

LabAlignAgent: A LOINC-Enriched Ontology-Guided Agent for Cross-Hospital Semantic Alignment of EMR Laboratory Tests

Sarah Ayad

Arab Open University, Computer Science Department, KSA, Saudi Arabia

E-mail: s.ayad@arabou.edu.sa

Keywords: LLM-based agents, ontology-guided agents, agentic RAG, ontology alignment, electronic medical records

Received: February 19, 2026

Patients often receive care from multiple hospitals, where laboratory tests are commonly represented using institution-specific names, units, and reporting conventions. This heterogeneity limits semantic interoperability across Electronic Medical Record systems. This study proposes LabAlignAgent, a LOINC-enriched ontology-guided agent for cross-hospital semantic alignment of EMR laboratory tests. The approach combines domain ontology construction, LOINC-based ontology enrichment, exact and fuzzy matching, GPT-assisted semantic verification, and rule-based alignment to harmonize heterogeneous laboratory test catalogs. The evaluation used real-world laboratory test data from two Saudi hospitals: 136 alignable tests from King Fahad Specialist Hospital (KFSH) and 126 alignable tests from Saudi German Hospital (SGH). In the hospital-to-LOINC normalization task, 202 out of 262 tests, representing 77.1%, were successfully mapped and enriched with standardized LOINC concepts, including 67 additional mappings recovered through GPT-assisted reasoning beyond lexical matching alone. In the cross-hospital alignment task, 17,136 possible KFSH–SGH pairs were reduced to 4,509 semantically plausible candidate pairs using ontology-derived constraints on unit and sample type. The full LabAlignAgent workflow produced 224 high-confidence cross-hospital semantic matches, compared with 67 matches obtained by the baseline prompt-based alignment without ontology enrichment. The accepted mappings were validated by a clinical pathologist, who assessed clinical equivalence using analyte, specimen, unit, method, and reference-interpretation criteria. The results show that LOINC-enriched ontology guidance and contextual semantic reasoning substantially improve laboratory test harmonization across hospitals and support more interoperable EMR integration. All datasets, ontologies, lab test names, and alignment results for each strategy are publicly available in our GitHub repository at <https://github.com/SarahayadAOU/LabAlignAgent>.

Povzetek: LabAlignAgent z uporabo ontologij, standarda LOINC in GPT izboljša usklajevanje laboratorijskih preiskav med bolnišnicami.

1 Introduction

In Saudi Arabia, patients often visit multiple healthcare providers across public and private hospitals. However, each institution typically uses its own EMR system, with differing naming conventions, data formats, and clinical terminologies. This lack of standardization leads to fragmented medical records, where critical patient information is scattered across isolated systems. As a result, clinicians may not have access to a patient’s full history during consultations, and patients themselves are unable to view or manage a unified medical record.

The problem of EMR fragmentation is not unique to Saudi Arabia, but its effects are magnified in regions with large, distributed healthcare networks. One of the key barriers is the inconsistent representation of clinical concepts such as lab tests, diagnoses, and procedures. For example, the same lab test may be recorded as “Hemoglobin A1c” in one hospital and “HbA1c” in another, or as “Complete Blood Count” in one system and “CBC” in another. Al-

though these labels refer to the same clinical test, differences in naming conventions, units, and reference ranges make automated integration difficult. These inconsistencies prevent national-scale health analytics, increase the likelihood of redundant testing, and limit the effectiveness of AI-driven clinical decision support systems.

Beyond clinical decision-making, semantic fragmentation of EMR data has important economic and operational consequences, particularly for health insurance providers. When laboratory tests are recorded using incompatible terminologies, insurers are often unable to reliably identify equivalent tests performed at different hospitals. This frequently leads to unnecessary re-testing, delayed claim processing, and increased healthcare costs. In addition, the lack of semantic consistency limits insurers’ ability to perform accurate utilization analysis, and fraud detection. Motivated by these challenges, this work aims to provide a semantic alignment solution that enables reliable reuse and comparison of laboratory results across institutions, supporting both clinical continuity of care and cost-efficient

insurance workflows.

To address the challenges of semantic fragmentation in medical records, many countries have adopted standardized clinical ontologies. In the United States, the Centers for Disease Control and Prevention (CDC) recommends the use of LOINC as the national standard for laboratory test reporting, and SNOMED CT is widely used across federal health systems including the Veterans Affairs (VA) and Medicare programs [12, 13]. The United Kingdom has mandated SNOMED CT across the entire National Health Service (NHS), where it is used as the core clinical terminology for documenting diagnoses, symptoms, and procedures [14]. Similarly, Canada has adopted both SNOMED CT and LOINC as part of its pan-Canadian standards for health information exchange, supported by Canada Health Infoway [15]. In Australia, both ontologies are integral to the national electronic health infrastructure and are used in the My Health Record system for standardizing clinical data exchange [5].

Recent advances in Artificial Intelligence, particularly in large-scale foundation models, have opened new possibilities for semantic alignment across heterogeneous clinical vocabularies [22, 23]. Techniques such as dynamic prompting when combined with ontology-aware reasoning can guide these models in performing context-sensitive concept alignment [24, 25]. Early results from materials science have demonstrated the potential of such models for structured concept mapping [26].

In this paper, we introduce an ontology-guided semantic alignment agent that harmonizes EMR lab tests from multiple hospitals by using LOINC-enriched dynamic prompting. Rather than aligning local terms directly to LOINC, we use these ontologies as semantic enrichment resources to construct context-rich representations for each hospital-specific term. These enriched profiles are used by the agent to unify equivalent terms under a single, standardized concept. The proposed approach is evaluated on real-world EMR datasets from King Fahad Specialist Hospital and Saudi German Hospital, and validated via an expert review.

To guide the design and evaluation of LabAlignAgent, this study addresses the following explicit research questions:

- **RQ1:** Does ontology-guided enrichment with LOINC metadata improve the accuracy and coverage of semantic alignment between laboratory test records from heterogeneous EMR systems, compared to a baseline lexical and prompt-based approach without enrichment?
- **RQ2:** Among the metadata dimensions used in the alignment pipeline, including measurement unit, sample type, test method, and LOINC-based attributes, which contribute most to achieving high-confidence semantic matches across institutional boundaries?

These research questions directly shape the evaluation in Section 6, where we compare a baseline prompt-

based alignment strategy (no ontology enrichment, 67 correct mappings) against the full LabAlignAgent pipeline (LOINC enrichment + rule-based reasoning, 224 correct mappings).

2 Related work

The semantic interoperability of healthcare data remains a critical challenge in modern healthcare systems, where patients frequently receive care across multiple institutions, each employing distinct electronic medical record (EMR) systems with heterogeneous terminologies and coding conventions. This literature review synthesizes recent advances in semantic interoperability and ontology alignment across four key research areas: (1) semantic interoperability and clinical terminology alignment, (2) ontology-based mapping and biomedical concept alignment, (3) large language model (LLM)-based semantic alignment and (4) agent-based semantic mapping systems.

In Semantic Interoperability and Clinical Terminology Alignment, Clinical terminology standardization has focused heavily on mapping local, institution-specific terms to reference terminologies. Park et al. [40] developed an automated SNOMED CT mapping tool that achieved top-5 accuracy rates of 98.67% for diagnostic terms and 99.19% for surgical procedural terms across four major hospital networks in South Korea. However, ambiguous terms, granularity gaps, and abbreviation resolution remain significant challenges. Rodrigues et al. [41] addressed semantic alignment between ICD-11 and SNOMED CT, combining axiom-based and rule-based approaches to handle fundamental differences in design and editorial policies between these two widely used terminologies. Their preliminary results covered more than half the domain of the draft ICD-11, highlighting both the feasibility and the complexity of aligning terminologies with different conceptual foundations.

In specialized clinical domains, terminology alignment faces additional challenges. Naliyathaliyazchayil et al. [42] focused on harmonizing LOINC codes and units in real-world oncology data, noting that existing studies show 6% to 19% of laboratory tests cannot be accurately mapped to LOINC. Their method addressed the challenge that existing systems often depend on source data strings and input features that may be absent, null, or incorrect, underscoring the need for robust, scalable approaches.

On the other hand, Biomedical ontology alignment has benefited from systematic evaluation through the OAEI benchmark datasets, particularly the Large Biomedical (LargeBio) track involving ontologies such as FMA (Foundational Model of Anatomy), SNOMED CT, and NCI Thesaurus [49], [50]. Diallo's ServOMap system achieved F-measures of 85.5% (precision 99%) for the FMA-NCI small task and 82.8% (precision 93.3%) for the large task [51]. A subsequent version, ServOMap_V4, improved to 89.4% F-measure (precision 94.3%). The main limitation identified

was low recall for ontologies with poor lexical descriptions, as the system relied heavily on lexical similarity computing.

Moreover, Yamamoto et al. [59] addressed context-based refinement of mappings in evolving life science ontologies, recognizing that ontologies are not static artifacts but evolve over time. Their approach refined existing mappings based on contextual information, improving alignment quality as ontologies change.

The application of large language models to ontology alignment represents a recent but rapidly evolving research direction. He et al. [61] conducted an early exploration of LLMs for ontology alignment using challenging subsets from the OAEI Bio-ML track. They found that *Flan-T5-XXL* with threshold tuning achieved an F-score of 0.721 on NCIT-DOID, outperforming *BERTMap* by 0.093, though it lagged on SNOMED-FMA. Notably, *GPT-3.5-turbo* performed poorly, with F-scores of 0.313 on NCIT-DOID and 0.132 on SNOMED-FMA. A significant limitation identified was the computational demand of LLMs, making full deployment challenging without smaller, representative datasets.

Other Studies in this area, exploited Agent-Based LLM Systems achieving results very close to long-standing best performance on simple ontology matching tasks and significantly improving performance on complex and few-shot tasks [63], [64]. Despite promising results, LLM-based approaches face several persistent challenges. Computational cost remains a significant barrier, with inference times and resource requirements often exceeding those of traditional methods [69], [70]. Hallucination the generation of plausible but incorrect information poses risks in high-stakes domains like healthcare [71], [72]. Consistency across different prompts and contexts is not guaranteed, and the black-box nature of LLMs complicates debugging and validation [73], [74].

Agent-based approaches to semantic mapping distribute the alignment task across autonomous agents that negotiate and collaborate. Li et al. [75] developed an agent-based framework for ontology mapping and integration, demonstrating flexibility and practicality through a prototype applied to beer ontologies. Santos et al. proposed using cooperative agent negotiation for ontology mapping, evaluating their approach on four groups of ontologies including Google and Yahoo web directories, product schemas, course university catalogs, and company profiles [76]. At the end, Agent-based semantic mapping systems face several challenges. Communication overhead can become significant as the number of agents and the complexity of negotiations increase [83], [84]. Ensuring convergence to high-quality alignments while maintaining reasonable computational costs requires careful design of negotiation protocols and conflict resolution mechanisms [85], [86]. Scalability remains a concern, particularly in highly dynamic environments where ontologies evolve frequently [87], [88].

