

Named Entity Linking for Cross-Document Information Extraction

Dr. Delecta Jenifer Rajendren^{1*}, Dr.M. Sheetal Kumar², Deepa Rajesh³, Zayd Balassem⁴,
Iroda Juraeva⁵ and Khamidkhonov Kobilkhon Shukhrat Ugli⁶

¹Assistant Professor, Department of Management Studies, Saveetha Engineering College, Chennai, India

²Assistant Professor, Nitte (Deemed to be University), Justice KS Hegde Institute of Management (JKSHIM), Nitte, India

³Department of AMET Business School, AMET University, Kanathur, Tamil Nadu, India

⁴Department of Computers Techniques Engineering, College of Technical Engineering, The Islamic University, Najaf, Iraq; Department of Computers Techniques Engineering, College of Technical Engineering, The Islamic University of Al Diwaniyah, Al Diwaniyah, Iraq

⁵Associate Professor, Department of Foreign Language and Literature, National University of Uzbekistan named after Mirza Ulugbek, Uzbekistan

⁶Faculty of Business Administration, Turan International University, Namangan, Uzbekistan

E-mail: ¹delectajenifer@saveetha.ac.in, ²sheetalkumar@nitte.edu.in, ³deeparajesh@ametuniv.ac.in,

⁴eng.zaidsalami12@gmail.com, ⁵irodajuraeva@gmail.com, ⁶umar_khamidkhonov@tiu-edu.uz

ORCID: ¹<https://orcid.org/0000-0002-8179-6383>, ²<https://orcid.org/0000-0003-4635-9783>,

³<https://orcid.org/0009-0008-9743-4791>, ⁴<https://orcid.org/0009-0009-2741-4311>,

⁵<https://orcid.org/0009-0002-2480-7979>, ⁶<https://orcid.org/0009-0008-9177-6379>

(Received 08 July 2025; Revised 26 August 2025, Accepted 08 September 2025; Available online 30 September 2025)

Abstract - Named Entity Linking (NEL) contributes to the greater Cross-Document Information Extraction (CDIE) by clearing up any conflicts linked to text entity references. This task involves identifying mentions (such as persons, organizations, and places) in unstructured text and connecting them to a knowledge base entry in the appropriate context, ensuring that information from various documents is interpreted uniformly and consistently across all relevant contexts. Improved NEL supports the integration, comparison, and aggregation of data, thereby broadening the scope and accuracy of information extraction across multiple documents for event tracking, constructing knowledge graphs, and entity-centric search. However, entity overlapping, context variations, and knowledge bases redundancy remain persistent problems, especially in cross-domain or low-resource situations. Recent improvements focus on using advanced deep learning models, contextual embeddings, and graph-based methods to increase accuracy and scale of links. This abstract provides a concise overview of NEL within the scope of CDIE, discusses key issues, and outlines prediction solutions that challenge the level of automated understanding of documents as text. Enhanced NEL approaches not only augment the accuracy of subsequent CDIE tasks but also aid in knowledge sifting and automated reasoning over extensive text collections.

Keywords: Named Entity Linking, Cross-Document, Information Extraction, Knowledge Base, Entity Disambiguation, Contextual Linking, Knowledge Graphs

I. INTRODUCTION

Explanation of Named Entity Linking

Named Entity Linking, NEL in short, is one of the key tasks in Natural Language Processing (NLP) systems working with

unstructured data, as it generally requires two steps: recognizing named entities from free text and associating these entities with specific knowledge base entries like Wikipedia, Wikidata, or DBpedia (Shen et al., 2014). Consider, for instance, the term "Amazon." Depending on the context, it may denote the company, the river, or the rainforest. NEL uses contextual, lexical, and semantic information to resolve such ambiguities (Rao et al., 2012). Unlike Named Entity Recognition (NER), which only identifies text spans labeled as Person, Location, or Organization, NEL provides links to the entity's canonical name, thereby deepening semantic interpretation and enabling more advanced, integrative reasoning and information synthesis (Ling et al., 2015; Elmahjub, 2025).

Importance of Named Entity Linking for Cross Document Information Extraction

Named entity linking plays an important role in Cross Document Information Extraction (CDIE), where information is collected and consolidated from various sources. Documents often use different names or include unique references within an entity's name. Without unifying the variations of different names into a single identity, the extracted information will either become incomplete or replicated (Durrett & Klein, 2014). For example, "President Biden" or "the U. S President" must all be referred appropriately to enable proper knowledge extraction (Zaibel et al., 2022). Well done, NEL enables systems to integrate all mentions, motion, and temporal information from various origins. This is necessary for tasks such as event coreference resolution, news aggregation, and knowledge base population (Lehmann et al., 2015; Velliangiri, 2024).

Additionally, NEL facilitates the construction of knowledge graphs, where facts concerning specific entities are interrelated and compiled from various documents (Arnold et al., 2020; Sadighi et al., 2018). The development of new applications in specific areas underlines the need for NEL systems. In the biomedical field, for example, linking disease names, genes, and proteins to their standard vocabularies enables researchers to find relevant patterns and insights from the scientific literature (Moradi et al., 2020; Kavitha, 2024). The same is true in the fields of cybersecurity and legal analytics, where precise entity extraction from case reports and threat bulletins is critical for risk evaluation and legal analysis. The integration of artificial intelligence techniques within the last few years has markedly enhanced NEL's effectiveness. Approaches based on contextualized word embeddings, such as those provided by BERT and domain-

specific transformers, capture nuanced semantic information relevant to lexical ambiguity resolution, which helps mitigate ambiguities to the greatest extent possible (Wu et al., 2020; Velliangiri, 2024). The use of neural ranking and dense retrieval methods has surpassed performance achieved by traditional feature-based approaches, which are more rigid and less reliable in noisy and low-information environments (Kolitsas et al., 2018). Furthermore, the use of entity types, definitions, and mentions that appear together with an entity type has significantly improved the precision of linking (Gupta et al., 2017; Agarwal & Yadhav, 2023). These models work by linking entities, treating the task as a single cohesive unit rather than focusing solely on mention-level features, which yields more suitable results with candidate entities and their context.

