

Indonesian named entity recognition using BiLSTM-CRF with domain-specific word embeddings

Joko Santoso

STKIP PGRI Metro, Lampung, Indonesia

Article Info

Article history:

Received Feb 15, 2026

Revised Mar 10, 2026

Accepted Mar 20, 2026

Keywords:

BiLSTM-CRF;
domain-specific embeddings;
Indonesian language;
named entity recognition;
sequence labeling.

ABSTRACT

Background: The rapid growth of Indonesian digital news content increases the demand for accurate Named Entity Recognition (NER), yet linguistic complexity and limited annotated data remain key challenges. **Objective:** This study evaluates the effectiveness of a BiLSTM-CRF model enhanced with domain-specific word embeddings for Indonesian NER across person, location, and organization entities. **Method:** A supervised sequence-labeling approach was applied using annotated news data, with embeddings trained on large-scale political and economic corpora and evaluated via precision, recall, and F1-score. **Results:** The model shows stable performance for person and location entities, while domain-specific embeddings improve all categories, especially organization entities with high variation; remaining errors relate to boundary detection and semantic ambiguity. **Implication:** These findings highlight the importance of domain-adaptive representations for improving NER systems in low-resource languages and support more reliable information extraction in Indonesian digital media. **Novelty:** This study demonstrates that embedding-level domain adaptation significantly enhances Indonesian NER without increasing model complexity, while clarifying distinctions between corpus resources and annotated data for better reproducibility.

This is an open access article under the [CC BY-SA](#) license.



Corresponding Author:

Joko Santoso

STKIP PGRI Metro, Lampung, Indonesia

Jl. Ki Hajar Dewantara, Banjar Rejo, Kec. Batanghari, Lampung Timur, Lampung 34381, Indonesia

Email: joko.spbsi@gmail.com

1. INTRODUCTION

The growth of digital textual data in Bahasa Indonesia has significantly intensified the demand for robust natural language processing (NLP) systems capable of extracting structured information from unstructured texts. Indonesia is among the world's largest producers of online news and social media content, with more than 215 million internet users and an estimated millions of news articles published annually, particularly in political and economic domains. Named Entity Recognition (NER) plays a pivotal role in transforming this vast textual landscape into actionable knowledge by identifying entities such as persons, locations, and organizations. However, Indonesian presents distinctive linguistic challenges, including rich morphology, flexible word order, inconsistent capitalization, and extensive use of local names and abbreviations. These characteristics render direct adoption of NER models developed for high-resource languages ineffective. Consequently, advancing Indonesian NER is not merely a technical endeavor but a crucial step toward improving information retrieval, media monitoring, policy analysis, and digital governance in one of the world's most linguistically dynamic societies.

From a scholarly perspective, NER research has evolved from rule-based and statistical approaches to deep learning architectures, particularly recurrent neural networks and conditional random fields [1], [2], [3]. Previous studies demonstrate that BiLSTM-CRF models consistently outperform traditional machine-

learning methods by capturing bidirectional contextual dependencies and enforcing global label consistency [4], [5], [6]. In the Indonesian context, several works have applied BiLSTM-CRF using general-purpose embeddings or multilingual pretrained vectors, yielding moderate performance gains [4], [7], [8]. More recent studies have explored transformer-based architectures [9], [10]; however, these models often require extensive computational resources and large annotated corpora that remain scarce for low-resource languages. Critically, existing Indonesian NER research has paid limited attention to domain-specific word embeddings, despite evidence from other languages that embeddings trained on in-domain corpora substantially improve semantic representation [2], [6], [7]. As a result, the interaction between BiLSTM-CRF architectures and domain-specific embeddings for Indonesian NER remains underexplored, creating a methodological gap that this study seeks to address.

This study aims to investigate how domain-specific word embeddings enhance the performance of a BiLSTM-CRF model for Indonesian Named Entity Recognition. Specifically, the research addresses the following questions: (1) To what extent can a BiLSTM-CRF model effectively recognize PER, LOC, and ORG entities in Indonesian news texts? (2) How does the use of domain-specific embeddings trained on Indonesian political and economic news compare to general embeddings in improving NER performance? (3) What types of errors persist despite architectural and representational improvements? To answer these questions, the study employs a labeled Indonesian NER corpus as token-level data while leveraging large-scale, domain-specific news corpora to train word embeddings. By distinguishing between corpus-level textual resources and annotated token-level data, this research ensures methodological clarity and reproducibility, addressing long-standing concerns in Indonesian NLP research regarding dataset transparency.

The central argument advanced in this study is that domain-specific word embeddings significantly enhance the capacity of BiLSTM-CRF models to capture semantic and contextual nuances in Indonesian NER tasks. By embedding lexical knowledge derived from Indonesian news discourse, the model is expected to better resolve local name variations, institutional abbreviations, and morphologically complex entities. This research hypothesizes that such embeddings will yield measurable improvements in precision, recall, and F1-score, particularly for challenging entity types such as ORG. Beyond performance gains, the study also contributes theoretically by demonstrating that representational adaptation at the embedding level can partially compensate for limited annotated data in low-resource languages. Practically, the findings offer implications for scalable Indonesian NLP applications, including automated news analysis, knowledge graph construction, and domain-aware information extraction systems, reinforcing the strategic importance of linguistically grounded model design.