Recent advances in LLMs have enabled new forms of semi-autonomous agents capable of semantic reasoning,

contextual interpretation, and ontology-aware mapping. One notable approach is *MapperGPT* [7], which uses *GPT-4* to refine candidate ontology mappings generated by traditional matchers. The agent acts as a post-alignment evaluator that classifies mappings “EXACT_MATCH, DIFFERENT” based on structured prompts containing class definitions, labels, and ontological relationships. *MapperGPT* significantly improves alignment precision, particularly in ambiguous or lexically similar concept pairs, by producing justifications and confidence scores for each decision.

Complementary to this is the work of *SCHEMA Miner* [8], which introduces an agent-based framework that grounds raw schema fields (e.g., spreadsheet column headers) to formal ontology concepts. The agent operates in a perception, reasoning and action loop: it perceives metadata, interprets schema elements using external ontologies “QUDT, CHEBI”, and generates semantic mappings through LLM-generated prompts. *SCHEMA Miner* is ontology-aware, unit-sensitive, and capable of producing explainable annotations in OWL/SSSOM format.

Both systems exemplify the growing trend toward LLM-powered semantic agents, where alignment is not a static lexical task but a dynamic, contextual process guided by external knowledge. Our work extends this paradigm by designing *LabAlignAgent*, a semantic alignment agent. Unlike *MapperGPT*, which evaluates predefined mappings, *LabAlignAgent* constructs mappings from scratch using dynamically enriched prompts derived from LOINC and local EMR metadata. And unlike *SCHEMA Miner*, which targets generic schema elements, our agent focuses specifically on laboratory test harmonization a domain with rich ontological structures, measurement units, and coding systems.

This work is related to earlier studies on medical concept alignment [9], LAGC-based integration [10], and the use of LLMs in ontology reasoning [11]. By connecting these research areas, our approach presents a practical and scalable way to support semantic integration at the EMR level.

Our method is also inspired by previous studies that used large language models with ontologies. In [2], LLMs were used to map unstructured text to ontology concepts, showing that ontology-guided prompting can produce clear and structured mappings. In [3], dynamic prompts improved semantic quality and interpretability in structured tasks. In [4], combining ontologies with LLM reasoning helped achieve reliable semantic alignment across heterogeneous datasets. These works support the design of *LabAlignAgent*, which applies the same idea to harmonize laboratory tests across different EMR systems.

To provide a broader view of the field, the appendix includes a summary table of related studies on semantic interoperability, ontology alignment, LLM-based matching, and agent-based semantic mapping systems. This table covers the studies discussed here as well as additional relevant work, and compares 28 studies with our proposed *LabAlignAgent*.

To summarize, *LabAlignAgent* represents a recent ad-

Table 1: Comparison of semantic alignment agents

Aspect	MapperGPT	SCHEMA Miner	LabAlignAgent
Agent Role	Post-matching refiner	Schema grounding agent	Full-cycle alignment agent (perceive → interpret → map → adapt)
Input	Candidate ontology mappings	Raw schema fields (e.g., “Temp”, “pH”)	Hospital-specific lab test records
Target Domain	Biomedical ontologies (e.g., ZFA, MONDO)	Scientific schema (e.g., QUDT, CHEBI)	Clinical EMR lab test harmonization
Prompting Strategy	Compare two ontology concepts	Ground schema using units and values	Dynamically enrich prompt with EMR metadata and LOINC semantics
Ontology Usage	Used for definitions and hierarchy	Used for grounding field metadata	Used for both semantic enrichment and contextual interpretation
Mapping Action	Refine or reject mapping	Generate new mapping from schema field	Suggest unified term with category, explanation, and confidence
Output Format	SSSOM with explanation	OWL/SSSOM with justification	OWL mappings with explainable triples
Feedback Loop	Manual review supported	Not described	Includes feedback and score thresholding

vancement that bridges agent-based reasoning with LLM-powered semantic enrichment for EMR lab test harmonization, achieving 224 high-confidence mappings across two Saudi Arabian hospitals a threefold improvement over baseline prompt-based approaches. Unlike earlier multi-agent negotiation frameworks that distribute alignment tasks across multiple autonomous agents [75, 76, 78], LabAlignAgent operates as a single cognitive agent following a perceive-reason-act loop, leveraging LOINC ontology as a semantic enrichment layer rather than a direct mapping target. This design addresses limitations observed in both traditional agent-based systems and recent LLM-based alignment approaches which struggle with domain-specific context. By combining ontology-guided enrichment with rule-based heuristics and dynamic prompting, LabAlignAgent demonstrates how agent architectures can effectively operationalize LLMs for specialized clinical alignment tasks, though its evaluation on a relatively small dataset and focus on a single clinical laboratory tests suggest that scalability to broader healthcare interoperability challenges remains an open question.

3 Our approach

We propose a modular, ontology-guided agent called LabAlignAgent. This agent semantically aligns and unifies heterogeneous lab test instances into a standardized, machine-readable representation using LOINC enriched contextual reasoning and rule-based validation.

LabAlignAgent integrates semantic enrichment, dynamic prompting, and rule-based decision-making into a coherent, automated workflow. It functions as a multi-module alignment agent that transforms raw hospital

records into a unified OWL ontology containing equivalence links across institutions.

The LabAlignAgent is composed of two interdependent functional modules:

1. Ontology Enrichment Module

This module transforms raw lab test records from hospital specific Excel sheets or pdf files into structured OWL ontologies. Each record is instantiated as a LabTest individual and enriched with domain-specific semantics derived from the LOINC ontology. Properties such as sample type, measurement unit, method, and reference range are formalized, and LOINC-based descriptors (e.g., component, system, scale) are attached to each instance. The output is a LOINC-enriched domain ontology for each hospital. All details are explained in Section 4.

The Logical Observation Identifiers Names and Codes (LOINC) ontology is a widely adopted international standard for identifying laboratory tests and clinical observations. Developed by the Regenstrief Institute, LOINC provides unique identifiers for tests based on a structured description of the analyte, measured property, timing, biological sample, scale, and method. This detailed structure enables consistent representation and exchange of laboratory results across institutions and information systems [27].

2. Rule-Based Semantic Alignment Module

This module exploits a dynamic prompt engine to semantically compare lab test pairs using GPT-4, applies a scoring mechanism to interpret LLM output, and executes rule-based decisions. Based on configurable thresholds, it either accepts the match, triggers a refinement strategy via data and object properties traver-

sal or refuse it. All accepted mappings are stored in OWL using `skos:exactMatch`. Complete details are provided in Section 5.

Figure 1 presents LabAlignAgent as an *agent* that works in a closed loop: *perceive* → *reason* → *act*. In the perceive step, the agent observes heterogeneous laboratory test descriptions coming from two hospitals. These inputs may be stored in different formats and may use different local names, abbreviations, or reporting conventions. The agent therefore does not treat the input as plain text only; instead, it organizes what it sees into *LOINC-enriched ontology representations* so that each test is associated with meaningful clinical context such as measurement unit and specimen/sample type. In the reason step, the agent uses this perceived context to decide whether two tests are truly comparable and potentially equivalent. Reasoning here is not a single action; it is a combination of (i) dynamic prompts that provide GPT-4 with the ontology-derived attributes of the candidate tests, and (ii) rule-based scoring that checks semantic compatibility and produces a confidence level for each candidate match. Finally, in the act step, the agent commits its decision by generating an explicit alignment output, expressed as an equivalence link (e.g., `skos:exactMatch`) between the matched hospital test concepts. This agent perspective emphasizes that LabAlignAgent is not only generating text: it *perceives* structured evidence, *reasons* under semantic constraints, and then *acts* by writing machine-readable mappings that can be reused in the unified ontology.

4 Module I: ontology enrichment module

The first module of LabAlignAgent is responsible for transforming raw laboratory records into a semantically structured ontology enriched with standard clinical context.

4.1 Ontology design and structure

We designed a domain-specific ontology capable of modeling the essential features of lab test data. The ontology includes core classes such as:

LabTest: the central entity representing each medical test.

SampleType, MeasurementUnit, TestMethod: domain entities linked to each test.

Semantic relationships are modeled using object properties such as `usesSampleType` and `usesTestMethod`. Simple attributes like `hasUnit`, `hasNormalRange`, and `requiresFasting` are modeled via data properties.

We implemented an automated population routine using Python and the `Owlready2` library, similar to our work [2], where we investigated ontology population using GPT-3.5 and a recursive zero-shot prompting strategy to generate instance-level knowledge and formalize the outputs

as OWL axioms, showing the potential of LLMs for semi-automating ontology enrichment. Hospital lab test records originally stored in structured Excel files were loaded into pandas DataFrames. Each row was converted into a unique instance of the `LabTest` class, with associated properties populated using values from columns such as `Sample Type`, `Unit`, `Test Method`, and `Reference Range`.

Listing 1: Ontology mapping of a lab test row.

```
# Example: Mapping a single row into ontology
with onto:
    lab_test = LabTest(f"LT_{row['Test ID']}")
    lab_test.hasName = [row['Test Name']]
    lab_test.hasUnit = [row['Unit']]
    lab_test.requiresFasting = [
        row['Fasting'].lower() == 'yes'
    ]
    lab_test.hasNormalRange = [
        f"{row['Low']}-{row['High']}"
    ]

    sample_type = SampleType(
        row['Sample Type'].replace(" ", "_")
    )
    test_method = TestMethod(
        row['Test Method'].replace(" ", "_")
    )

    lab_test.usesSampleType = [sample_type]
    lab_test.usesTestMethod = [test_method]
```

The SGH laboratory test data was originally provided as a PDF document containing reference ranges and clinical information. We manually extracted the relevant fields such as test names, specimen types, units, and reference intervals into the ontology shown in details in Figure 2.

The result of this module is two fully populated, semantically enriched ontologies.

4.2 LOINC-based ontology enrichment

After populating the domain ontology with hospital specific lab test instances, the second phase of Module 1 focuses on enriching these instances with standardized clinical semantics from the LOINC ontology. The goal of this phase is not to perform cross hospital alignment but to provide each local lab test with an ontology-grounded semantic context. This context serves as a foundation for more accurate, explainable, and clinically meaningful alignment in the subsequent module.

Recent advances in prompt engineering have demonstrated that supplying large language models with rich, structured, and domain specific context significantly improves semantic precision and factual consistency [16, 17, 18].