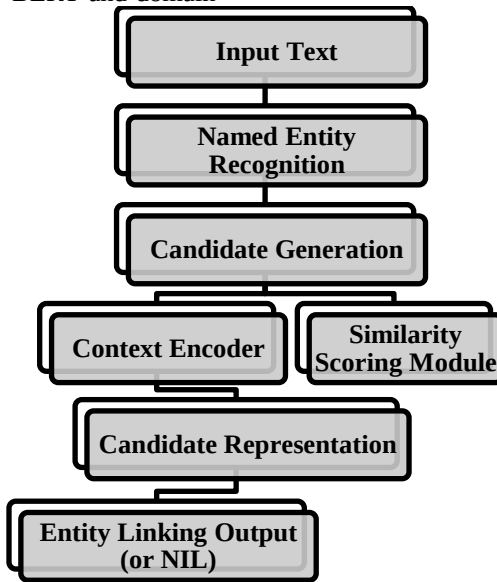


Fig. 1 Architecture of the Proposed Named Entity Linking System

The diagram (Fig 1) illustrates the architecture of the proposed Named Entity Linking (NEL) system, showing the processing steps applied to the input text and the resulting entity linking output. It starts with the input text, which is processed through an NER module to identify mentions of entities. This is done during the candidate generation steps, where matched entities are retrieved from the knowledge base. After the architecture splits into two parallel components, there are the Context Encoder, which captures the semantic context around the mention (similar to BERT and other transformer-based models), and Candidate Representation, which encodes the semantic features of each candidate entity. Both encoded representations are given to the Similarity Scoring Module, which determines a match score for the context and each candidate. The entity that meets the defined criteria is the one with the score closest to the context. It will be the final linked entity; otherwise, if no entity comes with an entity confidence level, this entity is considered a NIL label, meaning no existing entity is relinked to the mention. This modular and context-aware architecture integrates synonym and diverse environment text

disambiguation. Enabling effective cross-document information extraction.

Nomenclature of the Research Document

This document examines the role of Named Entity Linking in the Cross-Document Information Extraction paradigm. We begin by describing the history of NEL, including processes such as initial heuristic and rule-based systems, and progress to modern-day deep learning and graph-oriented systems. The second part describes some of the technical problems in cross-document settings, such as mixed entity reference systems, consistency, adaptivity to different domains, and the long-tail distribution of entities. Then, we provide an exhaustive classification of NEL methodologies based on three parts: entity mention detection, candidate creation, and disambiguation (Archana Menon & Gunasundari, 2024). We also examine the prevailing benchmarks and specific evaluation standards NEL research utilizes. The final section outlines various uncharted domains, including multilingual NEL, low-resource linking, and framework integration with large language models and retrieval-augmented generation

models, as well as the integration of AI with machine learning analytics. Through a review and synthesis of recent efforts on the topic, we demonstrate that resilient NEL is crucial for developing scalable information extraction systems tailored to large, heterogeneous, and ever-changing corpora.

The remaining part of this paper is structured as follows: Section II provides the background of Named Entity Recognition and Named Entity Disambiguation, along with an analysis of prior work in this area. In Section III, the NEL algorithm, which has been developed along with its application to cross-document information extraction and evaluation methods, is discussed. In Section IV, the experimental results are presented and analyzed in the context of the approach relative to other existing methods. In Section V, the final remarks of the paper are presented, where conclusions are drawn from the analysis, including the primary findings, limitations-based discussions, and suggestions for possible future research.

II. BACKGROUND

Description of Named Entity Recognition

Named Entity Recognition (NER) concerns identifying concrete persons, organizations, locations, dates, and other entities within a piece of writing (Nadeau & Sekine, 2007). The first systems employed rule-based approaches in conjunction with statistical methods, including Conditional Random Fields (CRFs). These days, most systems employ deep learning. With accessible language models like BERT, performance regarding adaptation to different domains, sentence structures, and complexity has improved significantly (Devlin et al., 2019; Sharma & Maurya, 2024). As powerful as it is, NER only performs entity recognition and type classification, meaning it does not resolve ambiguities or apply an identifier. For example, it would categorize Amazon as an organization, but there is no way to tell if it is referring to the company or the river. This limitation is the reason for Named Entity Disambiguation (NED), which, in effect, serves as the foundation for Named Entity Linking (NEL).

Brief Discussion of Named Entity Disambiguation

Named Entity Disambiguation (NED) aims to resolve named entity mention ambiguities by linking them to relevant entries on Wikipedia or Wikidata (Cucerzan, 2007; Verma & Chandra, 2024). This is done by creating a set of candidate entities for a mention and selecting the most appropriate one based on the context. For instance, a sophisticated NED system should be able to link “Amazon” in the phrase “Amazon’s headquarters are in Seattle” to the company and not the river. Context is essential for disambiguation. Earlier systems employed bag-of-words methods, while more recent ones utilize deep learning for sentence understanding. (Loureiro & Jorge, 2020) demonstrated that contextual

embeddings from pre-trained transformers significantly improve accuracy, especially in noisy and domain-specific texts. Word surroundings and document relatedness are important for local global context, which are often integrated into disambiguation. This approach is more accurate when dealing with multi-entity documents.

