Datatype	Loosened	0.630	0.465	0.525	0.356	0.297	0.304

Table 11: Results of sheXer under different matching criteria.

Settings	Matching Criteria		YAGOS			WES		
	Node Constraint	Cardinality	P	R	F1	P	R	F1
local	Exact	Exact	0.575	0.550	0.559	0.370	0.175	0.222
	Subclass	Exact	0.598	0.572	0.581	0.407	0.190	0.242
	Datatype	Exact	0.640	0.611	0.621	0.613	0.290	0.368
	Exact	Loosened	0.614	0.590	0.597	0.408	0.198	0.249
	Subclass	Loosened	0.640	0.616	0.623	0.456	0.217	0.275
	Datatype	Loosened	0.685	0.658	0.666	0.698	0.338	0.427
global	Exact	Exact	0.421	0.362	0.388	0.306	0.328	0.312
	Subclass	Exact	0.421	0.362	0.388	0.326	0.349	0.333
	Datatype	Exact	0.681	0.587	0.628	0.570	0.618	0.587
	Exact	Loosened	0.428	0.367	0.394	0.326	0.350	0.334
	Subclass	Loosened	0.428	0.367	0.394	0.346	0.372	0.354
	Datatype	Loosened	0.739	0.636	0.681	0.635	0.685	0.651
triples	Exact	Exact	0.631	0.564	0.591	0.335	0.200	0.242
	Subclass	Exact	0.650	0.581	0.608	0.391	0.232	0.281
	Datatype	Exact	0.698	0.621	0.652	0.643	0.395	0.472
	Exact	Loosened	0.678	0.607	0.634	0.377	0.231	0.276
	Subclass	Loosened	0.697	0.625	0.652	0.440	0.267	0.320
	Datatype	Loosened	0.744	0.664	0.695	0.754	0.477	0.563

Table 12: Results of GPT-4o mini under different matching criteria.

```

{
  "role": "system",
  "content": "You are a skilled knowledge engineer with deep expertise in writing
    ShEx (Shape Expressions) schemas. Carefully analyze the provided few-shot
    examples to understand the end-to-end generation process. Generate precise,
    well-structured ShEx scripts based on given example items and their related
    triples."
},
{
  "role": "user",
  "content": "Based on the information, generate the ShEx schema for the class 'http
    ://www.wikidata.org/entity/Q4220917 (film award)'. The provided list contains
    example instances of this class with the following fields: 'subject' (label),
    'predicate' (label), 'object' (label), and 'datatype'.
    Example instances:
    [
      wd:Q105447 (Saturn Award) wdt:P31 (instance of) [wd:Q107655869 (group of
        awards) (datatype: IRI), wd:Q4220917 (film award) (datatype: IRI)],
      wd:Q105447 (Saturn Award) wdt:P1027 (conferred by) [wd:Q2822376 (Academy of
        Science Fiction) (datatype: IRI)],
      wd:Q105447 (Saturn Award) wdt:P17 (country) [wd:Q30 (United States of
        America) (datatype: IRI)],
      wd:Q105447 (Saturn Award) wdt:P276 (location) [wd:Q34006 (Hollywood)(
        datatype: IRI)],
      wd:Q105447 (Saturn Award) wdt:P856 (official website) [<http://www.
        saturnawards.org/> (datatype: IRI)],
      ...
      wd:Q154590 (Golden Bear) wdt:P31 (instance of) [wd:Q4220917 (film award) (
        datatype: IRI)],
      wd:Q154590 (Golden Bear) wdt:P1027 (conferred by) [wd:Q130871 (Berlin
        International Film Festival) (datatype: IRI)],
      wd:Q154590 (Golden Bear) wdt:P17 (country) [wd:Q183 (Germany) (datatype: IRI
        )],
      wd:Q154590 (Golden Bear) wdt:P571 (inception) [1951-01-01T00:00:00Z (
        datatype: xsd:dateTime)],
      ...
      wd:Q182366 (Nordic Council Film Prize) wdt:P31 (instance of) [wd:Q4220917 (
        film award) (datatype: IRI)],
      wd:Q182366 (Nordic Council Film Prize) wdt:P1027 (conferred by) [wd:Q146165
        (Nordic Council) (datatype: IRI)],
      wd:Q182366 (Nordic Council Film Prize) wdt:P138 (named after) [wd:Q146165 (
        Nordic Council) (datatype: IRI)],
      ...
      wd:Q209459 (Golden Lion) wdt:P31 (instance of) [wd:Q4220917 (film award) (
        datatype: IRI)],
      wd:Q209459 (Golden Lion) wdt:P1027 (conferred by) [wd:Q49024 (Venice Film
        Festival) (datatype: IRI)],
      wd:Q209459 (Golden Lion) wdt:P159 (headquarters location) [wd:Q641 (Venice)
        (datatype: IRI)],
      wd:Q209459 (Golden Lion) wdt:P17 (country) [wd:Q38 (Italy) (datatype: IRI)],
      wd:Q209459 (Golden Lion) wdt:P276 (location) [wd:Q641 (Venice) (datatype:
        IRI)],
      ...
    ]"
}