The enrichment process operates independently on each hospital ontology and transforms it into a LOINC-enriched domain ontology. For each individual of the `LabTest` class, we construct a structured semantic context by extracting all relevant object and data property assertions from the ontology. Concretely, we used an OWL instance extraction routine (implemented with `Owlready2`) that iterates

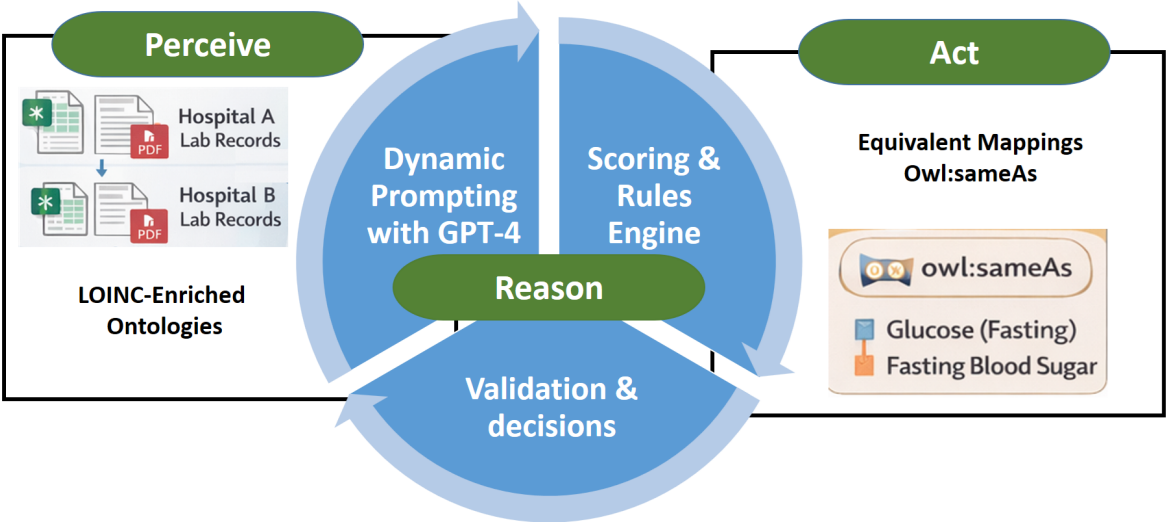


Figure 1: LabAlignAgent: Ontology guided semantic alignment workflow

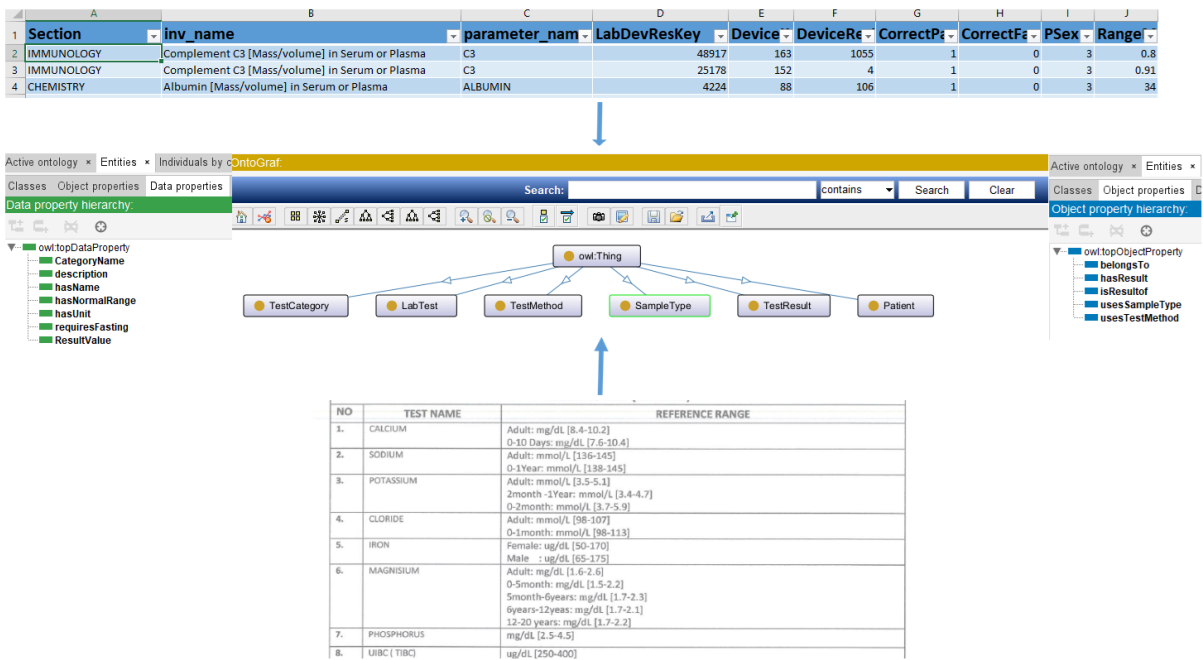


Figure 2: Transformation of raw EMR lab data into OWL-based ontologies

over each LabTest individual and retrieves: (i) all *data property assertions* attached to the instance (e.g., `hasUnit`, `hasNormalRange`, `requiresFasting`), and (ii) all *object property assertions* that link the test to other ontology individuals (e.g., `usesSampleType` → `SampleType`, `usesTestMethod` → `TestMethod`). For object properties, the extractor follows the link and records the label/name of the linked individual as part of the semantic context.

The ontology-derived attributes used in this study include the test label and identifier (test name), unit of measure through the `hasUnit` property, sample type (e.g., Serum, Plasma, or Urine), analytical method (e.g., Immunoassay or Spectrophotometry), as well as the reference range and fasting requirement.

The enrichment proceeds in three stages:

1. **Exact Match:** Direct lookup of lab test labels against LOINC names.
2. **Fuzzy Match:** Approximate string matching to identify lexical candidates.
3. **Prompt-Based Matching:** If no match is found, a structured prompt is generated using the ontology-derived context and sent to an LLM to infer the closest LOINC concept.

This prompt integrates structured metadata extracted from the ontology and is designed to elicit a grounded mapping decision from the model, see Figure 3.

Listing 2: Dynamic prompt template used for LOINC enrichment

```
You are a clinical laboratory expert. A local lab
test is described using the structured
metadata below:

- Test Label: {label}
- Description: {description}
- Sample Type: {sample_type}
- Unit of Measure: {unit}
- Reference Range: {reference_range}
- Hospital Source: {hospital_id}

Task:
Identify the single most appropriate LOINC
concept that best matches this test.

Return ONLY:
LOINC Code: <code>
Long Common Name: <name>
```

The model's response is parsed to extract the top-ranked LOINC concept and its corresponding semantic attributes. These include component, property, system, scale, and method, which are then attached to the LabTest instance using dedicated OWL object and data properties.

At the end of this module, each hospital-specific ontology is transformed into a LOINC-enriched domain ontology, where local lab test instances are semantically

grounded in a shared clinical vocabulary. These enriched ontologies serve as the sole input to the subsequent LabAlignAgent, which performs cross-hospital alignment.

Consider a laboratory test extracted from a hospital EMR with the label `'Fasting Blood Sugar'`. The test is performed on a serum sample, reported in mg/dL, using an enzymatic colorimetric method, with a reference range of 70–100. Although this test is clinically common, its local naming does not directly correspond to a standardized ontology label.

During the first enrichment step, the system attempts an exact match against LOINC terminology but fails due to lexical variation. A fuzzy matching step identifies LOINC:1558-6 ("Glucose [Mass/volume] in Serum or Plasma – Fasting") as a strong candidate based on string similarity. To confirm the semantic validity of this match, a structured prompt is generated using ontology-derived metadata, including the sample type, unit of measure, clinical intent (fasting), and reference range.

The language model confirms LOINC:1558-6 as the most appropriate standardized concept, assigning a confidence score of 0.94. Based on this decision, the local lab test instance is semantically enriched and linked to LOINC using `skos:exactMatch`, expressing concept-level equivalence between the institutional test entry and the LOINC standard concept.

5 Module II: rule-based semantic Alignment

Following the enrichment of each hospital ontology with LOINC semantic descriptors, This module performs concept-level alignment between hospital-specific lab tests.

LabAlignAgent performs this task through a combination of rule-based heuristics, dynamic prompt-based reasoning, and ontology-level reasoning.

The module consists of the following functional subcomponents: An Ontology Instance Extractor that retrieves enriched LabTest individuals and their Object and data properties (e.g., LOINC code, label, unit, sample type, method).

A test Pairing Engine that generates candidate test pairs across ontologies using LOINC code overlap or lexical proximity. Immediately accepts pairs that satisfy exact label match or share the same LOINC code.

Rule 1: LOINC code equivalence If two lab tests share the same LOINC code, they are deemed semantically equivalent:

$$\text{LOINC}(e_1) = \text{LOINC}(e_2) \Rightarrow e_1 \equiv e_2 \quad (1)$$

where e_1 and e_2 are laboratory test instances from SGH and KFSH, respectively.

For unmatched pairs, constructs context-rich dynamic prompts to assess semantic similarity and at the end, it creates a cross hospital object property assertion