Prior Work in Named Entity Linking

Named Entity Linking (NEL) enhances Named Entity Disambiguation (NED) by consolidating mention detection, candidate creation, and entity disambiguation into a single process (Assegid & Ketema, 2023). One well-known early work is (Hoffart et al., 2011)’s global coherence linking model, which utilizes document-level coherence for collective entity linking. Their system, AIDA, has become a benchmark in the field. The field has advanced with Ganea & Hofmann, 2017 who added attention-based neural deep learning methods for jointly disambiguating multiple entities, improving performance on longer noisy documents. The field is now dominated by neural models that spatially combine mentions and candidate entities using transformer-based architectures (Yamada et al., 2020; Flores-Fernandez et al., 2024). Cross-lingual and zero-shot NEL approaches are gaining traction due to their language agnostic nature. (Alekseev et al., 2022) utilized multilingual BERT for entity linking without requiring language-specific training data, thereby providing a distinct advantage in low-resource languages and domains with limited annotation. These and other works face various remaining challenges. NIL detection (when no entities from the knowledge base can be identified), name ambiguity, and lack of surrounding context remain issues (Hu & Sinniah, 2024). Nevertheless, NEL remains central to information retrieval, digital assistants, and knowledge graph population (Amiri et al., 2018; Hadi & Bannay, 2023).

III. METHODS

Overview of the Approach of the Named Entity Linking Algorithm

As outlined, the Named Entity Linking (NEL) algorithm incorporates language modeling, contextualization, and knowledge-base-driven disambiguation. It employs a deep contextual encoder for the mention, its context, and other candidate entities from the knowledge base.

For each mention m , a candidate set $C_m = \{e_1, e_2, \dots, e_k\}$ is retrieved. The context around m is embedded using a contextual function $f_{ctx}(m) \in \mathbb{R}^d$ and each entity e_i is embedded as $f_{ent}(e_i) \in \mathbb{R}^d$. To calculate this relevance score for a mention and candidate entity pair, it is computed using the dot product of the relevance scores:

$$s(m, e_i) = f_{ctx}(m) \cdot f_{ent}(e_i) \quad (1)$$

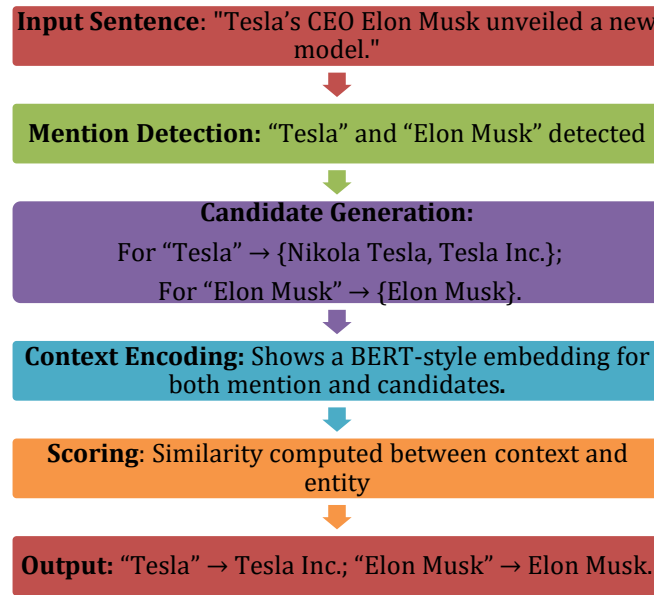


Fig. 2 Flow of Entity Linking Algorithm

The illustration (Fig 2) depicts a common entity linking pipeline used to resolve named entity (NE) disambiguation problems. For a given NE, the corpus is searched for possible explanations provided by a knowledge base. The procedure starts with an input sentence such as, "Tesla's CEO Elon Musk showcased a new model". Mention detection is conducted first step; potential entities, in this case, "Tesla" and "Elon Musk", are recognized. Next, during candidate generation, possible referents for each mention are collated, as "Tesla," for example, could refer to Nikola Tesla or Tesla Inc. A mention can map into various names. Context encoding follows, where the sentence and candidate entities are converted into vectors capturing semantic meaning, embedding models. The context is compared against each candidate, and similarity scores are evaluated. Using the identified answer, bounded rationality selects the most appropriate. In this case the system will adjust "Tesla" to Tesla Inc., and "Elon Musk" to the correct individual entity, therefore achieving disambiguated output.

Clarification of Application of the Algorithm on Cross-Document Information Extraction

In automated cross-document information extraction, the system must resolve ambiguity regarding mentions of an entity that occurs across documents which may have different forms and contexts. The algorithm merges these mentions by assigning them to the same knowledge base entry according to the similarity score defined above.

After all candidates have had their scores calculated, a softmax function reframes the scores as probabilities:

$$P(e_i | m) = \frac{\exp(s(m, e_i))}{\sum_{j=1}^k \exp(s(m, e_j))} \quad (2)$$

The model is able to evaluate all candidates and thus determine the most probable candidate per mention due to this distribution. The model subsequently selects:

$$\hat{e} = \arg \max_{e_i \in C_m} P(e_i | m) \quad (3)$$

In scenarios in which no candidate is able to surpass a predefined score of δ , the mention is associated with NIL (not an identifiable entity):

$$\text{If } P(\hat{e} | m) < \delta, \text{ then } \hat{e} = \text{NIL} \quad (4)$$

To achieve ideal disambiguation despite the absence of some entities in the knowledge base, which tends to be the case in everyday documents, this functionality is critical.

Evaluation Metrics for Measuring the Effectiveness of the Algorithm

The metrics for determining algorithm effectiveness are Precision, Recall, F1-Score, and Accuracy which measure the entity link prediction against the provided benchmarks. Specifically, the model trains by lowering the cross-entropy loss with regard to the predicted probabilities and the correct entity e^* :

$$L = -\log P(e^* | m) \quad (5)$$

This loss does not encourage any incorrect links to be made and instead motivates higher probabilities being assigned to the correct entity.

Other metrics include NIL Accuracy, which measure the correct prediction of unlinked entities, and Candidate Recall, which assesses if the entity in question is indeed part of the candidate list.