2. LITERATURE REVIEW

2.1. Named Entity Recognition (NER)

Named Entity Recognition (NER) is a foundational task in natural language processing concerned with identifying and classifying real-world entities such as persons, locations, and organizations within text [1], [11]. Conceptually, NER has been defined either as a sequence labeling problem or as a semantic information extraction task, depending on theoretical orientation [2], [3]. Linguistic-oriented definitions emphasize the role of syntax and namedness [4], [9], whereas computational perspectives frame NER as probabilistic inference over token sequences [11], [12]. In low-resource languages such as Indonesian, definitional boundaries become blurred due to morphological richness, flexible capitalization, and context-dependent entity realization. Recent literature highlights that NER accuracy is strongly influenced by how entities are conceptualized—whether as surface forms, semantic referents, or discourse-bound units—underscoring the importance of context-sensitive modeling [1].

Scholars have further operationalized NER through a set of well-established categories and annotation schemes [4], [11]. Most studies adopt entity taxonomies such as PER, LOC, ORG, and MISC [11], while others expand categories to capture geopolitical entities, facilities, or products [12]. Annotation strategies such as BIO and BILOU tagging schemes are widely used to encode entity boundaries at the token level, each offering different trade-offs between granularity and complexity. Evaluation indicators commonly include precision, recall, and F1-score, with some studies incorporating span-level accuracy and error typologies. Collectively, these categorizations shape how NER systems are trained and evaluated, particularly in languages where multi-token and nested entities are prevalent.

2.2. BiLSTM-CRF framework

Recurrent neural network architectures have become central to contemporary NER research, with bidirectional long short-term memory networks combined with conditional random fields emerging as a dominant approach [5], [6]. BiLSTM architectures are valued for their ability to capture both preceding and following contextual information, while CRF layers model label dependencies to enforce sequence-level

consistency. However, literature reveals differing interpretations of the BiLSTM-CRF framework: some view it as a purely discriminative sequence model [7], whereas others emphasize its implicit linguistic modeling capacity [8]. These differences influence architectural choices such as layer depth, hidden dimensions, and feature integration, particularly when adapting models to morphologically complex languages.

The effectiveness of BiLSTM-CRF models has been assessed across various dimensions, including contextual sensitivity, robustness to sparse data, and generalization across domains [6], [13]. Empirical studies consistently report strong performance gains over standalone BiLSTM or CRF models, especially for well-defined entity types [2], [4]. Nevertheless, limitations persist in handling long-range dependencies and rare or ambiguous entities. Performance disparities across entity categories—most notably for organizations—are frequently reported, suggesting that structural complexity and lexical variation remain unresolved challenges. These findings motivate further exploration of representational enhancements beyond architectural optimization alone.

2.3. Word embeddings

Word embeddings constitute another critical strand of NER research, serving as the primary mechanism for encoding lexical semantics [9]. Early approaches relied on static embeddings trained on large, general-purpose corpora, while more recent studies emphasize domain-specific embeddings tailored to particular textual genres. Definitions of domain-specific embeddings vary, ranging from narrowly trained vectors on homogeneous corpora [7] to hybrid embeddings fine-tuned from general models [10]. In Indonesian NLP research, debates persist regarding whether localized semantic representations sufficiently capture socio-cultural naming conventions and linguistic variation, highlighting ongoing conceptual divergence [9].

Literature increasingly demonstrates that embedding quality directly affects NER performance, particularly in low-resource settings [2], [6]. Key aspects examined include corpus size, domain relevance, subword modeling, and morphological sensitivity [14]. Domain-specific embeddings have been shown to improve recognition of local names, institutional acronyms, and compound entities that are poorly represented in general corpora. Indicators of embedding effectiveness include downstream task performance, semantic similarity coherence, and error reduction rates. Collectively, these studies suggest that embedding adaptation is not merely a technical refinement but a strategic intervention that reshapes how NER models internalize linguistic and contextual knowledge.

3. METHOD

The unit of analysis in this study consists of token sequences representing named entities in Indonesian-language texts. The material object is Indonesian news discourse, while the formal object is token-level named entity labeling using the BIO tagging scheme. Two types of textual resources were employed: a labeled NER corpus and unlabeled domain-specific corpora for embedding training. The NER corpus comprises annotated news texts containing PER, LOC, ORG, and optional MISC entities, while the domain corpus consists of large-scale political and economic news articles. This distinction ensures methodological clarity between raw textual corpora and processed token-level data, as summarized in Table 1. This unit of analysis enables fine-grained modeling of contextual and sequential dependencies in Indonesian NER.

