

Agent-OM: Leveraging LLM Agents for Ontology Matching

Zhangcheng Qiang
Australian National University
Canberra, ACT, Australia
qzc438@gmail.com

Weiying Wang
Monash University
Melbourne, VIC, Australia
teresa.wang@monash.edu

Kerry Taylor
Australian National University
Canberra, ACT, Australia
kerry.taylor@anu.edu.au

ABSTRACT

Ontology matching (OM) enables semantic interoperability between different ontologies and resolves their conceptual heterogeneity by aligning related entities. OM systems currently have two prevailing design paradigms: conventional knowledge-based expert systems and newer machine learning-based predictive systems. While large language models (LLMs) and LLM agents have revolutionised data engineering and have been applied creatively in many domains, their potential for OM remains underexplored. This study introduces a novel agent-powered LLM-based design paradigm for OM systems. With consideration of several specific challenges in leveraging LLM agents for OM, we propose a generic framework, namely Agent-OM (Agent for Ontology Matching), consisting of two Siamese agents for retrieval and matching, with a set of OM tools. Our framework is implemented in a proof-of-concept system. Evaluations of three Ontology Alignment Evaluation Initiative (OAEI) tracks over state-of-the-art OM systems show that our system can achieve results very close to the long-standing best performance on simple OM tasks and can significantly improve the performance on complex and few-shot OM tasks.

PVLDB Reference Format:

Zhangcheng Qiang, Weiying Wang, and Kerry Taylor. Agent-OM: Leveraging LLM Agents for Ontology Matching. PVLDB, 18(3): 516 - 529, 2024.
doi:10.14778/3712221.3712222

PVLDB Artifact Availability:

The source code, data, and/or other artifacts have been made available at <https://github.com/qzc438/ontology-llm>.

1 INTRODUCTION

Large language models (LLMs) are pre-trained with an enormous corpus of common knowledge and therefore have powerful generative capabilities. Despite the success of using LLMs in a wide range of applications, leveraging LLMs for downstream tasks still has several challenges. (1) LLMs are pre-trained models that do not capture late-breaking information. (2) LLM hallucinations are often observed in domain-specific tasks and hamper their reliability. LLMs often generate unsound responses that are syntactically sound but factually incorrect [33]. (3) LLMs are good models of linguistic competence, but have shown limited capabilities in non-linguistic tasks, such as planning and routing [72]. LLMs are originally designed

for sequential question-answering, but most real-world tasks are designed with complex logic rather than following a single path.

To overcome the limitations of LLM customisation for downstream tasks, LLM-based autonomous agents have become a prominent research area. In the field of artificial intelligence (AI), the notion of agents was first introduced in the famous Turing Test [71], referring to intelligent computational entities that can display human-like behaviours. Such AI agents have fallen short of human-level capabilities, as they can only act on simple and heuristic policy functions learnt from constrained environments and they lack efficient central control to simulate the human learning process [75]. LLMs, with remarkable success in demonstrating autonomy, reactivity, pro-activeness, and social ability, have attracted growing research efforts aiming to construct AI agents, so-called LLM agents [79].

The core concept of LLM agents is to employ the LLM as a controller or “brain” rather than as a predictive model only (a.k.a Model as a Service). LLM agents extend LLM’s capability with advanced planning, memory, and pluggable tools, and allow LLMs to communicate with open-world knowledge [78]. (1) Planning breaks down a complex task into simpler and more manageable subtasks. LLMs can also receive feedback on plans and perform reflection and refinement. The most practical technique used for LLM planning is chain of thought (CoT) [77]. (2) Tools allow LLMs to call external resources for additional information. They are often invoked by LLM actions. (3) Memory provides context to inherently stateless LLMs, including short-term memory and long-term memory. Short-term memory can be considered as context information obtained from planning and tools via in-context learning (ICL) [11]. Long-term memory often uses database storage with retrieval-augmented generation (RAG) [39] to retain information. Unlike fine-tuning, where models need to be retrained to learn new context data, ICL/RAG instead augments the LLM prompts with new information. ICL/RAG is more scalable for working with dynamic information. Almost 90% of use cases can be achieved by ICL/RAG-based search and retrieval [47]. A recent paper [58] demonstrates that RAG surpasses fine-tuning across a diversity of knowledge-intensive tasks.

Ontology matching (OM) is a classic alignment task, aiming to find possible correspondences between a pair of ontologies [20]. OM systems are developed to automate this matching process. There are two dominant design paradigms for OM systems: traditional knowledge-based OM systems that implement pre-defined logic and expert knowledge; and more recent machine learning-based OM systems that transform the matching task into machine-enhanced learning and prediction. The former expert systems require intensive expert knowledge, while the latter predictive systems need a large amount of high-quality data to train the model. The prevalence of LLMs and LLM agents has driven many successful domain-specific applications. However, in the context of OM, using LLMs and LLM agents is currently under-explored. Leveraging LLMs and

This work is licensed under the Creative Commons BY-NC-ND 4.0 International License. Visit <https://creativecommons.org/licenses/by-nc-nd/4.0/> to view a copy of this license. For any use beyond those covered by this license, obtain permission by emailing info@vldb.org. Copyright is held by the owner/author(s). Publication rights licensed to the VLDB Endowment.

Proceedings of the VLDB Endowment, Vol. 18, No. 3 ISSN 2150-8097.
doi:10.14778/3712221.3712222

LLM agents for OM tasks is not an intuitive task; the challenges will be presented in the Related Work section.

This paper introduces a novel agent-powered LLM-based design paradigm for OM systems. We propose a generic framework and implement it with a proof-of-concept system. The system extends the LLM capability beyond general question-answering, offering a powerful problem solver for OM tasks. The system includes tools that facilitate information retrieval, entity matching, and memory storage. The system is compared to state-of-the-art OM systems, achieving considerable matching performance improvements across three Alignment Evaluation Initiative (OAEI) tracks. Specifically, this paper makes the following contributions:

- We introduce a new agent-powered LLM-based design paradigm for OM systems and propose a novel Agent-OM framework. It consists of the following key components:
 - A LLM acts as a central “brain” to link different modules and instruct their functions via prompt engineering;
 - A pair of planning modules use CoT for OM decomposition;
 - A set of OM tools use ICL/RAG to mitigate LLM hallucinations;
 - A shared memory module uses dialogue and hybrid data storage to support the search and retrieval of entity mappings.
- We implement our proposed Agent-OM framework in a proof-of-concept system. The system deals with several critical downstream challenges in leveraging LLM agents for OM, such as cost-effective entity information retrieval, matching candidate selection, and search-based matching functions.
- The experimental results of the system show that Agent-OM achieves results very close to the best long-standing performance on simple OM tasks and significantly improves matching performance on complex and few-shot OM tasks.

An ontology contains classes, properties, and individuals. In this study, we consider only classes and properties, and individuals are excluded. Possible logical relations between classes (respectively properties) can be equivalence ($\equiv$) and subsumption (either $\subseteq$ or $\supseteq$). In this study, we only consider the logical relation of equivalence ($\equiv$) between classes and properties.

The rest of the paper is organised as follows. Section 2 reviews related work. We illustrate the design of our agent-powered LLM-based OM framework in Section 3 and present implementation details in Section 4. Section 5 and 6 evaluate the system, with a discussion in Section 7. We discuss the limitations and future work in Section 8 and 9, respectively. Section 10 concludes the paper.

2 RELATED WORK

OM is typically a non-trivial but essential alignment task for data integration, information sharing, and knowledge discovery [68]. While matching is a prerequisite for interoperating applications with heterogeneous ontologies, OM systems are designed to partially or fully automate matching. The traditional approach using knowledge-based OM systems, such as LogMap [34, 35], AgreementMakerLight (AML) [22, 23], and FCA-Map [40, 86] has been shown to be precise and effective. However, it is resource-hungry and labour-intensive. It is often difficult to find domain experts to evaluate the matches, and any group of experts may not be able to cover all domain concepts that an expert system requires. A new approach uses machine learning (ML), implemented in systems

such as BERTMap [27] and LogMap extension LogMap-ML [13]. ML-based OM systems employ the concept of training and testing in ML, using ontology entities as features for model training or for fine-tuning and then using the model to predict additional correspondences. Specifically, the leading system BERTMap uses a common language model originating from natural language processing (NLP) (i.e. BERT [16]).

Although ML-based approaches have shown a significant improvement in matching performance, their training-testing paradigm is not feasible for LLMs. The number of parameters used in LLMs is much larger than those of common language models. This means re-training the entire LLM is usually infeasible, and fine-tuning such a large model requires a number of samples that may be infeasibly large for OM. A survey in [88] implies that 1000 is a reasonable number of training samples to fine-tune GPT-3, but generally speaking, a domain ontology has only around 100-200 entities. Furthermore, currently some LLMs are only accessible through a web service. This means that training or fine-tuning LLMs risks leaking sensitive information, while synthetic data, on the other hand, makes it difficult to ensure training quality.

Early studies using LLMs for OM can be found in [28] and [52]. Both works use a purely prompt-based approach. The prompts are structured as a binary question: given an entity from the source ontology and an entity from the target ontology, the LLMs perform a classification task to determine whether these two entities are identical or not. A similar approach is also used in OLala [31] and LLMs4OM [24], but their candidate generation is integrated with text embedding extractor models. The authors of [2] explore the potential of using LLMs for complex ontology OM challenges.

LLM agents were introduced in AutoGPT [66] and BabyAGI [50]. The recent release of *OpenAI GPTs* [54], *Microsoft Copilot* [45], and *Copilot Studio* [46] sparked an interest in “the world of LLM agents”. Building applications with LLM agents allows users to build their own custom GPTs to support custom business scenarios. In ontology-related tasks, LLM-driven agents have shown impressive performance in automating manual activities in the broader task of ontology engineering. These works pay attention to the use of conversational dialogue to enhance the agent’s capability with human feedback. While this is suitable for tasks that require humans to be in the loop, such as collecting competency questions in ontology engineering [84] or validating extended terms in ontology learning [7], modern OM seeks to automate a complex task with minimal human intervention. Contrasting with these works, our aim is to design a new infrastructure that is able to instruct LLM agents to use planning to decompose a complex task into steps and to use tools to facilitate automated matching (a.k.a. function calling), instead of purely using agent-based conversational dialogue, even when specialised as ontology-oriented dialogue like [59, 85].

We introduce our novel agent-powered LLM-based design paradigm for OM systems. We have two generic agents; each one is self-contained and designed to instruct LLMs to use extensive planning, memory, and tools, thus unlocking their generative capacity to handle various types of OM tasks in different contexts. Meanwhile, as a key enabler for precise decision-making, we also limit the current LLM’s flaws in hallucination, context understanding, and non-linguistic reasoning. Several OM-related tools have been created for this purpose. These tools enable LLM agents to simulate

a traditional OM system, automating the entire matching process without human intervention. The overall infrastructure offers high scalability and allows extensive customisation. To the best of our knowledge, this study is the first to introduce an LLM-agent-based framework for OM tasks.

