

Explainable Models for Justifying Search Result Rankings in Ambiguous Knowledge Base Queries

Karim Elkhoury¹, Tarek Mansour²

Abstract

This paper explores the design and justification of explainable models for ranking search results in situations where knowledge base queries are ambiguous. Ambiguity often arises when user queries contain limited context, leading to multiple plausible interpretations within large-scale knowledge repositories. The objective is to develop strategies that systematically identify and rank relevant records while providing transparent rationales for why certain results appear in higher positions. Our approach merges algorithmic ranking methods, intuitive interpretability components, and structured knowledge base representations to enhance user trust and understanding of the retrieval process. We examine the interplay between latent semantic structures, data-driven features, and explicit logical constraints that can reconcile ambiguous query terms. The key elements we discuss include the integration of domain-agnostic feature extraction mechanisms, the incorporation of human-understandable rules for interpretability, and the formal modeling of query-to-result relationships. This research expands on foundational work in semantic search and explainable artificial intelligence by focusing on methods prevalent before and around 2019, highlighting the methodological gaps that remain in rendering accurate but justifiable rankings. We conduct an in-depth analysis of the interplay between uncertainty in query interpretation and the algorithmic processes that prioritize relevant knowledge base entries. Our findings aim to advance the creation of robust, interpretable ranking mechanisms that address both performance and user-oriented transparency in search systems.

¹ Lebanese International University, Department of Computer Science, Beirut Campus, Salim Salam, Beirut, Lebanon

² Université Antonine, Department of Information Technology, Hadat-Baabda Campus, Hadat, Lebanon

Contents

	16
1 Introduction	16
2 Background and Related Efforts	19
3 Methodological Framework	20
3.1 Representation of Entities and Queries	20
3.2 Retrieval Algorithm	20
3.3 Interpretability Layer	20
3.4 Logic Statements and Example Rules	21
4 Empirical Evaluation and Illustrative Scenarios	21
4.1 Metrics for Disambiguation and Ranking Quality	21
4.2 Experimental Methodology	21
4.3 Illustrative Scenario: Scientific Article Retrieval	21
4.4 Illustrative Scenario: Enterprise Data Warehouse Queries	22
5 Discussion of Results and Broader Implications	22

6 Conclusion	23
References	23

1. Introduction

Ambiguous knowledge base queries represent a persistent challenge within the domain of information retrieval [1]. When a user initiates a search request but provides minimal context or uses polysemous terms, it can be difficult for an automated system to identify the specific subtopic or concept area intended by the user [2]. The problem intensifies in large-scale knowledge bases, where concept overlap and varied terminology usage often lead to multiple plausible answer sets. Conventional ranking algorithms frequently maximize relevance based on statistical features while overlooking the necessity of detailed justifications, especially crucial in professional and high-stakes environments such as medical, legal, or scientific research settings [3]. The capability to articulate the reasoning behind search result ordering is thus essential for fostering user trust and mitigating the risks of misinformation.

The ambiguity inherent in query formulation stems from

both linguistic complexity and domain-specific challenges [4]. In natural language, homonyms, synonyms, and syntactic variations contribute to multiple potential interpretations of a single query [5]. Additionally, in specialized fields, terminological inconsistency further complicates retrieval efforts. For instance, in the biomedical domain, the term “cold” could reference a viral infection, a physiological temperature state, or a cryogenic preservation method [6]. Knowledge bases populated with heterogeneous datasets from multiple sources exacerbate this issue by encoding overlapping yet distinct representations of similar concepts. Without appropriate disambiguation mechanisms, retrieval systems risk returning an extensive yet imprecise set of results, thereby increasing cognitive load on the user to manually filter and interpret the information. [7]

Another complicating factor is the evolving nature of knowledge and the contextual dependencies of queries. Temporal shifts in meaning, newly introduced terminology, and variations in conceptual framing across disciplines all contribute to inconsistencies in query interpretation [8]. Legal documents, for example, frequently undergo amendments that modify the scope of key terms, leading to discrepancies in how past and present documents relate to a given query [9]. Similarly, scientific paradigms evolve, and terms that were once definitive may acquire new connotations over time. In dynamic knowledge bases, ensuring that retrieval mechanisms adapt to these changes without introducing inconsistencies remains an ongoing challenge. [10]

Furthermore, the structural characteristics of knowledge bases significantly impact the effectiveness of disambiguation strategies. Hierarchical taxonomies, ontological mappings, and graph-based knowledge representations all offer different mechanisms for structuring information; however, they also introduce constraints in how ambiguous terms are resolved [11]. Some systems rely on explicit entity linking, where terms are mapped to predefined concepts in an ontology, but this approach can be limited when user queries contain novel or underrepresented terms. Others incorporate statistical co-occurrence models, which infer meaning based on contextual usage patterns, but such models may fail to capture nuanced semantic relationships that require deeper contextual understanding. [12]

Another pressing concern is the interaction between user intent and query ambiguity [13]. Users may express a query with an implicit expectation that the system understands their intended meaning, often relying on prior knowledge or domain expertise that the retrieval system does not possess. This is particularly evident in technical domains where expert users employ shorthand notation or field-specific jargon [14]. In contrast, non-expert users might inadvertently use broad or vague descriptors, further complicating disambiguation. In both cases, there exists a fundamental gap between how users conceptualize their queries and how retrieval systems process them [15]. Bridging this gap necessitates techniques that dynamically infer user intent based on additional contextual

cues, such as search history, user profile data, or interactive clarification mechanisms.

Moreover, ambiguous queries pose a significant challenge in high-stakes applications where information retrieval accuracy is critical [16]. In medical knowledge bases, for instance, a query for “diabetes treatment” could yield a wide array of results ranging from pharmaceutical interventions to lifestyle modifications and experimental therapies [17]. Misinterpretation or misranking of results in such contexts could lead to misinformation, potentially influencing medical decisions with adverse consequences. Legal and regulatory databases similarly demand precision, as ambiguous retrieval results could affect case law interpretations or compliance assessments [18]. In scientific research, retrieval ambiguity may lead to improper citation of sources, misalignment of hypotheses, or duplication of efforts due to incomplete awareness of prior work.

The challenge is further compounded by differences in retrieval performance across various knowledge base architectures [19]. Traditional keyword-based search engines rely heavily on term-matching heuristics, which are susceptible to ambiguity in polysemous terms. Semantic search models, in contrast, attempt to map queries to conceptual representations, yet they too face challenges when dealing with multi-faceted or overlapping meanings [20]. Graph-based knowledge bases leverage structured relationships between entities, but they depend on the completeness and accuracy of the underlying ontology, which may not always reflect the most current or comprehensive state of knowledge [21]. Thus, the efficacy of retrieval depends on both the technical framework employed and the richness of the knowledge base itself.

Additionally, the problem of ambiguous knowledge base queries extends beyond individual retrieval instances and influences broader aspects of knowledge management [22]. When users frequently encounter ambiguous or misleading results, it erodes confidence in the system, leading to reduced adoption and reliance on alternative sources that may lack credibility. In enterprise environments, poor retrieval performance can hinder decision-making processes, reduce operational efficiency, and contribute to knowledge silos where employees resort to informal or localized repositories of information rather than centralized, organization-wide knowledge bases [23, 24]. In research and academic settings, inadequate disambiguation can distort literature reviews, skew citation patterns, and impede interdisciplinary collaboration by failing to align related concepts across fields.

Table 1 provides illustrative examples of ambiguous queries across different domains, highlighting the challenges they pose and the potential implications of incorrect retrieval.

