

Enriching Source for English-to-Urdu Machine Translation

Bushra Jawaid and Amir Kamran and Ondřej Bojar

Charles University in Prague, Faculty of Mathematics and Physics

Institute of Formal and Applied Linguistics

Malostranské nám. 25, Praha 1, CZ-118 00, Czech Republic

`{jawaid,kamran,bojar}@ufal.mff.cuni.cz`

Abstract

This paper focuses on the generation of case markers for free word order languages that use case markers as phrasal clitics for marking the relationship between the dependent-noun and its head. The generation of such clitics becomes essential task especially when translating from fixed word order languages where syntactic relations are identified by the positions of the dependent-nouns. To address the problem of missing markers on source-side, artificial markers are added in source to improve alignments with its target counterparts. Up to 1 BLEU point increase is observed over the baseline on different test sets for English-to-Urdu.

1 Introduction

Phrase-based statistical machine translation (SMT) systems encounter many challenges when translating from morphologically poor to morphologically rich languages. One main challenge is the correct identification of the grammatical structure of a sentence when the required information lies outside the phrasal boundaries. In fixed word order languages such as English, syntactic structure of a sentence follows a fixed subject-verb-object (SVO) pattern; hence, it omits the need of marking the grammatical roles of words. On the contrary, in free word order languages syntactic roles are either embedded as noun inflections or added as a separate token before or after the head noun. In either case, the generation of morphologically complex language becomes difficult task for SMT systems.

In Urdu, a separate token is added after head noun to identify the case such as nominative, accusative, dative etc. The existence of separate case markers not only introduces errors in alignment due to missing source counterparts but it also directly effects the selection of noun forms, which can either be “oblique” if followed by a case marker or “direct” otherwise.

Several approaches have been explored for the enrichment of the source corpus while dealing with the agreement phenomenon on target side. This work focuses on pre-processing the source corpus by adding pseudo-words that can improve alignments with their target counterparts. The experiments are carried out on the phrase-based English-to-Urdu SMT, a language pair that exemplifies the lack of information on source side for the generation of case markers on the target side.

This work is licenced under a Creative Commons Attribution 4.0 International License. License details: <http://creativecommons.org/licenses/by/4.0/>

2 Related Work

Several attempts have been made for the integration of the linguistics information to the existing phrase-based SMT systems. Few models that pre-process source corpus for dealing with the agreement phenomenon on target side are discussed below:

The method of source pseudo-words insertion to generate the target words is not novel. We build upon the work of Kamran (2011) who exploited the use of pseudo-words for generating the target case markers for the English-Urdu language pair. Kamran (2011) used preliminary set of linguistic rules to add case markers for subject, object, indirect object and additionally for verb auxiliaries. We refine the oversimplified linguistic rules for adding pseudo-words by first identifying the various syntactic and morphological features such as transitivity and animacy.

Avramidis and Koehn (2008) model case agreement phenomenon for English-to-Greek by adding case information as factor on source side. This approach uses source CFG parses to identify the grammatical roles of words, whereas we use the dependency parses. Also, due to the fact that Greek noun inflections depend on their role, information is added in the form of factors, whereas we use the single-factored setup with the assumption that pseudo-words will play a role in the selection of the correct noun forms.

Goldwater and McClosky (2005) aim at overcoming the data sparseness issue by increasing the similarity between languages using source morphological analysis for Czech-to-English MT. In this approach, the source input is first lemmatized and then extra tokens are added for the information that is stripped off during the lemmatization process, such as for negation words.

Birch et al. (2007) have shown the use of Combinatorial Categorical Grammar (CCG) supertags on source sentence, for German-to-English translation, in an attempt to capture the syntactic structure of the source language in factored SMT models. Recently, Dungarwal et al. (2014) have used CCG supertags as an additional factor on source for English-to-Hindi SMT system.

