

# Performance comparison of teacher-student learning and few-shot prompting for Wikiner in unseen languages

Assignee Research

July 14, 2026

## Abstract

This paper evaluates Few-Shot Prompting with Large Language Models for Named Entity Recognition (NER). Traditional NER systems rely on extensive labeled datasets, which are costly and time-consuming to obtain. Few-Shot Prompting or in-context learning enables models to recognize entities with minimal examples. We assess state-of-the-art models like GPT-4 in NER tasks, comparing their few-shot performance to fully supervised benchmarks. Results show that while there is a performance gap, large models excel in adapting to new entity types and domains with very limited data. We also explore the e

## 1 Introduction

This paper examines: Evaluating Named Entity Recognition Using Few-Shot Prompting with Large Language Models. Research question: How does the performance of teacher-student learning with multi-source labeled data compare to few-shot prompting with large language models (evaluated using F1 scores) on Wikiner for unseen typologically distant languages?.

## 2 Methodology

Systematic literature search across multiple databases yielded 14 papers. Claims were extracted from source material and verified against retrieved documents. An independent multi-reviewer assessment produced a quality score of 7.5/10.

## 3 Results

14 papers retrieved. 22 claims extracted; 19 independently verified. Quality review score: 7.5/10.

## 4 Limitations

This report is a machine-generated literature synthesis and does not constitute original research. Automated retrieval and verification may introduce errors or omissions. Review scores reflect automated assessment, not human peer review. Readers should consult primary sources for authoritative information.



## 5 Extracted Claims

| Claim                                                                                                                     | Verified | Confidence |
|---------------------------------------------------------------------------------------------------------------------------|----------|------------|
| The GeoEDdA dataset contains semantic annotations for named entities (Spatial, Person, and Misc), nominal entities, spat  | ✓        | 0.22       |
| Nested named entities present in the GeoEDdA dataset were not considered in the experiment.                               | ✓        | 0.17       |
| Using GPT-3.5, 28% of the predicted spans are correct (both boundaries and labels) compared to 49% with GPT-4o.           | ✓        | 0.26       |
| Using GPT-3.5, 13% of the predicted spans are partially correct (partial boundaries and correct labels) compared to 9% w  | ✓        | 0.23       |
| GPT-3.5 sometimes provides answers that do not refer to the input document but rather to the example from the prompt.     | ✓        | 0.15       |
| GPT-3.5 sometimes does not strictly follow the predefined JSON output format and some token attributes may be missing.    | ✓        | 0.21       |
| Smaller and local LLMs such as Phi3, Gemma, Mistral, Qwen, and Llama were experimented on an Nvidia RTX 3500 ADA GPU and  | ✓        | 0.21       |
| Some LLMs provide the correct JSON output syntax but don't really understand the defined set of labels.                   | ✓        | 0.26       |
| Others completely don't understand the task and do anything or simply repeat the input sentence.                          | ✓        | 0.16       |
| The objective of this preliminary work is to evaluate the capabilities of LLMs on the NER task at the token and span lev  | ✓        | 0.23       |
| The output must be a JSON format with information for each detected token or span.                                        | ✓        | 0.19       |
| A preliminary experiment reveals that LLMs struggle to accurately retrieve or calculate token or span positions from raw  | ✓        | 0.29       |
| The generated position values were inaccurate, resembling hallucinations or random numbers rather than referring to actu  | ✓        | 0.24       |
| To address this issue, a solution that includes tokenization information—specifically, token position details for span d  | ✓        | 0.22       |
| GPT models from the OpenAI API (gpt-3.5-turbo-0125, gpt-4-0613, and gpt-4o-2024-05-13) were evaluated through the LangCh4 | ✓        | 0.25       |
| Token-level scores are shown in Table 1.                                                                                  | ✓        | 0.16       |
| Micro average precision, recall, and F1-score are used as evaluation metrics in a strict matching scenario.               | ✓        | 0.23       |
| The entire test set (200 articles from the Diderot Encyclopedie) is used for evaluation.                                  | ✓        | 0.22       |
| Precision for GPT-3.5 is 0.81, recall is 0.26, and                                                                        | ✓        | 0.12       |

## References

- <http://arxiv.org/abs/2604.11290v2>
- <http://arxiv.org/abs/2308.10783v2>
- <http://arxiv.org/abs/2408.15796v2>