Table 1. Corpus and data description

No	Resource Type	Description	Size
1	General NER Corpus	Indonesian news texts labeled with PER, LOC, ORG	~12,000 sentences
2	Domain-Specific Corpus	Political & economic news articles	~5 million tokens
3	Validation/Test Data	Manually annotated token sequences	~2,000 sentences

This research adopts a quantitative experimental design within a supervised machine-learning framework. The core objective is to evaluate the effectiveness of BiLSTM-CRF models combined with domain-specific word embeddings for Indonesian NER tasks. Comparative experiments were conducted by training identical model architectures with different embedding configurations, allowing controlled assessment of embedding impact. Such experimental designs are widely employed in NER research to isolate representational effects from architectural influences. This study follows a train–validation–test split protocol to ensure generalizability and prevent data leakage. This design supports systematic performance comparison across entity categories and embedding strategies [4].

Information sources for this study were derived from both primary and secondary linguistic data. Primary data consist of manually annotated Indonesian NER datasets, serving as gold-standard labels for supervised learning and evaluation. Secondary data include large-scale, publicly accessible Indonesian news

articles used to train domain-specific word embeddings. These texts were sourced from reputable national media outlets to ensure linguistic authenticity and domain relevance. Previous studies emphasize that embedding corpora must reflect the discourse domain of the target task to maximize semantic alignment [15], [16]. By integrating authoritative news sources and expert-annotated labels, the study ensures both empirical rigor and linguistic validity.

Data collection followed a multi-stage preprocessing and annotation workflow. Unlabeled news articles were collected through web crawling and subsequently cleaned through tokenization, lowercasing normalization, and punctuation handling. For labeled data, existing NER datasets were supplemented with additional manual annotations performed by trained annotators following standardized annotation guidelines [17]. Inter-annotator agreement was calculated to ensure label consistency and reliability. The annotated data were then converted into BIO-formatted sequences compatible with sequence-labeling models. This structured collection process ensures that both corpus-level embeddings and token-level annotations meet reproducibility and quality standards.

Data analysis was conducted through sequential modeling and quantitative evaluation. Word embeddings were first trained using Word2Vec and FastText algorithms on the domain-specific corpus [9]. These embeddings were then integrated into a BiLSTM-CRF architecture to perform NER classification. Model performance was evaluated using precision, recall, and F1-score at the entity level. Comparative analysis was conducted between models using general-purpose and domain-specific embeddings, followed by error analysis focusing on entity confusion patterns and boundary detection. This analytical procedure enables robust assessment of both model effectiveness and linguistic limitations.

4. RESULTS

4.1. Performance of BiLSTM-CRF on Indonesian NER tasks

The performance of the BiLSTM-CRF model on Indonesian Named Entity Recognition tasks is summarized in Table 2, which reports entity-level precision, recall, and F1-scores across PER, LOC, ORG, and overall categories. The table is constructed based on evaluation results from the held-out test set derived from Indonesian news texts, consistent with the annotated corpus described in the Method section. Similar evaluation protocols are widely adopted in sequence-labeling studies to ensure comparability across models and datasets. By presenting disaggregated scores for each entity type, the table enables a fine-grained assessment of model behavior rather than relying solely on aggregate performance. This visualization provides an empirical foundation for identifying strengths and weaknesses of the BiLSTM-CRF architecture when applied to linguistically complex, low-resource languages such as Indonesian.

Table 2. Performance of BiLSTM-CRF model on Indonesian named entity recognition tasks

Entity Type	Precision (%)	Recall (%)	F1-score (%)	Support (Entity Count)	Avg. Entity Length (Tokens)	Boundary Accuracy (%)	Confusion Rate (%)
PER	91.84	90.27	91.05	3,842	2.1	93.12	4.6
LOC	89.76	88.43	89.09	2,965	2.4	91.48	6.1
ORG	84.21	79.63	81.86	2,417	3.7	86.02	11.8
MISC	82.35	80.14	81.23	1,128	2.9	87.55	9.4
Micro Average	88.74	86.92	87.82	10,352	–	89.63	7.5
Macro Average	87.04	84.62	85.81	–	–	89.54	8.0

Notes:

- Precision, recall, and F1-score are computed at the entity level using exact span matching.
- *Support* indicates the number of gold-standard entities per category in the test set.
- *Average entity length* reflects the mean number of tokens per entity span.
- *Boundary accuracy* measures correct identification of entity start–end positions.
- *Confusion rate* indicates misclassification frequency with other entity types, particularly ORG–LOC overlaps.

Table 2 reveals a clear performance gradient across entity types. The BiLSTM-CRF model achieves its highest F1-score on PER entities, followed closely by LOC, while ORG entities exhibit comparatively lower scores. Specifically, person names demonstrate strong precision and recall, indicating stable recognition of personal naming patterns in Indonesian news discourse. Location entities also show robust performance, reflecting relatively consistent lexical and contextual cues. In contrast, organization entities display reduced recall and more variable precision, suggesting higher classification difficulty. The overall

micro-averaged F1-score indicates competitive performance relative to previously reported Indonesian NER benchmarks. These patterns align with earlier findings that entity complexity and lexical variability substantially influence NER accuracy in news-based corpora [2], [11].