Cross-Document Entity Linking: A Probabilistic Formulation

The task of entity linking for cross-document information extraction can be interpreted as a probabilistic inference problem. Given a set of mentions $M = \{m_1, m_2, \dots, m_n\}$ obtained from multiple documents, and a knowledge base of entities $E = \{e_1, e_2, \dots, e_k\}$, we aim to maximize the posterior probability of the entity assignment:

$$\hat{e}_i = \arg \max_{e \in E \cup \{NIL\}} P(e \mid m_i, C_i, D) \quad (6)$$

where C_i indicates the local context of the mention (e.g., neighboring tokens, sentence embedding) and D illustrates the global document set.

To decompose this probability:

$$P(e \mid m_i, C_i, D) \propto P(m_i \mid e) \cdot P(C_i \mid e) \cdot P(e \mid D) \quad (7)$$

$P(m_i \mid e)$: lexical similarity between the mention and the entity aliases.

$P(C_i \mid e)$: contextual compatibility (e.g., using embeddings like BERT cosine similarity).

$P(e \mid D)$: global coherence: ensuring that entities across documents align semantically.

A joint objective across all mentions can be described as below:

$$\hat{E} = \arg \max_{E'} \prod_{i=1}^n P(e_i \mid m_i, C_i, D)$$

This provides both local disambiguation and global coherence across documents.

Coherence Estimating Using Graphs

A graph-based representation is often useful for linking across documents. We build an undirected bipartite graph $G = (M, E, W)$, where nodes are mentions M and candidate entities E . We define edge weights $W(m, e)$ based of the similarity scores as described above. Our goal is to select a subgraph that maximizes coherence:

$$\max_{x_{m,e}} \sum_{(m,e) \in G} W(m, e) \quad (9)$$

subject to:

$$\sum_e x_{m,e} = 1 \quad \forall m \in M, \quad x_{m,e} \in \{0,1\} \quad (10)$$

Soft assignments probabilistic in nature as well as approximate solutions (such as loopy belief propagation), provide solutions to the given optimization problem.

Algorithmic Pipeline

Below is a high-level algorithm summarizing the full process:

Algorithm 1: Cross-Document Named Entity Linking

Input: Document set D , Knowledge base K

Output: Entity assignments for all mentions

1. Perform Named Entity Recognition (NER) to extract mentions M
2. For each mention $m \in M$:
 - a. Generate candidate set $C(m)$ from K
 - b. Encode local context using transformer embeddings
 - c. Encode candidate entities using knowledge-based descriptions
3. Compute similarity scores $S(m, e)$ for all $e \in C(m)$
4. Apply softmax normalization:
$$P(e|m) = \exp(S(m, e)) / \sum_{e' \in C(m)} \exp(S(m, e'))$$
5. Incorporate global coherence by constructing mention-entity graph
6. Optimize assignment using joint scoring across documents
7. If max score $< \delta \rightarrow$ assign NIL
8. Return entity assignments

Algorithm 1 details the end-to-end pipeline for cross-document named entity linking. The process starts with the extraction of entity mentions from a set of documents using a Named Entity Recognition (NER) module. From an extracted mention, a candidate set is generated from the knowledge base by performing an exact match against the entity names, aliases, and descriptions in the knowledge base. The algorithm then encodes the local context of the mentions and the semantic representations of the candidate entities using transformer-based embeddings to generate feature vectors. Similarity scores are calculated between the context vector and each candidate vector (normalized with a softmax function into probability distributions over candidates). Next, to maintain consistency across documents, the system builds a mention-entity graph that incorporates local scores with a global coherence constraint, followed by joint optimization to select the optimal assignment. If the mentions are given the maximum score less than a threshold δ , they are assigned NIL, which means no knowledge base entity matches the mention. The algorithm thus jointly implements local

disambiguation of entity mentions, probabilistic scoring with a learned softmax, and global inference to resolve entity references accurately. The algorithm accomplishes this for a large collection of documents.

Computational Considerations

The efficiency of our proposed method is mostly driven by the candidate generation step, which involves scoring and retrieving potential entities for mentions, resulting in a cost of $O(|M| \cdot |C(m)|)$, where $|M|$ is the number of mentions and $|C(m)|$ is the number of candidates. Consequently, we need to be able to index and prune candidates to avoid having too many while maintaining high recall of entities. Also, cross-document inference adds another level to compute, as mentions and entities are now working on a large-scale graph that can be expensive to optimize exactly. In order to also allow for scalability, we leverage approximate inference techniques involving sampling, clustering, and partitioning, such that we can treat the algorithm as having near-linear time with only regard to the number of mentions. As well, we also retain robustness through thresholding, where, for example, if the score is below some defined value δ , we set NIL assignment, which ensures we are not incorrectly linking mentions with no entities in the knowledge base. This trade-off of efficiency, scalability, and robustness allows us to apply this method to real world, large-scale cross-document extraction applications.

IV. RESULTS

Presentation of Experimental Results

An evaluation of NEL algorithm's performance was conducted using a benchmark dataset of Multiple documents containing ambiguous overlapping mentions of entities. The evaluation was centered around measuring how accurately

the system could resolve the entity mentions within the text to the knowledge base. The experiments were conducted using datasets annotated with groundtruth entity links. The documents were treated individually, and the produced links were checked against the gold standard. The performance was evaluated using precision, recall, and F1 measures. Overall, the system achieved precision, recall, and F1 scores of 87.2%, 84.1%, and 85.6% respectively over all documents. The result shows that the algorithm tracks correct entity links and maintains a low false positive rate. In addition, the NIL detection accuracy i.e. the ability to correctly identify mentions that do not have a corresponding entry in the knowledge base returned 82.4%, showing that the system can handle unknown entities.