Given these complexities, the necessity of refining query interpretation mechanisms extends beyond simple keyword expansion or lexical disambiguation [25]. Effective handling of ambiguous queries demands an integrated approach that considers linguistic variation, contextual dependencies, user intent inference, and domain-specific knowledge representa-

Table 1. Examples of Ambiguous Queries and Their Contextual Challenges

Query	Potential Interpretations	Contextual Challenges
“Cold”	Viral infection, temperature state, cryogenic preservation	Requires domain knowledge to differentiate medical vs. physical meaning
“Apple”	Technology company, fruit, record label	Disambiguation needed based on user intent (e.g., tech-related search vs. nutrition-related search)
“Jaguar”	Animal, luxury car brand, sports team	Context-dependent retrieval necessary to avoid irrelevant results
“Mercury”	Element, planet, automobile brand, Roman deity	Knowledge graph linking or contextual keyword expansion needed
“Java”	Programming language, Indonesian island, type of coffee	Must infer from surrounding text or user history

tions [26]. It also requires advancements in explainability to ensure that retrieval outputs can be understood, trusted, and refined based on user feedback. Table 2 summarizes key performance metrics that must be considered when evaluating the effectiveness of retrieval systems in handling ambiguous queries.

The increasing reliance on large-scale knowledge bases across disciplines necessitates continual refinement of information retrieval methodologies [27]. As knowledge repositories grow in size and complexity, the risks posed by ambiguous queries become more pronounced, demanding robust strategies to ensure accurate and transparent retrieval. The future of information retrieval hinges not only on improving precision but also on fostering user trust through intelligible and context-aware search mechanisms. [28]

Traditional approaches to ranking documents in information retrieval involve constructing relevance scores by comparing query terms to index representations. Early strategies relied on term frequency-inverse document frequency (TF-IDF) techniques, which provided numerical estimates of term importance [29]. More recent approaches, emerging before 2019, leverage neural embeddings and structured knowledge graph embeddings to capture semantic relationships [30]. However, most methods in practice still provide a ranking output with limited or no interpretability, leaving practitioners and end-users uncertain about the mechanics behind retrieval decisions. Specifically, while neural embedding models may improve retrieval performance, their latent representations are opaque, leading to inherent difficulties when attempting to provide interpretable rationales. [31]

In ambiguous query scenarios, the situation is further complicated by the fact that multiple potential senses, context angles, or subdomains of a query term could be equally valid. Without sophisticated disambiguation techniques, a system might either return a heterogeneous set of results or fail to address certain interpretations altogether [32]. This dual risk highlights the necessity of robust, explainable frameworks that effectively handle the multiple candidate meanings of a single query. Such frameworks must incorporate interpretable

logic structures to guide users toward correct interpretations of system outputs, especially if query expansions or concept disambiguation steps are performed automatically. [33]

Effective explanation mechanisms should extend beyond merely attributing a query’s result to the presence of matching keywords [34]. Instead, they should elucidate the logical connections and semantic relationships that cause some documents, or entities, to appear more relevant than others. By aligning the internal representation of content with interpretable logic rules, retrieval systems can transparently reveal the steps taken to rank and filter the search results [35]. This not only informs end-users of how a particular document was selected, but also allows researchers and developers to diagnose potential pitfalls and biases in the retrieval pipeline.

In this paper, we undertake a comprehensive exploration of how interpretability and explainability can be systematically introduced into the ranking pipeline for ambiguous knowledge base queries [36]. We examine the roles of structured embeddings, symbolic logic, and uncertainty quantification in reconciling multiple possible interpretations, focusing on design choices that support traceable rationales. Specifically, we present approaches for handling uncertain or conflicting evidence, discuss the interplay between high-dimensional feature representations and symbolic constraints, and explore the architectural frameworks that integrate these components into an operational pipeline [37]. Throughout the discussion, we assume a broad range of application domains, such as scientific literature indexing, enterprise data warehouses, and large encyclopedic knowledge graphs, each requiring robust solutions for ambiguous queries [38, 39].

The remainder of this paper is organized into four main sections. First, we provide a background of related efforts addressing ambiguity resolution in knowledge base search [40]. Next, we detail our methodological framework, presenting core formal statements and data structures that enable explainable ranking. We then discuss empirical insights and illustrative scenarios where these methods have been tested, describing the metrics and assessment strategies used for performance evaluation [41]. Finally, we examine the

Table 2. Key Performance Metrics for Evaluating Retrieval Systems on Ambiguous Queries

Metric	Description
Disambiguation Accuracy	Measures how often the system correctly identifies the intended meaning of a query.
Relevance Precision	Evaluates the proportion of retrieved results that are contextually relevant to the user’s intended query meaning.
Explainability Score	Assesses the system’s ability to provide human-understandable justifications for retrieved results.
User Correction Rate	Tracks how frequently users need to refine their query or manually adjust search parameters due to ambiguity.
Query Resolution Time	Measures the time required for the system to disambiguate and return meaningful results.

broader implications, including potential limitations, practical deployment concerns, and the importance of standardized interpretability benchmarks. In the concluding section, we summarize our findings and propose directions for future work on interpretable ranking for ambiguous queries. [42]

2. Background and Related Efforts

Research on handling ambiguous queries in knowledge bases and information retrieval systems spans multiple disciplines, including semantic web technologies, natural language processing, machine learning, and logic-based reasoning frameworks [43]. Since the early stages of information retrieval, ambiguity has been a recognized obstacle, prompting efforts to refine indexing and search algorithms. The classical vector space model, as well as probabilistic retrieval approaches, sought to compute relevance scores based on lexical overlap and document distributions [44]. Despite the documented success of these models, they provided limited means of explaining why certain documents ranked higher than others, restricting the diagnostic capabilities available to both system designers and end-users.

The field witnessed a growing interest in embedding-based techniques that mapped query and document content into dense vector spaces [45]. Word2Vec, GloVe, and other neural embedding frameworks provided semantic representations capable of capturing contextual relationships. Extensions of these approaches led to knowledge graph embeddings, which allowed for structured representations of entities and their relations [46]. Although these techniques improved retrieval performance, their interpretability was typically not a primary focus [47]. Researchers recognized that while embedding-based algorithms could address certain forms of lexical ambiguity by inferring semantic distance, the latent nature of these models hindered any transparent explanation process.

Simultaneously, the semantic web community made strides in employing ontology-based queries, leveraging formal logic representations such as Description Logics [48]. Query languages like SPARQL provided a means to pose structured queries against knowledge graphs, yet these methods were still limited when user queries were vague or under-specified. To address such gaps, some efforts combined ontology-based

reasoning with query expansion strategies driven by lexical resources (for instance, synonym lists and taxonomic hierarchies) [49]. Although these approaches improved retrieval coverage, the question of how to systematically communicate the chosen expansion paths or disambiguation decisions to non-expert users remained unresolved.

Explaining result rankings involves not only clarifying the source of ambiguities but also exposing the internal reasoning used by the system [50]. The field of explainable artificial intelligence, especially in machine learning contexts, witnessed notable progress with techniques that highlight influential features in classification and regression tasks [51]. Methods such as Layer-Wise Relevance Propagation (LRP), Local Interpretable Model-Agnostic Explanations (LIME), and gradient-based saliency maps provided ways to measure local feature contributions. Yet, these methods were primarily suited to classification problems and did not directly address the unique challenges of ranking tasks, especially when the queries themselves were ambiguous. [52]

An additional strand of research pertinent to interpretability is the formalization of logic-based explanation frameworks. In these paradigms, a user might receive an explicit logical derivation or proof for why a particular result is relevant [53]. In knowledge representation, some investigators introduced justification systems for ontology-based queries, generating minimal sets of axioms or inference chains that support a query result. However, these were often geared toward domain experts capable of parsing logical constructs, leaving open questions about how to adapt the same clarity to broader user bases and large-scale, multi-domain knowledge bases. [54]

In parallel, user studies in information retrieval indicated that trust in system outputs significantly increases when users can ascertain the rationale behind the ranking [55]. Thus, there is a strong motivation for bridging the gap between robust, high-performance retrieval techniques and transparent, intelligible explanations. Advances in probabilistic and fuzzy logic-based frameworks offered partial solutions, enabling some measure of uncertainty quantification [56]. Still, ambiguous queries complicate even these methods, as multiple plausible inferences might exist for different interpretations

of the query.