3 MATCHING WITH LLM AGENTS

Given a source ontology (O_s) and a target ontology (O_t), OM aims to find an alignment (A) that contains a set of pairs-matched entities $\{(e_i, e_j) | e_i \in O_s, e_j \in O_t\}$. A classical matching process has two main steps: retrieval and matching. The retrieval step involves retrieving internal information from the ontology itself (R_{int}) and external information from a domain-specific thesaurus (R_{ext}). The matching step involves selecting the matching candidates (M_{sel}), running the matching algorithms (M_{alg}), and refining the matching results (M_{ref}). A classical OM can be formulated as:

$$R_{int} \Rightarrow R_{ext} \Rightarrow M_{sel} \Rightarrow M_{alg} \Rightarrow M_{ref} \quad (1)$$

Figure 1 shows the architecture of *Agent-OM*, our LLM-based agent-powered OM framework. It retains the original input and output of the classical OM but modularises the two main steps with autonomous LLM agents, namely Retrieval Agent (*Agent_R*) and Matching Agent (*Agent_M*). We call these two LLM agents “Siamese” because they have their own planning modules and related tools but share memory. The memory is responsible for storing the information retrieved from the Retrieval Agent (R_{sto}) and facilitating the search of the stored information by the Matching Agent (M_{sea}). Therefore, an agent-based OM is formulated as:

$$Agent_R_{(R_{int} \Rightarrow R_{ext} \Rightarrow R_{sto})} \Rightarrow Agent_M_{(M_{sea} \Rightarrow M_{sel} \Rightarrow M_{alg} \Rightarrow M_{ref})} \quad (2)$$

For each autonomous agent, the workflow is described as follows. The planning module decomposes a complex task into several subtasks and defines the order of subtasks and tools to be invoked. The plan is stored in the dialogue and passed to the LLMs. LLMs then invoke the tools to perform the subtasks. The tools may communicate with each other, with intermediate results stored in the dialogue. The tools can also access the database via the CRUD (create, read, update, and delete) functions provided. The entire workflow is driven by LLM prompts. We use solid lines to show the actual workflow controlled by the LLMs, and dotted lines to show the implicit link between a subtask and its corresponding tool activated by the LLMs.

3.1 Retrieval Agent

The Retrieval Agent is responsible for extracting entities from the ontologies, eliciting their metadata and ontology context information, and storing them in the hybrid database. For each entity extracted from the source and target ontologies, the planning module generates the instruction for retrieving the relevant information and feeding it into the LLMs to invoke the corresponding retrieval tools. The tools used in the Retrieval Agent include Metadata Retriever, Syntactic & Lexical & Semantic Retriever, and Hybrid Database Store with Content Embedder, described as follows.

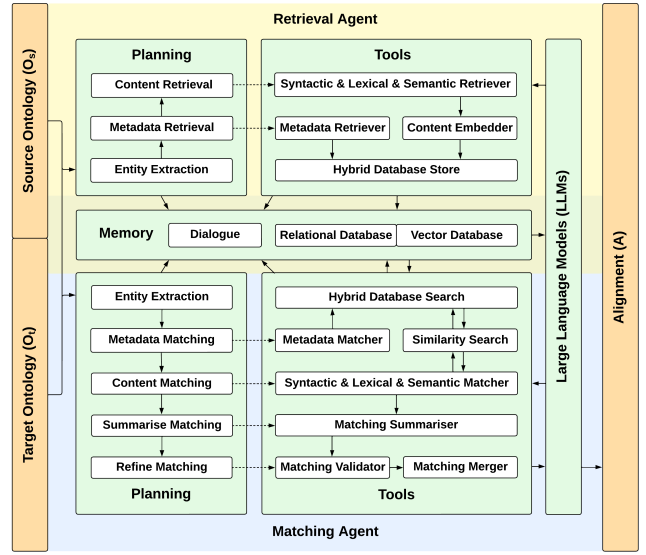


Figure 1: Architecture of Agent-OM. All components are executed twice, once for each of the source and target ontologies, apart from the matching merger tool that combines the results from each pass.

- Metadata Retriever (R_{int}): The metadata retriever collects the metadata of the input entity from the ontology, including its *category* (i.e. either from the *source* ontology or from the *target* ontology) and *type* (i.e. *named class* or *property*).
- Syntactic Retriever (R_{int}): The syntactic retriever is responsible for providing a unified text preprocessing result. A common text preprocessing pipeline consists of tokenisation, normalisation, stop words removal, and stemming/lemmatisation [43]. According to our prior study in [63], only tokenisation and normalisation help both matching completeness and correctness. The other two pipeline methods, stop words removal and stemming/lemmatisation, could cause unwanted false mappings. For this reason, our syntactic retriever considers only tokenisation and normalisation. We select white spaces to separate the words so that the outputs are short sentences that are easier for LLMs to interpret.
- Lexical Retriever (R_{int} & R_{ext}): We consider three key aspects of the entity’s lexical information: the general meaning (R_{ext}), the context meaning (R_{ext}), and the content meaning (R_{int}). In the context of OM, the general meaning is traditionally generated from Wikidata [73] or similar corpus-based knowledge bases (KBs). As LLMs are trained from these KBs, we use the prompt “What is the meaning of {entity_name}?” for the same function. However, using only the general meaning is not sufficient. Using the context constraint “Context: {context}” is effective in domain-specific tasks. Popular GPT-based domain applications, such as LawGPT [87] and MedicalGPT [74], use similar approaches. Additionally, we also retrieve content information from `rdfs:label`, `rdfs:comment`, `skos:prefLabel`, and `skos:definition` properties, where the ontology creators may add comments or explanations. These are also useful for retrieving the meaning of the entity.

- **Semantic Retriever (R_{int}):** The entity’s semantic information includes its basic triple-based relations and more complex logic-based axioms. In this study, we only consider triple-based relations that can be verbalised into a more natural language-like presentation via a prompt-based verbalisation tool. Such verbalisation tools are not capable of handling complex logic-based axioms. These functions can only be achieved with external packages, such as OWL Verbaliser [36], Sydney OWL Syntax [15], and the DeepOnto [29] verbalisation module [30].
- **Hybrid Database Store with Content Embedder (R_{sto}):** We use a hybrid database system consisting of a traditional relational database and an advanced vector database. Entity metadata, such as the entity’s category and type, are stored in the traditional relational database. In contrast, natural language-based content information, such as the entity’s syntactic, lexical, and semantic information, is vectorised via an embedding model and then stored in the vector database to enable similarity search based on relative distance in the embedding space. A unique key links these two databases.

3.2 Matching Agent

The Matching Agent is responsible for finding possible correspondences, ranking and refining the results according to different criteria, and selecting the most relevant candidate. For each entity extracted from the ontologies, the planning module generates the instruction for the matching types to be considered and feeds it into the LLMs to invoke the corresponding matching tools. The planning module first selects the source ontology as a starting point, extracting the entities from the ontology. Then, different matchers perform syntactic, lexical, or semantic matching functions to find the most relevant mappings to the input entity, using a hybrid database search across the relational and vector databases. A predicted mapping is based on a summarised profile measure of syntactic matching, lexical matching, and semantic matching, with matching validation. The same procedure applies to the target ontology as a starting point, and the results of the common candidates are combined. The tools used in the Matching Agent include Hybrid Database Search, Metadata Matcher, Syntactic & Lexical & Semantic Matcher, Matching Summariser, Matching Validator, and Matching Merger, described as follows.

- **Hybrid Database Search (M_{sea}):** The hybrid database search serves as an interface for the database accessible by the Metadata Matcher and Syntactic & Lexical & Semantic Matcher.
- **Metadata Matcher (M_{sel}):** Given an input entity, the metadata matcher collects the type and category of the input entity from the relational database.
- **Syntactic & Lexical & Semantic Matcher (M_{sel}):** Given an input entity, the syntactic & lexical & semantic matchers search for similar syntactic/lexical/semantic information respectively in the vector database using cosine similarity, defined for entities **A** and **B** as:

$$S_C(\mathbf{A}, \mathbf{B}) = \frac{\mathbf{A} \cdot \mathbf{B}}{\|\mathbf{A}\| \|\mathbf{B}\|} = \frac{\sum_{i=1}^n \mathbf{A}_i \mathbf{B}_i}{\sqrt{\sum_{i=1}^n \mathbf{A}_i^2} \cdot \sqrt{\sum_{i=1}^n \mathbf{B}_i^2}} \quad (3)$$

An extended search in the relational database is then used to filter the results based on the entity’s metadata.

- **Matching Summariser (M_{alg}):** We use reciprocal rank fusion (RRF) [14] to summarise the matching results. Viewing each of the

syntactic, lexical, and semantic descriptions as a document, the purpose of reciprocal rank is to accumulate the inverse of the ranks r of documents d over three ranking results from syntactic, lexical, and semantic matching, defined as:

$$RRF(d \in D) = \sum_{r \in R} \frac{1}{k + r(d)} \quad (4)$$

k is a constant parameter that is conventionally set to 0 as we do here. This ensures that the formula most highly rewards highly ranked entities. In our case, we are evaluating each entity that occurs in the top@ k of each of the three rankings (i.e. syntactic matching, lexical matching, and semantic matching), and combining their results as an overall matching summary.

- **Matching Validator (M_{ref}):** Validation is a critical step in minimising LLM hallucinations, as illustrated in SelfCheckGPT [42] and Self-RAG [5]. We also apply this method to the summarised results. We ask the LLM a binary question “Question: Is {entity_name} equivalent to {predicted_entity_name}? Context: {context} Answer the question within the context. Answer yes or no. Give a short explanation.” to check whether the predicted entity is equivalent to the input entity in the provided context. For computational efficiency, we iterate the comparison from rank 1 to n and select the highest-ranked match with a “yes” answer for the matching merger.
- **Matching Merger (M_{ref}):** The matching merger is responsible for combining the results from a search of the source ontology and a search of the target ontology. In this study, we select only the correspondences found on both sides. As an agent-based system, this can be extended to use multi-agent negotiation via the correspondence inclusion dialogue [59].

4 IMPLEMENTATION DETAILS

We implement our design of the framework in a proof-of-concept system. The components and their implementation are as follows:

- **LLMs:** Our system supports a wide range of LLMs, including OpenAI GPT [56], Anthropic Claude [3], Meta Llama [44], Alibaba Qwen [1], Google Gemma [26], and ChatGLM [83]. We select 10 models for this study. 4 models are API-accessed commercial LLMs, while the other 6 are open-source LLMs. Table 1 gives the details. For API-accessed LLMs, we include two models of different sizes for each category. For open-source LLMs, we select models with similar sizes (7-9 billion parameters) from different categories. They are accessed via the Ollama library [53].
- **Planning:** We select the LangChain library [38]. The library provides a wide range of agents. We select the tool calling agent (a.k.a function calling agent). At the time of writing, the LangChain package only supports this type of agent used with commercial API-accessed LLMs. To extend our framework to open-source LLMs, we employ the similar concept of “chain” to simulate tool calling agents for open-source LLMs.
- **Memory:** (1) Short-term memory: We use a conversational dialogue to store the original intermediate output of each operating process, with no map-reduce applied. (2) Long-term memory: We select a hybrid database consisting of a traditional relational database and an advanced vector database. PostgreSQL [62] supports a standalone integration of the traditional relational database and the extended vector database using pgvector [61]. We select OpenAI

Table 1: Details of LLMs used in the study.