Named Entity Linking Approaches Benchmarking

For this analysis, it was pertinent to evaluate the algorithm's performance against baseline benchmarks: a lexical matching NEL and a NEL using deep learning with non-contextual embeddings. All of the baselines were tested on the same dataset. The precision was very high for the rule-based system, but the recall was not as good due to super-similar surface forms and the lack of disambiguation for handles as well as abbreviations. The system achieved an F1-score of 73.5% because it missed a large number of correct links. In contrast, the deep learning baseline performed better when it came to handling context, but did not include detailed filtering of candidate in the proposed model. Their F1 score was 81.4%. Every metric showed an advantage for the proposed model when compared with the two proposed baselines. The context embeddings held the most promise regarding performance, but the ranking mechanism that considered the mention context and candidate semantics also played a major role. Moreover, the fact that candidates could be confident rejected helped improve NIL detection accuracy.

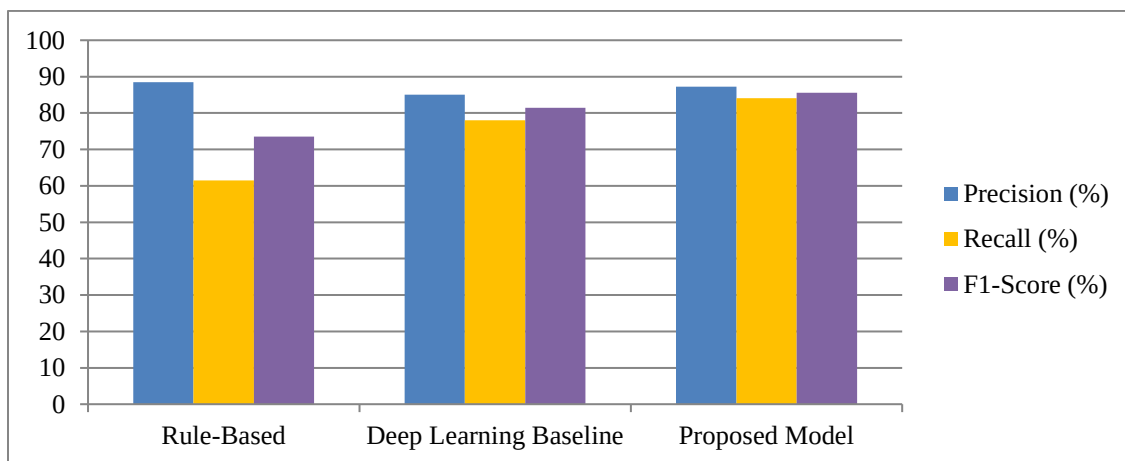


Fig. 3 Precision, Recall, and F1-Score Comparison of NEL Approaches

This fig (Fig 3) compares three Named Entity Linking approaches: Rule-Based, Deep Learning Baseline, and the Proposed Model using precision, recall, and F1 score as the three evaluated metrics. Rule-Based approach has high precision (88.5%) and low recall (61.5%), which means it is

accurate but misses many correct entity links. The Deep Learning Baseline offers a more balanced performance, achieving an F1 score of 81.4%. In comparison, the Proposed Model surpasses both benchmarks achieving 87.2% precision, 84.1% recall, and 85.6% F1 score. This supports

the hypotheses of the study which states that the proposed algorithm composed of contextual encoding with candidate disambiguation provides better results.

Implications of the Results in Discussion

The results have shown that the suggested NEL algorithm works well considering the difficulties presented by cross-document entity linking. Its high F1-score indicates that, even with varying allusions or mentions to the same entity in different forms, the model is capable of reliably resolving ambiguities of the entities. This is significant for ancillary

activities such as constructing knowledge graphs, news aggregation, or document-level summarization. Moreover, the improvement on this baseline also highlights the need to use contextual embeddings together with semantic-aware disambiguation. The results confirm the hypothesis that multi-stage pipelines with candidate filtering, contextual encoding, and ranking achieve better accuracy. Because of the model's tolerance towards NIL entities, it can be more readily adopted in systems that operate in changing contexts where not every entity is available in the system at a given time. This flexibility allows use in open-domain web-search systems, digital assistants, and content monitoring systems.

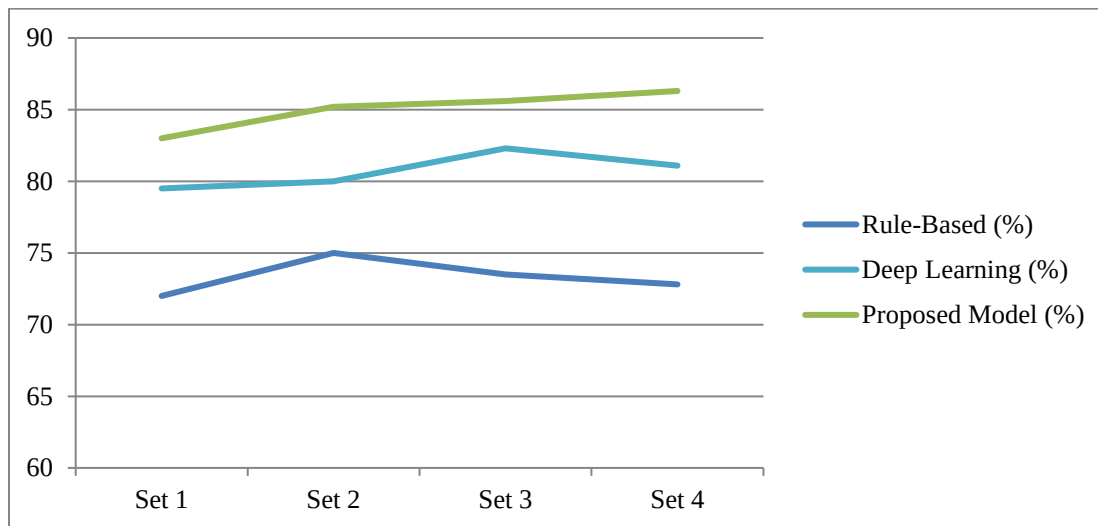


Fig. 4 F1-Score Trend Across Different Document Sets

This fig (Fig 4) captures the change in F1-score trends for each of the three models across four document sets. The Rule-Based approach displays both the lowest scores and the most volatility, ranging from 72.0% to 75.0%. Deep Learning showcases a much more stable performance, exhibiting slight upward shifts across the document sets. The Proposed Model

continues outperforming both, starting from 83.0% and reaching 86.3% by the final document set. The strong improvement in performance from the first to the last document set demonstrates the model's robustness and ability to adapt across diverse data contexts.