In summary, while many relevant lines of research have contributed building blocks—ranging from advanced semantic embeddings to logic-based explanations—the integration of these components into a coherent, unified approach for ambiguous queries remains a challenge [57]. A major objective in the proposed work is to address this gap by designing an end-to-end pipeline that provides understandable justifications for search result rankings. This pipeline would incorporate modern data-driven embeddings, symbolic representations, and model-agnostic explanation layers, tailored specifically for the disambiguation problem in knowledge base retrieval tasks. [58]

3. Methodological Framework

The proposed framework for delivering explainable rankings in the context of ambiguous queries rests on three key components: a structured representation of knowledge, a multi-phase retrieval algorithm, and an interpretability layer grounded in formal logic [59]. Central to this methodology is the assumption that knowledge base entities can be represented by vectors in a semantic embedding space and simultaneously be described by symbolic annotations that reflect domain-specific or general ontological constraints. Such dual representations facilitate both broad coverage of potential ambiguities via data-driven techniques and explicit reasoning over definitional constraints. [60]

3.1 Representation of Entities and Queries

Let $\mathcal{E}$ denote the set of entities in the knowledge base. Each entity $e \in \mathcal{E}$ is associated with two forms of representation:

$$e \mapsto (\mathbf{v}(e), \sigma(e))$$

where $\mathbf{v}(e) \in \mathbb{R}^d$ is a vector embedding capturing semantic information, and $\sigma(e)$ is a symbolic descriptor containing attributes, relevant ontological classes, and relationships. For instance, an entity representing a scientific article might have $\mathbf{v}(e)$ derived from distributed text representations, while $\sigma(e)$ enumerates authors, keywords, or associated concepts.

Similarly, an incoming user query q can be represented by:

$$q \mapsto (\mathbf{v}(q), \sigma(q))$$

where $\mathbf{v}(q)$ is a vector representation, often computed by aggregating or encoding the user's query text, and $\sigma(q)$ is a partially formed symbolic descriptor. In ambiguous queries, $\sigma(q)$ may be incomplete or contain placeholders indicating uncertain aspects of the query's meaning. [61]

3.2 Retrieval Algorithm

The retrieval process proceeds in two stages, aiming first to capture all potentially relevant entities and then to refine this list based on disambiguation signals: [62, 63]

Stage 1: Broad Candidate Selection A broad set of candidate entities $\mathcal{C} \subset \mathcal{E}$ is selected by comparing $\mathbf{v}(q)$ to $\mathbf{v}(e)$ for each entity. A common approach is to compute a similarity score using the cosine distance, or a bilinear form $\mathbf{v}(q)^\top M \mathbf{v}(e)$ for some learnable matrix M . Entities whose similarity score exceeds a threshold θ form the candidate set $\mathcal{C}$.

$$\mathcal{C} = \{e \in \mathcal{E} \mid \text{sim}(\mathbf{v}(q), \mathbf{v}(e)) \geq \theta\}$$

Stage 2: Disambiguation-Guided Ranking For each candidate $e \in \mathcal{C}$, the symbolic representations $\sigma(q)$ and $\sigma(e)$ are used to refine the ranking. Logic rules capture interpretive constraints, such as:

$$\Gamma : (\sigma(q), \sigma(e)) \models \tau(q, e) [64]$$

where $\tau(q, e)$ is a statement indicating how well e matches the intended meaning of q . Possible rules can include hierarchical constraints, entity-type matching, or specialized domain constraints [65]. Each rule in Γ is assigned a weight w_i indicating its importance. The total symbolic compatibility score is:

$$\text{Score}_{\text{sym}}(q, e) = \sum_{i=1}^k w_i \cdot \mathbf{1}\{\gamma_i \in \Gamma \wedge \gamma_i(\sigma(q), \sigma(e)) = \text{true}\}$$

The final ranking metric for each candidate combines the embedding-based similarity and symbolic compatibility: [66]

$$\text{RankScore}(q, e) = \alpha \cdot \text{sim}(\mathbf{v}(q), \mathbf{v}(e)) + \beta \cdot \text{Score}_{\text{sym}}(q, e)$$

where α and β are hyperparameters [67]. The ranked list is then generated by sorting $\mathcal{C}$ in descending order of RankScore . In ambiguous scenarios, the symbolic checks help discriminate among multiple plausible meanings of the query.

3.3 Interpretability Layer

The interpretability component is designed to elucidate the role of both the embedding similarity and the logical constraints in shaping final rankings [68]. We introduce a function:

$$\Lambda(q, e) = \{(r, c_r) \mid r \in \Gamma, c_r \in \{0, 1\}\}$$

which enumerates whether each rule r in Γ was satisfied ($c_r = 1$) or not ($c_r = 0$) for the query-entity pair [69]. This set indicates precisely which symbolic conditions contributed to a high rank. In addition, a subfunction: [70]

$$\Theta(\mathbf{v}(q), \mathbf{v}(e)) = \{\delta \mid \delta = \text{sim}(\mathbf{v}(q), \mathbf{v}(e))\}$$

exposes the embedding-based similarity score [71]. The explanation for why e was ranked in a particular position can thus be formulated as:

$$\Xi(q, e) = \langle \Theta(\mathbf{v}(q), \mathbf{v}(e)), \Lambda(q, e) \rangle$$

The system can present $\Xi(q, e)$ either as a short textual summary or as a structured justification, depending on the user interface [72]. In essence, Ξ clarifies how the embedding-based similarity and the relevant logic rules mutually influence the final ranking.

3.4 Logic Statements and Example Rules

To illustrate how logic statements might be employed, consider a scenario in which a query references an ambiguous term like “jaguar,” which could refer to an animal or a vehicle brand [73]. We define a logic statement:

IsAnimal($\sigma(e)$) $\wedge$ NotVehicle($\sigma(e)$) $\Rightarrow$ AnimalContextScore
and

IsVehicle($\sigma(e)$) $\wedge$ NotAnimal($\sigma(e)$) $\Rightarrow$ VehicleContextScore

The user’s query may provide partial signals about the context (for instance, mention of species, environment, or mechanical attributes) [74]. Those signals are encoded in $\sigma(q)$ [75]. If $\sigma(q)$ strongly correlates with the concept “species,” then the rule IsAnimal($\sigma(e)$) $\wedge$ NotVehicle($\sigma(e)$) might have a higher weight w_{animal} . Conversely, if the textual representation of the query suggests a mechanical context, the rule IsVehicle($\sigma(e)$) $\wedge$ NotAnimal($\sigma(e)$) would be more relevant. By enumerating which rules were triggered, the system conveys to users that it considered the animal sense or the vehicle sense, thereby explaining the final ranking.

4. Empirical Evaluation and Illustrative Scenarios

Evaluating the proposed framework involves two main components: measuring how effectively the system handles ambiguous queries (disambiguation performance) and assessing the interpretability or explainability of the ranking decisions [76]. We outline below a prototypical procedure for such an empirical evaluation, along with simplified scenarios to illustrate the outcomes.