Family	Model	Size	Version
GPT	gpt-4o	N/A	gpt-4o-2024-05-13
	gpt-4o-mini	N/A	gpt-4o-mini-2024-07-18
Claude	claude-3-sonnet	N/A	claude-3-sonnet-20240229
	claude-3-haiku	N/A	claude-3-haiku-20240307
Llama	llama-3-8b*	4.7 GB	Ollama Model ID: 365c0bd3c000
	llama-3.1-8b*	4.9 GB	Ollama Model ID: 46e0c10c039e
Qwen	qwen-2-7b*	4.4 GB	Ollama Model ID: dd314f039b9d
	qwen-2.5-7b*	4.7 GB	Ollama Model ID: 845dbda0ea48
Gemma	gemma-2-9b*	5.4 GB	Ollama Model ID: ff02c3702f32
GLM	glm-4-9b*	5.5 GB	Ollama Model ID: 5b699761eca5

* Open-source LLM (retrieved December 1, 2024).

Embedding [55] for the content embedding in the vector database. Alternatives are VertexAIEmbeddings [25] or Sentence-BERT [65], but the dimension of the embedding changes between different embedding models.

- Tools: To demonstrate the flexibility of our framework, we present the usage of prompt-based tools and programming-based tools, as well as the tools that combine a mixture of prompt-based and programming-based tools.

4.1 Ontology Naming Conventions

In this work, the term *entity* is a general expression for ontology classes or properties (without specifying which). We use *entity uri* to mean a fully expanded class name or property name with respect to its prefix. We use *entity name* to mean a class name or property name without its prefix. For example, the *entity uri* is “http://cmt#ProgramCommitteeChair” and the *entity name* is “ProgramCommitteeChair”.

Naming conventions for entity names fall into two categories: the name has a natural language meaning (Type 1); or the name is a code (Type 2). We observe that LLMs can perform well with meaningful entity names (e.g. ProgramCommitteeChair and Chair_PC). Often in larger biomedical ontologies each entity name is a unique identifier or code, and the meaningful description of the entity is in its label or comment (e.g. MA_0000270 and NCI_C33736). For this type of naming convention, current LLMs tend to generate the wrong synthesised label or comment corresponding to the code. For example, LLMs can mistakenly interpret the codes “MA_0000270” and “NCI_C33736” both to be “liver”, while the intended meanings of these two codes are “eyelid tarsus” and “Tarsal Plate”.

To handle the variety of ontology naming conventions and standardise their usage in LLM-based OM, we use a unified naming convention in this study. For an entity with a unique identifier or code, we use its label or comment instead. For example, we use “eyelid tarsus” and “Tarsal Plate” instead of “MA_0000270” and “NCI_C33736” respectively. In case the two ontologies reuse the same entity name, we add the type as a prefix for each entity. For example, if it were the case that “ProgramCommitteeChair” appears in both the source ontology and the target ontology, the unique identifier for each entity would be like “026-Source-Class-ProgramCommitteeChair” and “042-Target-Class-ProgramCommitteeChair”.

4.2 Running Example

To demonstrate the usability of our framework, we choose the CMT-ConfOf alignment as a sample alignment. The CMT Ontology is the source and the ConfOf Ontology is the target. These two ontologies contain similar concepts related to conference organisation. The running example aims to find the best matching entity in the target corresponding to the entity “http://cmt#ProgramCommitteeChair” in the source.

4.2.1 Retrieval Agent. Table 2 illustrates the tool calling in the Retrieval Agent. For the input “http://cmt#ProgramCommitteeChair” from the source ontology, the output of the Syntactic Retriever is “program committee chair”. The agent further invokes the Lexical Retriever to generate a detailed description: “In the context of a conference, ‘ProgramCommitteeChair’ refers to ...”. The Semantic Retriever generates related triple relations, such as “ProgramCommitteeChair rdfs:subClassOf ProgramCommitteeMember”. These triples are verbalised using natural language: “ProgramCommitteeChair” is a subclass of “ProgramCommitteeMember”.

While each entity has its own syntactic, lexical, and semantic information, a naive approach to deciding if two entities are the same is to generate a binary question for every pair of entities as a prompt to the LLM: “Is Entity1 equivalent to Entity2? Consider the following: The syntactic information of Entity1 is... The lexical information of Entity1 is... The semantic information of Entity1 is... The syntactic information of Entity2 is... The lexical information of Entity2 is... The semantic information of Entity2 is...” This approach has two limitations. (1) LLMs have token limits that restrict the number of tokens processed for each interaction. Combining all the retrieved information may exceed token limits. (2) The binary comparison is costly because the complexity of the comparison is the Cartesian product of the number of entities in the source ontology and the target ontology. We bypass these limitations by (1) using an open question instead and (2) storing useful information in a searchable database. Figure 2 shows the entity metadata and content information stored in the relational database and the vector database, respectively. On one hand, the entity’s metadata is needed to find an exact match. In other words, the matched entity needs to have the same metadata as the target entity. For example, “http://cmt#ProgramCommitteeChair” is a class in the source ontology, so the matched entity should be a class in the target ontology. On the other hand, content information including an entity’s syntactic, lexical, and semantic information is used for a similarity-based match because they are usually retrieved as natural language, which can be more ambiguous than metadata. Similarity between natural language terms is commonly based on embedding vectors, for which the vector database enables fast similarity searches.

4.2.2 Matching Agent. Table 3 demonstrates the tool calling in the Matching Agent. Given “http://cmt#ProgramCommitteeChair” from the source ontology, the most relevant entity found by each matcher is stored in short-term memory and combined using the RRF function. The result of the Matching Summariser is a list of predicted mappings. The last step is to refine the predicted mappings. The Matching Validator asks a binary question to compare whether the given entity is the same or different to the predicted relevant entity in RRF-ascending order. Because the validator receives a “yes”

Table 2: Tool calling in the Retrieval Agent.

Tool: Metadata Retriever
Input: {entity_uri} = "http://cmt#ProgramCommitteeChair"
Extract: {source_or_target} = "Source", {entity_type} = "Class"
Tool: Syntactic Retriever
Input: {entity_uri} = "http://cmt#ProgramCommitteeChair"
Extract: {entity_name} = "ProgramCommitteeChair" based on {entity_uri}.
Method: cleaning(entity_name)
Output: "program committee chair" (AI-generated content)
Tool: Lexical Retriever
Input: {entity_uri} = "http://cmt#ProgramCommitteeChair", {context} = "conference"
Extract: {entity_name} = "ProgramCommitteeChair" based on {entity_uri}.
{extra_information} based on {entity_uri} related rdfs:label, rdfs:comment, skos:prefLabel, and skos:definition.
Prompt: <i>Question: What is the meaning of {entity_name}? Context: {context} Extra Information: {extra_information}</i>
<i>Answer the question within the context and using the extra information.</i>
Output: "In the context of a conference, 'ProgramCommitteeChair' refers to..." (AI-generated content)
Tool: Semantic Retriever
Input: {entity_uri} = "http://cmt#ProgramCommitteeChair"
Method: generate_subgraph(entity_uri)
Output: subgraph
Prompt: <i>Verbalise triples into phrases: {subgraph}</i>
Output: "The class 'ProgramCommitteeChair' is a subclass of 'ProgramCommitteeMember'..." (AI-generated content)
Tool: Hybrid Database Store with Content Embedder
Input: {entity_uri} = "http://cmt#ProgramCommitteeChair", {source_or_target} = "Source", {entity_type} = "Class"
Extract: {entity_id} = "023-Source-Class-ProgramCommitteeChair"
Query: Create a relational database and store entity's metadata
<pre> DROP TABLE IF EXISTS ontology_matching CASCADE; CREATE TABLE ontology_matching (entity_id VARCHAR(1024) PRIMARY KEY, entity TEXT, source_or_target TEXT, entity_type TEXT); INSERT INTO ontology_matching (entity_id, entity, source_or_target, entity_type) VALUES ({entity_id}, {entity_uri}, {source_or_target}, {entity_type}); </pre>
Input: {entity_syntactic} = "program committee chair", {entity_lexical} = "In the context of conference, 'ProgramCommitteeChair' refers to...", {entity_semantic} = "The class 'ProgramCommitteeChair' is a subclass of 'ProgramCommitteeMember'..."
{matching_table} = "syntactic_matching/lexical_matching/semantic_matching"
Extract: {entity_embedding} based on {entity_syntactic}/{entity_lexical}/{entity_semantic}.
Query: Create a vector database and store entity's syntactic, lexical, and semantic information
<pre> CREATE EXTENSION IF NOT EXISTS vector; DROP TABLE IF EXISTS {matching_table}; CREATE TABLE {matching_table} (entity_id VARCHAR(1024) NOT NULL REFERENCES ontology_matching(entity_id), content TEXT, embedding vector(1536)); INSERT INTO {matching_table} (entity_id, content, embedding) VALUES ({entity_id}, {entity_syntactic}/{entity_lexical}/{entity_semantic}, {entity_embedding}); </pre>
Output: One relational database table (ontology_matching) and three vector database tables (syntactic_matching, lexical_matching, and semantic_matching).

answer for the first iteration of the entity "http://confOf#Chair_PC", the Matching Agent outputs "http://confOf#Chair_PC" as the best matching entity found in the target ontology. The Matching Merger combines the results from the same procedure applied in the search from "http://confOf#Chair_PC" in the target ontology. These two terms are considered as matched entities only if the mapping can be found bidirectionally (i.e. "http://cmt#ProgramCommitteeChair" is also found to be the best matching entity in the source ontology for the search from "http://confOf#Chair_PC" in the target ontology).

5 EVALUATION

5.1 Evaluation Criteria

In information retrieval, a common assessment for matching tasks is based on comparing predicted results with expected output. Precision and recall are used to measure the correctness and completeness of the matching, respectively. When adapting these measures to OM, the predicted results generated by the system are denoted Alignment (A), and the expected results provided by the domain

Table 3: Tool calling in the Matching Agent.

Tool: Metadata Matcher Input: {entity_uri} = "http://cmt#ProgramCommitteeChair", {source_or_target} = "Source" {matching_table} = "syntactic_matching/lexical_matching/semantic_matching" Query: Get entity id <pre> SELECT o.entity_id FROM ontology_matching o WHERE o.entity = {entity_uri} and o.source_or_target = {source_or_target}; </pre> Output: {entity_id} = "023-Source-Class-ProgramCommitteeChair" Query: Get entity metadata <pre> SELECT o.entity_type, m.embedding From ontology_matching o, {matching_table} m WHERE o.entity_id = m.entity_id AND o.entity_id = {entity_id}; </pre> Output: {entity_type} = "Class", {content_embedding} = [...]
Tool: Syntactic & Lexical & Semantic Matcher Input: {content_embedding} = [...], {matching_table} = "syntactic_matching/lexical_matching/semantic_matching", {similarity_threshold} = 0.90, {top_k} = 3, {source_or_target} = "Source", {entity_type} = "Class" Query: Search for similar entities <pre> WITH vector_matches AS (SELECT entity_id, 1 - (embedding <=> '{content_embedding}') AS similarity FROM {matching_table} WHERE 1 - (embedding <=> '{content_embedding}') >= {similarity_threshold}) SELECT o.entity_id, v.similarity as similarity FROM ontology_matching o, vector_matches v WHERE o.entity_id IN (SELECT entity_id FROM vector_matches) AND o.entity_id = v.entity_id AND o.source_or_target != {source_or_target} AND o.entity_type = {entity_type} ORDER BY v.similarity DESC LIMIT {top_k}; </pre> Output: {rankings} = [syntactic_matching: [], lexical_matching: ['095-Target-Class-Chair_PC'], semantic_matching: ['103-Target-Class-Member_PC', '092-Target-Class-Author', '123-Target-Class-University']] (AI-generated content)
Tool: Matching Summariser Method: reciprocal_rank_fusion(rankings) Output: {matching_summary} = [(1.0, ["095-Target-Class-Chair_PC", "103-Target-Class-Member_PC"]), (0.5, ["092-Target-Class-Author"]), (0.33, ["123-Target-Class-University"])] (AI-generated content)
Tool: Matching Validator Input: {context} = "conference", {entity_uri} = "http://cmt#ProgramCommitteeChair", {matching_summary} = [...] Extract: {entity_name} from {entity_uri}, each {predicted_entity_name} from {matching_summary}. Prompt: <i>Question: Is {entity_name} equivalent to {predicted_entity_name}? Context: {context}</i> <i>Answer the question within the context. Answer yes or no. Give a short explanation.</i> Output: {entity_id} = "095-Target-Class-Chair_PC" "095-Target-Class-Chair_PC": "Yes. In the context of a conference, the term "program committee chair" is equivalent to "chair PC." Both refer to the individual responsible for leading the program committee, which is in charge of organizing and overseeing the review and selection of conference submissions." "103-Target-Class-Member_PC": "No. The program committee chair is not equivalent to a member of the program committee (PC). The chair is responsible for overseeing the entire review process, coordinating the activities of the PC members, and making final decisions on the conference program. In contrast, a PC member primarily reviews and evaluates submitted papers." (AI-generated content) Query: Get entity uri <pre> SELECT o.entity FROM ontology_matching o WHERE o.entity_id = {entity_id}; </pre> Output: {entity_uri} = "http://confOf#Chair_PC"
Tool: Matching Merger Output: Merge the entity found in the source ontology that is equivalent to "http://confOf#Chair_PC" in the target ontology.