Analytically, the observed performance differences can be attributed to linguistic and structural properties inherent to each entity category. PER entities often follow predictable syntactic patterns and capitalization conventions, enabling the BiLSTM to capture contextual dependencies effectively. LOC entities similarly benefit from recurring geographic markers and prepositional cues. ORG entities, however, tend to exhibit longer multi-token spans, frequent abbreviations, and semantic overlap with LOC, particularly in Indonesian institutional naming practices. Prior studies have noted that such ambiguity challenges sequence-labeling models, even with CRF-based label optimization [4]. Consequently, while the BiLSTM-CRF architecture demonstrates strong baseline effectiveness for Indonesian NER, the results indicate that representational enhancements are necessary to address entity categories characterized by high lexical and structural variation.

Table 3. Comparative performance of BiLSTM-CRF models using general vs. domain-specific word embeddings

Entity Type	Embedding Type	Precision (%)	Recall (%)	F1-score (%)	Δ F1 (Abs.)	Δ F1 (%)	Boundary Accuracy (%)	OOV Handling Gain (%)
PER	General Embedding	90.62	89.11	89.86	–	–	92.08	–
	Domain-Specific Embedding	91.84	90.27	91.05	+1.19	+1.32	93.12	+8.4
LOC	General Embedding	88.01	86.72	87.36	–	–	90.12	–
	Domain-Specific Embedding	89.76	88.43	89.09	+1.73	+1.98	91.48	+10.1
ORG	General Embedding	81.34	75.92	78.54	–	–	82.87	–
	Domain-Specific Embedding	84.21	79.63	81.86	+3.32	+4.23	86.02	+17.6
MISC	General Embedding	80.27	77.65	78.94	–	–	84.31	–
	Domain-Specific Embedding	82.35	80.14	81.23	+2.29	+2.90	87.55	+13.9
Micro Avg.	General Embedding	86.71	84.09	85.38	–	–	87.46	–
	Domain-Specific Embedding	88.74	86.92	87.82	+2.44	+2.86	89.63	+12.7
Macro Avg.	General Embedding	85.56	82.85	84.17	–	–	87.35	–
	Domain-Specific Embedding	87.04	84.62	85.81	+1.64	+1.95	89.54	+12.5

Notes:

- Δ F1 (Abs.): the absolute difference in F1-score between the domain-specific embedding and the general embedding.
- Δ F1 (%): the percentage improvement relative to the general embedding baseline.
- Boundary Accuracy: the accuracy of entity boundary detection (start–end token).
- OOV Handling Gain: the improvement in the recognition success of out-of-vocabulary words, especially local names, institutional abbreviations, and news domain terms.

4.2. Impact of domain-specific word embeddings

The impact of domain-specific word embeddings on Indonesian NER performance is presented in Table 3, which compares the BiLSTM-CRF model trained with general-purpose embeddings and domain-specific embeddings derived from Indonesian political and economic news corpora. The visualization follows established evaluation practices in NER research, where embedding configurations are treated as controlled experimental variables. By holding the model architecture constant, the table isolates the contribution of embedding quality to downstream sequence-labeling performance. Similar comparative tables are commonly used in embedding-sensitive studies to demonstrate representational effects across entity categories [10]. This structured presentation allows for direct assessment of how in-domain semantic representations influence recognition accuracy for PER, LOC, and ORG entities within Indonesian news discourse.

Table 3 demonstrates a consistent performance improvement when domain-specific embeddings are employed. Across all entity types, models using in-domain embeddings achieve higher precision, recall, and F1-scores compared to their general-embedding counterparts. The most pronounced gains are observed for ORG entities, where F1-score improvements exceed those for PER and LOC. Moderate yet stable improvements are also evident for LOC, while PER shows marginal but positive gains, reflecting its already strong baseline performance. The overall micro-averaged F1-score increases noticeably, indicating that embedding adaptation contributes to global model robustness. These numerical patterns suggest that embeddings trained on Indonesian news corpora better encode lexical variations, institutional naming conventions, and contextual semantics specific to the target domain [9].

From an analytical perspective, the observed improvements can be attributed to enhanced semantic alignment between the embedding space and the linguistic characteristics of the evaluation corpus. Domain-specific embeddings capture distributional regularities of Indonesian news discourse, including abbreviations, compound organization names, and morphologically inflected forms that are underrepresented in general corpora. Prior studies have shown that such embeddings reduce semantic sparsity and improve contextual discrimination in low-resource settings [12]. The disproportionate improvement for ORG entities further suggests that embedding-level knowledge plays a critical role in resolving structurally complex and lexically diverse entity spans. These findings support the argument that representational adaptation, rather than architectural complexity alone, is a key driver of NER performance gains in Indonesian-language contexts.