```

Listing 5: An example of local setting input information. Only a subset of instances is shown for brevity.

```

{
  "role": "system",
  "content": "You are a skilled knowledge engineer with deep expertise in writing
  ShEx (Shape Expressions) schemas. Carefully analyze the provided few-shot
  examples to understand the end-to-end generation process. Generate precise,
  well-structured ShEx scripts based on given example items and their related
  triples."
},
{
  "role": "user",
  "content": "Generate a ShEx schema for the class 'http://www.wikidata.org/entity/
  Q4220917 (film award)' based on the provided information. The provided list
  contains example triples of instances of this class, with the following fields
  : 'subject' (label), 'predicate' (label), 'object' (label), and 'datatype'.
  Each predicate used by instances of this class is represented with triples
  from 5 instances.
  Example triples of predicates:
  [
    wd:Q105447 (Saturn Award) wdt:P31 (instance of) [wd:Q107655869 (group of
      awards) (datatype: IRI), wd:Q4220917 (film award) (datatype: IRI)],
    wd:Q154590 (Golden Bear) wdt:P31 (instance of) [wd:Q4220917 (film award) (
      datatype: IRI)],
    wd:Q209459 (Golden Lion) wdt:P31 (instance of) [wd:Q4220917 (film award) (
      datatype: IRI)],
    ...
    wd:Q105447 (Saturn Award) wdt:P17 (country) [wd:Q30 (United States of
      America) (datatype: IRI)],
    wd:Q154590 (Golden Bear) wdt:P17 (country) [wd:Q183 (Germany) (datatype: IRI
      )],
    wd:Q209459 (Golden Lion) wdt:P17 (country) [wd:Q38 (Italy) (datatype: IRI)],
    wd:Q290627 (Golden stars of French cinema) wdt:P17 (country) [wd:Q142 (
      France) (datatype: IRI)],
    ...
    wd:Q105447 (Saturn Award) wdt:P1027 (conferred by) [wd:Q2822376 (Academy of
      Science Fiction) (datatype: IRI)],
    wd:Q154590 (Golden Bear) wdt:P1027 (conferred by) [wd:Q130871 (Berlin
      International Film Festival) (datatype: IRI)],
    wd:Q182366 (Nordic Council Film Prize) wdt:P1027 (conferred by) [wd:Q146165
      (Nordic Council) (datatype: IRI)],
    wd:Q209459 (Golden Lion) wdt:P1027 (conferred by) [wd:Q49024 (Venice Film
      Festival) (datatype: IRI)],
    ...
    wd:Q105447 (Saturn Award) wdt:P571 (inception) [1972-01-01T00:00:00Z (
      datatype: xsd:dateTime)],
    wd:Q154590 (Golden Bear) wdt:P571 (inception) [1951-01-01T00:00:00Z (
      datatype: xsd:dateTime)],
    wd:Q182366 (Nordic Council Film Prize) wdt:P571 (inception) [2002-01-01T00
      :00:00Z (datatype: xsd:dateTime)],
    wd:Q209459 (Golden Lion) wdt:P571 (inception) [1949-01-01T00:00:00Z (
      datatype: xsd:dateTime)],
    ...
    wd:Q182366 (Nordic Council Film Prize) wdt:P138 (named after) [wd:Q146165 (
      Nordic Council) (datatype: IRI)],
    wd:Q321207 (Alfred Bauer Prize) wdt:P138 (named after) [wd:Q1686200 (Alfred
      Bauer) (datatype: IRI)],
    wd:Q630018 (Bambi Award) wdt:P138 (named after) [wd:Q43051 (Bambi) (datatype
      : IRI)],
    wd:Q734335 (Louis Delluc Prize) wdt:P138 (named after) [wd:Q1091635 (Louis
      Delluc) (datatype: IRI)],
    ...
  ]"
}