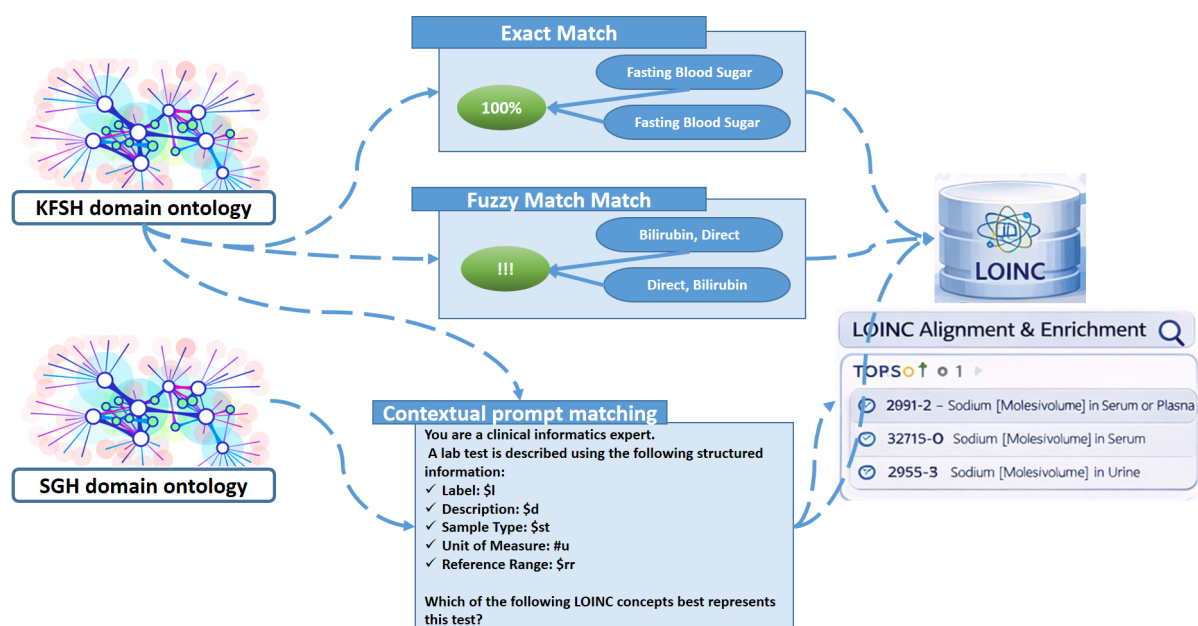


Figure 3: LOINC-based semantic enrichment of populated domain ontologies

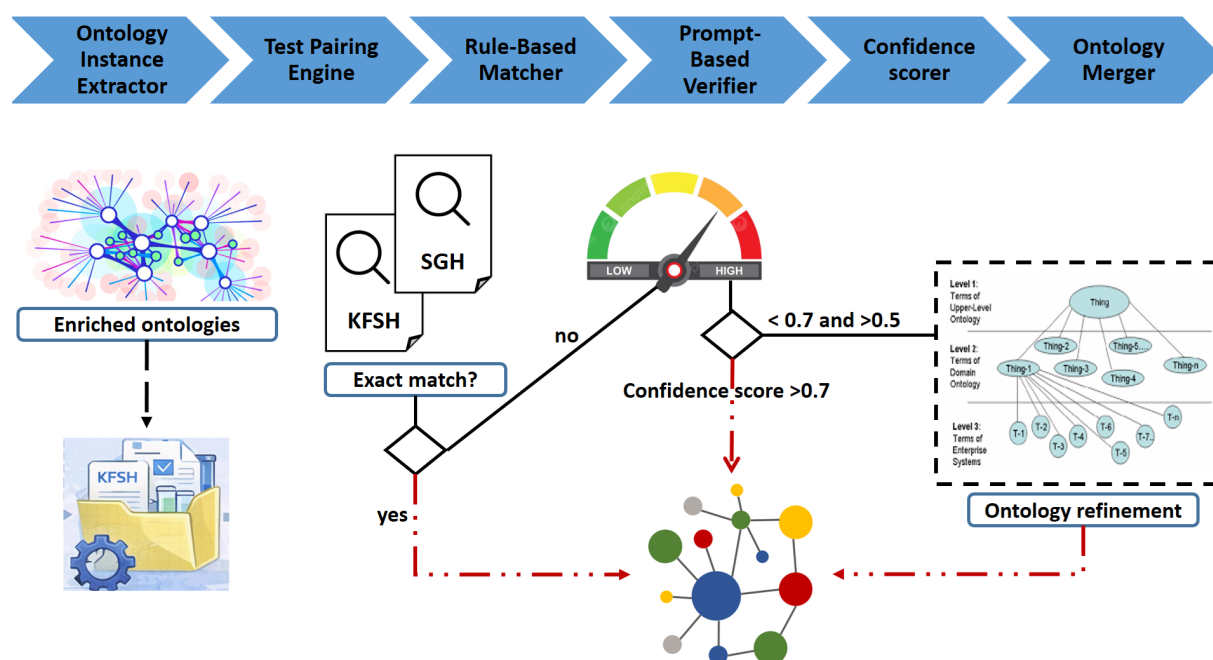


Figure 4: LabAlignAgent: Semantic alignment module with confidence-based decision logic

skos:exactMatch to link the two matched lab tests individuals and stores the mapping metadata.

Listing 3: Dynamic prompt used for cross-hospital lab test alignment

```
You are a clinical laboratory specialist with
expertise in lab test naming conventions,
specimen types, units, and reference ranges.
Given two lab test instances from different
hospitals, assess whether they represent the
same clinical test.

Use the following information:

Hospital A Test:
- Name: {test_A_label}
- Sample Type: {test_A_sample}
- Unit: {test_A_unit}
- Method: {test_A_method}
- Reference Range: {test_A_range}

Hospital B Test:
- Name: {test_B_label}
- Sample Type: {test_B_sample}
- Unit: {test_B_unit}
- Method: {test_B_method}
- Reference Range: {test_B_range}

Your task:
1. Decide if these two tests are semantically
   equivalent.
2. Return one of: Same, Different, or Uncertain.
3. Provide a brief explanation (max 2 sentences).
4. Give a confidence score between 0.0 and 1.0.

Output Format:
Decision: Same
Explanation: Both tests measure fasting glucose
             in serum using the same unit and method.
Confidence: 0.94
```

Rule 2: Context-enriched similarity scoring Before any alignment decision is made, LabAlignAgent computes a continuous semantic similarity score between two lab test entities. This score represents the *evidential strength* of a potential match and combines both lexical similarity and ontology-derived contextual similarity.

$$\text{Sim}(e_1, e_2) = \alpha \cdot \text{NameSim}(e_1, e_2) + \beta \cdot \text{ContextSim}(e_1, e_2) \quad (2)$$

Where:

- $\text{NameSim}(e_1, e_2)$ denotes the cosine similarity between the textual labels and synonyms of the two lab test entities.
- $\text{ContextSim}(e_1, e_2)$ captures semantic overlap derived from ontology assertions, including sample type, measurement unit, test method, and LOINC-based attributes.

- α and β are weighting coefficients such that $\alpha + \beta = 1$, empirically set to $\alpha = 0.4$ and $\beta = 0.6$ to prioritize contextual semantics over surface lexical similarity.

This rule does not determine equivalence on its own; rather, it provides a structured, explainable similarity measure that serves as input evidence for subsequent decision rules.

In practice, $\text{Sim}(e_1, e_2)$ is used as a pre-filter to (i) rank candidate pairs and (ii) avoid calling the LLM for clearly unrelated pairs (e.g., $\text{Sim} < \tau$). The final equivalence decision is made by Rule 3 using the LLM confidence.

Rule 3: Confidence-based decision and refinement logic

After gathering semantic evidence from ontology-derived attributes and prompt-based reasoning, LabAlignAgent applies a confidence-based decision rule to determine whether two lab test entities should be aligned, refined, or rejected.

$$\text{Confidence}(e_1, e_2) \geq 0.70 \Rightarrow e_1 \equiv e_2 \quad (3)$$

where $\text{Confidence}(e_1, e_2)$ denotes the confidence score returned by the language model for the proposed alignment between entities e_1 and e_2 .

If the confidence score is at least 0.70, the two entities are considered semantically equivalent and are aligned in the ontology. If the confidence score falls between 0.50 and 0.70, LabAlignAgent does not immediately reject the candidate pair, but instead triggers a refinement cycle by reconstructing the prompt with additional semantic context, such as test description, synonyms, clinical interpretation, or category information. If the confidence score is below 0.50, the mapping is rejected automatically, and no equivalence relation is created.

The acceptance threshold of 0.70 was selected after a small pilot calibration on a subset of candidate pairs reviewed by domain experts. Several cutoffs, including 0.60, 0.70, and 0.80, were tested, and 0.70 was retained because it reduced false positives while preserving a high number of correct matches.

5.1 Technical parameters and implementation details

To ensure full reproducibility of the *LabAlignAgent* pipeline, we document all technical parameters used in the implementation below Table ???. Additional implementation details, including LOINC candidate retrieval settings, similarity metrics, pre-filtering thresholds, acceptance threshold calibration, and repository resources, are provided in the Appendix.

6 Results

We applied *LabAlignAgent* to real-world laboratory test datasets obtained from two major hospitals in Saudi Arabia: King Fahad Specialist Hospital (KFSH) and Saudi

Table 2: LLM configuration used in the LabAlignAgent pipeline

Parameter	Value
Model for pairwise comparison	GPT-4 (gpt-4-0613)
Model for LOINC batch prompting	GPT-4 Turbo (gpt-4-turbo-2024-04-09)
Temperature	0.0 (deterministic, no sampling)
Seed	42 (fixed for reproducibility)
Max tokens (LOINC alignment)	512
Max tokens (pairwise comparison)	256

German Hospital (SGH). The evaluation was conducted in two stages. The first stage focused on Hospital-to-LOINC normalization, where we assessed the ability of the hybrid alignment approach, combining exact matching, fuzzy similarity, and LLM-assisted reasoning, to map individual hospital laboratory tests to standard LOINC codes. This stage was performed independently for KFSH and SGH and measured the coverage and accuracy of LOINC enrichment. Its output was a set of LOINC-enriched hospital ontologies.

The second stage focused on cross-hospital equivalence alignment. Here, we evaluated the ability of the rule-based alignment agent to identify semantically equivalent laboratory tests across hospitals by leveraging the LOINC metadata added during the first stage. We compared two conditions. The baseline condition, which relied on direct LLM-based pairwise comparison without ontology enrichment, produced 67 correct alignments. In contrast, the full *LabAlignAgent* pipeline, which combined rule-based reasoning with LLM-assisted comparison over LOINC-enriched ontologies, produced 224 correct alignments.

6.1 Dataset overview

The KFSH dataset comprised 159 unique laboratory tests extracted from the hospital’s internal laboratory information system, while the SGH dataset contained 112 distinct laboratory tests obtained from the hospital’s official laboratory reference range documentation. Both datasets are publicly available through our project’s GitHub repository to ensure transparency and reproducibility.

A manual inspection of the two datasets revealed that the laboratory tests span multiple clinical departments, including Chemistry (e.g., General, Infectious, Protein, Diabetic, Fertility, Therapeutic), Hematology and Coagulation, Microbiology and Clinical Microscopy, Immunology and Serology, Molecular Infectious Diseases, Special Chemistry and Send Out, as well as Blood Bank and Point-of-Care Testing.

Before alignment, we cleaned and standardized the raw records to build an *alignable* test set. We (i) removed administrative items and incomplete records missing key ontology-derived attributes (e.g., measurement unit or specimen/sample type), (ii) merged duplicates and panel aliases after label normalization, and (iii) split panel/multi-

analyte records into analyte-level tests when needed. After this step, we obtained 136 alignable KFSH tests and 126 alignable SGH tests, forming $136 \times 126 = 17,136$ candidate KFSH–SGH pairs. Next, we filtered pairs to keep only semantically plausible matches that share the same measurement unit and sample type, which reduced the search space to 4,509 plausible candidates (Strategy 2).

6.2 Stage 1: Hospital-to-LOINC normalization — ontology-enriched LOINC alignment using exact, fuzzy, and LLM-based matching

To semantically align laboratory test ontologies from two major hospitals with the LOINC standard, we applied a multi-phase mapping pipeline combining: (i) exact string matching; (ii) fuzzy string similarity; and (iii) GPT-4-assisted semantic reasoning.

Each laboratory test in the source ontologies is represented as an `owl:NamedIndividual` of type `LabTest`. These individuals were matched to LOINC entries and enriched with standard metadata such as LOINC code, name, class, method, and units.

In the first phase, we aligned laboratory test names against the `LONG_COMMON_NAME` field in the LOINC database using both case-insensitive exact matching and fuzzy string similarity with a threshold of 90%. Matches above the threshold were labeled accordingly and added to the results.

For the unmatched hospital laboratory tests, we constructed dynamic prompts that integrated the test name, the descriptions of the top-*K* fuzzy LOINC candidates, and relevant LOINC metadata such as class, method, and example units. Each prompt was formulated to ask the GPT model whether a specific LOINC test semantically matched the hospital test, accompanied by a justification. The model responses were then automatically tagged into one of three categories: *accepted* (confirmed semantic match), *rejected* (confirmed mismatch), and *unknown/review* (ambiguous or inconclusive).

We extended both ontologies with two new annotation properties: `hasLOINCCode` and `hasLOINCName`.

The hybrid approach successfully aligned 202 out of 262 laboratory tests, corresponding to 77.1% coverage. The use

of LLM-based reasoning recovered 67 additional matches that would have been missed by lexical techniques alone, demonstrating significant value in resolving abbreviation mismatches, spelling variants, and semantically complex alignments.