4.1 Metrics for Disambiguation and Ranking Quality

We adopt standard ranking metrics—such as Mean Reciprocal Rank (MRR), Normalized Discounted Cumulative Gain (nDCG), and precision/recall at various cutoffs—to quantify retrieval performance [77]. For ambiguous queries, these metrics are typically computed with respect to multiple valid ground truths, each corresponding to a distinct interpretation

of the query. For instance, if a query like “mercury” has multiple relevant result sets—planets, chemical elements, automotive references—each is treated as a possible correct interpretation. [78, 79]

Alongside these classical metrics, we incorporate an *Interpretability Score* designed to measure the clarity of the system’s explanations. While measuring interpretability is an inherently subjective task, we can use surrogate methods such as user surveys or proxy metrics like the fraction of logic rules satisfied that match an expert-defined gold standard [80]. For the latter, domain experts might specify which rules should logically apply if the query were intended to reference a certain context, and the system’s alignment with these rules can be tallied.

4.2 Experimental Methodology

The empirical methodology typically involves the following steps: [81]

1. **Dataset Preparation:** Construct or select a knowledge base containing entities that exhibit overlaps in meaning or usage. Ensure that annotations and embedding vectors are available. Curate a set of ambiguous queries, each with multiple valid interpretation labels. [82]
2. **Baseline Methods:** Implement or select baseline retrieval methods, such as pure embedding-based rankers, and possibly naive expansions of the query that do not incorporate logic-based constraints.
3. **Implementation of Proposed Framework:** Integrate the symbolic logic rule set with the embedding-based retrieval pipeline. Adjust weighting parameters α and β to balance the contribution of vector-based similarity against symbolic constraints.
4. **Evaluation Protocol:** For each ambiguous query, retrieve a ranked list of entities using both the baseline and the proposed framework. Compute the ranking metrics for each distinct interpretation [83]. Assess the interpretability by analyzing the system’s logic-based justifications, either through an automated alignment measure or a user study. [84]
5. **Statistical Analysis:** Compare the performance gains across different methods via significance tests, ensuring that improvements in ranking quality are robust. Similarly, compile user feedback on interpretability if a user study is conducted, to gauge the clarity and usefulness of the logic-based explanations.

4.3 Illustrative Scenario: Scientific Article Retrieval

Consider a knowledge base of scientific articles in the biomedical domain, where terms like “gene expression,” “cell line,” and “model organism” may appear within thousands of publications [85]. An ambiguous query could be “viral model,” which might refer to in vitro systems, animal models, or computational simulations of viral behavior. The system extracts

a vector $\mathbf{v}(q)$ from the textual representation of the user’s query, then computes similarity scores with candidate articles. The symbolic descriptor $\sigma(q)$ might specify the domain as “biology,” with partial indications that the user is interested in experimental setups [86]. The rule set can include:

$$\begin{aligned} &\text{IsExperimentalModel}(\sigma(e)) \wedge \text{MentionVirus}(\sigma(e)) \\ &\Rightarrow \text{ExperimentalContextRelevance} \end{aligned} \quad (1)$$

$$\begin{aligned} &\text{IsComputationalSimulation}(\sigma(e)) \wedge \text{MentionVirus}(\sigma(e)) \\ &\Rightarrow \text{SimulationContextRelevance} \end{aligned} \quad (2)$$

Articles that mention laboratory experiments and viral cultures would score highly on [87] `ExperimentalContextRelevance`, while purely computational studies would satisfy `(SimulationContextRelevance)`.

Depending on the weighting, the system might rank actual experimental studies higher if the embedding similarity also aligns with that context [88, 89]. The interpretability layer would then indicate that the experimental context rule was triggered, providing the user with a concise explanation for why a certain set of articles was prioritized.

4.4 Illustrative Scenario: Enterprise Data Warehouse Queries

In an enterprise context, knowledge bases often comprise a variety of semi-structured records: sales reports, inventory logs, and HR documentation [90]. An ambiguous query, such as “employee turnover,” can refer to an HR policy document, a financial report describing turnover costs, or a data analytics presentation on retention rates. The system’s logic rules might consider the domain categories (finance, HR, analytics), checking for domain alignment between the query’s symbolic descriptor $\sigma(q)$ and each entity’s $\sigma(e)$ [91]. If $\sigma(q)$ denotes a focus on policy guidelines, the rule might be:

$$\begin{aligned} &\text{IsPolicyDocument}(\sigma(e)) \wedge \text{InvolvesEmployees}(\sigma(e)) \\ &\Rightarrow \text{PolicyFocusScore} \end{aligned} \quad (3)$$

Entities describing purely numeric turnover trends might earn a “`numericAnalysisScore`” instead [92]. The final explanation to the user highlights that the chosen ranked item is indeed an HR policy file, matching the user’s interest in guidelines rather than statistical data, thus reinforcing trust in the retrieval system. [93]

A broad retrieval module identifies potential matches, which are then scored by a logic reasoning component to yield the final ranked list. An interpretability layer sits atop these components, extracting a set of justifications for the final ordering of results [94]. The synergy between vector-based matching and symbolic rules forms the crux of the explainable mechanism.

5. Discussion of Results and Broader Implications

The integration of symbolic logic with embedding-based retrieval for ambiguous queries yields multiple benefits [95]. First, the availability of rule-based justifications allows domain specialists and casual users alike to trace the system’s reasoning, bridging the gap between opaque neural models and actionable insights. The synergy of embedding distances with explicit constraints provides a means to systematically handle multi-faceted queries where lexical overlaps may be insufficient to fully capture the user’s intent. [96]

Nevertheless, the framework carries several practical caveats [97]. One concern is the manual curation of logic rules, which can be labor-intensive, especially for highly specialized domains. While adaptive or data-driven rule induction could mitigate some of these efforts, the correctness and completeness of automatically extracted rules remain an open problem [98, 99]. Moreover, the vector embedding module may drift or degrade if the underlying text corpus evolves over time, forcing periodic retraining or updating of the embedding space. In domains with dynamic content, both the symbolic annotations and the embedding models must be systematically maintained to preserve the reliability of the explanation mechanism. [100]

Another consideration relates to computational scalability. As knowledge bases grow in size, the two-stage retrieval process can become computationally heavy [101]. Pre-filtering with efficient approximate nearest neighbor searches can mitigate some of these challenges [102]. Nonetheless, the final disambiguation stage, which applies multiple logic constraints, may still require significant processing. The tuning of hyperparameters α , β , and rule weights must be performed with care to ensure robust performance across a variety of query types. [103]

In terms of user interaction, the interpretability layer’s format is critical for acceptance. Logical proofs or symbolic trace statements can overwhelm users lacking domain or technical expertise [104]. Therefore, creating succinct and user-friendly explanations that faithfully represent the underlying reasoning is essential. Strategies might include natural language generation systems that condense symbolic logic matches into readable sentences, or well-designed graphical interfaces that highlight the relevant portions of an ontology. [105]

The question of standardizing interpretability metrics remains a pressing issue [106]. While user studies offer valuable subjective feedback, the field lacks uniformly accepted quantitative measures for interpretability, particularly in the context of retrieval tasks. Initiatives to develop shared benchmarks or standardized explanation tasks are thus highly relevant [107]. Another area that deserves further exploration is the potential synergy with active learning or human-in-the-loop paradigms, where real-time user feedback can refine both the logic rule sets and the system’s overall ranking strategy.