3 Enriching Source

3.1 Stanford Parser

Stanford parser¹ is a toolkit that contains java implementation for both probabilistic context-free grammar (PCFG) and dependency parsers. The dependency parser extracts the typed dependency parse (de Marneffe et al., 2006) using the phrase structure parse of the sentence. Typed dependencies – such as subject, direct object etc – represent the grammatical relations between the individual words. The Stanford dependencies are represented as triplets consist of the name of the dependency relation, the dependent and the governor (also known as the “head”).

The Stanford CoreNLP framework² (Manning et al., 2014) is used for applying the NLP pipeline on the input sentence. The framework uses “annotators” for linguistic processing of input text. We use following annotators to process a sentence: tokenize, ssplit, pos, lemma, ner, parse and dcoref. Additionally, we set splitting of sentence (ssplit) to one sentence per input and tokenization is restricted to white space only.

¹<http://nlp.stanford.edu/software/lex-parser.shtml>

²<http://nlp.stanford.edu/software/corenlp.shtml>

Stanford CoreNLP provides the dependency parse in three graphical representations: basic, collapsed and cc-processed (collapsed and propagated) dependencies. The collapsed and cc-processed dependencies are used to extract the typed dependencies. Example 1 shows the Stanford’s collapsed typed dependencies³ where each triplet begins with the name of a dependency relation followed by the head and the dependent consecutively.

(1) My dog also likes eating sausage.

poss(dog-2, My-1) nsubj(likes-4, dog-2) advmod(likes-4, also-3) root(ROOT-0, likes-4)
xcomp(likes-4, eating-5) dobj(eating-5, sausage-6)

3.2 Case Markers

There are seven cases in Urdu that are morphologically realized by seven markers (Butt and King, 2004). Table 1 shows the list of cases with their respective markers and grammatical functions, adapted from Butt and King (2004).

Case	Marker	Grammatical Function
Nominative	ϕ	subj/obj
Ergative	ne	subj
Accusative	ko	obj
Dative	ko	subj/ind. obj
Instrumental	se	subj/obl/adjunct
Genitive	k-	subj/specifier
Locative	mē/par/tak/ ϕ	obl/adjunct

Table 1: Case Markers in Urdu

Absence of marker with subject or object roles marks the nominative case, while accusative and dative share the marker “ko”. Due to the fact that nominative lacks the marker, we only add pseudo-words for ergative, accusative and dative markers. Rest of the three cases are not considered in this work.

4 Common Settings

For the training of our translation system, the standard training pipeline of Moses is used along with the GIZA++ (Och and Ney, 2000) alignment toolkit and a 5-gram SRILM language model (Stolcke, 2002). The source texts were processed using the Treex platform (Popel and Žabokrtský, 2010)⁴, which included tokenization and lemmatization.

The target side of the corpus is tokenized using a simple tokenization script⁵ by Dan Zeman and it is lemmatized using the Urdu Shallow Parser⁶ developed by Language Technologies Research Center of IIIT Hyderabad.

The alignments are learnt from the lemmatized version of the corpus. For the rest of the SMT pipeline, word forms (i.e. no morphological decomposition) in their true case (i.e. names capitalized but sentence starts lowercased) are used. The lexicalized word-based reordering model (Koehn et al., 2005) is trained using *msd* orientation in both forward and backward direction, with model conditioned on both the source and the target languages (*msd-bidirectional-fe*).

³<http://nlp.stanford.edu:8080/parser/index.jsp>

⁴<http://ufal.mff.cuni.cz/treex/>

⁵The tokenization script can be downloaded from: <http://hdl.handle.net/11858/00-097C-0000-0023-65A9-5>

⁶http://ltrc.iiit.ac.in/showfile.php?filename=downloads/shallow_parser.php

The parallel and monolingual data is summarized in Table 2. The parallel data reported in Jawaid et al. (2014a) (called “ALL”) is used for training, development and test with the similar data splits. Jawaid et al. (2014b) released large plain and annotated Urdu monolingual data from mix of several domains. The plain text monolingual data is used to build the language model.