4.3. Error analysis and model limitations

The distribution and typology of model errors are summarized in Table 4, which presents a detailed breakdown of misclassification patterns observed in the test set. The table categorizes errors based on entity confusion (e.g., ORG–LOC), boundary detection failures, and orthographic inconsistencies. This form of error visualization is commonly adopted in NER studies to complement aggregate performance metrics with qualitative insights into model behavior. By aligning error categories with entity types and embedding configurations, the table provides a structured overview of where and how the BiLSTM-CRF model fails. Such visualization is particularly important for low-resource languages, where performance bottlenecks often stem from linguistic ambiguity rather than architectural inadequacy.

The error distribution reveals several dominant patterns. The most frequent error type involves confusion between ORG and LOC entities, accounting for a substantial proportion of total misclassifications. Boundary-related errors are also prevalent, especially for long, multi-token organization names. Additionally, inconsistencies in capitalization and the presence of foreign loanwords contribute to incorrect labeling, particularly in headlines and abbreviated news texts. While domain-specific embeddings reduce overall error rates—especially for OOV terms—certain categories remain problematic. Notably, MISC entities exhibit heterogeneous error patterns due to their semantic diversity. These descriptive findings suggest that while representational improvements mitigate lexical sparsity, structural and annotation-related challenges persist [3], [16].

From an interpretive standpoint, these error patterns reflect deeper linguistic and methodological constraints. ORG–LOC ambiguity is rooted in Indonesian naming conventions, where institutions frequently incorporate geographic identifiers, blurring semantic boundaries. Boundary detection errors arise from the model's limited capacity to represent long-span dependencies, despite CRF-based label optimization. Capitalization errors indicate sensitivity to orthographic variation, a known issue in Indonesian news discourse. Prior literature suggests that such limitations may be addressed through character-level modeling, syntactic features, or contextualized embeddings [2], [8], [18], [19]. Therefore, while the BiLSTM-CRF with domain-specific embeddings demonstrates strong overall performance, the error analysis underscores the need for richer linguistic features and expanded annotation schemes to further advance Indonesian NER robustness.

Table 4. Error typology and distribution in Indonesian NER using BiLSTM-CRF with domain-specific embeddings

Error Category	Entity Type Affected	Error Description	Frequency (Count)	Proportion (%)	Avg. Entity Length (Tokens)	Reduced by Domain Embedding (%)
ORG to LOC Confusion	ORG	Organization names misclassified as locations	312	23.6	3.9	18.4
LOC to ORG Confusion	LOC	Location entities labeled as organizations	147	11.1	3.1	12.7
Boundary Truncation	ORG	Partial recognition of multi-token entities	221	16.7	4.6	14.2
Boundary Expansion	ORG / MISC	Over-extended entity spans	103	7.8	4.2	9.6
Capitalization Error	PER / ORG	Incorrect labeling due to inconsistent casing	96	7.3	2.3	6.1
Foreign Loanword Mislabeling	ORG / MISC	Foreign terms misclassified or ignored	88	6.6	3.5	15.9
OOV Token Failure	ORG	Unrecognized rare or local entity names	79	6.0	3.8	22.5
Nested Entity Collapse	ORG / LOC	Failure to detect nested entity structures	64	4.8	4.1	5.4
Annotation Ambiguity	MISC	Inconsistent or vague category boundaries	113	8.6	2.7	3.2
Other Errors	All	Residual misclassification patterns	99	7.5	–	4.1
Total Errors	–	–	1,322	100.0	–	–

Notes:

- Reduced by Domain Embedding (%) indicates the percentage of errors reduced compared to the general embedding model (see Table 3).
- ORG–LOC errors were the largest contributor to the ORG F1-score reduction (Table 2).
- Boundary-related errors (truncation + expansion) accounted for 24.5% of the total error.
- Errors that were not significantly reduced by embedding indicate structural limitations of the model, not representational ones.

5. DISCUSSION

The empirical evidence demonstrates that the BiLSTM-CRF architecture provides a reliable functional baseline for Indonesian Named Entity Recognition, particularly for person and location entities. This finding carries important implications for applied NLP in low-resource languages, where stability and interpretability often outweigh marginal accuracy gains. The model's consistent performance across these entity types suggests that sequence-based contextual learning can effectively capture recurrent syntactic and semantic patterns in Indonesian news discourse [9], [20], [21], [22]. At the same time, the observed

performance gap for organization entities highlights a structural dysfunction: standard architectures remain unevenly effective across entity categories with varying linguistic complexity. This imbalance signals that baseline NER systems, while operationally viable, may reproduce systematic blind spots when deployed in real-world applications such as media monitoring or knowledge extraction. Consequently, the results underscore the necessity of complementing architectural robustness with representational and linguistic adaptations.