From an ethical standpoint, providing transparent explanations can reduce the risk of biased or incorrect retrieval results

going unchecked [108]. Nevertheless, the system must be carefully designed to ensure that partial or overly simplistic explanations do not mislead users into a false sense of security. There is also the risk that malicious actors could exploit the interpretability layer to reverse-engineer the system's logic for harmful purposes, such as spamming or misinformation campaigns [109]. Addressing these adversarial aspects calls for robust safeguard mechanisms, such as anomaly detection in user behavior and masked release of certain rule sets in sensitive domains. [110]

In conclusion, while the integration of explainable ranking for ambiguous queries requires careful methodological considerations, it promises to significantly enhance trust, clarity, and effectiveness in knowledge base retrieval. By combining vector-based similarities, symbolic logic rules, and user-centric explanation interfaces, the approach provides a balanced means of reconciling multiple query interpretations [111]. This alignment of interpretability and performance stands as a crucial step forward in making large-scale knowledge systems more transparent and user-friendly.

6. Conclusion

This paper has presented a robust, integrative approach for delivering explainable search result rankings in the context of ambiguous knowledge base queries [112]. By unifying embedding-based representations with symbolic logic constraints, the framework allows multiple potential interpretations to be systematically identified and then weighed according to domain-specific rules. The explanation layer renders these processes transparent, exposing both the contribution of semantic similarity scores and the satisfaction of logic-based constraints to end-users [113]. Our discussion has explored the major challenges—ranging from the manual creation of rule sets to the computational complexity of multi-stage ranking—and provided examples illustrating how this methodology can be applied in scientific, enterprise, or general-purpose knowledge bases. [114]

Empirical evaluations emphasize not only the accuracy of result lists for ambiguous queries but also the quality of interpretability that is crucial for trust and validation in knowledge retrieval systems. We have highlighted metrics, user studies, and logical alignment checks that can guide the development of standardized benchmarks in this evolving field [115]. While various aspects of the approach remain open for future research, such as automated rule discovery, improved explanation formats, and the interplay with active learning, the present study underscores the value of merging symbolic and sub-symbolic paradigms for transparent information retrieval.

Ultimately, incorporating explainable ranking mechanisms not only improves user satisfaction but also mitigates risks linked to misinterpretation or hidden biases in large-scale systems [116]. These considerations remain pivotal for many domains, including healthcare, finance, and scientific information retrieval, where accountability and verifiability are paramount. With a careful balance of interpretability and per-

formance, the proposed framework lays the groundwork for more trustworthy and comprehensible knowledge base search services, particularly in cases where the user's initial query may be ambiguous or prone to multiple meanings. [117]

References

- [1] B. G. Arsintescu, S. Deo, and W. Harris, "Pathquery pregel: high-performance graph query with bulk synchronous processing," *Pattern Analysis and Applications*, vol. 23, pp. 1493–1504, August 2019.
- [2] J. Hellings, M. Gyssens, D. V. Gucht, and Y. Wu, "First-order definable counting-only queries," *Annals of Mathematics and Artificial Intelligence*, vol. 87, pp. 109–136, July 2019.
- [3] B. B. Nandi, S. C. Ghosh, A. Banerjee, and N. Banerjee, "Customer on-boarding strategies for cloud computing services with dynamic service-level agreements," *Service Oriented Computing and Applications*, vol. 11, pp. 47–63, June 2016.
- [4] L. Müller, R. Gangadharaiah, S. C. Klein, J. C. Perry, G. Bernstein, D. Nurkse, D. H. Wailes, R. Graham, R. El-Kareh, S. Mehta, S. A. Vinterbo, and E. Aronoff-Spencer, "An open access medical knowledge base for community driven diagnostic decision support system development.," *BMC medical informatics and decision making*, vol. 19, pp. 93–, April 2019.
- [5] B. Müller, C. Poley, J. Pössel, A. Hagelstein, and T. Gübitz, "Livivo – the vertical search engine for life sciences," *Datenbank-Spektrum : Zeitschrift für Datenbanktechnologie : Organ der Fachgruppe Datenbanken der Gesellschaft für Informatik e.V.*, vol. 17, pp. 29–34, January 2017.
- [6] T. Xie, B. Wu, B. Jia, and B. Wang, "Graph-ranking collective chinese entity linking algorithm," *Frontiers of Computer Science*, vol. 14, pp. 291–303, August 2019.
- [7] Y. Guo, J. Hu, and Y. Peng, "Research of new strategies for improving cbr system," *Artificial Intelligence Review*, vol. 42, pp. 1–20, March 2012.
- [8] X. Liu and H. Fang, "Latent entity space: a novel retrieval approach for entity-bearing queries," *Information Retrieval Journal*, vol. 18, pp. 473–503, September 2015.
- [9] A. Kamyabniya, M. Lotfi, M. Naderpour, and Y. Yih, "Robust platelet logistics planning in disaster relief operations under uncertainty: a coordinated approach," *Information Systems Frontiers*, vol. 20, pp. 759–782, August 2017.
- [10] A. Wong and H. Shatkay, "Protein function prediction using text-based features extracted from the biomedical literature: The cafa challenge," *BMC bioinformatics*, vol. 14, pp. 1–14, February 2013.