	Dataset	Sents (en/ur)	Tokens (en/ur)
Parallel	Train	74.9k	1.5M/1.7M
	Dev	2K	41.5K/45.2K
	Test	2K	41.8K/45.6K
Mono	-	5.4M	95.4M

Table 2: Summary of training data.

Final BLEU scores (Papineni et al., 2002) are reported on the test set called “PTEST” in the following and also on the three independent official test sets briefly explained by Jawaid et al. (2014a).

5 Experiments

The experiments are conducted with the insertion of pseudo-words on the un-preordered source side as well as after preordering the source corpus. For preordering of the English corpus, we use the transformation module of Jawaid and Zeman (2011) that utilizes the Stanford PCFG parse trees to first parse the input sentences and afterwards applies the hand-written rules to transform the English sentences to closely match the syntactic structure of Urdu sentences.

For preordered system with pseudo-words, the pseudo-words are added to the input that also contains the index of each word as an additional information. After generating the case markers, words are printed in the order of the reordered indexes together with pseudo-words.

In the following section, the Stanford dependencies that are used to generate the pseudo-words as well as the process of generating the case markers are briefly explained.

5.1 Case Marker Generation

Table 1 shows that ergative, accusative and dative cases take the roles of either subject, object or indirect object. Stanford dependency parser identifies these roles as: nominal subject (nsubj), direct object (dobj) and indirect object (iobj). The name of the dependencies are used to add the respective pseudo-words. Ergative and accusative cases take the nsubj and dobj pseudo-words respectively, whereas for dative case iobj marker is used to mark both subjects and indirect objects. We only use passive subjects (nsubjpass) for marking the subject role of the dative case. Only those relations are contemplated that hold verb as a governor of a relation unless stated explicitly.

5.1.1 Ergative Case

In Urdu, noun represents ergative case for transitive head verbs with perfective aspect. If verbs are tagged with “VBD” or “VBN” tags⁷, they are considered as perfective, whereas verb take the transitivity feature if it also hold dobj relation. There are cases where transitivity feature requires to deal with few exceptions.

⁷https://www.ling.upenn.edu/courses/Fall_2003/ling001/penn_treebank_pos.html

In case of a missing dobj dependency of a head verb, verb is marked transitive if followed by a prepositional phrase⁸.

Exception is also made for intransitive verbs or verbs with missing objects that contain the clausal complement (ccomp) relation with other verbs. In ccomp relations the internal subject of a dependent of ccomp relation acts as an object of a governor that qualifies the nsubj dependency relation to take the transitivity attribute.

Perfective attribute is ignored for question sentences in past indefinite tense where head noun of nsubj relation is tagged with either “VBP” or “VB”.

If a subject of a relation is Wh-determiner then the case marking process is only followed for “which”, “who” or “that” determiners. We don’t add erg marker if cardinals are dependent of a nsubj relation and also if auxiliary verbs are governor of a relation.

The dependency relation of reduced non-finite verbal modifier (vmod) is also used for marking the ergative case of nouns. In vmod relation, dependent modifies the meaning of a governor that can either be a verb or a noun. We only deal with cases where noun is a governor of relation and its also not a dependent of dobj, iobj or nsubjpass relations. Rest of the checks for adding an erg marker are similar to the way we deal with nsubj relation. vmod relations are used only after looking at few training examples but they need to be further investigated.

5.1.2 Accusative Case

The assignment of a dobj marker for an accusative case is not always straight forward. Butt and King (2004) show examples where “ko” alternates with null marker of nominative on direct objects.