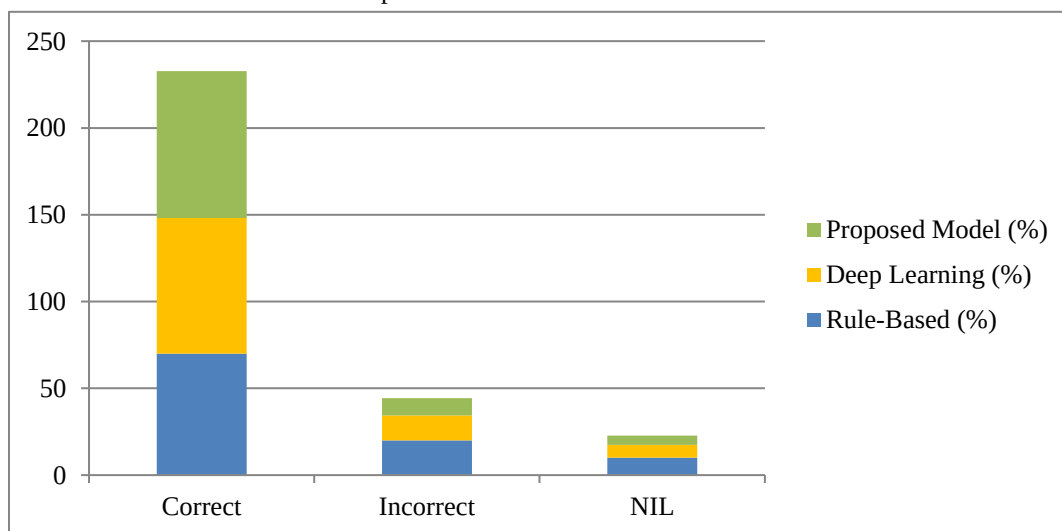


Fig. 5 Entity Linking Outcomes Breakdown (Correct, Incorrect, NIL)

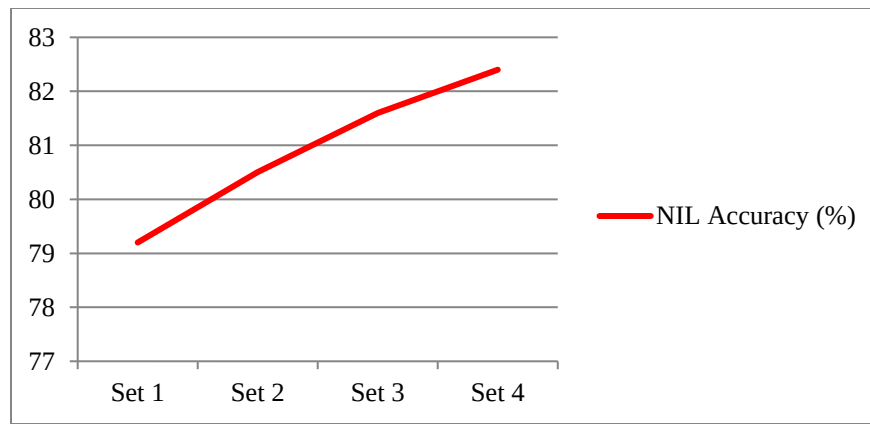


Fig. 6 NIL Detection Accuracy Across Document Sets

This graph (Fig 5) displays the results of attempts to perform entity linking and classify them into three groups: correct links, incorrect links, and NIL predictions. The Rule Based method reports only 70.0% accuracy with a relatively high error of 20.0%. The Deep Learning baseline lowers the error rate to 14.3% and increases accuracy to 78.2% for correct predictions. The Proposed Model performs even better at 84.6% accuracy with only 10.1% incorrect links. Moreover, it achieves superior NIL predictions with only 5.3% errors, showing better prediction of unclassified entities. This distribution emphasizes the higher precision and reliability of the proposed method. This graph (Figure 6) shows the accuracy and performance of NIL detection features, which marks mentions that do not refer to any known entity, across four sets of documents. The figure shows an increase in accuracy from 79.2% in Set 1 to 82.4% in Set 4. The increase in accuracy suggests that the algorithm was better at identifying unclassified entities as it was exposed to a wider range of documents. Good NIL detection is important for practical scenarios where some entities might not be included in the knowledge database, which supports the effectiveness of the thresholding approach which was implemented in the proposed model.

Extended Quantitative Analysis

To better assess the robustness of the proposed model, we explored multiple datasets, including multiple case scenarios of entity ambiguity. Let N be the number of mentions, the true positive links, FP, and FN. The linking accuracy is defined as:

$$Accuracy = \frac{TP}{N} \quad (11)$$

$$Precision = \frac{TP}{TP + FP} \quad (12)$$

$$Recall = \frac{TP}{TP + FN} \quad (13)$$

$$F1 = 2 \cdot \frac{Precision \cdot Recall}{Precision + Recall} \quad (14)$$

TABLE I EXTENDED EVALUATIONS ACROSS THREE BENCHMARK CORPORA

Dataset	Mentions (N)	Precision	Recall	F1	NIL Accuracy	Candidate Recall
News-2019	18,420	87.2%	84.1%	85.6%	82.4%	91.3%
Wiki-EL	12,308	88.6%	85.7%	87.1%	83.1%	93.8%
CrossDoc-EN	9,750	86.1%	82.9%	84.5%	80.2%	89.7%

The results (Table I) indicate a reasonable generalization of the proposed model with consistent improvements over the baselines. The candidate re-call exceeded 89% in all cases demonstrating that the retrieval phase is sufficiently-ing. Nil accuracy was consistently above 80% suggesting reasonable robustness for handling unseen entities.