- [11] J. W. Selsky, R. Ramírez, and O. N. Babüroğlu, “Collaborative capability design: Redundancy of potentialities,” *Systemic Practice and Action Research*, vol. 26, pp. 377–395, November 2012.
- [12] O. Choudhury, A. Chakrabarty, and S. J. Emrich, “Hecil: A hybrid error correction algorithm for long reads with iterative learning,” *Scientific reports*, vol. 8, pp. 9936–9936, July 2018.
- [13] J. S. Saltz and N. I. Dewar, “Data science ethical considerations: a systematic literature review and proposed project framework,” *Ethics and Information Technology*, vol. 21, pp. 197–208, March 2019.
- [14] C. Jin, J. Chen, and H. Liu, “Mapreduce-based entity matching with multiple blocking functions,” *Frontiers of Computer Science*, vol. 11, pp. 895–911, August 2016.
- [15] W. Qu, D. Wang, S. Feng, Y. Zhang, and G. Yu, “A novel cross-modal hashing algorithm based on multi-modal deep learning,” *Science China Information Sciences*, vol. 60, pp. 092104–, March 2017.
- [16] H. Ziaimatin, T. Groza, T. Tudorache, and J. Hunter, “Modelling expertise at different levels of granularity using semantic similarity measures in the context of collaborative knowledge-curation platforms,” *Journal of intelligent information systems*, vol. 47, pp. 469–490, August 2015.
- [17] Z. Yang and J. V. Sweedler, “Application of capillary electrophoresis for the early diagnosis of cancer,” *Analytical and bioanalytical chemistry*, vol. 406, pp. 4013–4031, March 2014.
- [18] L. A. McNamara, “Biologically-inspired analysis in the real world: computing, informatics, and ecologies of use,” *Security Informatics*, vol. 1, pp. 17–, November 2012.
- [19] I. Jeon, E. E. Papalexakis, C. Faloutsos, L. Sael, and U. Kang, “Mining billion-scale tensors: algorithms and discoveries,” *The VLDB Journal*, vol. 25, pp. 519–544, March 2016.
- [20] Y. Lai, D. Liu, and S. Wang, “Reduced ordered binary decision diagram with implied literals: a new knowledge compilation approach,” *Knowledge and Information Systems*, vol. 35, pp. 665–712, August 2012.
- [21] A. V. Gelder, “Producing and verifying extremely large propositional refutations,” *Annals of Mathematics and Artificial Intelligence*, vol. 65, pp. 329–372, November 2012.
- [22] S. McGrath and D. Gherzi, “Building towards precision medicine: empowering medical professionals for the next revolution,” *BMC medical genomics*, vol. 9, pp. 23–23, May 2016.
- [23] P. Cremaschi, S. Rovidia, L. Sacchi, A. Lisa, F. Calvi, A. Montecucco, G. Biamonti, S. Bione, and G. Sacchi, “Correlagenes: a new tool for the interpretation of the human transcriptome,” *BMC bioinformatics*, vol. 15, pp. 1–10, January 2014.
- [24] Abhishek and V. Rajaraman, “A computer aided short-hand expander,” *IETE Technical Review*, vol. 22, no. 4, pp. 267–272, 2005.
- [25] X. Xue and J. Chen, “Optimizing ontology alignment through hybrid population-based incremental learning algorithm,” *Memetic Computing*, vol. 11, pp. 209–217, March 2018.
- [26] H. N. Tran, E. Cambria, and A. Hussain, “Towards gpu-based common-sense reasoning: Using fast subgraph matching,” *Cognitive Computation*, vol. 8, pp. 1074–1086, August 2016.
- [27] J. I. Seeman, “From “multiple simultaneous independent discoveries” to the theory of “multiple simultaneous independent errors”: a conduit in science,” *Foundations of Chemistry*, vol. 20, pp. 219–249, February 2018.
- [28] A. Koutsoukas, K. J. Monaghan, X. Li, and J. Huan, “Deep-learning: investigating deep neural networks hyper-parameters and comparison of performance to shallow methods for modeling bioactivity data,” *Journal of cheminformatics*, vol. 9, pp. 42–42, June 2017.
- [29] Y. Lin and Y. He, “The ontology of genetic susceptibility factors (ogsf) and its application in modeling genetic susceptibility to vaccine adverse events,” *Journal of biomedical semantics*, vol. 5, pp. 19–19, April 2014.
- [30] S. R. Johnson, X. Qi, A. E. Cicek, and G. Ozsoyoglu, “ipathcase kegg : An ipad interface for kegg metabolic pathways,” *Health information science and systems*, vol. 1, pp. 4–4, January 2013.
- [31] S. Zapolsky and E. Drumwright, “Inverse dynamics with rigid contact and friction,” *Autonomous Robots*, vol. 41, pp. 831–863, November 2016.
- [32] Y. qiang Yang, “Research on the single image super-resolution method based on sparse bayesian estimation,” *Cluster Computing*, vol. 22, pp. 1505–1513, February 2018.
- [33] A. Hall, P. M. Cox, C. Huntingford, and S. A. Klein, “Progressing emergent constraints on future climate change,” *Nature Climate Change*, vol. 9, pp. 269–278, March 2019.
- [34] H. R. Strasberg, G. D. Fiol, and J. J. Cimino, “Terminology challenges implementing the hl7 context-aware knowledge retrieval (‘infobutton’) standard,” *Journal of the American Medical Informatics Association : JAMIA*, vol. 20, pp. 218–223, October 2012.
- [35] Y. Weinstein, C. R. Madan, and M. A. Sumeracki, “Teaching the science of learning,” *Cognitive research: principles and implications*, vol. 3, pp. 2–2, January 2018.

- [36] J. P. F. Bai, J. C. Earp, and V. C. Pillai, “Translational quantitative systems pharmacology in drug development: from current landscape to good practices.,” *The AAPS journal*, vol. 21, pp. 72–, June 2019.
- [37] X. Liu, F. Chen, H. Fang, and M. Wang, “Exploiting entity relationship for query expansion in enterprise search,” *Information Retrieval*, vol. 17, pp. 265–294, January 2014.
- [38] H. Li, X. Feng, L. Cao, E. Li, H. Liang, and X. Chen, “A new ecg signal classification based on wpd and apen feature extraction,” *Circuits, Systems, and Signal Processing*, vol. 35, pp. 339–352, May 2015.
- [39] A. Sharma, K. M. Goolsbey, and D. Schneider, “Disambiguation for semi-supervised extraction of complex relations in large commonsense knowledge bases,” in *7th Annual Conference on Advances in Cognitive Systems*, 2019.
- [40] M. S. Ackerman, J. Dachtera, V. Pipek, and V. Wulf, “Sharing knowledge and expertise: The cscw view of knowledge management,” *Computer Supported Cooperative Work (CSCW)*, vol. 22, pp. 531–573, August 2013.
- [41] F. C. Constantino and J. Z. Simon, “Dynamic cortical representations of perceptual filling-in for missing acoustic rhythm,” *Scientific reports*, vol. 7, pp. 17536–17536, December 2017.
- [42] J. Lai, Y. Mu, F. Guo, W. Susilo, and R. Chen, “Fully privacy-preserving and revocable id-based broadcast encryption for data access control in smart city,” *Personal and Ubiquitous Computing*, vol. 21, pp. 855–868, July 2017.
- [43] M. Zhang, J. Zhu, Z. Wang, and Y. Chen, “Providing personalized learning guidance in moocs by multi-source data analysis,” *World Wide Web*, vol. 22, pp. 1189–1219, May 2018.
- [44] Y. Yan and J. Guan, “How multiple networks help in creating knowledge: evidence from alternative energy patents,” *Scientometrics*, vol. 115, pp. 51–77, January 2018.
- [45] A. Szabóová, O. Kuželka, F. Železný, and J. Tolar, “Prediction of dna-binding propensity of proteins by the ball-histogram method using automatic template search,” *BMC bioinformatics*, vol. 13, pp. 1–11, June 2012.
- [46] Q. Qiu, Z. Li, K. Ahmed, W. Liu, S. F. Habib, H. Li, and M. Hu, “A neuromorphic architecture for context aware text image recognition,” *Journal of Signal Processing Systems*, vol. 84, pp. 355–369, November 2015.
- [47] Y. Zhu, P. Pan, S. Fang, L. Xu, J. Song, J. Zhang, and M. Feng, “The development and application of e-geoscience in china,” *Information Systems Frontiers*, vol. 18, pp. 1217–1231, June 2015.
- [48] E. Mizraji and J. Lin, “Modeling spatial-temporal operations with context-dependent associative memories,” *Cognitive neurodynamics*, vol. 9, pp. 523–534, May 2015.
- [49] D. C. Chen, S. Liu, P. R. Kingsbury, S. Sohn, C. B. Storlie, E. B. Habermann, J. M. Naessens, D. W. Larson, and H. Liu, “Deep learning and alternative learning strategies for retrospective real-world clinical data,” *NPJ digital medicine*, vol. 2, pp. 1–5, May 2019.
- [50] X. Yan, Q. Gu, F. Lu, J. Li, and J. Xu, “Gsa: a gpu-accelerated structure similarity algorithm and its application in progressive virtual screening,” *Molecular diversity*, vol. 16, pp. 759–769, October 2012.
- [51] Y. Xie and J. Wu, “Multi-objective constraint task scheduling algorithm for multi-core processors,” *Cluster Computing*, vol. 22, pp. 953–964, December 2018.
- [52] X. Xue and Y. Li, “Robust particle tracking via spatio-temporal context learning and multi-task joint local sparse representation,” *Multimedia Tools and Applications*, vol. 78, pp. 21187–21204, March 2019.
- [53] A. Huang, “A risk detection system of e-commerce: researches based on soft information extracted by affective computing web texts,” *Electronic Commerce Research*, vol. 18, pp. 143–157, July 2017.
- [54] E. A. Keshner, P. T. Weiss, D. Geifman, and D. R. Raban, “Tracking the evolution of virtual reality applications to rehabilitation as a field of study,” *Journal of neuroengineering and rehabilitation*, vol. 16, pp. 76–76, June 2019.
- [55] H. Jia, J. Xiong, J. Xu, and J. Wang, “Research and application of role theory in ocean carbon cycle ontology construction,” *Journal of Ocean University of China*, vol. 13, pp. 979–984, November 2014.
- [56] S. H. Rubin, T. Bouabana-Tebibel, and Y. Hoadjli, “On the empirical justification of theoretical heuristic transference and learning,” *Information Systems Frontiers*, vol. 18, pp. 981–994, June 2016.
- [57] K. Yang, X. Ding, Y. Zhang, L. Chen, B. Zheng, and Y. Gao, “Distributed similarity queries in metric spaces,” *Data Science and Engineering*, vol. 4, pp. 93–108, June 2019.
- [58] M. Fan, Q. Zhou, A. Abel, T. F. Zheng, and R. Grishman, “Probabilistic belief embedding for large-scale knowledge population,” *Cognitive Computation*, vol. 8, pp. 1087–1102, August 2016.
- [59] S. Hunt, Q. Meng, C. J. Hinde, and T. Huang, “A consensus-based grouping algorithm for multi-agent cooperative task allocation with complex requirements,” *Cognitive computation*, vol. 6, pp. 338–350, April 2014.
- [60] G. A. Mensah, W. Yu, W. Barfield, M. Clyne, M. M. Engलगau, and M. J. Khoury, “Hlbs-popomics: an on-line knowledge base to accelerate dissemination and