The underlying causes of this performance differentiation can be traced to the interaction between Indonesian linguistic structures and sequence-model assumptions. Person and location entities tend to exhibit stable surface cues, such as capitalization, honorifics, or geographic markers, which align well with BiLSTM contextual encoding and CRF transition constraints. Organization entities, however, frequently involve longer spans, embedded nominal phrases, and semantic overlap with locations, introducing ambiguity that challenges linear sequence modeling. Prior studies have shown that such entities strain models that rely primarily on local context windows and token-level dependencies [10], [15]. This suggests that the observed performance pattern is not incidental but structurally rooted in how BiLSTM-CRF models abstract linguistic information. As a result, the architecture's strengths and limitations are intrinsically tied to the syntactic regularity and semantic coherence of entity categories.

The introduction of domain-specific word embeddings yields substantive functional gains, particularly in addressing representational gaps inherent in general-purpose embeddings. The improvement in recognition accuracy, most notably for organization entities, indicates that embedding adaptation enhances the model's ability to internalize domain-bound lexical semantics. This has broader implications for Indonesian NLP research, as it demonstrates that performance gains need not rely solely on more complex architectures or larger annotated datasets. Instead, aligning the embedding space with the discourse domain offers a scalable and computationally efficient pathway to improvement. Practically, this finding supports the deployment of domain-aware NER systems in specialized contexts such as political analysis, economic reporting, and institutional document processing [3], [5]. The functional contribution of embedding specialization thus extends beyond numerical improvement, reinforcing the importance of contextualized lexical knowledge.

The causal mechanisms behind these improvements lie in the distributional properties captured by domain-specific embeddings. By training on large-scale Indonesian news corpora, the embeddings encode frequent co-occurrence patterns of institutional names, abbreviations, and morphologically inflected forms that are underrepresented in general corpora. This richer semantic grounding reduces out-of-vocabulary effects and sharpens contextual discrimination during sequence labeling. The disproportionate improvement observed for organization entities further suggests that embedding-level semantic cohesion plays a critical role in resolving structurally complex spans. These findings resonate with prior literature emphasizing that representation quality often exerts a stronger influence on NER performance than architectural depth [9]. Thus, the observed gains are best understood as the outcome of structural alignment between lexical representation and task-specific discourse.

The analysis of error patterns provides further insight into the functional boundaries of the proposed approach. While overall error rates decrease with domain-specific embeddings, certain error types persist, particularly those involving entity boundary detection and semantic ambiguity. These limitations have practical implications for downstream applications that require high precision in entity span extraction, such as relation extraction or knowledge graph construction. The persistence of boundary-related errors indicates that sequence-labeling models, even when enhanced with improved embeddings, struggle with long-range dependencies and nested structures. Moreover, the limited reduction of errors associated with annotation ambiguity highlights that model performance is partially constrained by dataset design rather than algorithmic capacity [12], [23]. These dysfunctions suggest that performance ceilings are shaped by both linguistic complexity and methodological choices.

At a structural level, the remaining errors reflect deeper constraints inherent in flat sequence-labeling frameworks. Ambiguities between organizations and locations are embedded in Indonesian naming conventions, where geographic references frequently function as institutional identifiers [24]. Boundary errors arise from the model's reliance on token-local transitions, which inadequately represent hierarchical or nested entity structures. Orthographic inconsistencies further expose sensitivity to surface-level variation, a known challenge in Indonesian news texts. These patterns indicate that while embedding adaptation mitigates lexical sparsity, it cannot fully compensate for structural limitations in model design. Addressing these issues likely requires integrating character-level features, syntactic representations, or contextualized embeddings capable of modeling hierarchical semantics [25], [26], [27]. Accordingly, the discussion situates the present findings within a broader trajectory of Indonesian NER research, emphasizing incremental yet strategically significant advances rather than exhaustive solutions.

6. CONCLUSION

This study demonstrates that integrating a BiLSTM-CRF architecture with domain-specific word embeddings constitutes an effective and methodologically sound approach for Indonesian Named Entity Recognition. The findings reveal that sequence-based contextual modeling reliably captures syntactic and semantic regularities in Indonesian news discourse, while embedding specialization significantly enhances representational alignment, particularly for structurally complex entity categories. Beyond empirical performance gains, the study contributes theoretically by reaffirming the centrality of representation quality over architectural complexity in low-resource language settings. Methodologically, it advances Indonesian NER research by clearly distinguishing between corpus-level textual resources and token-level annotated data, thereby strengthening reproducibility and analytical transparency.

Despite these contributions, the study is subject to several limitations that open avenues for future research. The reliance on news-domain corpora constrains generalizability to other genres such as social media or conversational text. In addition, the flat sequence-labeling framework limits the model's ability to represent nested or hierarchical entities, contributing to persistent boundary-related errors. Future studies should explore the integration of character-level representations, syntactic features, or contextualized embeddings to address these structural constraints. Expanding annotated datasets across diverse domains and refining annotation schemes would further enhance model adaptability.

ACKNOWLEDGMENTS

The author sincerely thanks colleagues and academic peers who provided technical support and constructive suggestions for this study.

FUNDING INFORMATION

Authors state no funding involved.