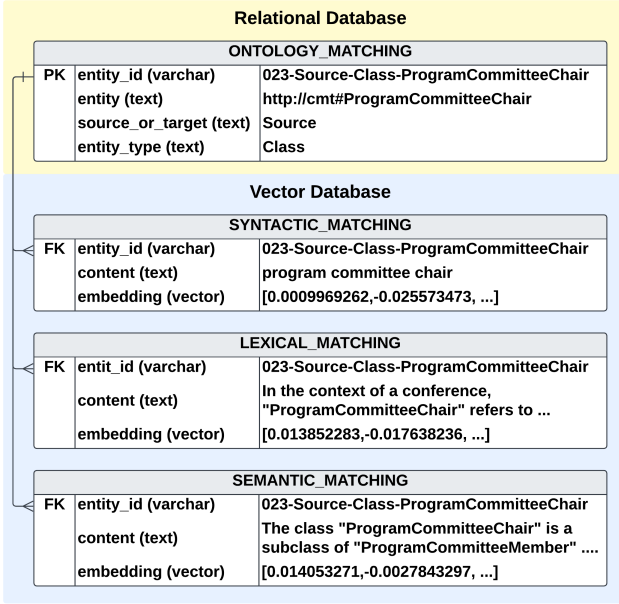


Figure 2: Storing “http://cmt#ProgramCommitteeChair”.

experts are denoted Reference (R) [17]. Therefore, precision and recall for OM tasks are defined as:

$$Precision = \frac{|A \cap R|}{|A|} \quad Recall = \frac{|A \cap R|}{|R|} \quad (5)$$

Precision and recall are commonly combined into a single measure F1 score, defined as:

$$F_1 \text{ Score} = \frac{2}{Precision^{-1} + Recall^{-1}} \quad (6)$$

5.2 Evaluation of Three OAEI Tracks

In this section, we test our proof-of-concept system with three OAEI tracks containing different types of OM tasks. These include few-shot tasks with a small proportion of trivial correspondences (5.2.1 Test Case), simple tasks with a large proportion of trivial correspondences (5.2.2 Test Case 1 and 5.2.3 Test Case 3), complex tasks with a large proportion of non-trivial correspondences (5.2.2 Test Case 2), with complex references (5.2.3 Test Case 1), or requiring domain-specific knowledge (5.2.3 Test Case 2). We report the evaluation metrics for the best-performing singular model gpt-4o and its hyperparameter settings over a single run. We ran multiple trials and found slight differences in the results due to the non-determinism of LLMs, but these differences are not significant with respect to the precision of the results we report. For all test cases in the three OAEI tracks, we select the hyperparameter settings of similarity_threshold = 0.90 and top@k = 3. See Section 6.2 for discussion on the hyperparameter settings of Agent-OM.

5.2.1 OAEI Conference Track. The OAEI conference track contains a pairwise alignment of 7 small and medium-sized conference-related ontologies with a total of 21 matching tasks [12, 69, 70, 82]. In each alignment, the trial correspondences that can be used to train

the models are very limited (commonly less than 10). All conference ontologies in this track use the Type 1 naming convention, where the names of classes and properties have meanings. In this study, we use the publicly available reference ra1-M3 as the reference (R), including class and property mappings.

Figure 3 compares Agent-OM with the 15 OM systems in OAEI 2022 and OAEI 2023. Agent-OM achieves above-average performance. Its overall F1 score ranks 3/13 in 2022 and 5/12 in 2023. We note that ra1-M3 is known to be missing valid equivalence mappings. We believe that Agent-OM could achieve better performance over a complete reference such as ra2-M3 or rar2-M3. These are not publicly available at the time of writing.

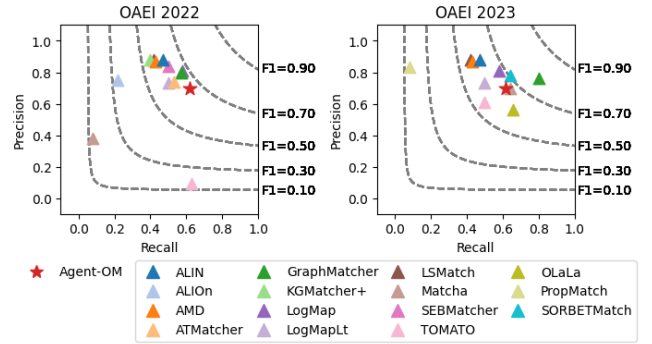


Figure 3: OAEI Conference Track Test Case.

5.2.2 OAEI Anatomy Track. The OAEI anatomy track contains only a reference alignment of human and mouse anatomy, created and evolved from [8, 10, 18, 21]. Both ontologies use the Type 2 naming convention, where the names of classes and properties are biomedical codes. We report the results of our evaluation in two test cases: alignment with trivial correspondences and alignment with non-trivial correspondences.

(1) **Test Case 1:** The track originally contains a large proportion of trivial correspondences that have the same standardised labels (e.g. “femoral artery” and “Femoral_Artery”). Figure 4 compares Agent-OM with the results of the 12 OM systems in OAEI 2022 and OAEI 2023 for alignment with trivial correspondences. Almost all OM systems achieve relatively high precision and recall when matching with trivial correspondences. For Agent-OM, its overall F1 score ranks the second highest in 2022 and 2023.

(2) **Test Case 2:** We remove these trivial correspondences from both the reference (R) and alignment (A) to focus the matching performance comparison on non-trivial correspondences. Figure 5 compares Agent-OM with the results of the 12 OM systems in OAEI 2022 and OAEI 2023 for alignment with non-trivial correspondences. We observe better performance by Agent-OM. The overall F1 score of Agent-OM is superior to the 9 traditional OM systems, better than the two LLM-based systems (OLaLa and SORBETMatch), and only behind one deep learning (DL)-based OM system (Matcha), which could have benefited from an unusual large training set available for this case.

5.2.3 OAEI MSE Track. The OAEI MSE track provides a reference alignment of ontologies in materials science and engineering [51].

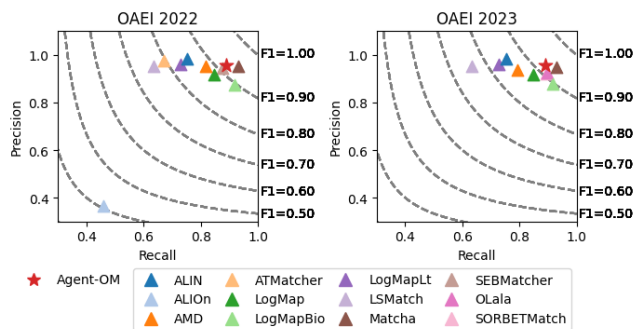


Figure 4: OAEI Anatomy Track Test Case 1.

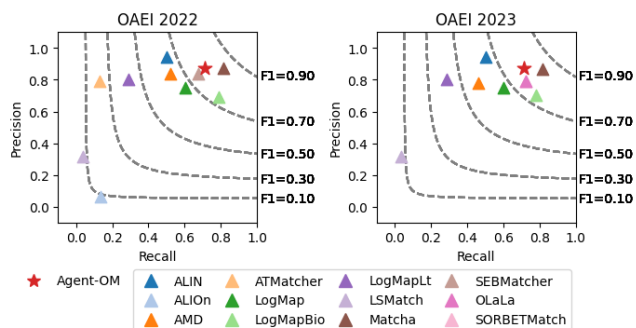


Figure 5: OAEI Anatomy Track Test Case 2.

The track contains three test cases aligning MaterialInformation [6], MatOnto [32], and EMMO [19]. MaterialInformation and MatOnto use the Type 1 naming convention, while EMMO uses the Type 2 naming convention.

(1) Test Case 1: This case provides a reference alignment for a small version of MaterialInformation and the medium-sized MatOnto. The challenge of this task arises due to the reference intentionally including several subsumption correspondences, when OM systems may mistakenly map subsumptions into equivalence relations. Figure 6 compares Agent-OM with the results of the 4 OM systems in OAEI 2022 and OAEI 2023. Agent-OM achieves the best performance, with the highest F1 score.

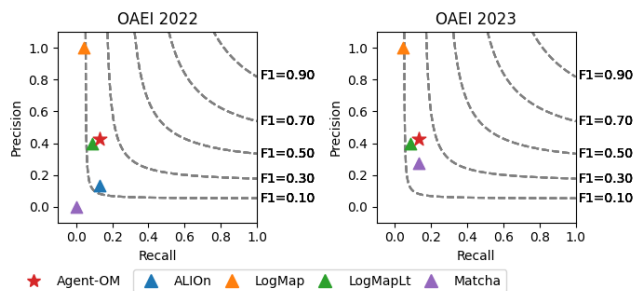


Figure 6: OAEI MSE Track Test Case 1.

(2) Test Case 2: This case provides the reference alignment of the full version of the large MaterialInformation and the medium-sized MatOnto. The challenge of this task is to align many examples of specific terminology, abbreviations, and acronyms used in materials science. For example, “Au” stands for “Gold”, “Ag” stands for “Silver”, and “Cu” stands for “Copper”. Figure 7 compares Agent-OM with the results of the 4 OM systems in OAEI 2022 and OAEI 2023. Agent-OM achieves the best performance in precision, recall, and overall F1 score across the OAEI 2022 and OAEI 2023 results. We should expect an LLM-based matcher to have high recall on this test case due to its access to extensive training literature.