implementation of research advances in population genomics to reduce the burden of heart, lung, blood, and sleep disorders,” *Genetics in medicine : official journal of the American College of Medical Genetics*, vol. 21, pp. 519–524, September 2018.

- [61] M. T. Bahadori, Y. Liu, and D. Zhang, “A general framework for scalable transductive transfer learning,” *Knowledge and Information Systems*, vol. 38, pp. 61–83, April 2013.
- [62] A. G. Baydin, R. L. de Mántaras, and S. Ontañón, “A semantic network-based evolutionary algorithm for computational creativity,” *Evolutionary Intelligence*, vol. 8, pp. 3–21, November 2014.
- [63] A. Sharma and K. M. Goolsbey, “Learning search policies in large commonsense knowledge bases by randomized exploration,” 2018.
- [64] A. Mahmoud and N. Niu, “On the role of semantics in automated requirements tracing,” *Requirements Engineering*, vol. 20, pp. 281–300, January 2014.
- [65] J. N. K. Liu, Y.-L. He, E. H. Y. Lim, and X. Wang, “Domain ontology graph model and its application in chinese text classification,” *Neural Computing and Applications*, vol. 24, pp. 779–798, December 2012.
- [66] X. Zhang, D. Zhang, Y. L. Traon, Q. Wang, and L. Zhang, “Roundtable: Research opportunities and challenges for emerging software systems,” *Journal of Computer Science and Technology*, vol. 30, pp. 935–941, September 2015.
- [67] S. Gong, W. Hu, W. Ge, and Y. Qu, “Modeling topic-based human expertise for crowd entity resolution,” *Journal of Computer Science and Technology*, vol. 33, pp. 1204–1218, November 2018.
- [68] L. F. T. Moraes, S. Baki, R. M. Verma, and D. Lee, “Identifying reference spans: topic modeling and word embeddings help ir,” *International Journal on Digital Libraries*, vol. 19, pp. 191–202, June 2017.
- [69] D. Wu, H. Zhou, J. Shi, and N. Mamoulis, “Top-k relevant semantic place retrieval on spatiotemporal rdf data,” *The VLDB Journal*, vol. 29, pp. 893–917, November 2019.
- [70] H.-Y. Kwon, H. Wang, and K.-Y. Whang, “G-index model: A generic model of index schemes for top-k spatial-keyword queries,” *World Wide Web*, vol. 18, pp. 969–995, June 2014.
- [71] A. Charlesworth, “The comprehensibility theorem and the foundations of artificial intelligence,” *Minds and Machines*, vol. 24, pp. 439–476, November 2014.
- [72] C. Peltier and M. W. Becker, “Eye movement feedback fails to improve visual search performance,” *Cognitive research: principles and implications*, vol. 2, pp. 47–47, November 2017.
- [73] C. P. Walmsley, S. A. Williams, T. L. Grisbrook, C. Elliott, C. Imms, and A. Campbell, “Measurement of upper limb range of motion using wearable sensors: A systematic review,” *Sports medicine - open*, vol. 4, pp. 53–53, November 2018.
- [74] K. Su, G. Yang, L. Yang, P. Su, and Y. Yin, “Non-negative locality-constrained vocabulary tree for finger vein image retrieval,” *Frontiers of Computer Science*, vol. 13, pp. 318–332, April 2019.
- [75] R. W. L. So, S. W. Chung, H. H. C. Lau, J. J. Watts, E. Gaudette, Z. A. Al-Azzawi, J. Bishay, L. T.-W. Lin, J. Joung, X. Wang, and G. Schmitt-Ulms, “Application of crispr genetic screens to investigate neurological diseases,” *Molecular neurodegeneration*, vol. 14, pp. 1–16, November 2019.
- [76] J. Chen, L. Wang, and E. Charbon, “A quantum-implementable neural network model,” *Quantum Information Processing*, vol. 16, pp. 245–, August 2017.
- [77] K. Arlitsch, “Semantic web identity of academic libraries,” *Journal of Library Administration*, vol. 57, pp. 346–358, April 2017.
- [78] J. Wang, B. Gong, H. Liu, and S. Li, “Intelligent scheduling with deep fusion of hardware-software energy-saving principles for greening stochastic nonlinear heterogeneous super-systems,” *Applied Intelligence*, vol. 49, pp. 3159–3172, February 2019.
- [79] A. Abhishek and A. Basu, “A framework for disambiguation in ambiguous iconic environments,” in *AI 2004: Advances in Artificial Intelligence: 17th Australian Joint Conference on Artificial Intelligence, Cairns, Australia, December 4-6, 2004. Proceedings 17*, pp. 1135–1140, Springer, 2005.
- [80] W. V. Li and J. J. Li, “Modeling and analysis of rna-seq data: a review from a statistical perspective,” *Quantitative biology (Beijing, China)*, vol. 6, pp. 195–209, August 2018.
- [81] M. Arif and G. Wang, “Fast curvelet transform through genetic algorithm for multimodal medical image fusion,” *Soft Computing*, vol. 24, pp. 1815–1836, April 2019.
- [82] X. Chen and D. Wang, “Formalization and specification of geometric knowledge objects,” *Mathematics in Computer Science*, vol. 7, pp. 439–454, November 2013.
- [83] Y. Han, S. Chen, and Z. Feng, “Mining integration patterns of programmable ecosystem with social tags,” *Journal of Grid Computing*, vol. 12, pp. 265–283, January 2014.
- [84] A. Santella, R. Catena, I. Kovacevic, P. K. Shah, Z. Yu, J. Marquina-Solis, A. Kumar, Y. Wu, J. C. Schaff, D. A. Colón-Ramos, H. Shroff, W. A. Mohler, and Z. Bao, “Wormguides: an interactive single cell developmental atlas and tool for collaborative multidimensional data