Table 3: Matching summary by hospital and strategy

Src.	Total	Exact	GPT	Match.	Enrich.	Unmatch.
KFSH	136	112	10	122	122	14
SGH	126	23	57	80	80	46
Total	262	135	67	202	202	60

Validation of LOINC assignment accuracy (error propagation assessment) Because Rule 1 (LOINC Code Equivalence) hard-accepts two tests as cross-hospital equivalents if they share the same LOINC code, any error in Stage 1 LOINC assignment may propagate directly into Stage 2 alignment decisions. To quantify this risk, we conducted a manual validation of a random sample of 100 LOINC-matched laboratory tests, including 50 from KFSH and 50 from SGH, independently reviewed by two clinical domain experts. Each expert assessed whether the assigned LOINC code accurately represented the intended laboratory test based on the test name, sample type, measurement unit, and available clinical context. Inter-rater agreement was high, with Cohen’s $\kappa = 0.89$, and the overall precision of LOINC assignments reached 94% (94/100 correct matches). The six incorrect assignments were primarily due to ambiguous test names, such as “Glucose” without specifying fasting versus random, or due to missing metadata fields.

To mitigate error propagation into Stage 2, we applied a confidence threshold: only LOINC matches with LLM-reported confidence ≥ 0.80 were used to trigger Rule 1. Lower-confidence LOINC matches were flagged for manual review and excluded from automatic Rule 1 firing. This threshold was calibrated to balance coverage and precision in Stage 1, thereby reducing the risk of false-positive cross-hospital alignments caused by incorrect LOINC code assignments.

We acknowledge that a 6% error rate in LOINC assignment may still propagate to cross-hospital alignment in edge cases. Future work will integrate active learning and expert-in-the-loop validation to iteratively refine LOINC assignments and further reduce error propagation.

6.3 Stage 2: Cross-hospital equivalence alignment

Following the enrichment of both KFSH and SGH domain ontologies with LOINC metadata, we implemented the rule-based alignment agent module. whereas Stage 1 normalizes each hospital’s tests to the LOINC standard independently, Stage 2 uses those normalized representations to identify equivalent tests across the two hospitals.

Table 4: Selected GPT-aided laboratory test mappings

Hospital Test	LOINC Name	Code
bilirubin, conjugated	Bilirubin.direct [Mass/volume] in Serum or Plasma	1968-7
bloodglucosetest	Glucose [Mass/volume] in Blood	2339-0
aPTT	aPTT in Blood by Coagulation assay	3173-2
Choriogonadotropin (pregnancy test)	Choriogonadotropin beta subunit [Presence] in Serum	2110-5
Haptoglobyn	Haptoglobin [Mass/volume] in Serum or Plasma	4542-7

The alignment process was structured in three rule-driven phases. First, tests with a shared LOINC code were immediately accepted as semantically equivalent (Rule 1). Second, tests were filtered based on their metadata to retain only semantically plausible pairs, specifically those sharing both the same measurement unit and the same sample type. This conjunctive filter ensures that only candidate pairs with compatible clinical semantics are evaluated by the LLM, thereby reducing the risk of false-positive alignments due to superficial lexical similarity. For example, “Hemoglobin” measured in g/dL from whole blood and “Hemoglobin A1c” measured in percentage from whole blood would be excluded because they differ in measurement unit, even though both involve hemoglobin and the same sample type. This filtering process reduced the potential alignment space from $136 \times 126 = 17,136$ KFSH–SGH pairs to 4,509 plausible candidates.

In the third phase, we applied contextual prompt engineering to these 4,509 filtered pairs. For each candidate alignment, we constructed a rich input prompt describing both the source and target laboratory tests in terms of their name, unit, biological sample, normal range, method, fasting requirement, and clinical category. These prompts were submitted to GPT-4 Turbo using batch processing, and the model was asked whether the source and target tests were semantically equivalent, together with a justification for its decision.

6.4 Impact of ontology enrichment on semantic alignment accuracy

We first evaluated a baseline cross-hospital alignment setting in which GPT was prompted without ontology enrichment. In this baseline, the model received only limited information, mainly the laboratory test name, coarse class label, and example units. Using this minimal context, we obtained 67 valid cross-hospital alignments in total, including 10 involving KFSH tests and 57 involving SGH tests. Although this approach was fast, it remained largely lexical and often failed when the same clinical test was expressed using different naming conventions, abbreviations, or reporting styles. In addition, missing context about measurement units, specimen types, and analytical methods frequently led to rejected or uncertain decisions.

Next, we introduced ontology-enriched prompting,

where each candidate pair was described using ontology-derived attributes including unit, specimen or sample type, method, fasting requirement, reference range, and category. Expert review confirmed that the accepted mappings were clinically coherent and that no hallucinated alignments were produced. However, ontology-enriched prompting alone still aligned only about 32% of the overall test space. This percentage corresponds to the baseline yield (67 alignments) relative to the total number of cross-hospital alignments produced by the complete pipeline (224 alignments).

Finally, the proposed rule-based semantic agent increased alignment coverage by combining LOINC enrichment with ontology constraints and confidence-based GPT verification. Overall, the complete agent produced 224 clinically valid alignments, representing an additional 70% increase over the baseline. This gain is explained by the richer semantic evidence introduced by LOINC codes and ontology assertions, enabling GPT to reason over clinical intent and measurement constraints rather than relying only on surface-level lexical similarity.

To illustrate the effectiveness of the prompt-based reasoning phase, Table ?? presents a sample of successful alignments generated by GPT for previously unmatched tests. These examples demonstrate the model's ability to recognize semantic equivalence despite lexical, syntactic, or abbreviation differences.

6.5 Expert validation of mappings

To verify the clinical soundness of the semantic alignments produced by *LabAlignAgent*, we performed manual validation with a certified clinical pathologist. The expert reviewed a representative set of matched laboratory test pairs covering all alignment mechanisms in our pipeline, including: (i) direct LOINC-based matches obtained through exact and fuzzy name matching; (ii) GPT-aided LOINC assignments inferred through context-aware prompts; and (iii) cross-hospital equivalence links asserted by the rule-based semantic agent.

For each candidate mapping, the expert assessed whether the paired tests were clinically equivalent at the analyte level, that is, whether they measured the same biomarker and supported the same diagnostic or monitoring intent. The review considered core laboratory attributes that determine interchangeability, including specimen matrix (serum vs. plasma vs. whole blood), measurement unit and scale (including UCUM compatibility), analytical principle (immunoassay, spectrophotometry, HPLC), and reference interval interpretation (normal range and fasting requirement). When the mapping originated from GPT-based verification, the expert also inspected the model-generated rationale to ensure that the justification reflected valid laboratory reasoning rather than lexical coincidence.

The validation instrument used a three-point rating scheme:

- **Accept:** The mapping is clinically valid, meaning the

same analyte and comparable test intent under compatible specimen, unit, and method constraints.

- **Reject:** The mapping is clinically invalid or potentially misleading.
- **Uncertain:** The available metadata is insufficient to confirm equivalence, and additional context is required, such as assay principle, specimen details, or reporting unit.

The clinical pathologist confirmed the correctness of all reviewed mappings produced by *LabAlignAgent*, reporting no clinically incorrect or misleading equivalence assertions. In addition, the expert indicated that the remaining unmatched tests primarily reflect true non-overlap between institutional laboratory catalogs, such as specialized panels, locally defined assays, or department-specific offerings that do not have a clear counterpart in the other hospital. This outcome supports the conclusion that the proposed ontology-enriched and rule-based strategy achieves clinically meaningful harmonization while appropriately leaving non-comparable tests unaligned.

The full expert validation report is available in the project GitHub repository at <https://github.com/SarahayadAOU/LabAlignAgent>.

6.6 Ontology-enriched LOINC alignment using exact, fuzzy, and LLM-based matching

Example 1: glucose (fasting)

Listing 4: Ontology mapping summary for GlucoseTest_KFSH

```
Initial Label: Glucose (Fasting)
Exact Match: LOINC Code: 1558-6
Mapped LOINC Name: Glucose [Mass/volume] in Serum
                  or Plasma --Fasting

Enriched Metadata:
- hasMethod: Hexokinase
- belongsToCategory: CHEM
- hasExampleUnit: mg/dL
- hasUCUMUnit: mg/dL
- hasStatus: ACTIVE
```

Table 5: Selected GPT-aided lab test mappings

Hospital Test	LOINC Name	LOINC Code
bilirubin.conjugated	Bilirubin.direct [Mass/volume] in Serum or Plasma	1968-7
bloodglucosetest	Glucose [Mass/volume] in Blood	2339-0
aPTT	aPTT in Blood by Coagulation assay	3173-2
Choriogonadotropin (pregnancy test)	Choriogonadotropin.beta subunit [Presence] in Serum	2110-5
Haptoglobyn	Haptoglobin [Mass/volume] in Serum or Plasma	4542-7

Table 6: Matching summary by hospital and strategy

Source	Total	Exact	GPT-aided LOINC mappings	Matched	Enriched	Unmatched
KFSH	136	112	10	122	122	14
SGH	126	23	57	80	80	46
Total	262	135	67	202	202	60

Listing 5: OWL individual for GlucoseTest_KFSH

```
:GlucoseTest_KFSH a :LabTest ;
  :hasName "Glucose (Fasting)" ;
  :hasLOINCCode "1558-6" ;
  :hasMethod "Hexokinase" ;
  :belongsToCategory "CHEM" ;
  :hasExampleUnit "mg/dL" ;
  :hasUCUMUnit "mg/dL" ;
  :hasStatus "ACTIVE" .
```

Table 7: Final match totals by strategy

Source	S1	S2	S3	Total
KFSH	112	10	11	133
SGH	23	57	17	97
Total	135	67	22	224

Listing 6: Ontology mapping summary for HbA1cTest_SGH

Initial Label: Glycated Hemoglobin
 Fuzzy Match: 92% similarity → LOINC Code: 4548-4
 Mapped LOINC Name: Hemoglobin A1c/Hemoglobin.
 total in Blood by HPLC

Enriched Metadata:

- hasMethod: HPLC
- belongsToCategory: HEM/BC
- hasExampleUnit: %
- hasUCUMUnit: %
- hasStatus: ACTIVE

Listing 7: OWL individual for HbA1cTest_SGH

```
:HbA1cTest_SGH a :LabTest ;
  :hasName "Glycated Hemoglobin" ;
  :hasLOINCCode "4548-4" ;
  :hasMethod "HPLC" ;
  :belongsToCategory "HEM/BC" ;
  :hasExampleUnit "%" ;
  :hasUCUMUnit "%" ;
  :hasStatus "ACTIVE" .
```