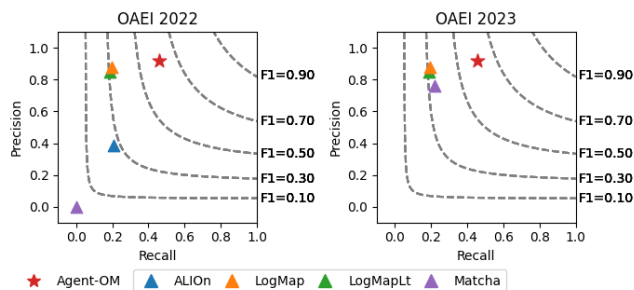


Figure 7: OAEI MSE Track Test Case 2.

(3) Test Case 3: This case provides the reference alignment of the full version of the large MaterialInformation and medium-sized EMMO. Unlike MatOnto used in Test Case 1 and Test Case 2, EMMO extends the upper ontology called Basic Formal Ontology (BFO). This means that the classes in EMMO are somewhat standardised according to the BFO classes. Figure 8 compares Agent-OM with the results of the 4 OM systems in OAEI 2022 and OAEI 2023. The performance of Agent-OM is competitive with the best of the OAEI 2022 results and in OAEI 2023 is bettered by the DL-based OM system Matcha.

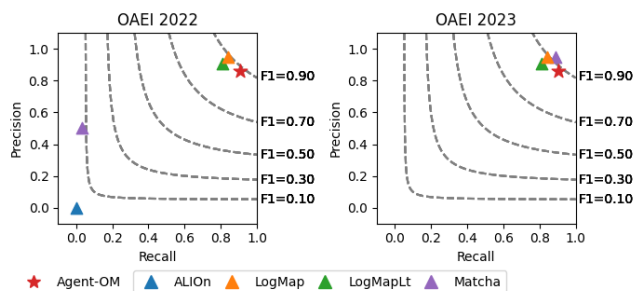


Figure 8: OAEI MSE Track Test Case 3.

6 ABLATION STUDY

6.1 System Components

6.1.1 *Architectures.* We compare Agent-OM with two simpler architectures where the OM is much more reliant on straightforward LLM use. (1) *LLM only:* Given O_s and O_t , this approach extracts

each $e_1 \in O_s$ and $e_2 \in O_t$. The matching decision is purely based on LLMs without any additional information. (2) *LLM with context*: Given O_s and O_t , this approach extracts each $e_1 \in O_s, e_2 \in O_t$ and their syntactic, lexical, and semantic information. The matching decision uses LLMs to determine whether two concepts are identical or not based on the information provided.

The experiment is run on the CMT-ConfOf alignment demonstrated in Section 4.2. We use the GPT model gpt-4o with similarity_threshold = 0.90. Figure 9 compares Agent-OM with *LLM Only* and *LLM with Context*. *LLM Only* shows low precision and recall. *LLM with Context* partially overcomes this deficiency by providing additional information, but token consumption is extremely high without optimising the matching candidate selection. The performance of *LLM with Context* is unstable over multiple runs due to the effects of stochastic LLM results. Our Agent-OM architecture handles these challenges with tool calling agents and hybrid database searches. Note that the CMT-ConfOf alignment demonstrated here is a small alignment task, while Agent-OM is expected to be relatively more effective and efficient in large-scale OM tasks.

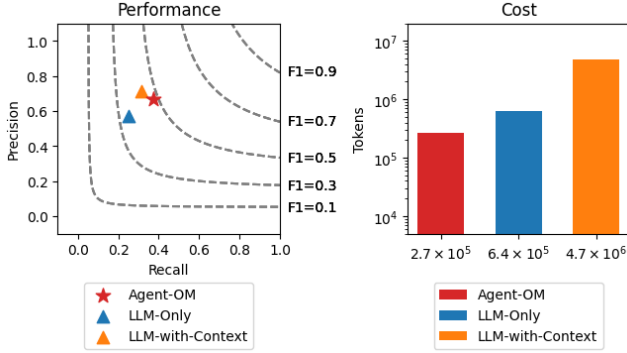


Figure 9: Comparison with LLM-based architectures.

6.1.2 LLMs. Figure 10 varies LLMs in Agent-OM on the OAEI Anatomy Track. In general, API-accessed models perform better than open-source models. The leading models, gpt-4o and claude-3-sonnet are both large API-accessed models. Among open-source models, gemma-2-9b achieves the best performance, while llama-3-8b is relatively poor. Curiously, although we see improved performance with llama-3.1-8b over its previous version, qwen-2.5-7b does not show an advantage over its previous version. This may be a side-effect of LLM developers optimising for tasks other than OM. We experimented with other LLMs derived from the Llama and Qwen families and generally found poor performance, possibly due to the fine-tuning for specific tasks.

6.1.3 Text Embedding Models. We also test three different text embeddings in [55] on the OAEI Anatomy Track. The default length of the embedding vector is 1536 for text-embedding-ada-002 and text-embedding-3-small, and 3072 for text-embedding-3-large. We do not observe a significant difference arising from varying the text embeddings from text-embedding-3-small to text-embedding-3-large. We do not see that text-embedding-3-small and text-embedding-3-large perform better than text-embedding-ada-002.

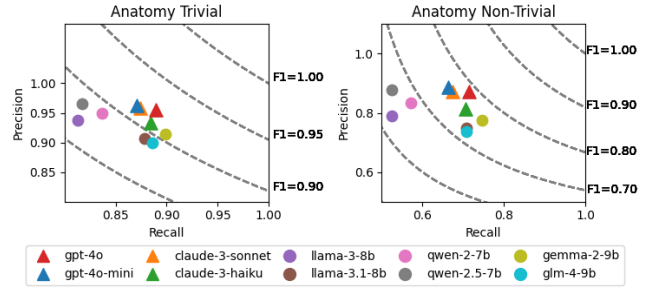


Figure 10: Comparison of different LLMs. API-accessed models are shown as triangles and open-source models as circles.

6.1.4 Hybrid Database. The use of a hybrid database unlocks the potential for search-based OM. We define N_s and N_t as the number of entities extracted from the source ontology (O_s) and the target ontology (O_t) respectively. In naive LLM-based OM, the matching is based on binary questions to compare each pair of entities from O_s and O_t based on their relevant information. The number of tokens consumed is $N_s * N_t$ (for retrieval and matching). In search-based OM, we first retrieve entity information from O_s and O_t and store it in a hybrid database. Following our methods in Section 4, the number of tokens consumed is $N_s + N_t$ (for retrieval) + 0 (for search) + $k(N_s + N_t)$ (for validation) + 0 (for merge). Search-based OM is cost-effective because the inequality $(k + 1)(N_s + N_t) < N_s * N_t$ always holds in common OM task settings where $N_s, N_t \gg k + 1$. We also apply two approaches to reduce LLM hallucinations (i.e. false positive mappings in search results): validator and merger.

6.1.5 Matching Validator. We employ a validator by asking the LLM to self-check the candidate correspondences. It is helpful in detecting two common types of false positive mappings: (1) non-existent mappings and (2) counter-intuitive mappings. Figure 11 compares precision, recall, and F1 score with and without validation in the three OAEI tracks we analysed. The matching results with validation generally achieved an improvement in precision and F1 score, with a slight decrease in recall. This is in line with the findings in CoT with self-consistency (CoT-SC) [76], where the provision of a self-check can reduce LLM hallucinations. Note that our experiment set the similarity_threshold = 0.90, so the improvement in matching performance is not statistically significant, but the matching validator will have a great impact on performance with lower similarity thresholds.

6.1.6 Matching Merger. We apply a merge function $O_s \Leftrightarrow O_t$ combining the results of $O_s \Rightarrow O_t$ and $O_s \Leftarrow O_t$ to improve the matching performance. Figure 12 shows the comparison of precision, recall, and F1 score in $O_s \Rightarrow O_t$, $O_t \Leftarrow O_s$, and $O_s \Leftrightarrow O_t$ in the three OAEI tracks we analysed. The merged matching results generally achieved a significant improvement in precision and F1 score, with a slight decrease in recall. The results are in line with the findings in RAG-Fusion [64], where providing two different paths to perform the same matching task can reduce LLM hallucinations.

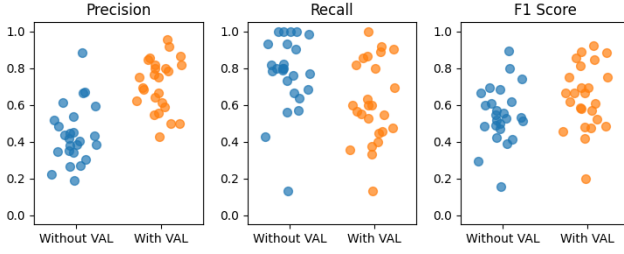


Figure 11: Comparison of without/with matching validator.

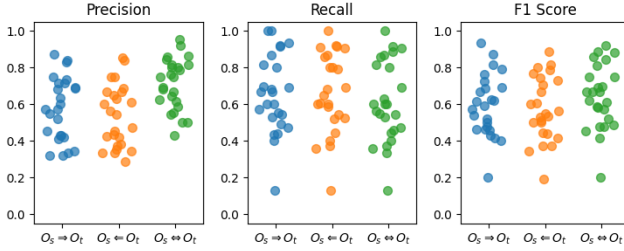


Figure 12: Comparison of $O_s \Rightarrow O_t$, $O_t \Leftarrow O_s$, and $O_s \Leftrightarrow O_t$

6.2 Hyperparameter Settings

6.2.1 Similarity Threshold. We test the similarity threshold $T \in [0.50, 0.55, 0.60, \dots, 0.90, 0.95, 1.00]$ in the three OAEI tracks we analysed. The optimal similarity threshold appears to be $T \in [0.90, 0.95]$, balancing the trade-off between precision and recall, thus achieving a higher overall F1 score. If we consider the similarity threshold as the required confidence interval (CI) for candidates for equivalence matching, this setting reflects the convention of accepting a 5-10% probability of observing values outside the estimation. The insensitivity of the results to the threshold value is more apparent in large-scale OM tasks, such as the Anatomy Track.

6.2.2 Top@k. We also test the top@k values $k \in [1, 2, 3, \dots, 8, 9, 10]$ in the three OAEI tracks we analysed. We observe that $k = 1$ and $k = 2$ do not provide enough candidates for the LLM to select, while appropriate correspondences are rarely found where $k > 5$. We recommend setting $k \in [3, 4, 5]$ to balance the computational complexity and precision of the results. Note that we deal with the tie-break case where multiple entities have the same RRF scores. In such cases, the total number of entities tested may be greater than k because these equally-scored entities share the same ranking.

6.2.3 How to choose Similarity Threshold and Top@k? Ontologies are context-dependent conceptual models that follow different conventions and restrictions to reflect different application-level requirements [67]. The hyperparameter settings can be adjusted for each specific OM task using reference mapping to achieve an optimal result. We observe that higher similarity thresholds and lower top@k values could result in high precision where most of the trivial mappings can be found, but some more obscure true mappings may be missed, lowering recall. On the other hand, lower

similarity thresholds and higher top@k values could result in high recall, but the precision may be low as more false mappings are generated during the matching process. This indicates that the applied matching refinements (validator and merger) would be more powerful in settings with lower similarity thresholds and higher top@k values. We found in our extensive OAEI experiments that threshold $T \in [0.90, 0.95]$ and top@k $k \in [3, 4, 5]$ was optimal. In the real-world application of Agent-OM to a matching problem with no reference, we advise choosing T and k within these ranges.

7 DISCUSSION

Google DeepMind classifies AI autonomy into 6 levels [49]. We believe that the potential of LLM is not only as a consultant, a collaborator, or an expert to answer binary classification questions in OM tasks, but it can also be an agent to simulate human behaviour in performing OM tasks, including data preprocessing, data preparation, data analysis, and data validation. Higher autonomy reduces barriers to accessing LLMs in OM tasks.