(2) Nadya has driven a car

nādyh ne gārī člāi he

Nadya=Erg car=Nom drive=perfective be=present

nādyh ne gārī ko člāyā he

Nadya=Erg car=Acc drive=perfective be=present

To avoid the complexity, we do not add markers for “inanimate” objects. “dobj” marker is added for accusative cases that satisfy following conditions: governing verb is transitive, it does not contain the iobj relation and the dependent of dobj relation is “animate” object.

Similar to the ergative case, there are few exception for checking the transitivity of verb before adding the dobj marker. If the head verb has missing nsubj and iobj relation then we search for prepositional clausal modifier (prepc) and conjunct (conj) dependency relations that contains head verb either as a governor or a dependent. If binding of head verb is found in any of prepc or conj relation then dependent noun is marked as accusative and get the dobj marker.

5.1.3 Dative Case

“iobj” marker is added for all iobj dependencies without any constraints and exceptions.

⁸we ignore following prepositions for transitivity check: in, into, of, on, by, from, since, until, behind, between, beyond, but, with, near, inside, after, at, before, within, without, under, underneath, up, upon, opposite.

For nsubjpass relations transitivity and perfective features are validated before adding iobj marker. Verb of a nsubjpass relation is attributed transitive if it either contains direct object or prepositional phrase following the verb. Similar to the ergative case, perfective aspect of verb is verified using VBD and VBN POS tags.

5.2 Markers Positioning

The placement of pseudo-words play crucial role due to the word order differences in English-Urdu language pair. We look for conj, appos, dep and prep_of dependencies of the nouns before adding the erg marker and only conj dependency incase of obj marker, if these dependencies exist then markers are added only with the dependents of these dependencies. Example 3 shows the movement of erg marker from head of prep_of dependency to the dependent of a relation, whereas Example 4 shows the deletion of obj marker from the head of conj relation when both governor and dependent are acting as an object.

(3) Before: The savagery **erg** of the attack has shocked the government and observers.

After: The savagery of the attack **erg** has shocked the government and observers.

(4) Before: Prime Minister Gilani erg brought his penchant for consensus politics to bear upon the problem recently by bringing together top federal **obj** and provincial leaders **obj** for a two-day conference to develop consensus.

After: Prime Minister Gilani erg brought his penchant for consensus politics to bear upon the problem recently by bringing together top federal and provincial leaders **obj** for a two-day conference to develop consensus.

With preordered source corpus, we don't reposition markers of prep_of dependency because they are automatically repositioned after reordering the source corpus.

5.3 Results

Table 3 shows the source preordering and psuedo-words insertion results on all four test sets. Baseline results of phrase-based and hierarchical systems are also reported from Jawaid et al. (2014a) to see the relative gain in BLEU scores. All results reported in Table 3 were tested with MultEval⁹ for statistical significance of the improvement over the baseline. Based on 3 independent MERT runs of both the baseline and the experiment in question, • marks the 100% confidence on improvement over the baseline. Similarly, † and ‡ marks 96% and 90% confidence and * shows 80% confidence on gain in systems performance over the baseline setup.

The preordering of source corpus, PBR system, brings minimum 1 point (on PTEST) to maximum 2.8 point (on CLE) gain in BLEU scores. The phrase-based system with case markers (PBC) bring 0.6 to 1 point increase in BLEU on all independent test sets except PTEST that did not gain any improvements over the baseline with the additional pseudo-words in source corpus. On the other hand, hierarchical system with pseudo-words also shows minimum 0.2 (again on PTEST) to maximum 1 point gain in BLEU on all test sets. CLE shows maximum performance gain in all setups due to the availability of multiple reference translations.

⁹<https://github.com/jhclark/multeval>

	PTEST	CLE	IPC	NIST2008
	1 refs	3 ref	1 ref	1 ref
Phrase-based Baseline (PB)	19.3	18.2	15.8	15.0
With-Markers (PBC)	‡ 19.3	• 19.1	• 16.5	• 15.6
Preordered (PBR)	• 20.1	• 21.0	• 17.9	• 16.5
Preordered-with-Markers (PBCR)	• 20.5	• 21.1	• 18.8	• 16.7
PBCR without definite article	• 20.7	• 21.3	• 18.6	• 17.1
Hierarchical Baseline	21.4	19.4	18.7	16.7
With-Markers	† 21.6	• 20.4	* 19.0	• 17.1

Table 3: Results of Phrase-based and Hierarchical MT with and without case markers.