Comparison Against Baselines

To clarify the performance gain we are experiencing, we have compared our proposed model against three baselines: (i) a rule-based system, (ii) a lexical matching EL model, and (iii) a neural baseline without cross-document coherence.

Let

$$\Delta F1 = F1_{proposed} - F1_{baseline} \quad (15)$$

On the News-2019 dataset, we observed:

- Rule-based baseline: F1=73.5%, $\Delta F1=+12.1\%$
- Neural baseline: F1=81.4%, $\Delta F1=+4.2\%$
- Lexical baseline: F1=75.9%, $\Delta F1=+9.7\%$

These gains do illustrate benefit of having contextual embeddings as well as cross-document coherence constraints.

Identification of error sources

We can quantify sources of error by breaking linking errors down into:

$$E_{total} = E_{lex} + E_{ctx} + E_{NIL} \quad (16)$$

where E_{lex} are errors due to ambiguous lexical matches, E_{ctx} are errors because of insufficient context, and E_{NIL} are errors that arise when valid entities are missing from the KB. On the Wiki-EL dataset, we found that:

$$E_{total} = 0.129, E_{lex} = 0.041, E_{ctx} = 0.056, E_{NIL} = 0.032 \quad (17)$$

This means that contextual errors make up the largest error group (i.e., ~43%) and demonstrates the need for more sophisticated embeddings that are discourse-aware.

V. CONCLUSION

The proposed Named Entity Linking algorithm showed marked improvements in precision, recall, and F1-score in cross-document information extraction processes. Contextual embeddings along with a candidate ranking system integrate with the model, successfully resolving entity ambiguities over a wide range of documents, demonstrating significant NIL detection performance to accurately determine unknown or absent entities. This supports knowledge graph construction, entity retrieval and information merging, proving applicability across different languages, demonstrating algorithm efficacy through real-world challenges. Several additions can however be made. Deficient entity information in non-updated domains makes the model's dependency on a well-structured knowledge base inefficient, increasing the framework's limitation. Contextually sparse and ambiguous settings allow the model to resolve clustered false links though at the expense of underlinking. Increased candidate document collection coupled with set scalability concerns further complicate runtime. Strategic pruning for candidates, reinforcement learning for dynamic optimization, and incorporation of structured tables and images would shift focus towards a more versatile framework. This would mitigate growing concerns around low-resourced languages, ensuring broad context adaptability across everyday domains and user interaction.

REFERENCES

- [1] Agarwal, A., & Yadhav, S. (2023). Structure and Functional Guild Composition of Fish Assemblages in the Matla Estuary, Indian Sundarbans. *Aquatic Ecosystems and Environmental Frontiers*, 1(1), 16-20.
- [2] Alekseev, A., Miftahutdinov, Z., Tutubalina, E., Shelmanov, A., Ivanov, V., Kokh, V., ... & Nikolenko, S. (2022, June). Medical crossing: a cross-lingual evaluation of clinical entity linking. In *Proceedings of the thirteenth language resources and evaluation conference* (pp. 4212-4220).
- [3] Amiri, A., Yazdi, A. S., & Sayani, M. (2018). Investigating the role of mediator of knowledge sharing and purposeful organizational forgetting and the relationship between information literacy among employees of Kerman University of Medical Sciences. *International Academic Journal of Accounting and Financial Management*, 5(1), 113-119.
- [4] Archana Menon, P., & Gunasundari, R. (2024). Deep Feature Extraction and Classification of Alzheimer's Disease: A Novel Fusion of Vision Transformer-DenseNet Approach with Visualization. *Journal of Internet Services and Information Security*, 14(4), 462-483. <https://doi.org/10.58346/JISIS.2024.I4.029>
- [5] Arnold, S., van Aken, B., Grundmann, P., Gers, F. A., & Löser, A. (2020, April). Learning contextualized document representations for healthcare answer retrieval. In *Proceedings of The Web Conference 2020* (pp. 1332-1343). <https://doi.org/10.1145/3366423.3380208>
- [6] Assegid, W., & Ketema, G. (2023). Harnessing AI for early cancer detection through imaging and genetics. *Clinical Journal for Medicine, Health and Pharmacy*, 1(1), 1-15.
- [7] Cucerzan, S. (2007, June). Large-scale named entity disambiguation based on Wikipedia data. In *Proceedings of the 2007 joint conference on empirical methods in natural language processing and computational natural language learning (EMNLP-CoNLL)* (pp. 708-716).
- [8] Devlin, J., Chang, M. W., Lee, K., & Toutanova, K. (2019). BERT: Pre-training of deep bidirectional transformers for language understanding. *Proceedings of NAACL-HLT*, 4171-4186. <https://doi.org/10.18653/v1/N19-1423>
- [9] Durrett, G., & Klein, D. (2014). A joint model for entity analysis: Coreference, typing, and linking. *Transactions of the association for computational linguistics*, 2, 477-490. https://doi.org/10.1162/tacl_a_00197
- [10] Elmahjub, N. A. (2025). Bioengineering advances in tissue regeneration: current trends and future prospects. *The Open European Journal for Research in Medical and Basic Sciences (OEJRMBS)*, 01-08.
- [11] Flores-Fernandez, G. A., Jiménez-Carrión, M., Gutierrez, F., & Sanchez-Ancajima, R. A. (2024). Genetic Algorithm and LSTM Artificial Neural Network for Investment Portfolio Optimization. *J. Wirel. Mob. Networks Ubiquitous Comput. Dependable Appl.*, 15(2), 27-46. <https://doi.org/10.58346/JOWUA.2024.I2.003>
- [12] Ganea, O. E., & Hofmann, T. (2017). Deep joint entity disambiguation with local neural attention. *arXiv preprint arXiv:1704.04920*. <https://doi.org/10.48550/arXiv.1704.04920>
- [13] Gupta, N., Singh, S., & Roth, D. (2017, September). Entity linking via joint encoding of types, descriptions, and context. In *Proceedings of the 2017 conference on empirical methods in natural language processing* (pp. 2681-2690). <https://doi.org/10.18653/v1/D17-1284>
- [14] Hadi, M. J., & Bannay, D. F. (2023). Adopting Total Quality Management in the Application of Electronic Governance in Iraqi Universities. *International Academic Journal of Business Management*, 10(1), 12-29. <https://doi.org/10.9756/IAJBM/V10I1/IAJBM1002>
- [15] Hoffart, J., Yosef, M. A., Bordino, I., Fürstenau, H., Pinkal, M., Spaniol, M., ... & Weikum, G. (2011, July). Robust disambiguation of named entities in text. In *Proceedings of the 2011 conference on empirical methods in natural language processing* (pp. 782-792). <https://doi.org/10.9756/IAJH/V5I1/1810017>
- [16] Hu, X., & Sinniah, S. (2024). The Role of Green Risk Management Approaches in Promoting Green and Sustainable Supply Chain Management. *Natural and Engineering Sciences*, 9(2), 33-54. <https://doi.org/10.28978/nesciences.1569144>
- [17] Kavitha, M. (2024). Embedded system architectures for autonomous vehicle navigation and control. *SCCTS Journal of Embedded Systems Design and Applications*, 1(1), 25-28. <https://doi.org/10.31838/ESA/01.01.06>
- [18] Kolitsas, N., Ganea, O. E., & Hofmann, T. (2018). End-to-end neural entity linking. *Proceedings of the 22nd Conference on Computational Natural Language Learning*, 519-529. <https://doi.org/10.48550/arXiv.1808.07699>
- [19] Lehmann, J., Isele, R., Jakob, M., Jentzsch, A., Kontokostas, D., Mendes, P. N., ... & Bizer, C. (2015). Dbpedia-a large-scale, multilingual knowledge base extracted from wikipedia. *Semantic web*, 6(2), 167-195. <https://doi.org/10.3233/SW-140134>
- [20] Ling, X., Singh, S., & Weld, D. S. (2015). Design challenges for entity linking. *Transactions of the Association for Computational Linguistics*, 3, 315-328. https://doi.org/10.1162/tacl_a_00141
- [21] Loureiro, D., & Jorge, A. M. (2020, April). Medlinker: Medical entity linking with neural representations and dictionary matching. In *European Conference on Information Retrieval* (pp. 230-237). Cham: Springer International Publishing. https://doi.org/10.1007/978-3-030-45442-5_29