(1) Agent-based OM is more efficient than LLM-based OM. LLMs are computationally expensive. While LLM-based OM using binary classification questions has repetitive LLM prompting, agent-based OM leverages the vector database to store ontology retrieval results, reducing the financial cost of token consumption.

(2) Agent-based OM is also more effective than LLM-based OM. LLM-based OM is commonly observed to have limitations. Due to the nature of its large knowledge base, it is possible to discover positive correspondences but also to find unavoidable false mappings or missing true mappings. LLMs are zero-shot reasoners [37], but they are also few-shot learners [11]. Their capacity for reasoning depends on the richness of the information provided. With the assistance of autonomous agents for extensive planning, tools, and memory, agent-based OM can access more detailed information and therefore feature the following advantages:

- (a) Context Learning: LLMs have a large corpus of background knowledge. Given a context, LLMs can select relevant background knowledge and therefore perform better in lexical matching.
- (b) Transitive Reasoning: LLMs can reason on transitive relationships. They can also understand general and domain-specific scenarios and apply lexical validation when necessary.
- (c) Self Correction: LLMs have a strong capacity for self correction. Even given a wrong statement, LLMs have good judgment to automatically remove false mappings. For example, semantic matching can cause false mappings because it considers only the data structure and ignores the linguistic meaning of the entity. However, such a small piece of false information does not influence the correct truth that LLMs nevertheless learn.

Despite the success of agent-powered LLMs for OM, there may be further opportunities for improvement as follows.

(1) A matching process could be more complex. Although CoT may simulate how humans plan and perform tasks, it is still an incomplete model of human thought. Human reasoning employs a more complex network of thoughts, as humans tend to try different isolated paths (i.e. ToT, tree-of-thoughts [41, 80]), explore multiple paths (i.e. GoT, graph-of-thoughts [9, 81]), and backtrack, split, or merge to find the optimal solution to the problem. For example, people may use discovered mappings as input to the next iteration.

- (2) Prompt engineering is the key to instructing efficient LLM agents. These prompts are currently hand-crafted. For prompt-based tools, different LLMs may have varying default chat templates. Finding generic hard prompts across all LLMs is almost impossible. However, we provide the simplest standardised version of the prompts from our experiments. The prompts used in our system currently support mainstream LLMs, such as OpenAI GPT models, Anthropic Claude models, Meta Llama 3, Alibaba Qwen 2, Google Gemma 2, and ChatGLM 4. For those models not included in the list, we also provide an interface to add new LLMs to our system, but it may require minor code customisation to fit the LLMs used. We expect that our system will support more models via open-source community efforts in the future. We seek automatic prompt engineering and will consider using soft prompts in future versions.
- (3) LLM hallucinations can be mitigated, but cannot be eliminated. The accuracy of the RAG remains an open question. Human-in-the-loop may remain necessary [57]. Advanced RAG techniques, such as including the explanatory context in the RAG process [4], are promising directions.
- (4) There is a trade-off between precision and recall. Strict rules could result in high precision where most of the implicit true mappings can be found, but they may miss some explicit true mappings. On the other hand, loose rules could result in a high recall score, but the precision score may become very low as more mappings are generated during the matching process.
- (5) For LLMs used for OM, we find Moravec’s paradox [48]: “the hard problems are easy and the easy problems are hard” (p192) [60]. Although Agent-OM performs well in complex and few-shot OM tasks, it is not outstanding on simple OM tasks. We will also consider integrating the LLM-based approach with traditional knowledge-based and machine learning-based approaches.

8 LIMITATIONS

We evaluate only the T-Box matching datasets that match named classes, object properties, and data type properties. A-Box matching datasets (including individual data instances) are not considered due to privacy concerns in our targeted application domain. Additional data engineering (e.g. data de-identification and fuzzing) may be required to apply LLMs to A-Box matching datasets to avoid personal and sensitive information exposure.

Due to the high cost of API calls for API-accessed commercial LLMs, experiments with the newest commercial models (e.g. gpt-o1-preview and claude-3-opus) are not included in this study. According to our findings in Section 6.1.2, we hypothesise that these models could achieve better performance in OM tasks. There may be additional resource requirements for running open-source LLMs locally. The run time for API-accessed commercial LLMs is controlled by the LLM providers.

9 FUTURE WORK

- (1) Multimodal OM: We have packaged our system into several natural language-based commands. It could be integrated with advanced LLM functions to support multimodal input, such as

ontology diagrams and online seminars. The richer information sources might improve OM performance.

- (2) Multilingual OM: Agent-OM supports ontologies in multiple languages. We have tested it in the OAEI 2022 and 2023 Multifarm Track, a modified conference dataset translated into nine different languages. Results are not included here because there are few benchmarks available.

- (3) Small language models (SLM) for OM: SLM (e.g. gemma-2-2b) are useful in resource-constrained devices, but they have problematic tool interfaces at present.

10 CONCLUSION

In this paper, we introduce a new design paradigm for OM systems. Agent-OM, an agent-powered LLM-based framework, is proposed and implemented with a proof-of-concept system. We compare our systems with state-of-the-art OM systems to perform different types of OM tasks. The system has shown a powerful capability to perform OM tasks at different levels of complexity, leveraging the potential of using LLM agents for OM. We also discuss our observations on the advantages and current limitations of using LLMs and LLM agents for OM tasks.

Our work focuses on pre-trained large foundation models that are impossible to retrain and hard to fine-tune. Our approach yields good results on LLMs for OM tasks without changing the LLM model itself, but by utilising CoT, ICL/RAG, and prompt engineering techniques. It is a simple, lightweight, and natural language-driven approach with high scalability. Agent-OM is all you need. While OM has been studied for two decades or more, we are now at a point where the goal of 100% accurate, fully-automated, and domain-independent OM seems to be within reach.

ACKNOWLEDGMENTS

The authors thank the reviewers for providing insightful comments. The authors thank Sven Hertling for curating the Ontology Alignment Evaluation Initiative (OAEI) datasets stored in the Matching Evaluation Toolkit (MELT). The authors thank the organisers of Conference Track (Ondřej Zamazal, Jana Vataščinová, and Lu Zhou), Anatomy Track (Mina Abd Nikooie Pour, Huanyu Li, Ying Li, and Patrick Lambrix), and MSE Track (Engy Nasr and Martin Huschka), for helpful advice on reproducing benchmarks. The authors thank Jing Jiang from the Australian National University (ANU) for helpful advice on the verbalisation tool used in the semantic retriever. The authors thank Alice Richardson from the ANU Statistical Support Network for helpful advice on the statistical analysis. The authors thank the Commonwealth Scientific and Industrial Research Organisation (CSIRO) for supporting this project.

Agent-OM did not participate in the OAEI 2022 and 2023 campaigns. According to the OAEI data policy (retrieved December 1, 2024), “OAEI results and datasets, are publicly available, but subject to a use policy similar to the one defined by NIST for TREC. These rules apply to anyone using these data.” Please find more details from the official website: <https://oaei.ontologymatching.org/doc/oaei-deontology.2.html>.