We also report results of phrase-based system together with preordered source corpus and added case markers (PBCR) to achieve the maximum performance gain in terms of BLEU. Over the PBR system, this system brings approximately 1 point gain on IPC to minimum 0.1 increase on CLE test set. The PBCR system did not bring significant improvements on all test sets (except IPC) compared to PBR system. It is not evident from the results, whether PBCR system has performed better than hierarchical system with case markers or vice versa. Even though, except PTEST, results of PBCR system always exceed (remain same for NIST test set) the hierarchical baseline results.

Figure 1 shows the impact of average source phrase length used during decoding on BLEU scores for all four phrase-based systems. The results verify that the systems perform better when the longer source phrases are matched during decoding. Figure 1 also shows the significance of preordering the source corpus that allows the MT engine to extract the longer matching phrases.

In Table 4, alignment statistics of baseline setup and our best performing phrase-based system (PBCR) is provided. In baseline system, case marker ‘*nay*’ gets mostly aligned to auxiliary ‘*have*’, followed by alignments with verbs and definite article. Interestingly, ‘*nay*’ remains unaligned 2.7K times out of 23K occurrences in reference. Furthermore, ‘*ko*’ aligns to ‘*the*’ most of the time, followed by alignments with prepositions. Out of 25K total occurrences, it remains unaligned 4.3K times.

In PBCR system, the statistics of most frequent alignment pairs change drastically for both markers. ‘*nay*’ gets aligned to ‘*erg*’ marker on source side 16K times, whereas the unaligned count reduces by 48.5%. The ‘*erg*’ marker remains unaligned around 6.6K times, which suggests that there might be an over generation of the ‘*erg*’ marker. This speculation can be confirmed

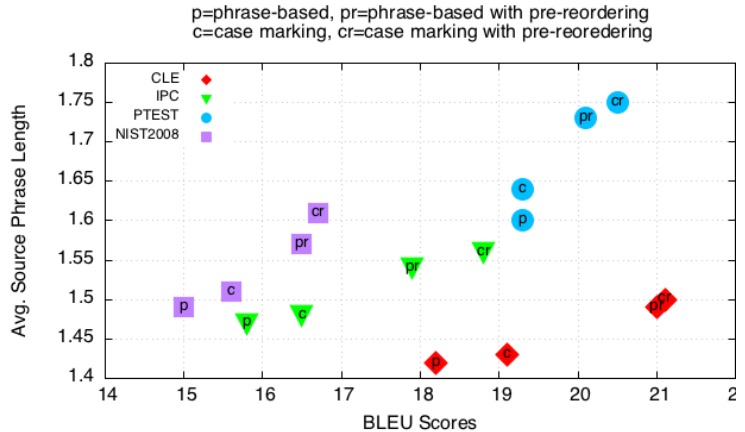


Figure 1: Plot of BLEU vs average source phrase length of each experimental setting indicated in “p”, “pr”, “c” and “cr” for all four test sets.