- exploration.,” *BMC bioinformatics*, vol. 16, pp. 189–189, June 2015.
- [85] B. D. Nye, “Its, the end of the world as we know it: Transitioning aied into a service-oriented ecosystem,” *International Journal of Artificial Intelligence in Education*, vol. 26, pp. 756–770, February 2016.
- [86] H. Jiang, J. Zhang, H. Ma, N. Nazar, and Z. Ren, “Mining authorship characteristics in bug repositories,” *Science China Information Sciences*, vol. 60, pp. 012107–, November 2016.
- [87] C. W. Rosenbrock, E. R. Homer, G. Csányi, and G. L. W. Hart, “Discovering the building blocks of atomic systems using machine learning,” *npj Computational Materials*, vol. 3, pp. 29–, August 2017.
- [88] Z. Zhang, X. Qi, Y. Wang, C. Jin, J. Mao, and A. Zhou, “Distributed top- k similarity query on big trajectory streams,” *Frontiers of Computer Science*, vol. 13, pp. 647–664, November 2018.
- [89] A. Basu *et al.*, “Iconic interfaces for assistive communication,” in *Encyclopedia of Human Computer Interaction*, pp. 295–302, IGI Global, 2006.
- [90] L. Wu, “Student model construction of intelligent teaching system based on bayesian network,” *Personal and Ubiquitous Computing*, vol. 24, pp. 419–428, October 2019.
- [91] K.-W. Lim, R. Kawamoto, E. Andò, G. Viggiani, and J. E. Andrade, “Multiscale characterization and modeling of granular materials through a computational mechanics avatar: a case study with experiment,” *Acta Geotechnica*, vol. 11, pp. 243–253, August 2015.
- [92] R. Xu, L. Gui, Q. Lu, S. Wang, and J. Xu, “Incorporating multi-kernel function and internet verification for chinese person name disambiguation,” *Frontiers of Computer Science*, vol. 10, pp. 1026–1038, May 2016.
- [93] J. Wang and M. Gribskov, “Irespy: an xgboost model for prediction of internal ribosome entry sites,” *BMC bioinformatics*, vol. 20, pp. 409–409, July 2019.
- [94] Y. Huo, “Analysis of intelligent evaluation algorithm based on english diagnostic system,” *Cluster Computing*, vol. 22, pp. 13821–13826, February 2018.
- [95] R. C. Tillquist and M. E. Lladser, “Low-dimensional representation of genomic sequences,” *Journal of mathematical biology*, vol. 79, pp. 1–29, March 2019.
- [96] K. D. Onal, Y. Zhang, I. S. Altingovde, M. Rahman, P. Karagoz, A. Braylan, B. Dang, H.-L. Chang, H. Kim, Q. McNamara, A. Angert, E. Banner, V. Khetan, T. McDonnell, A. T. Nguyen, D. Xu, B. C. Wallace, M. de Rijke, and M. Lease, “Neural information retrieval: at the end of the early years,” *Information Retrieval Journal*, vol. 21, pp. 111–182, November 2017.
- [97] T. Ebesu and Y. Fang, “Neural semantic personalized ranking for item cold-start recommendation,” *Information Retrieval Journal*, vol. 20, pp. 109–131, March 2017.
- [98] M. Sharifi, D. A. Buzatu, S. C. Harris, and J. G. Wilkes, “Development of models for predicting torsade de pointes cardiac arrhythmias using perceptron neural networks,” *BMC bioinformatics*, vol. 18, pp. 497–497, December 2017.
- [99] A. Sharma, M. Witbrock, and K. Goolsbey, “Controlling search in very large commonsense knowledge bases: a machine learning approach,” *arXiv preprint arXiv:1603.04402*, 2016.
- [100] G. Zacharewicz, S. Y. Diallo, Y. Ducq, C. Agostinho, R. Jardim-Goncalves, H. Bazoun, Z. Wang, and G. Doumeingts, “Model-based approaches for interoperability of next generation enterprise information systems: state of the art and future challenges,” *Information Systems and e-Business Management*, vol. 15, pp. 229–256, May 2016.
- [101] A. S. Gibbons and M. B. Langton, “The application of layer theory to design: the control layer,” *Journal of Computing in Higher Education*, vol. 28, pp. 97–135, March 2016.
- [102] V. Vapnik, “Complete statistical theory of learning,” *Automation and Remote Control*, vol. 80, pp. 1949–1975, November 2019.
- [103] J. M. Homan, J. J. Cimino, S. G. Peters, B. Pickering, V. Herasevich, and M. A. Ellsworth, “A survey from a large academic medical center,” *Applied clinical informatics*, vol. 06, pp. 305–317, May 2015.
- [104] A. Jourabloo and X. Liu, “Pose-invariant face alignment via cnn-based dense 3d model fitting,” *International Journal of Computer Vision*, vol. 124, pp. 187–203, April 2017.
- [105] R. Zuech, T. M. Khoshgoftaar, and R. Wald, “Intrusion detection and big heterogeneous data: a survey,” *Journal of Big Data*, vol. 2, pp. 3–, February 2015.
- [106] W. Wallach and C. Allen, “Framing robot arms control,” *Ethics and Information Technology*, vol. 15, pp. 125–135, October 2012.
- [107] H. Ogoe, S. Visweswaran, X. Lu, and V. Gopalakrishnan, “Knowledge transfer via classification rules using functional mapping for integrative modeling of gene expression data,” *BMC bioinformatics*, vol. 16, pp. 226–226, July 2015.
- [108] A. Schroeder, B. P. Trease, and A. Arsie, “Balancing robot swarm cost and interference effects by varying robot quantity and size,” *Swarm Intelligence*, vol. 13, pp. 1–19, December 2018.
- [109] L. Ma, F. Ju, J. Wan, and X. Shen, “Emotional computing based on cross-modal fusion and edge network data

incentive,” *Personal and Ubiquitous Computing*, vol. 23, pp. 363–372, May 2019.

- [110] B. Ge, H. Wang, P. Wang, Y. Tian, X. Zhang, and T. Liu, “Discovering and characterizing dynamic functional brain networks in task fmri.,” *Brain imaging and behavior*, vol. 14, pp. 1660–1673, April 2019.
- [111] F. Sewell, N. Gellatly, M. Beaumont, N. Burden, R. A. Currie, L. de Haan, T. H. Hutchinson, M. N. Jacobs, C. Mahony, I. Malcomber, J. Mehta, G. Whale, and I. Kimber, “The future trajectory of adverse outcome pathways: a commentary.,” *Archives of toxicology*, vol. 92, pp. 1657–1661, March 2018.
- [112] S. Xie and P. S. Yu, “Active zero-shot learning: a novel approach to extreme multi-labeled classification,” *International Journal of Data Science and Analytics*, vol. 3, pp. 151–160, February 2017.
- [113] Z. Chunjie, X. Wang, Z. Zhiwang, Z. Zhang, and H. Qu, “The time model for event processing in internet of things,” *Frontiers of Computer Science*, vol. 13, pp. 471–488, November 2018.
- [114] L. van de Wijngaert, H. Bouwman, and N. Contractor, “A network approach toward literature review,” *Quality & Quantity*, vol. 48, pp. 623–643, November 2012.
- [115] M. R. Lakin and A. Phillips, “Automated analysis of tethered dna nanostructures using constraint solving,” *Natural Computing*, vol. 17, pp. 709–722, June 2018.
- [116] X. Wu, Z. Yang, Z. Li, H. Lin, and J. Wang, “Disease related knowledge summarization based on deep graph search.,” *BioMed research international*, vol. 2015, pp. 428195–428195, August 2015.
- [117] C. Dai, X. H. Cai, Y. Cai, Q. Huo, Y. Lv, and G. Huang, “An interval-parameter mean-cvar two-stage stochastic programming approach for waste management under uncertainty,” *Stochastic Environmental Research and Risk Assessment*, vol. 28, pp. 167–187, May 2013.