AUTHOR CONTRIBUTIONS STATEMENT

Joko Santoso: conceptualization (lead), methodology (lead), software (lead), formal analysis (lead), writing – original draft (lead), writing – review and editing (lead).

CONFLICT OF INTEREST STATEMENT

Authors state no conflict of interest.

INFORMED CONSENT

We have obtained informed consent from all individuals included in this study.

ETHICAL APPROVAL

This research related to human use has been complied with all the relevant national regulations and institutional policies in accordance with the tenets of the Helsinki Declaration and has been approved by the authors' institutional review board or equivalent committee.

DATA AVAILABILITY

Data availability is not applicable to this article as no new data were created or analyzed in this study.

REFERENCES

- [1] Y. Qiu, L. Dong, W. Zhang, H. Xing, and J. Huang, "A diffusion enhanced CRF and BiLSTM framework for accurate entity recognition," *Scientific Reports*, vol. 15, Jun. 2025, doi: 10.1038/s41598-025-04036-x.
- [2] E. Yulianti, N. Bhary, J. Abdurrohman, F. W. Dwitilas, E. Q. Nuranti, and H. S. Husin, "Named entity recognition on Indonesian legal documents: a dataset and study using transformer-based models," *International Journal of Electrical and Computer Engineering (IJECE)*, Oct. 2024, doi: 10.11591/ijece.v14i5.pp5489-5501.
- [3] E. Subowo, I. Bukhori, and Wanto, "Corpus Development and NER Model for Identification of Legal Entities (Articles, Laws, and Sanctions) in Corruption Court Decisions in Indonesia," *Transactions on Informatics and Data Science*, Jun. 2025, doi: 10.24090/tids.v2i1.13592.
- [4] I. M. K. Karo, S. Dewi, and A. S. Nasution, "Named Entity Recognition on Indonesian Online News Based on Bidirectional LSTM-CRF," *2025 4th International Conference on Electronics Representation and Algorithm (ICERA)*, pp. 251–256, Jun. 2025, doi: 10.1109/icera66156.2025.11087367.
- [5] K. L. Thant, K. Nongpong, Y. K. Thu, T. Aung, K. Wai, and T. Oo, "myNER: Contextualized Burmese Named Entity Recognition with Bidirectional LSTM and fastText Embeddings via Joint Training with POS Tagging," *2025 IEEE International Conference on Cybernetics and Innovations (ICCI)*, pp. 1–6, Apr. 2025, doi: 10.1109/icci64209.2025.10987237.
- [6] P. Chen, M. Zhang, X. Yu, and S. Li, "Named entity recognition of Chinese electronic medical records based on a hybrid neural network and medical MC-BERT," *BMC Medical Informatics and Decision Making*, vol. 22, Dec. 2022, doi: 10.1186/s12911-022-02059-2.
- [7] S. Arslan, "Application of BiLSTM-CRF model with different embeddings for product name extraction in unstructured Turkish text," *Neural Computing and Applications*, vol. 36, pp. 8371–8382, Feb. 2024, doi: 10.1007/s00521-024-09532-1.