from the total number of ‘*erg*’ occurrences in source text that are 3.8K times more than its target counterpart. The stats of ‘*ko*’ marker does not show the same amount of improvement as ‘*nay*’. Out of 25K ‘*obj*’ markers only 5K aligned to ‘*ko*’ and out of 1.3K ‘*iobj*’ markers only 380 aligned to ‘*ko*’. The count of unaligned ‘*ko*’ markers only reduced by 10.6% compare to the baseline unaligned frequency. Even though, compared to the baseline setup, alignment count of ‘*ko*’ with definite article reduces by 35% but still 3K ‘*ko*’ markers aligned with the definite article. Our initial hypothesis was that due to the unavailability of the definite article in Urdu, the alignment between ‘*ko*’ and ‘*obj*’ was not learnt properly. To investigate this issue, we stripped off definite article from the source side and then re-ran the PBCR system. The result of this system is also reported in Table 3; small gains in terms of BLEU is observed on most test sets over the PBCR system but unfortunately improvements in alignment count of ‘*obj*’ and ‘*ko*’ markers are not up to the expectations, instead alignment count of ‘*the-ko*’ pair shifts to the unaligned ‘*ko*’ count, raising it to 5.3K. It is hard to predict why the large number of ‘*obj*’ markers remain unaligned; by looking at the total count of ‘*obj*’ marker in source, it can not be attributed to the over generation problem. Perhaps, it is added to the places where there was no matching marker on the target side exists. One simple solution would be (only for training) to add the ‘*obj*’ or ‘*iobj*’ marker in source when there exists at least one occurrence of ‘*ko*’ marker on target side. This way, it is possible to avoid the addition of the marker to unwanted places. The in-depth analysis of ‘*ko*’ is needed to investigate this issue further.

Markers	erg		لے ne		obj		iobj		کے ko	
Count in Refer.	–		23,747		–		–		25,095	
Count in Source	27,574		–		25,238		1341		–	
Baseline system	–	–	5588	have	–	–	–	–	6147	the
	–	–	4046	say	–	–	–	–	5348	to
	–	–	2727	unalign	–	–	–	–	4379	unalign
	–	–	2696	the	–	–	–	–	895	on
	–	–	456	do	–	–	–	–	500	as
PBCR system	16,664	لے(ne)	16,664	erg	7904	unalign	380	کے(ko)	5382	obj
	6676	unalign	1404	unalign	5788	کا(ka)	360	unalign	3915	unalign
	492	کے(ko)	1043	the	5382	کے(ko)	69	انہیں(*)	3700	to
	406	میں(meñ)	641	say	1760	سے(se)	26	تمہیں(*)	2979	the
	356	سے(se)	624	by	855	پر(per)	23	سے(se)	953	on

Table 4: Most frequent word alignments for source artificial markers and target case markers in training corpus for baseline and PBCR experiments.

6 Conclusion

The approach of introducing artificial source marking for phrasal clitics in Urdu (target side) shows significant improvements over baseline (PB vs PBC) except for one test set i.e., PTEST. In order to encounter target-side reordering problems, experiments are also carried out with preordered source sentences together with artificial markers. Due to the fact that reordering helps phrasal SMT to match longer phrases, it eventually helps to produce missing case markers due to longer matches. Hence, less improvements have been observed between PBR and PBCR

* انہیں = inheñ, تمہیں = tümheñ

systems with one exception being the IPC test set that shows significant gain over PBR system. The problem of over-generation of markers might have caused the inconsistent improvements over different test sets; however, it is still an open question and needs further investigation.

Acknowledgments

This work was supported by European Union’s innovation programme under grant agreement no. 645452 (QT21).