REFERENCES

- [1] Alibaba Qwen Team. [n.d.]. Qwen Models. Retrieved December 1, 2024 from <https://qwenlm.github.io>
- [2] Reihaneh Amini, Sanaz Saki Norouzi, Pascal Hitzler, and Reza Amini. 2024. Towards Complex Ontology Alignment using Large Language Models. arXiv:2404.10329 [cs.AI] <https://arxiv.org/abs/2404.10329>
- [3] Anthropic. [n.d.]. Claude Models. Retrieved December 1, 2024 from <https://docs.anthropic.com/en/docs/about-claude/models>
- [4] Anthropic. [n.d.]. Introducing Contextual Retrieval. Retrieved December 1, 2024 from <https://www.anthropic.com/news/contextual-retrieval>
- [5] Akari Asai, Zequi Wu, Yizhong Wang, Avirup Sil, and Hannaneh Hajishirzi. 2024. Self-RAG: Learning to Retrieve, Generate, and Critique through Self-Reflection. In *The Twelfth International Conference on Learning Representations*. OpenReview.net, Vienna, Austria.
- [6] Toshihiro Ashino. 2010. Materials Ontology: An Infrastructure for Exchanging Materials Information and Knowledge. *Data Science Journal* 9 (2010), 54–61. <https://doi.org/10.2481/dsj.008-041>
- [7] Hamed Babaei Giglou, Jennifer D'Souza, and Sören Auer. 2023. LLMs4OL: Large Language Models for Ontology Learning. In *The Semantic Web – ISWC 2023*. Springer, Athens, Greece, 408–427. https://doi.org/10.1007/978-3-031-47240-4_22
- [8] Elena Beisswanger and Udo Hahn. 2012. Towards Valid and Reusable Reference Alignments—Ten Basic Quality Checks for Ontology Alignments and Their Application to Three Different Reference Data Sets. *Journal of biomedical semantics* 3, Article S4 (2012), 14 pages. <https://doi.org/10.1186/2041-1480-3-S1-S4>
- [9] Maciej Besta, Nils Blach, Ales Kubicek, Robert Gerstenberger, Michal Podstawski, Lukas Gianinazzi, Joanna Gajda, Tomasz Lehmann, Hubert Niewiadomski, Piotr Nyczyk, and Torsten Hoefler. 2024. Graph of Thoughts: Solving Elaborate Problems with Large Language Models. In *Proceedings of the AAAI Conference on Artificial Intelligence*, Vol. 38. AAAI Press, Washington, DC, USA, 17682–17690. <https://doi.org/10.1609/aaai.v38i16.29720>
- [10] Olivier Bodenreider, Terry F Hayamizu, Martin Ringwald, Sherri De Coronado, and Songmao Zhang. 2005. Of Mice and Men: Aligning Mouse and Human Anatomies. In *American Medical Informatics Association Annual Symposium – AMIA 2005*. AMIA, Washington, DC, USA, 61–65.
- [11] Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D 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 Ziegler, Jeffrey Wu, Clemens Winter, Chris Hesse, Mark Chen, Eric Sigler, Mateusz Litwin, Scott Gray, Benjamin Chess, Jack Clark, Christopher Berner, Sam McCandlish, Alec Radford, Ilya Sutskever, and Dario Amodei. 2020. Language Models are Few-Shot Learners. In *Proceedings of the 34th Annual Conference on Neural Information Processing Systems*, Vol. 33. Curran Associates, Inc., Vancouver, BC, Canada, Article 159, 25 pages.
- [12] Michelle Cheatham and Pascal Hitzler. 2014. Conference v2.0: An Uncertain Version of the OAEI Conference Benchmark. In *The Semantic Web – ISWC 2014*. Springer, Riva del Garda, Italy, 33–48. https://doi.org/10.1007/978-3-319-11915-1_3
- [13] Jiaoyan Chen, Ernesto Jiménez-Ruiz, Ian Horrocks, Denvar Antonyrajah, Ali Hadian, and Jaehun Lee. 2021. Augmenting Ontology Alignment by Semantic Embedding and Distant Supervision. In *The Semantic Web – ESWC 2021*. Springer, Virtual Conference, 392–408. https://doi.org/10.1007/978-3-030-77385-4_23
- [14] Gordon V. Cormack, Charles L A Clarke, and Stefan Buettcher. 2009. Reciprocal Rank Fusion Outperforms Condorcet and Individual Rank Learning Methods. In *Proceedings of the 32nd international ACM SIGIR conference on Research and Development in Information Retrieval*. ACM, Boston, MA, USA, 758–759. <https://doi.org/10.1145/1571941.1572114>
- [15] Anne Cregan, Rolf Schwitter, and Thomas Meyer. 2007. Sydney OWL Syntax-towards a Controlled Natural Language Syntax for OWL 1.1. In *Proceedings of the OWLED 2007 Workshop on OWL: Experiences and Directions*, Vol. 258. CEUR-WS.org, Innsbruck, Austria.
- [16] Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*. ACL, Minneapolis, MN, USA, 4171–4186. <https://doi.org/10.18653/v1/N19-1423>
- [17] Hong-Hai Do, Sergey Melnik, and Erhard Rahm. 2003. Comparison of Schema Matching Evaluations. In *Web, Web-Services, and Database Systems*. Springer, Erfurt, Germany, 221–237. https://doi.org/10.1007/3-540-36560-5_17
- [18] Zlatan Dragisic, Valentina Ivanova, Huan-yu Li, and Patrick Lambrix. 2017. Experiences from the Anatomy Track in the Ontology Alignment Evaluation Initiative. *Journal of biomedical semantics* 8, Article 56 (2017), 28 pages. <https://doi.org/10.1186/s13326-017-0166-5>
- [19] European Materials Modelling Council. [n.d.]. EMMO. Retrieved December 1, 2024 from <https://github.com/emmo-repo/EMMO>
- [20] Jérôme Euzenat and Pavel Shvaiko. 2013. *Ontology Matching (2nd ed.)*. Springer, Berlin, Heidelberg, Germany. <https://doi.org/10.1007/978-3-642-38721-0>
- [21] Euzenat, Jérôme and Meilicke, Christian and Stuckenschmidt, Heiner and Shvaiko, Pavel and Trojahn, Cássia. 2011. *Ontology Alignment Evaluation Initiative: Six Years of Experience*. Springer, Berlin, Heidelberg, Germany, 158–192. https://doi.org/10.1007/978-3-642-22630-4_6
- [22] Daniel Faria, Catia Pesquita, Emanuel Santos, Isabel F Cruz, and Francisco M Couto. 2014. AgreementMakerLight 2.0: Towards Efficient Large-Scale Ontology Matching. In *Proceedings of the ISWC 2014 Posters and Demonstrations Track*, Vol. 1272. CEUR-WS.org, Riva del Garda, Italy, 457–460.
- [23] Daniel Faria, Catia Pesquita, Emanuel Santos, Matteo Palmonari, Isabel F. Cruz, and Francisco M Couto. 2013. The AgreementMakerLight Ontology Matching System. In *On the Move to Meaningful Internet Systems: OTM 2013 Conferences*. Springer, Graz, Austria, 527–541. https://doi.org/10.1007/978-3-642-41030-7_38
- [24] Hamed Babaei Giglou, Jennifer D'Souza, Felix Engel, and Sören Auer. 2024. LLMs4OM: Matching Ontologies with Large Language Models. arXiv:2404.10317 [cs.AI] <https://arxiv.org/abs/2404.10317>
- [25] Google Cloud. [n.d.]. Google Vertex AI Embeddings. Retrieved December 1, 2024 from <https://cloud.google.com/vertex-ai/docs/generative-ai/embeddings/get-text-embeddings>
- [26] Google Gemma Team and Google DeepMind. [n.d.]. Google Gemma Models. Retrieved December 1, 2024 from <https://ai.google.dev/gemma>
- [27] Yuan He, Jiaoyan Chen, Denvar Antonyrajah, and Ian Horrocks. 2022. BERTMap: A BERT-based Ontology Alignment System. In *Proceedings of the 36th AAAI Conference on Artificial Intelligence*, Vol. 36. AAAI Press, Virtual Conference, 5684–5691. <https://doi.org/10.1609/aaai.v36i5.20510>
- [28] Yuan He, Jiaoyan Chen, Hang Dong, and Ian Horrocks. 2023. Exploring Large Language Models for Ontology Alignment. In *Proceedings of the ISWC 2023 Posters, Demos and Industry Tracks*, Vol. 3632. CEUR-WS.org, Athens, Greece.
- [29] Yuan He, Jiaoyan Chen, Hang Dong, Ian Horrocks, Carlo Allocca, Taehun Kim, and Brahmananda Sapkota. 2024. DeepOnto: A Python Package for Ontology Engineering with Deep Learning. *Semantic Web* 15, 5 (2024), 1991–2004. <https://doi.org/10.3233/SW-243568>
- [30] Yuan He, Jiaoyan Chen, Ernesto Jiménez-Ruiz, Hang Dong, and Ian Horrocks. 2023. Language Model Analysis for Ontology Subsumption Inference. In *Findings of the Association for Computational Linguistics: ACL 2023*. ACL, Toronto, Canada, 3439–3453. <https://doi.org/10.18653/v1/2023.findings-acl.213>
- [31] Sven Hertling and Heiko Paulheim. 2023. OLaLa: Ontology Matching with Large Language Models. In *Proceedings of the 12th Knowledge Capture Conference 2023*. ACM, Pensacola, FL, USA, 131–139. <https://doi.org/10.1145/3587259.3627571>
- [32] iNovex IRAD. [n.d.]. Ontologies for the MatOnto Project. Retrieved December 1, 2024 from <https://github.com/inovexcorp/MatOnto-Ontologies>
- [33] Ziwei Ji, Nayeon Lee, Rita Frieske, Tiezheng Yu, Dan Su, Yan Xu, Etsuko Ishii, Ye Jin Bang, Andrea Madotto, and Pascale Fung. 2023. Survey of Hallucination in Natural Language Generation. *ACM Computing Surveys* 55, 12, Article 248 (2023), 38 pages. <https://doi.org/10.1145/3571730>
- [34] Ernesto Jiménez-Ruiz and Bernardo Cuenca Grau. 2011. LogMap: Logic-Based and Scalable Ontology Matching. In *The Semantic Web – ISWC 2011*. Springer, Bonn, Germany, 273–288. https://doi.org/10.1007/978-3-642-25073-6_18
- [35] Ernesto Jiménez-Ruiz, Bernardo Cuenca Grau, and Yujiao Zhou. 2011. LogMap 2.0: Towards Logic-Based, Scalable and Interactive Ontology Matching. In *Proceedings of the 4th International Workshop on Semantic Web Applications and Tools for the Life Sciences*. ACM, London, UK, 45–46. <https://doi.org/10.1145/2166896.2166911>
- [36] Kaarel Kaljurand. 2007. OWL Verbalizer. Retrieved December 1, 2024 from <https://github.com/Kaljurand/owl-verbalizer>
- [37] Takeshi Kojima, Shixiang (Shane) Gu, Machel Reid, Yutaka Matsuo, and Yusuke Iwasawa. 2022. Large Language Models are Zero-Shot Reasoners. In *Proceedings of the 36th Annual Conference on Neural Information Processing Systems*, Vol. 35. Curran Associates, Inc., New Orleans, LA, USA, Article 1613, 15 pages.
- [38] LangChain, Inc. [n.d.]. LangChain. Retrieved December 1, 2024 from <https://www.langchain.com>
- [39] Patrick Lewis, Ethan Perez, Aleksandra Piktus, Fabio Petroni, Vladimir Karpukhin, Naman Goyal, Heinrich Küttler, Mike Lewis, Wen-tau Yih, Tim Rocktäschel, Sebastian Riedel, and Douwe Kiela. 2020. Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks. In *Proceedings of the 34th Annual Conference on Neural Information Processing Systems*, Vol. 33. Curran Associates Inc., Vancouver, BC, Canada, Article 793, 16 pages.
- [40] Guoxuan Li, Songmao Zhang, Jiayi Wei, and Wenqian Ye. 2021. Combining FCA-Map with representation learning for aligning large biomedical ontologies. In *Proceedings of the 16th International Workshop on Ontology Matching – ISWC 2021*, Vol. 3063. CEUR-WS.org, Virtual Conference, 207–208.
- [41] Jieyi Long. 2023. Large Language Model Guided Tree-of-Thought. arXiv:2305.08291 [cs.AI] <https://arxiv.org/abs/2305.08291>
- [42] Potsawee Manakul, Adian Liusie, and Mark Gales. 2023. SelfCheckGPT: Zero-Resource Black-Box Hallucination Detection for Generative Large Language Models. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*. ACL, Singapore, 9004–9017. <https://doi.org/10.18653/v1/2023.emnlp-main.557>
- [43] Christopher D. Manning, Prabhakar Raghavan, and Hinrich Schütze. 2008. *Introduction to Information Retrieval*. Cambridge University Press, Cambridge, UK. <https://doi.org/10.1017/CBO9780511809071>