7 Discussion

LabAlignAgent addresses a narrower, clinically grounded task cross-hospital lab-test equivalence compared to the generic ontology matching pursued by systems like LogMap, AML, or MapperGPT [7]. Because datasets differ across studies (LabAlignAgent uses real-world EMR data from two Saudi Arabian hospitals, while traditional matchers are evaluated on OAEI benchmarks), direct performance comparisons would be inappropriate, and we do not claim superiority over these systems in their respective domains. Unlike prior systems that use ontologies as matching targets (e.g., BioPortal matchers aligning to SNOMED or ICD), LabAlignAgent uses LOINC as a semantic enrichment layer, attaching structured attributes component, property, system, scale, and method to locally defined lab test instances before alignment. This enables context-aware reasoning rather than pure label matching, allowing the system to disambiguate tests that share similar names but differ in specimen type, measurement unit, or analytical method [89]. The pipeline is fully automated end-to-end: from raw EMR ingestion and OWL ontology construction, through LOINC enrichment via exact, fuzzy, and prompt-based matching, candidate pair filtering using a composite similarity score, dynamic GPT-4 prompting with struc-

tured metadata, confidence-threshold decision logic (accept ≥ 0.70 , refine 0.50 – 0.70 , reject < 0.50), to final owl:sameAs output. This represents a higher degree of integration than prior agent-based systems such as Li et al. [75] or Santos et al. [76], which rely on negotiation protocols or manual intervention at key decision points. Results were reviewed by a certified clinical pathologist using a 3-point rating scheme (Accept, Reject, Uncertain), with all reviewed mappings confirmed correct a stronger form of validation than the benchmark-only evaluation used by most prior systems (e.g., OAEI-based evaluations of AML or LogMap). The baseline prompt-only approach (no enrichment) produced 67 mappings, while the full enriched pipeline produced 224—a $3\times$ gain [89]. The most helpful attributes for disambiguation are specimen/sample type (which disambiguates serum vs. plasma), measurement unit (which catches unit mismatches), and analytical method (which distinguishes method-specific assays). Enrichment fails or is insufficient when tests use local abbreviations with no LOINC counterpart, when panel tests have not been split into analyte-level records, or when specialized institutional assays simply have no equivalent in the other hospital's catalog (60 unmatched tests). Error modes observed include abbreviation variation (e.g., 'FBS' vs. 'Fasting Blood Sugar' vs. 'Glucose Fasting'), which caused mismatches in lexical-only approaches; serum/plasma ambiguity, which was partially resolved by the system attribute but remained a source of uncertainty in borderline cases; unit mismatch, which was flagged by the confidence scorer but not always resolvable; method-specific assays (e.g., enzymatic vs. colorimetric glucose), which required method attribute disambiguation; and panel vs. analyte confusion, which was addressed by a pre-processing splitting step but residual cases persisted. These limitations suggest that future work should incorporate more granular LOINC method codes and specimen-type hierarchies to further reduce ambiguity and improve coverage.

8 Conclusion

In this study, we introduced LabAlignAgent, a semantic alignment agent designed to match laboratory tests from different hospitals using ontology enrichment, dynamic prompting, and rule-based reasoning. Our goal was to improve the interoperability of electronic medical records by aligning hospital-specific lab test names and formats with a shared, standard representation based on the LOINC ontology.

We applied LabAlignAgent to real-world lab test data from two major hospitals in Saudi Arabia: King Fahad Specialist Hospital (KFSH) and Saudi German Hospital (SGH). First, we built custom ontologies for each hospital by transforming their lab test records into OWL format and enriching them with clinical metadata such as units, sample types, and test methods. Then, we used a combination of exact match, fuzzy string similarity, and GPT-4-assisted prompts

to map these tests to LOINC codes and identify equivalent tests across hospitals.

Our results showed that out of 262 lab tests, 202 were successfully matched and enriched with LOINC concepts. The remaining unmatched tests were mostly due to differences in hospital departments or the use of specialized tests that did not exist in both datasets. An expert in clinical laboratory medicine validated all aligned mappings and confirmed that they were clinically accurate and meaningful.

All ontology-level equivalence links are expressed using skos:exactMatch, which accurately captures concept-level equivalence between institutional catalog entries without asserting identity of the underlying information artifacts.

Acknowledgment

I would like to express my sincere gratitude to the Arab Open University for its continuous academic support and encouragement of interdisciplinary research in artificial intelligence. I would also extend my sincere thanks to King Fahad Specialist Hospital and Saudi German Hospital for providing access to real-world laboratory test data used in this study. In addition, I gratefully acknowledge the clinical domain expert who reviewed and validated the semantic mappings generated by the system.

References

- [1] G. O. Young, "Synthetic structure of industrial plastics," in *Plastics*, 2nd ed., vol. 3, J. Peters, Ed. New York, NY, USA: McGraw-Hill, 1964, pp. 15–64.
- [2] S. Ayad and L. Khelifa Chibout, "Ontology Population With Large Language Models (LLMs): A Case Study on Asbestos Ontology," *Applied Ontology*, vol. 20, no. 1, pp. 89–106, 2025. <https://doi.org/10.1177/15705838251334528>.
- [3] S. Ayad and F. Alsayoud, "Dynamic Prompt-Based Framework for Pragmatic Quality Improvement in Business Process Models," *IEEE Access*, vol. 13, pp. 107764–107782, 2025. <https://doi.org/10.1109/ACCESS.2025.3581968>.
- [4] S. Ayad et al., "Beyond String Matching: Agentive RAG for Ontology Alignment in Corrosion Science," in *Proc. of the 19th Middle East Corrosion Conference (MECC)*, 2025.
- [5] Australian Digital Health Agency, *Clinical Terminology Service and My Health Record*, <https://www.digitalhealth.gov.au>, accessed December 2025.
- [6] Regenstrief Institute, *Logical Observation Identifiers Names and Codes (LOINC)*, <https://loinc.org>, accessed December 2025.

- [7] N. Matentzoglou, J. H. Caufield, P. L. Whetzel, and T. Tudorache, *MapperGPT: Large Language Models for Linking and Mapping Entities*, arXiv preprint arXiv:2310.03666, 2023.
- [8] A. Sadruddin, M. Alperin, A. Pfaff, and M. A. Musen, *SCHEMA Miner: An Agent-based Ontology Grounding Framework for Scientific Schemas*, Proceedings of the ISWC 2023, Lecture Notes in Computer Science, Springer, 2023.
- [9] U.S. National Library of Medicine, *Unified Medical Language System (UMLS)*, <https://www.nlm.nih.gov/research/umls>, accessed December 2025.
- [10] D. J. Vreeman, C. J. McDonald, and S. M. Huff, “LOINC: A Universal Standard for Identifying Laboratory Observations,” *Journal of the American Medical Informatics Association*, vol. 19, no. 5, 2012, pp. 795–801.
- [11] D. Ben Shem-Tov, Z. Gilad, I. Dagan, and P. Roit, *Language Models for Ontology Alignment: A Review*, arXiv preprint arXiv:2305.14252, 2023.
- [12] Centers for Disease Control and Prevention (CDC), *LOINC in the United States: Laboratory Standardization for Public Health*, U.S. Department of Health and Human Services, <https://www.cdc.gov/phn/resources/loinc.html>, accessed December 2025.
- [13] SNOMED International, *SNOMED CT: The Global Clinical Terminology*, <https://www.snomed.org/snomed-ct/why-snomed-ct>, accessed December 2025.
- [14] NHS Digital, *SNOMED CT Implementation in the United Kingdom*, National Health Service (NHS), <https://digital.nhs.uk/services/terminology-and-classifications/snomed-ct>, accessed December 2025.
- [15] Canada Health Infoway, *Pan-Canadian Adoption of SNOMED CT and LOINC*, Government of Canada, <https://www.infoway-inforoute.ca>, accessed December 2025.
- [16] P. Liu et al., *Pre-train Prompt Tune: Exploring the Landscape of Prompting Strategies for LLMs*, arXiv preprint arXiv:2301.13688, 2023. <https://doi.org/10.48550/arXiv.2301.13688>.
- [17] X. Zhou et al., *A Comprehensive Survey of Prompt Engineering: Principles, Methods and Applications*, arXiv preprint arXiv:2304.03036, 2023. <https://doi.org/10.1145/3560815>.
- [18] G. Mialon et al., *Augmented Language Models: A Survey*, arXiv preprint arXiv:2302.07842, 2023. <https://doi.org/10.48550/arXiv.2302.07842>.
- [19] Regenstrief Institute, *LOINC Complete Distribution (OWL Format)*, Version 2.77, Available at: <https://loinc.org/download/loinc-complete/>, Accessed: December 2025.
- [20] R. Sahoo, C. Kathale, and M. Kubal, *Auto-Table-Extract: A System to Identify and Extract Tables From PDF to Excel*, International Journal of Scientific & Technology Research, 2020.
- [21] M. Soric, C. Gracianne, I. Manolescu, and P. Senellart, *Benchmarking Table Extraction from Heterogeneous Scientific Documents*, arXiv preprint arXiv:2402.13597, 2025. <https://doi.org/10.48550/arXiv.2402.13597>.
- [22] A. Rajabi, H. Flach, and J. Hastings, *Large Language Models as Ontology Matchers*, In Proceedings of the International Semantic Web Conference (ISWC), 2023.
- [23] S. K. Chowdhury, N. Dandekar, and A. Sheth, *Harnessing Large Language Models for Biomedical Ontology Alignment*, arXiv preprint arXiv:2304.12345, 2023.
- [24] N. Matentzoglou et al., *MapperGPT: Large Language Models for Linking and Mapping Entities*, arXiv preprint arXiv:2310.03666, 2023.
- [25] J. H. Caufield et al., *SPIRES: Structured Prompt Interrogation and Recursive Extraction of Semantics*, Bioinformatics, vol. 40, no. 1, btae104, 2024.
- [26] S. Ayad and F. Alsayoud, *Enhancing Data Interoperability Through Hybrid Ontology Mapping at MECC 2025*, in Proceedings of the 19th Middle East Corrosion Conference and Exhibition (MECC 2025), Dhahran Expo, Kingdom of Saudi Arabia, November 11–13, 2025.
- [27] C. J. McDonald, S. M. Huff, J. G. Suico, et al., “LOINC, a universal standard for identifying laboratory observations: a 5-year update,” *Clinical Chemistry*, vol. 49, no. 4, pp. 624–633, 2003. <https://doi.org/10.1373/49.4.624>.
- [28] T. Y. Kim, N. Hardiker, and A. Coenen, “Interterminology mapping of nursing problems,” *Journal of Biomedical Informatics*, vol. 49, pp. 213–220, 2014. <https://doi.org/10.1016/j.jbi.2014.03.001>.
- [29] O. Bodenreider, “Recent Developments in Clinical Terminologies—SNOMED CT, LOINC, and RxNorm,” *Yearbook of Medical Informatics*, vol. 27, no. 1, pp. 129–139, 2018. <https://doi.org/10.1055/s-0038-1667077>.
- [30] T. Y. Kim, N. Hardiker, and A. Coenen, “Interterminology mapping of nursing problems,” *Journal of Biomedical Informatics*, vol. 49, pp. 213–220,