References

- Eleftherios Avramidis and Philipp Koehn. 2008. Enriching morphologically poor languages for statistical machine translation. In *Proceedings of ACL-08: HLT*, pages 763–770, Columbus, Ohio, June. Association for Computational Linguistics.
- Alexandra Birch, Miles Osborne, and Philipp Koehn. 2007. CCG Supertags in Factored Statistical Machine Translation. In *Proceedings of the Second Workshop on Statistical Machine Translation*, pages 9–16, Prague, Czech Republic, June. Association for Computational Linguistics.
- Miriam Butt and Tracy Holloway King. 2004. The status of case. In *Clause structure in South Asian languages*, pages 153–198. Springer Netherlands.
- Marie-Catherine de Marneffe, Bill MacCartney, and Christopher D. Manning. 2006. Generating typed dependency parses from phrase structure parses. In *Proceedings of Language Resources and Evaluation (LREC)*, pages 449–454.
- Piyush Dungarwal, Rajen Chatterjee, Abhijit Mishra, Anoop Kunchukuttan, Ritesh Shah, and Pushpak Bhattacharyya. 2014. The iit bombay hindi-english translation system at wmt 2014. In *Proceedings of the Ninth Workshop on Statistical Machine Translation*, pages 90–96, Baltimore, Maryland, USA, June. Association for Computational Linguistics.
- Sharon Goldwater and David McClosky. 2005. Improving statistical mt through morphological analysis. In *Proceedings of the Conference on Human Language Technology and Empirical Methods in Natural Language Processing, HLT ’05*, pages 676–683, Stroudsburg, PA, USA. Association for Computational Linguistics.
- Bushra Jawaid and Daniel Zeman. 2011. Word-order issues in english-to-urdu statistical machine translation. *The Prague Bulletin of Mathematical Linguistics*, (95):87–106.
- Bushra Jawaid, Amir Kamran, and Ondřej Bojar. 2014a. English to urdu statistical machine translation: Establishing a baseline. pages 1–6, Dublin, Ireland. Dublin City University, Dublin City University.
- Bushra Jawaid, Amir Kamran, and Ondřej Bojar. 2014b. A tagged corpus and a tagger for urdu. In Nicoletta Calzolari, Khalid Choukri, Thierry Declerck, Hrafn Loftsson, Bente Maegaard, and Joseph Mariani, editors, *Proceedings of the 9th International Conference on Language Resources and Evaluation (LREC 2014)*, pages 2938–2943, Reykjavík, Iceland. European Language Resources Association.
- Amir Kamran. 2011. Hybrid Machine Translation Approaches for Low-Resource Languages. Master’s thesis, UFAL, Prague, September.
- Philipp Koehn, Amittai Axelrod, Alexandra B. Mayne, Chris Callison-Burch, Miles Osborne, and David Talbot. 2005. Edinburgh System Description for the 2005 IWSLT Speech Translation Evaluation. In *Proceedings of International Workshop on Spoken Language Translation*.
- Christopher D. Manning, Mihai Surdeanu, John Bauer, Jenny Finkel, Steven J. Bethard, and David McClosky. 2014. The Stanford CoreNLP natural language processing toolkit. In *Proceedings of 52nd Annual Meeting of the Association for Computational Linguistics: System Demonstrations*, pages 55–60, Baltimore, Maryland, June. Association for Computational Linguistics.
- Franz Josef Och and Hermann Ney. 2000. A Comparison of Alignment Models for Statistical Machine Translation. In *Proceedings of the 17th conference on Computational linguistics*, pages 1086–1090. Association for Computational Linguistics.

- Kishore Papineni, Salim Roukos, Todd Ward, and Wei-Jing Zhu. 2002. BLEU: a Method for Automatic Evaluation of Machine Translation. In *ACL 2002, Proceedings of the 40th Annual Meeting of the Association for Computational Linguistics*, pages 311–318, Philadelphia, Pennsylvania.
- Martin Popel and Zdeněk Žabokrtský. 2010. TectoMT: Modular NLP Framework. In Hrafn Loftsson, Eiríkur Rögnvaldsson, and Sigrún Helgadóttir, editors, *Lecture Notes in Artificial Intelligence, Proceedings of the 7th International Conference on Advances in Natural Language Processing (IceTAL 2010)*, volume 6233 of *Lecture Notes in Computer Science*, pages 293–304, Berlin / Heidelberg. Iceland Centre for Language Technology (ICLT), Springer.
- Andreas Stolcke. 2002. Srilm - an extensible language modeling toolkit. In *In Proceedings of the 7th International Conference on Spoken Language Processing (ICSLP) 2002*, pages 901–904.