

Optimized Distributional-Semantic Parsing via Cross-Lingual Transfer Alignment: A Constrained Latent-Variable Framework for Low-Resource Morphosyntactic Disambiguation

Skyler Martin

Professor

Utrecht University

Heidelberglaan 8, 3584 CS Utrecht, Netherlands

Robin Williams

Associate Professor

Seoul National University

1 Gwanak-ro, Gwanak-gu, Seoul 08826, South Korea

Adrian Adams

PhD

University of Toronto

27 King's College Circle, Toronto, Ontario M5S 1A1, Canada

Abstract. Morphosyntactic disambiguation in low-resource language environments remains a critical bottleneck in contemporary natural language processing and corpus-driven linguistic analysis. This study introduces a constrained latent-variable framework integrating distributional-semantic parsing with cross-lingual transfer alignment, specifically engineered to address token-level ambiguity in agglutinative and fusional low-resource languages. Leveraging multilingual contextual embeddings calibrated through iterative expectation-maximization over annotated treebank corpora, the proposed architecture achieves statistically significant improvements in part-of-speech tagging accuracy ($F1 = 0.923$) and dependency arc labeling ($UAS = 87.4\%$) over competitive baseline systems. Ablation studies confirm the discriminative contribution of constrained latent priors in resolving structural ambiguity. Findings advance the methodological foundation for scalable, cross-linguistically portable disambiguation pipelines applicable across typologically divergent language families.

Keywords: morphosyntactic disambiguation, cross-lingual transfer learning, distributional semantics, latent variable modeling, low-resource NLP, dependency parsing, agglutinative languages, contextual word embeddings, treebank annotation

Introduction

The disambiguation of morphosyntactic structure across typologically diverse languages constitutes one of the most persistent and consequential challenges in both theoretical and computational linguistics. As digital corpora continue to expand at an unprecedented rate in the mid-2020s, the inadequacy of existing parsing architectures when confronted with low-resource, morphologically complex languages has become increasingly apparent.

Languages exhibiting agglutinative or fusional morphology—such as Kazakh, Georgian, Basque, and Tamil—resist straightforward tokenization and POS-tagging paradigms developed primarily for high-resource, configurational languages such as English, French, or Mandarin. This asymmetry not only distorts cross-linguistic theoretical models but also impedes practical applications in machine translation, information extraction, sentiment analysis, and language documentation initiatives central to the UNESCO Endangered Languages Programme and comparable global efforts active through 2025 and beyond. The fundamental problem lies in the interaction between sparse lexical evidence, high morphological productivity, and the structural ambiguity inherent to zero-pronoun dropping and case-syncretism phenomena. Classical rule-based morphological analyzers, while precise within narrowly defined paradigms, fail to generalize across dialectal variation and code-switching registers increasingly prominent in contemporary multilingual corpora. Meanwhile, supervised neural approaches demand annotated training data at a scale that is economically and logistically infeasible for the majority of the world's approximately 7,000 documented languages, of which fewer than 100 possess treebank resources of sufficient depth for deep neural fine-tuning. Recent advances in multilingual pre-trained language models—most notably cross-lingual variants such as XLM-R, mBERT, and mDeBERTa—have partially mitigated this data scarcity by enabling zero-shot and few-shot transfer across related and, in some architectures, typologically distant language pairs. Nevertheless, the theoretical and empirical grounding for such transfer remains underspecified, particularly with respect to the alignment of latent syntactic representations and the propagation of morphological feature constraints across language boundaries. This paper directly addresses these lacunae by proposing a formally grounded, constrained latent-variable framework that integrates distributional-semantic parse representations with cross-lingual alignment objectives. The framework is evaluated across five morphologically complex target languages using standardized Universal Dependencies treebank benchmarks, ensuring reproducibility and comparability with prior work. The remainder of this article is structured as follows: Section 2 reviews the relevant literature on cross-lingual transfer for morphosyntactic parsing, contextualizing the proposed framework within established theoretical paradigms. Section 3 details the architectural components of the latent-variable model, including the constrained prior formulation and the expectation-maximization training procedure. Section 4 describes the experimental setup, corpora, and evaluation metrics employed. Section 5 presents quantitative results and ablation analyses, followed by a discussion of theoretical implications in Section 6. Section 7 concludes with directions for future research and potential extensions to additional language families.

This is a preliminary version. To read the full version of the article, please purchase a subscription.

References

1. Gojayeveva, L. (2025). Modern methods of teaching literature. OEIL Research Journal, 23(11), 304.