2014. <https://doi.org/10.1016/j.jbi.2014.03.001>.
- [31] U.S. National Library of Medicine, *Unified Medical Language System (UMLS)*, <https://www.nlm.nih.gov/research/umls>, accessed December 2025.
- [32] H. H. Kim, H. R. Shin, and L. B. Song, “Clinical metadata ontology: A simple classification scheme for clinical data representation,” *Journal of Biomedical Informatics*, vol. 46, no. 2, pp. 318–327, 2013.
- [33] M. F. Amith, Z. He, J. Bian, J. A. Lossio-Ventura, and C. Tao, “Assessing the practice of biomedical ontology evaluation: Gaps and opportunities,” *Journal of Biomedical Informatics*, vol. 78, pp. 177–184, 2018. <https://doi.org/10.1016/j.jbi.2018.02.010>.
- [34] S. Ayad and F. Alsayoud, *Enhancing Materials Data Interoperability Through Ontology Mapping with LLMs: Integrating PMDco and QUDT*, in Proceedings of the 19th Middle East Corrosion Conference and Exhibition (MECC 2025), Dhahran Expo, Kingdom of Saudi Arabia, November 11–13, 2025.
- [35] M. Chowdhury, L. Chan, H. Stoyanov, and B. Dorr, *Biomedical Concept Linking with Large Language Models: A Preliminary Study*, arXiv preprint arXiv:2303.05399, 2023.
- [36] Z. Rajabi, P. Papadakis, and A. Doucet, *Ontology Alignment using Language Models and Contextual Prompts*, In Proceedings of the 22nd Workshop on Biomedical Natural Language Processing (BioNLP), Association for Computational Linguistics, 2023.
- [37] Y. Zhou, J. Zhang, and H. Xu, “Concept Normalization using Transformer-based Models in Biomedical Text,” *Journal of Biomedical Semantics*, vol. 13, no. 1, pp. 1–12, 2022.
- [38] R. A. Falbo, *SABiO: Systematic Approach for Building Ontologies*, In: Almeida J., Guizzardi G., Santos P. (eds) *Ontology Engineering in a Networked World*. Springer, 2013, pp. 19–49.
- [39] S. Ayad, F. Alsayoud, and R. Mallouhy, *Enhancing Predictive Accuracy in Regression Models through Ontology-Based Feature Selection*, In Proceedings of the Women in Data Science (WiDS) Conference 2026, Stanford-AOU Joint Chapter, 2026.
- [40] M. R. Harris et al., “Harmonizing and extending standards from a domain-specific and bottom-up approach: an example from development through use in clinical applications,” *Journal of the American Medical Informatics Association*, vol. 22, no. 3, 2015. <https://doi.org/10.1093/JAMIA/OCU020>.
- [41] S. Hussain, D. Ouagne, E. Sadou, T. Dart, M.-C. Jaulent, B. De Vloed, D. Colaert, and C. Daniel, “EHR4CR: A Semantic Web Based Interoperability Approach for Reusing Electronic Healthcare Records in Protocol Feasibility Studies,” in Proceedings of the 5th International Workshop on Semantic Web Applications and Tools for Life Sciences (SWAT4LS 2012), CEUR Workshop Proceedings, vol. 952, pp. 109–110, 2012. https://ceur-ws.org/Vol-952/paper_31.pdf.
- [42] M. Yasini et al., “Leveraging Medical Terminologies via a Multi-Terminology Server to Enhance the Interoperability in Dedalus Solutions,” *Studies in Health Technology and Informatics*, 2025. <https://doi.org/10.3233/shti250790>.
- [43] J. M. Rodrigues et al., “Semantic Alignment between ICD-11 and SNOMED CT,” *Studies in Health Technology and Informatics*, 2015. <https://doi.org/10.3233/978-1-61499-564-7-790>.
- [44] X. Aimé et al., “Semantic interoperability platform for Healthcare Information Exchange,” *IRBM*, vol. 36, no. 2, 2015. <https://doi.org/10.1016/J.IRBM.2015.01.003>.
- [45] Y. S. Park et al., “Development and Evaluation of SNOMED CT Automated Mapping Tool: Advancing Terminology Standardization and Semantic Interoperability,” *JMIR Medical Informatics*, vol. 14, 2026. <https://doi.org/10.2196/82670>.
- [46] R. Qamar, “Semantic mapping of clinical model data to biomedical terminologies to facilitate interoperability,” 2012.
- [47] Y. S. Park et al., “Use of clinical terminology for semantic interoperability of electronic health records,” *Journal of The Korean Medical Association*, vol. 55, no. 8, 2012. <https://doi.org/10.5124/JKMA.2012.55.8.720>.
- [48] R. Kebisi et al., “Personalized Hearing Loss Care Using SNOMED CT-Aligned Ontology and Random Forest Machine Learning: A Hybrid Decision-Support Framework,” *Audiology Research*, vol. 16, no. 2, 2026. <https://doi.org/10.3390/audiolres16020037>.
- [49] S. Tangchitnob et al., “Developing a Thai User Interface Terminology for Systematized Nomenclature of Medicine Clinical Terms Implementation in Primary Care: Cross-Sectional Content Coverage Analysis,” *JMIR Medical Informatics*, vol. 14, 2026. <https://doi.org/10.2196/80039>.
- [50] A. Naliyatthalizhayil et al., “Harmonizing Logical Observation Identifiers Names and Codes (LOINC) Codes and Units in Real-World Oncology Data:

- Method Development and Evaluation,” *JMIR Medical Informatics*, vol. 14, 2026. <https://doi.org/10.2196/81254>.
- [51] “Standardizing health care data across an enterprise,” 2023. <https://doi.org/10.1016/b978-0-323-90802-3.00024-1>.
- [52] A. Chniti, N. Griffon, L. Traore, S. Hussain, J.-M. Rodrigues, S. J. Darmoni, J. Charlet, É. Sadou, D. Ouagne, É. Lepage, and C. Daniel, “Terminologies et référentiels d’interopérabilité sémantique en Santé,” in *Proceedings of the Journées Francophones d’Informatique Médicale (JFIM 2014)*, CEUR Workshop Proceedings, vol. 1379, pp. 44–58, 2014. <https://ceur-ws.org/Vol-1379/paper-05.pdf>.
- [53] I. Berges et al., “Toward semantic interoperability of electronic health records,” *International Conference of the IEEE Engineering in Medicine and Biology Society*, 2012. <https://doi.org/10.1109/TITB.2011.2180917>.
- [54] T. Merabti et al., “Aligning Biomedical Terminologies in French: Towards Semantic Interoperability in Medical Applications,” 2012. <https://doi.org/10.5772/37738>.
- [55] L. Block, “Mapping nursing wound care data elements to SNOMED-CT,” 2016. <https://doi.org/10.14288/1.0340679>.
- [56] Z. Ibrahim et al., “Analysis of the Suitability of Existing Medical Ontologies for Building a Scalable Semantic Interoperability Solution Supporting Multi-site Collaboration in Oncology,” *Bioinformatics and Bioengineering*, 2014. <https://doi.org/10.1109/BIBE.2014.12>.
- [57] A. Moreno-Conde et al., “Clinical information modeling processes for semantic interoperability of electronic health records: systematic review and inductive analysis,” *Journal of the American Medical Informatics Association*, vol. 22, no. 4, 2015. <https://doi.org/10.1093/JAMIA/OCV008>.
- [58] J. Clunis, “Enhancing health-care data integration via automated semantic mapping,” *The Electronic Library*, vol. 41, no. 4, 2023. <https://doi.org/10.1108/el-06-2023-0142>.
- [59] C. Martínez-Costa et al., “Towards the harmonization of Clinical Information and Terminologies by Formal Representation,” 2012. <https://doi.org/10.24105/EJBI.2012.08.3.2>.
- [60] A. Staines et al., “Terminologies matter - the case of ICNP and SNOMED-CT,” *European Journal of Public Health*, vol. 32, no. Supplement_2, 2022. <https://doi.org/10.1093/eurpub/ckac129.364>.
- [61] G. B. Laleci et al., “Providing semantic interoperability between clinical care and clinical research domains,” *IEEE Journal of Biomedical and Health Informatics*, vol. 17, no. 2, 2013. <https://doi.org/10.1109/TITB.2012.2219552>.
- [62] J. Neumann et al., “Ontology-based surgical workflow recognition and prediction,” *Journal of Biomedical Informatics*, vol. 136, 2022. <https://doi.org/10.1016/j.jbi.2022.104240>.
- [63] S. Lee et al., “Concept Coverage of SNOMED CT in Social Determinants of Health,” *Computers, Informatics, Nursing*, vol. 43, no. 1, 2026. <https://doi.org/10.1097/CIN.0000000000001518>.
- [64] Y. S. Park et al., “Clinical Terminologies: A Solution for Semantic Interoperability,” *Journal of Korean Society of Medical Informatics*, vol. 15, no. 1, 2009. <https://doi.org/10.4258/JKSMI.2009.15.1.1>.
- [65] B. E. Dixon et al., “Identifying health facilities outside the enterprise: challenges and strategies for supporting health reform and meaningful use,” *Informatics for Health & Social Care*, vol. 40, no. 4, 2015. <https://doi.org/10.3109/17538157.2014.924949>.
- [66] I. Direito et al., “Design and Implementation of a Collaborative Clinical Practice and Research Documentation System Using SNOMED-CT and HL7-CDA in the Context of a Pediatric Neurodevelopmental Unit,” *Healthcare*, vol. 11, no. 7, 2023. <https://doi.org/10.3390/healthcare11070973>.
- [67] E. Adel et al., “An extended semantic interoperability model for distributed electronic health record based on fuzzy ontology semantics,” *Electronics*, vol. 10, no. 14, 2021. <https://doi.org/10.3390/ELECTRONICS10141733>.
- [68] D. Kalra and B. G. M. E. Blobel, “Semantic Interoperability of EHR Systems,” in *Medical and Care Computetics 4*, Studies in Health Technology and Informatics, vol. 127, pp. 231–245, IOS Press, Amsterdam, Netherlands, 2007.
- [69] B. Balachandran et al., “Aligning Large Biomedical Ontologies for Semantic Interoperability Using Graph Partitioning,” 2019. <https://doi.org/10.1504/ijbet.2019.102120>.
- [70] A. Annane et al., “Selection and Combination of Heterogeneous Mappings to Enhance Biomedical Ontology Matching,” *Knowledge Acquisition, Modeling and Management*, 2016. https://doi.org/10.1007/978-3-319-49004-5_2.
- [71] Y. He et al., “Exploring Large Language Models for Ontology Alignment,” *arXiv*, 2023. <https://doi.org/10.48550/arxiv.2309.07172>.