```

Listing 6: An example of triples setting input prompt. Only a subset of instances is shown for brevity.

```

{
  "role": "system",
  "content": "You are a skilled knowledge engineer with deep expertise in writing
    ShEx (Shape Expressions) schemas. Carefully analyze the provided few-shot
    examples to understand the end-to-end generation process. Generate precise,
    well-structured ShEx scripts based on given example items and their related
    triples."
},
{
  "role": "user",
  "content": "Based on the following information, generate constraints in JSON:
    {
      'class_uri': 'http://www.wikidata.org/entity/Q4220917',
      'class_label': 'film award',
      'class_description': 'recognition for cinematic achievements',
      'predicate_uri': 'http://www.wikidata.org/prop/direct/P664',
      'predicate_label': 'organizer',
      'predicate_description': 'person or institution organizing an event',
      'triple_examples': [
        'wd:Q3910523 wdt:P664 (organizer) [wd:Q2288813 (Italian National Syndicate
          of Film Journalists)]',
        'wd:Q11624249 (Fujimoto Prize) wdt:P664 (organizer) [wd:Q114256803]',
        'wd:Q18640780 (Florida Film Critics Circle Awards) wdt:P664 (organizer) [
          wd:Q3074282 (Florida Film Critics Circle)]',
        'wd:Q106867277 (NBR Freedom of Expression) wdt:P664 (organizer) [wd:
          Q1133614 (National Board of Review of Motion Pictures)]',
        'wd:Q109259295 (Gotham Independent Film Award for Breakthrough Nonfiction
          Series) wdt:P664 (organizer) [wd:Q892112 (Independent Feature Project)
        ]'
      ],
      'frequency': '1.73% instance(s) in the class use the predicate',
      'cardinality_distribution': '1.73% instances in the class have 1 object(s)
        when using the predicate',
      'datatype_of_objects': 'IRI',
      'object_class_distribution': '13.33% of subjects have objects in class wd:
        Q43229 (organization), 6.67% of subjects have objects in class wd:
        Q101007233 (film critics association), 6.67% of subjects have objects in
        class wd:Q10689397 (television production company)',
      'subject_type_constraint': 'Based on the subject type constraint of Wikidata
        , the item described by such predicates should be a subclass or instance
        of [wd:Q170584 (project), wd:Q288514 (fair), wd:Q464980 (exhibition),
        wd:Q1190554 (occurrence), wd:Q14136353 (fictional occurrence), wd:
        Q15275719 (recurring event), wd:Q15900616 (event sequence), wd:
        Q107736918 (series of concerts)].',
      'value_type_constraint': 'Based on the value type constraint of Wikidata,
        the value item should be a subclass or instance of [wd:Q5 (human), wd:
        Q43229 (organization), wd:Q49773 (social movement), wd:Q4164871 (
        position), wd:Q14623646 (fictional organization), wd:Q15275719 (
        recurring event), wd:Q16334295 (group of humans), wd:Q30017383 (
        fictional organism)].'
    }
  "
}

```

Listing 7: An example of global setting input prompt.