- [8] R. Anam *et al.*, “A deep learning approach for Named Entity Recognition in Urdu language,” *PLOS ONE*, vol. 19, Mar. 2024, doi: 10.1371/journal.pone.0300725.
- [9] Z. Zainuddin, -. Mudassir, and Z. Tahir, “Entity Extraction in Indonesian Online News Using Named Entity Recognition (NER) with Hybrid Method Transformer, Word2Vec, Attention and Bi-LSTM,” *JOIV: International Journal on Informatics Visualization*, May 2025, doi: 10.62527/joiv.9.3.2902.
- [10] A. Trewartha *et al.*, “Quantifying the advantage of domain-specific pre-training on named entity recognition tasks in materials science,” *Patterns*, vol. 3, Apr. 2022, doi: 10.1016/j.patter.2022.100488.
- [11] G. Shidik *et al.*, “Indonesian Disaster Named Entity Recognition from Multi Source Information using Bidirectional LSTM (BiLSTM),” *Journal of Open Innovation: Technology, Market, and Complexity*, Aug. 2024, doi: 10.1016/j.joitmc.2024.100358.
- [12] E. Dave and A. Chowanda, “IPerFEX-2023: Indonesian personal financial entity extraction using indoBERT-BiGRU-CRF model,” *Journal of Big Data*, vol. 11, Sep. 2024, doi: 10.1186/s40537-024-00987-6.
- [13] R. Yan, X. Jiang, and D. Dang, “Named Entity Recognition by Using XLNet-BiLSTM-CRF,” *Neural Processing Letters*, vol. 53, pp. 3339–3356, Jun. 2021, doi: 10.1007/s11063-021-10547-1.
- [14] A. Fawaid, H. Baharun, M. Hamzah, Rohimah, I. Munawwaroh, and D. F. Putri, “AI-based career management to improve the quality of decision making in higher education,” in *2025 IEEE Integrated STEM Education Conference (ISEC)*, IEEE, Mar. 2025, pp. 1–8. doi: 10.1109/isec64801.2025.11147274.
- [15] N. Hussain, A. Qasim, G. Mehak, O. Kolesnikova, A. Gelbukh, and G. Sidorov, “ORUD-Detect: A Comprehensive Approach to Offensive Language Detection in Roman Urdu Using Hybrid Machine Learning-Deep Learning Models with Embedding Techniques,” *Inf.*, vol. 16, p. 139, Feb. 2025, doi: 10.3390/info16020139.
- [16] K. Li, H. Zha, Y. Su, and X. Yan, “Concept Mining via Embedding,” presented at the Proceedings - IEEE International Conference on Data Mining, ICDM, 2018, pp. 267–276. doi: 10.1109/ICDM.2018.00042.
- [17] J. Neumann, R. Lange, Y. Susanti, and M. Färber, “Optimizing Small Transformer-Based Language Models for Multi-Label Sentiment Analysis in Short Texts,” *ArXiv*, vol. abs/2509.04982, Sep. 2025, doi: 10.48550/arxiv.2509.04982.
- [18] N. Yogeesh *et al.*, “Modeling Lexical Ambiguity in English Literature Using Fuzzy Logic and Equations,” *Applied Mathematics and Information Sciences*, vol. 19, no. 4, pp. 873–889, 2025, doi: 10.18576/amis/190413.
- [19] S. Zad, M. Heidari, P. Hajibabae, and M. Malekzadeh, “A Survey of Deep Learning Methods on Semantic Similarity and Sentence Modeling,” presented at the 2021 IEEE 12th Annual Information Technology, Electronics and Mobile Communication Conference, IEMCON 2021, 2021, pp. 466–472. doi: 10.1109/IEMCON53756.2021.9623078.
- [20] R. Adyanthaya and R. S. P., “YenLP_CS@DravidianLangTech 2025: Sentiment Analysis on Code-Mixed Tamil-Tulu Data Using Machine Learning and Deep Learning Models,” *Proceedings of the Fifth Workshop on Speech, Vision, and Language Technologies for Dravidian Languages*, Jan. 2025, doi: 10.18653/v1/2025.dravidianlangtech-1.50.
- [21] C. Anderson, M. Nguyen, and R. Coto-Solano, “Unsupervised, Semi-Supervised and LLM-Based Morphological Segmentation for Bribri,” *Proceedings of the Fifth Workshop on NLP for Indigenous Languages of the Americas (AmericasNLP)*, Jan. 2025, doi: 10.18653/v1/2025.americasnlp-1.7.
- [22] A. Petropoulos and V. Siakoulis, “Can central bank speeches predict financial market turbulence? Evidence from an adaptive NLP sentiment index analysis using XGBoost machine learning technique,” *Central Bank Review*, Dec. 2021, doi: 10.1016/j.cbrev.2021.12.002.
- [23] X. Fang and J. Yan, “Named entity recognition for U.S. export control based on the ERNIE-BiLSTM-Attention-CRF model,” vol. 13646, pp. 136460–136460, May 2025, doi: 10.1117/12.3056127.
- [24] A. Fawaid, R. Assyabani, I. Abdullah, C. Muali, M. S. Itqan, and S. Islam, “Human Intelligence and Algorithmic Precision: An Experimental Study of Indonesian Translation Pedagogy in Higher Education,” *Asian Journal of University Education*, vol. 21, no. 3, pp. 779–792, 2025, doi: 10.24191/ajue.v21i3.53.
- [25] A. Ustun and B. Can, “Incorporating word embeddings in unsupervised morphological segmentation,” *Natural Language Engineering*, vol. 27, pp. 609–629, Jul. 2020, doi: 10.1017/s1351324920000406.
- [26] R. Chauhan, A. Mishra, M. Jena, E. Yafi, and M. F. Zuhairi, “Mental Health Diagnostics Using Machine Learning: A Clustering and Predictive Modeling Approach,” presented at the Proceedings of the 2025 19th International Conference on Ubiquitous Information Management and Communication, IMCOM 2025, 2025. doi: 10.1109/IMCOM64595.2025.10857500.
- [27] V. Mayik, N. Lotoshynska, L. Mayik, K. Bazylyuk, and M. Khariv, “Mathematical Modeling of the Braille Font Formation for the Information and Communication Environment Creation for Blind People,” presented at the CEUR Workshop Proceedings, 2023, pp. 248–254. [Online]. Available: <https://ceur-ws.org/Vol-3641/short5.pdf>

BIOGRAPHY OF AUTHOR



Joko Santoso    is affiliated with STKIP PGRI Metro, Indonesia. His academic studies are related to computational linguistics, computer science, language technology, or related interdisciplinary fields. His research include natural language processing, named entity recognition, machine learning for language, corpus-based analysis, and Indonesian language technology. His work focuses on developing computational models to support language analysis and digital text processing. He can be contacted at email: joko.spbsi@gmail.com