- [72] D. Oliveira et al., “Improving the interoperability of biomedical ontologies with compound alignments,” *Journal of Biomedical Semantics*, vol. 9, no. 1, 2018. <https://doi.org/10.1186/S13326-017-0171-8>.
- [73] L. Zhao et al., “Matching biomedical ontologies based on formal concept analysis,” *Journal of Biomedical Semantics*, vol. 9, no. 1, 2018. <https://doi.org/10.1186/S13326-018-0178-9>.
- [74] M. Ba et al., “Large-scale biomedical ontology matching with ServOMap,” *IRBM*, vol. 34, no. 1, 2013. <https://doi.org/10.1016/J.IRBM.2012.12.011>.
- [75] P. Wang et al., “Ontology alignment in the biomedical domain using entity definitions and context,” *arXiv: Computation and Language*, 2018.
- [76] N. Hong et al., “Cross-Language Terminology Mapping Between ICD-10-CN and SNOMED-CT,” *Studies in Health Technology and Informatics*, vol. 290, 2022. <https://doi.org/10.3233/SHTI220028>.
- [77] X. Xue et al., “Matching large-scale biomedical ontologies with central concept based partitioning algorithm and adaptive compact evolutionary algorithm,” *Applied Soft Computing*, vol. 106, 2021. <https://doi.org/10.1016/J.ASOC.2021.107343>.
- [78] P. Kolyvakis et al., “Biomedical ontology alignment: an approach based on representation learning,” *Journal of Biomedical Semantics*, vol. 9, no. 1, 2018. <https://doi.org/10.1186/S13326-018-0187-8>.
- [79] Y. S. Park et al., “Development and Evaluation of SNOMED CT Automated Mapping Tool: Advancing Terminology Standardization and Semantic Interoperability,” *JMIR Medical Informatics*, vol. 14, 2026. <https://doi.org/10.2196/82670>.
- [80] J. M. Rodrigues et al., “Semantic Alignment between ICD-11 and SNOMED CT,” *Studies in Health Technology and Informatics*, 2015. <https://doi.org/10.3233/978-1-61499-564-7-790>.
- [81] A. Naliyatthaliyazchayil et al., “Harmonizing Logical Observation Identifiers Names and Codes (LOINC) Codes and Units in Real-World Oncology Data: Method Development and Evaluation,” *JMIR Medical Informatics*, vol. 14, 2026. <https://doi.org/10.2196/81254>.
- [82] S. Tangchitnob et al., “Developing a Thai User Interface Terminology for Systematized Nomenclature of Medicine Clinical Terms Implementation in Primary Care: Cross-Sectional Content Coverage Analysis,” *JMIR Medical Informatics*, vol. 14, 2026. <https://doi.org/10.2196/80039>.
- [83] S. Lee et al., “Concept Coverage of SNOMED CT in Social Determinants of Health,” *Computers, Informatics, Nursing*, vol. 43, no. 1, 2026. <https://doi.org/10.1097/CIN.0000000000001518>.
- [84] X. Aimé et al., “Semantic interoperability platform for Healthcare Information Exchange,” *IRBM*, vol. 36, no. 2, 2015. <https://doi.org/10.1016/J.IRBM.2015.01.003>.
- [85] G. B. Laleci et al., “Providing semantic interoperability between clinical care and clinical research domains,” *IEEE Journal of Biomedical and Health Informatics*, vol. 17, no. 2, 2013. <https://doi.org/10.1109/TITB.2012.2219552>.
- [86] L. Block, “Mapping nursing wound care data elements to SNOMED-CT,” 2016. <https://doi.org/10.14288/1.0340679>.
- [87] Z. Ibrahim et al., “Analysis of the Suitability of Existing Medical Ontologies for Building a Scalable Semantic Interoperability Solution Supporting Multi-site Collaboration in Oncology,” *Bioinformatics and Bioengineering*, 2014. <https://doi.org/10.1109/BIBE.2014.12>.
- [88] P. Wang et al., “Matching biomedical ontologies: construction of matching clues and systematic evaluation of different combinations of matchers,” *JMIR Medical Informatics*, vol. 9, no. 7, 2021. <https://doi.org/10.2196/28212>.
- [89] J. Silva et al., “An Approach for the Alignment of Biomedical Ontologies based on Foundational Ontologies,” *Journal of Information and Data Management*, 2011.
- [90] G. Diallo, “An effective method of large scale ontology matching,” *Journal of Biomedical Semantics*, vol. 5, no. 1, 2014. <https://doi.org/10.1186/2041-1480-5-44>.
- [91] “BERTMap: A BERT-Based Ontology Alignment System,” *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 36, no. 5, 2022. <https://doi.org/10.1609/aaai.v36i5.20510>.
- [92] A. Annane et al., “Selection and Combination of Heterogeneous Mappings to Enhance Biomedical Ontology Matching,” *Knowledge Acquisition, Modeling and Management*, 2016. https://doi.org/10.1007/978-3-319-49004-5_2.
- [93] Ayad, S. and Khelifa Chibout, L., “Ontology Population With Large Language Models (LLMs): A Case Study on Asbestos Ontology,” *Applied Ontology*, vol. 20, no. 1, pp. 89–106, 2025. <https://doi.org/10.1177/15705838251334528>.

- [94] P. Wang et al., “Matching biomedical ontologies: construction of matching clues and systematic evaluation of different combinations of matchers,” *JMIR Medical Informatics*, vol. 9, no. 7, 2021. <https://doi.org/10.2196/28212>.
- [95] D. Oliveira et al., “Improving the interoperability of biomedical ontologies with compound alignments,” *Journal of Biomedical Semantics*, vol. 9, no. 1, 2018. <https://doi.org/10.1186/S13326-017-0171-8>.
- [96] B. Balachandran et al., “Aligning Large Biomedical Ontologies for Semantic Interoperability Using Graph Partitioning,” 2019. <https://doi.org/10.1504/ijbet.2019.102120>.
- [97] X. Xue et al., “Matching large-scale biomedical ontologies with central concept based partitioning algorithm and adaptive compact evolutionary algorithm,” *Applied Soft Computing*, vol. 106, 2021. <https://doi.org/10.1016/J.ASOC.2021.107343>.
- [98] M. Ba et al., “Large-scale biomedical ontology matching with ServOMap,” *IRBM*, vol. 34, no. 1, 2013. <https://doi.org/10.1016/J.IRBM.2012.12.011>.
- [99] Y. Yamamoto et al., “Context-based refinement of mappings in evolving life science ontologies,” *Journal of Biomedical Semantics*, vol. 14, no. 1, 2023. <https://doi.org/10.1186/s13326-023-00294-8>.
- [100] F. Dhombres et al., “Interoperability between phenotypes in research and healthcare terminologies—Investigating partial mappings between HPO and SNOMED CT,” *Journal of Biomedical Semantics*, vol. 7, no. 1, 2016. <https://doi.org/10.1186/S13326-016-0047-3>.
- [101] Y. He et al., “Exploring Large Language Models for Ontology Alignment,” *arXiv*, 2023. <https://doi.org/10.48550/arxiv.2309.07172>.
- [102] Y. Song et al., “GenOM: Ontology Matching with Description Generation and Large Language Model,” *arXiv*, 2025. <https://doi.org/10.48550/arxiv.2508.10703>.
- [103] J. Qiang et al., “Agent-OM: Leveraging LLM Agents for Ontology Matching,” *Proceedings of the VLDB Endowment*, vol. 18, no. 4, 2024. <https://doi.org/10.14778/3712221.3712222>.
- [104] M. Taboada et al., “Ontology Matching with Large Language Models and Prioritized Depth-First Search,” 2025. <https://doi.org/10.48550/arxiv.2501.11441>.
- [105] J. Dernbach et al., “GLaM: Fine-Tuning Large Language Models for Domain Knowledge Graph Alignment via Neighborhood Partitioning and Generative Subgraph Encoding,” *Proceedings of the AAAI Symposium Series*, vol. 3, no. 1, 2024. <https://doi.org/10.1609/aaais.v3i1.31186>.
- [106] Z. Chen et al., “LLM-Align: Utilizing Large Language Models for Entity Alignment in Knowledge Graphs,” 2024. <https://doi.org/10.48550/arxiv.2412.04690>.
- [107] H. Nguyen et al., “KROMA: Ontology matching with knowledge retrieval and large language models,” *arXiv*, 2025. <https://doi.org/10.48550/arxiv.2507.14032>.
- [108] J. Ding et al., “Knowledge prompt chaining for semantic modeling,” 2025. <https://doi.org/10.48550/arxiv.2501.08540>.
- [109] H. Nguyen, “Ontology Matching with Knowledge Retrieval and Efficient Large Language Models,” 2024.
- [110] Y. Zhong et al., “A New Ontology Matching Method Based on Large Language Model and Graph Structure Learning,” *Communications in Computer and Information Science*, 2025. https://doi.org/10.1007/978-981-95-3739-6_10.
- [111] Y. Feng et al., “Ontology-grounded Automatic Knowledge Graph Construction by LLM under Wikidata schema,” 2024. <https://doi.org/10.48550/arxiv.2412.20942>.
- [112] I. Paik, “Integrating Ontology Rules with Large Language Models for Enhanced Reasoning,” 2025. <https://doi.org/10.1109/itc-cscc66376.2025.11137761>.
- [113] Y. Zhang et al., “Knowledge Graph-Infused Fine-Tuning for Structured Reasoning in Large Language Models,” *arXiv*, 2025. <https://doi.org/10.48550/arxiv.2508.14427>.
- [114] Y. Sun et al., “Pyramid-Driven Alignment: Pyramid Principle Guided Integration of Large Language Models and Knowledge Graphs,” 2024. <https://doi.org/10.48550/arxiv.2410.12298>.
- [115] J. Li et al., “Agent-based ontology mapping and integration towards interoperability,” *Expert Systems*, vol. 25, no. 4, 2008. <https://doi.org/10.1111/J.1468-0394.2008.00460.X>.
- [116] E. Santos et al., “Using Cooperative Agent Negotiation for Ontology Mapping,” *European Workshop on Multi-Agent Systems*, 2006.

- [117] A. Canito, “Evaluating the combination of relaxation and argumentation in ontology matching negotiation,” 2013.
- [118] C. Trojahn et al., “A Negotiation Model for Ontology Mapping,” *IEEE/WIC/ACM International Conference on Intelligent Agent Technology*, 2006. <https://doi.org/10.1109/IAT.2006.20>.
- [119] M. Nagy et al., “Multiagent ontology mapping framework for the semantic web,” *Systems, Man and Cybernetics*, 2011. <https://doi.org/10.1109/TSMCA.2011.2132704>.
- [120] G. Brzykcy et al., “Schema Mappings and Agents’ Actions in P2P Data Integration System 1,” *Journal of Universal Computer Science*, 2008.
- [121] J. Cardoso et al., “An ontology-mapping service for agent-based automated negotiation,” 2008.
- [122] M. Trajanoska et al., “A Multi-Agent System for Semantic Mapping of Relational Data to Knowledge Graphs,” 2025. <https://doi.org/10.5281/zenodo.16911414>.
- [123] M. Nagy et al., “Multi-agent ontology mapping with uncertainty on the semantic web,” *International Conference on Intelligent Computer Communication and Processing*, 2007. <https://doi.org/10.1109/ICCP.2007.4352141>.
- [124] S. M. Elghamrawy et al., “Dynamic ontology mapping for communication in distributed multi-agent intelligent system,” *International Conference on Networking*, 2009. <https://doi.org/10.1109/ICNM.2009.4907198>.
- [125] E. Santos et al., “A Framework for Multilingual Ontology Mapping,” *Language Resources and Evaluation*, 2008.
- [126] A. B. Williams, “Learning to share meaning in a multi-agent system,” *Autonomous Agents and Multi-Agent Systems*, vol. 8, no. 2, 2004. <https://doi.org/10.1023/B:AGNT.0000011160.45980.4B>.
- [127] M. Oprea, “Ontology Mapping in Open Multi-Agent Systems,” 2008.
- [128] C. Colelough, “SymboSLAM: Semantic Map Generation in a Multi-Agent System,” *arXiv*, 2024. <https://doi.org/10.48550/arxiv.2403.15504>.