- [44] Meta. [n.d.]. Meta Llama Models. Retrieved December 1, 2024 from <https://www.llama.com>
- [45] Microsoft. [n.d.]. Copilot. Retrieved December 1, 2024 from <https://www.microsoft.com/en-us/microsoft-copilot/personal-ai-assistant>
- [46] Microsoft. [n.d.]. Copilot Studio. Retrieved December 1, 2024 from <https://www.microsoft.com/en-us/microsoft-copilot/microsoft-copilot-studio>
- [47] Microsoft Azure. [n.d.]. GPT-RAG. Retrieved December 1, 2024 from https://github.com/Azure/GPT-RAG/blob/main/docs/RAG_CONCEPTS.md
- [48] Hans Moravec. 1988. *Mind Children: The Future of Robot and Human Intelligence*. Harvard University Press, Cambridge, MA, USA.
- [49] Meredith Ringel Morris, Jascha Sohl-Dickstein, Noah Fiedel, Tris Warkentin, Allan Dafoe, Aleksandra Faust, Clement Farabet, and Shane Legg. 2024. Levels of AGI for Operationalizing Progress on the Path to AGI. In *Proceedings of the 41st International Conference on Machine Learning*, Vol. 235. ML Research Press, Vienna, Austria, 36308–36321.
- [50] Yohei Nakajima. [n.d.]. BabyAGI. Retrieved December 1, 2024 from <https://github.com/yoheinakajima/babyagi>
- [51] Engy Nasr. 2020. *Evaluation of Automatic Ontology Matching for Materials Sciences and Engineering (Master's thesis)*. Master's thesis. Albert Ludwig University of Freiburg, Freiburg, Germany.
- [52] Sanaz Saki Norouzi, Mohammad Saeid Mahdavejad, and Pascal Hitzler. 2023. Conversational Ontology Alignment with ChatGPT. In *Proceedings of the 18th International Workshop on Ontology Matching – ISWC 2023*, Vol. 3591. CEUR-Ws.org, Athens, Greece, 61–66.
- [53] Ollama. [n.d.]. Ollama. Retrieved December 1, 2024 from <https://ollama.com/>
- [54] OpenAI. [n.d.]. GPTs. Retrieved December 1, 2024 from <https://openai.com/blog/introducing-gpts>
- [55] OpenAI. [n.d.]. OpenAI Embeddings. Retrieved December 1, 2024 from <https://platform.openai.com/docs/guides/embeddings/embedding-models>
- [56] OpenAI. [n.d.]. OpenAI Models. Retrieved December 1, 2024 from <https://platform.openai.com/docs/models>
- [57] Long Ouyang, Jeffrey Wu, Xu Jiang, Diogo Almeida, Carroll Wainwright, Pamela Mishkin, Chong Zhang, Sandhini Agarwal, Katarina Slama, Alex Ray, John Schulman, Jacob Hilton, Fraser Kelton, Luke Miller, Maddie Simens, Amanda Askell, Peter Welinder, Paul F Christiano, Jan Leike, and Ryan Lowe. 2022. Training Language Models to Follow Instructions with Human Feedback. In *Proceedings of the 36th Annual Conference on Neural Information Processing Systems*, Vol. 35. Curran Associates, Inc., New Orleans, LA, USA, 27730–27744.
- [58] Oded Ovadia, Menachem Brief, Moshik Mishaelli, and Oren Elisha. 2024. Fine-Tuning or Retrieval? Comparing Knowledge Injection in LLMs. In *Proceedings of the 24th Conference on Empirical Methods in Natural Language Processing*, ACL, Miami, FL, USA, 237–250. <https://doi.org/10.18653/v1/2024.emnlp-main.15>
- [59] Terry R. Payne and Valentina Tamma. 2014. Negotiating over Ontological Correspondences with Asymmetric and Incomplete Knowledge. In *Proceedings of the 2014 International Conference on Autonomous Agents and Multi-Agent Systems*. IFAAMAS, Paris, France, 517–524.
- [60] Steven Pinker. 1994. *The Language Instinct: How the Mind Creates Language*. William Morrow and Company, New York, NY, USA.
- [61] PostgreSQL Global Development Group. [n.d.]. pgvector. Retrieved December 1, 2024 from <https://github.com/pgvector/pgvector>
- [62] PostgreSQL Global Development Group. [n.d.]. PostgreSQL. Retrieved December 1, 2024 from <https://www.postgresql.org>
- [63] Zhangcheng Qiang, Kerry Taylor, and Weiqing Wang. 2024. How Does A Text Preprocessing Pipeline Affect Ontology Syntactic Matching? arXiv:2411.03962 [cs.CL] <https://arxiv.org/abs/2411.03962>
- [64] Adrian Raudaschl. [n.d.]. RAG-Fusion. Retrieved December 1, 2024 from <https://github.com/Raudaschl/rag-fusion>
- [65] Nils Reimers and Iryna Gurevych. 2019. Sentence-BERT: Sentence Embeddings using Siamese BERT-Networks. In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)*. ACL, Hong Kong, China, 3982–3992. <https://doi.org/10.18653/v1/D19-1410>
- [66] Toran Bruce Richards. [n.d.]. Auto-GPT. Retrieved December 1, 2024 from <https://github.com/Significant-Gravitas/AutoGPT>
- [67] Shenghui Wang. 2021. Linked Data and Semantic Web (2021-2A): Ontology Matching. Retrieved December 1, 2024 from <https://www.utwente.nl/en/ces/sal/exams/Linked-Data-and-Semantic-Web/ldsw-lecture6-om-2021-2a.pdf>
- [68] Pavel Shvaiko and Jérôme Euzenat. 2013. Ontology Matching: State of the Art and Future Challenges. *IEEE Transactions on Knowledge and Data Engineering* 25, 1 (2013), 158–176. <https://doi.org/10.1109/TKDE.2011.253>
- [69] Alessandro Solimando, Ernesto Jiménez-Ruiz, and Giovanna Guerrini. 2014. Detecting and Correcting Conservativity Principle Violations in Ontology-to-Ontology Mappings. In *The Semantic Web – ISWC 2014*. Springer, Riva del Garda, Italy, 1–16. https://doi.org/10.1007/978-3-319-11915-1_1
- [70] Alessandro Solimando, Ernesto Jiménez-Ruiz, and Giovanna Guerrini. 2017. Minimizing Conservativity Violations in Ontology Alignments: Algorithms and Evaluation. *Knowledge and Information Systems* 51 (2017), 775–819. <https://doi.org/10.1007/s10115-016-0983-3>
- [71] Alan M. Turing. 2009. *Computing Machinery and Intelligence*. Springer, Dordrecht, Netherlands, 23–65. https://doi.org/10.1007/978-1-4020-6710-5_3
- [72] Karthik Valmeekam, Alberto Olmo, Sarath Sreedharan, and Subbarao Kambhampati. 2022. Large Language Models Still Can't Plan (A Benchmark for LLMs on Planning and Reasoning about Change). In *NeurIPS 2022 Foundation Models for Decision Making Workshop*. Curran Associates, Inc., New Orleans, LA, USA.
- [73] Denny Vrandečić and Markus Krötzsch. 2014. Wikidata: A Free Collaborative Knowledge Base. *Communications of the ACM* 57, 10 (2014), 78–85. <https://doi.org/10.1145/2629489>
- [74] Haochun Wang, Chi Liu, Nuwa Xi, Zewen Qiang, Sendong Zhao, Bing Qin, and Ting Liu. 2023. HuaTuo: Tuning LLaMA Model with Chinese Medical Knowledge. arXiv:2304.06975 [cs.CL] <https://arxiv.org/abs/2304.06975>
- [75] Lei Wang, Chen Ma, Xueyang Feng, Zeyu Zhang, Hao Yang, Jingsen Zhang, Zhiyuan Chen, Jiakai Tang, Xu Chen, Yankai Lin, Wayne Xin Zhao, Zhewei Wei, and Jirong Wen. 2024. A Survey on Large Language Model Based Autonomous Agents. *Frontiers of Computer Science* 18, 6, Article 186345 (2024), 26 pages. <https://doi.org/10.1007/s11704-024-40231-1>
- [76] Xuezhi Wang, Jason Wei, Dale Schuurmans, Quoc V. Le, Ed H. Chi, Sharan Narang, Aakanksha Chowdhery, and Denny Zhou. 2023. Self-Consistency Improves Chain of Thought Reasoning in Language Models. In *The Eleventh International Conference on Learning Representations*. OpenReview.net, Kigali, Rwanda.
- [77] Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Brian Ichter, Fei Xia, Ed H. Chi, Quoc V. Le, and Denny Zhou. 2022. Chain-of-Thought Prompting Elicits Reasoning in Large Language Models. In *Proceedings of the 36th Annual Conference on Neural Information Processing Systems*, Vol. 35. Curran Associates, Inc., New Orleans, LA, USA, Article 1800, 14 pages.
- [78] Lilian Weng. 2023. LLM Powered Autonomous Agents. Retrieved December 1, 2024 from <https://lilianweng.github.io/posts/2023-06-23-agent/>
- [79] Zhiheng Xi, Wenxiang Chen, Xin Guo, Wei He, Yiwen Ding, Boyang Hong, Ming Zhang, Junze Wang, Senjie Jin, Enyu Zhou, Rui Zheng, Xiaoran Fan, Xiao Wang, Limao Xiong, Yuhao Zhou, Weiran Wang, Changhao Jiang, Yicheng Zou, Xiangyang Liu, Zhangyue Yin, Shihan Dou, Rongxiang Weng, Wensen Cheng, Qi Zhang, Wenjuan Qin, Yongyan Zheng, Xipeng Qiu, Xuanjing Huang, and Tao Gui. 2023. The Rise and Potential of Large Language Model Based Agents: A Survey. arXiv:2309.07864 [cs.AI] <https://arxiv.org/abs/2309.07864>
- [80] Shunyu Yao, Dian Yu, Jeffrey Zhao, Izhak Shafran, Tom Griffiths, Yuan Cao, and Karthik Narasimhan. 2023. Tree of Thoughts: Deliberate Problem Solving with Large Language Models. In *Proceedings of the 37th Annual Conference on Neural Information Processing Systems*, Vol. 36. Curran Associates, Inc., New Orleans, LA, USA, Article 517, 14 pages.
- [81] Yao Yao, Zuchao Li, and Hai Zhao. 2024. GoT: Effective Graph-of-Thought Reasoning in Language Models. In *Findings of the Association for Computational Linguistics: NAACL 2024*. ACL, Mexico City, Mexico, 2901–2921. <https://doi.org/10.18653/v1/2024.findings-naacl.183>
- [82] Ondřej Zamazal and Vojtěch Svátek. 2017. The Ten-Year OntoFarm and its Fertilization within the Onto-Sphere. *Journal of Web Semantics* 43 (2017), 46–53. <https://doi.org/10.1016/j.websem.2017.01.001>
- [83] Aohan Zeng, Bin Xu, Bowen Wang, Chenhui Zhang, Da Yin, Dan Zhang, Diego Rojas, Guanyu Feng, Hanlin Zhao, Hanyu Lai, Hao Yu, Hongning Wang, Jiada Sun, Jiajie Zhang, Jiale Cheng, Jiayi Gui, Jie Tang, Jing Zhang, Jingyu Sun, Juanzi Li, Lei Zhao, Lindong Wu, Lucen Zhong, Mingdao Liu, Minlie Huang, Peng Zhang, Qinkai Zheng, Rui Lu, Shuaiqi Duan, Shudan Zhang, Shulin Cao, Shuxun Yang, Weng Lam Tam, Wenyi Zhao, Xiao Liu, Xiao Xia, Xiaohan Zhang, Xiaotao Gu, Xin Lv, Xinghan Liu, Xinyi Liu, Xinyue Yang, Xixuan Song, Xunkai Zhang, Yifan An, Yifan Xu, Yilin Niu, Yuntao Yang, Yueyan Li, Yushi Bai, Yuxiao Dong, Zehan Qi, Zhaoyu Wang, Zhen Yang, Zhengxiao Du, Zhenyu Hou, and Zihan Wang. 2024. ChatGLM: A Family of Large Language Models from GLM-130B to GLM-4 All Tools. arXiv:2406.12793 [cs.CL] <https://arxiv.org/abs/2406.12793>
- [84] Bohui Zhang, Valentina Anita Carriero, Katrin Schreiberhuber, Stefani Tsaneva, Lucía Sánchez González, Jongmo Kim, and Jacopo de Berardinis. 2024. OntoChat: a Framework for Conversational Ontology Engineering using Language Models. arXiv:2403.05921 [cs.AI] <https://arxiv.org/abs/2403.05921>
- [85] Shiyao Zhang, Yuji Dong, Yichuan Zhang, Terry R. Payne, and Jie Zhang. 2024. Large Language Model Assisted Multi-Agent Dialogue for Ontology Alignment. In *Proceedings of the 2024 International Conference on Autonomous Agents and Multiagent Systems*. IFAAMAS, Auckland, New Zealand, 2594–2596.
- [86] Mengyi Zhao, Songmao Zhang, Weizhuo Li, and Guowei Chen. 2018. Matching Biomedical Ontologies Based on Formal Concept Analysis. *Journal of Biomedical Semantics* 9, Article 11 (2018), 27 pages. <https://doi.org/10.1186/s13326-018-0178-9>
- [87] Zhi Zhou, Jiang-Xin Shi, Peng-Xiao Song, Xiao-Wen Yang, Yi-Xuan Jin, Lan-Zhe Guo, and Yu-Feng Li. 2024. LawGPT: A Chinese Legal Knowledge-Enhanced Large Language Model. arXiv:2406.04614 [cs.CL] <https://arxiv.org/abs/2406.04614>
- [88] Mingyu Zong and Bhaskar Krishnamachari. 2022. A Survey on GPT-3. arXiv:2212.00857 [cs.CL] <https://arxiv.org/abs/2212.00857>