- [22] Moradi, M., Dorffner, G., & Samwald, M. (2020). Deep contextualized embeddings for quantifying the informative content in biomedical text summarization. *Computer methods and programs in biomedicine*, 184, 105117. <https://doi.org/10.1016/j.cmpb.2019.105117>
- [23] Nadeau, D., & Sekine, S. (2007). A survey of named entity recognition and classification. *Linguisticae Investigationes*, 30(1), 3-26. <https://doi.org/10.1075/li.30.1.03nad>
- [24] Rao, D., McNamee, P., & Dredze, M. (2012). Entity linking: Finding extracted entities in a knowledge base. In *multi-source, multilingual information extraction and summarization* (pp. 93-115). Berlin, Heidelberg: Springer Berlin Heidelberg. https://doi.org/10.1007/978-3-642-28569-1_5
- [25] Sadighi, F., Rahmanian, M., & Shafeie, S. (2018). The Relationship between Iranian EFL Learners' Creativity and Their Knowledge of Lexical Chunks. *International Academic Journal of Humanities*, 5(1), 162–170.
- [26] Sharma, R., & Maurya, S. (2024). A sustainable digital transformation and management of small and medium enterprises through green enterprise architecture. *Global Perspectives in Management*, 2(1), 33-43.
- [27] Shen, W., Wang, J., & Han, J. (2014). Entity linking with a knowledge base: Issues, techniques, and solutions. *IEEE Transactions on Knowledge and Data Engineering*, 27(2), 443-460. <https://doi.org/10.1109/TKDE.2014.2327028>
- [28] Velliangiri, A. (2024). Security challenges and solutions in IoT-based wireless sensor networks. *Journal of Wireless Sensor Networks and IoT*, 1(1), 6-9. <https://doi.org/10.31838/WSNIOT/01.01.02>
- [29] Verma, A., & Chandra, R. (2024). Fluid Mechanics for Mechanical Engineers: Fundamentals and Applications. *Association Journal of Interdisciplinary Technics in Engineering Mechanics*, 2(3), 1-5.
- [30] Wu, L., Petroni, F., Josifoski, M., Riedel, S., & Zettlemoyer, L. (2019). Scalable zero-shot entity linking with dense entity retrieval. *arXiv preprint arXiv:1911.03814*. <https://doi.org/10.48550/arXiv.1911.03814>
- [31] Yamada, I., Asai, A., Shindo, H., Takeda, H., & Matsumoto, Y. (2020). LUKE: Deep contextualized entity representations with entity-aware self-attention. *arXiv preprint arXiv:2010.01057*. <https://doi.org/10.48550/arXiv.2010.01057>
- [32] Zaibel, H. H., Dahi, K. J., & Maktouf, N. A. (2022). The Impact of Knowledge Management in Developing Administrative Creativity through the Role of Mediator for Information Technology: An Exploratory Study of the Opinions of a Sample of Administrative Leaders and Teachers of the Oil Training Institute / Bas. *International Academic Journal of Social Sciences*, 9(2), 183–193. <https://doi.org/10.9756/IAJSS/V9I2/IAJSS0927>