

Nízko-zdrojové filtrovanie korpusu pomocou viacjazyčných vložiek

Vishrav Chaudhary* Yuqing Tang* Francisco Guzman* Holger Schwenk* Philipp Koehn*

*Facebook AI, Johns Hopkins University

{vishrav,yuqtang, fguzman,schwenk}@fb.com phi@jhu.edu

Abstraktné

V tomto dokumente popisujeme naše podanie dozdĺanej úlohy filtrovania nízkozdrojového korpusu WMT19. Naš hlavný prístup je založený na súbore nástrojov LASER (Jazykovo-agnostické zastúpenia vety), ktorý používa encoder-dekoderovú architektúru vyškolenú na paralelnomkorpusu na získanie viacjazyčných vetových reprezentácií. Potom používame reprezentácie priamo na skóre a filtrovať hlučné paralelné vety bez ďalšieho tréningu funkcie bodovania. Porovnávame náš prístup k iným sľubným metódam a dokazujeme, že LASER prináša silné výsledky. Nakoniec vytvoríme súbor rôznych skórovacích metód a získame ďalšie zisky. V porovnaní s druhými najlepšimi systémami dosiahla naša ponuka najlepší celkový výkon tak pre nepálsko-anglické úlohy 1M, ako aj pre úlohy Sinhala-English 1M, resp. 1,4 BLEU. Navyše, naše experimenty ukazujú, že táto technika je sľubná pre scenáre s nízkymi a dokonca žiadnymi zdrojmi.

1 Úvod

Dostupnosť vysokokvalitných paralelných školiacích údajov je rozhodujúca pre dosiahnutie dobrého výkonu prekladu, keďže systémy NMT sú menej odolné voči hlučným paralelným údajom ako systémy štatistického strojového prekladu (Khayrallah a Koehn, 2018). Nedávno sa zvýšil záujem o filtrovanie hlučných paralelných korpusov (ako je Paracrawl1) o zvýšenie množstva dát, ktoré možno použiť na výcvik prekladateľských systémov (Koehn a kol., 2018).

Zatiaľ čo najmodernejšie metódy, ktoré používajú NMT modely, sa ukázali ako účinné pri ťažbe

Paralelné vety (Junczys-Dowmunt, 2018) pre jazyky s vysokým zdrojom, ich účinnosť nebola testovaná v jazykoch s nízkym zdrojom. Dôsledky nízkej dostupnosti údajov o odbornej príprave pre metódy paralelného označovania zatiaľ nie sú známe.

Za úlohu nízkozdrojového filtrovania (Koehn et al., 2019) sme dostali veľmi hlučné 40,6 milióna slov (počet anglických tokenov) Nepali- anglický korpus a 59,6 milióna slov Sinhala- anglický korpus vylínel z webu ako súčasť projektu Paracrawl. Problém spočíva v poskytovaní bodov pre každý pár viet v oboch hlučných paralelných súboroch. Skóre sa použije na podvzorky párov viet, ktoré predstavujú 1 miliónov a 5 miliónov anglických slov. Kvalita výsledných podmnožín je určená kvalitou štatistického strojového prekladu (Mozes, frázy (Koehn a kol., 2007)) a neurálneho strojového prekladu Fairseq (Ott a kol., 2019) vyškoleným na tieto údaje. Kvalita systému strojového prekladu bude meraná podľa BLEU skóre pomocou Sacrebleu (Post, 2018) navybranej testovacej sady Wikipédie preklady pre Sinhala-anglický a nepali-anglický z flores dataset (Guzman a kol., 2019).

V našom podaní pre túto zdieľanú úlohu používame viacjazyčné vkladanie vety získané od LASER, ktorý¹ používa encoder-dekoder architektúru na tréning viacjazyčného zastúpenia vety s použitím relatívne malého paralelného korpusu. Naše experimenty ukazujú, že navrhovaný prístup prekonáva iné existujúce prístupy. Okrem toho využívame súbor viacerých bodovacích funkcií na ďalšie zvýšenie filtračného výkonu.

¹<http://www.paracrawl.eu/>

2 Metodika

Zdieľaná úloha WMT 2018 pre paralelné filtrovanie korpusu (Koehn *et al.*, 2018)³ zaviedla niekoľko-metód na riešenie vysokozdrojového nemeckého a anglického dátového stavu. Zatiaľ čo mnohé z týchto metód boli úspešné pri filtrovaní hlučných prekladov, málo z nich bolo vyskúšaných za podmienok s nízkym zdrojom. V tomto dokumente riešime problém nízkozdrojového filtrovania viet pomocou reprezentácií na úrovni viet a porovnávame ich s inými populárnymi metódami používanými vo vysoko-zdrojových podmienkach.

Model LASER (Artetxe a Schwenk, 2018a) využíva viacjazyčné predstavy vety na meranie podobnosti medzi zdrojom a cieľovou vetou. Poskytuje najmodernejšie výkony na baníckej úlohe BuCC corpus a tiež účinne filtruje dáta WMT Paracrawl (Artetxe a Schwenk, 2018a). Tieto úlohy sa však týkali len jazykov s vysokým zdrojom, konkrétne francúzštiny, nemčiny, ruštiny a čínštiny. Našťastie, táto technika bola účinná aj v prípade nula-shot cross-lingual natural language indukcie v súbore dát XNLI (Artetxe a Schwenk, 2018b), čo sľubuje, že scenár s nízkymi zdrojmi bude zameraný na túto spoločnú úlohu. V tomto dokumente navrhujeme použiť adaptáciu LASER na nízke zdrojové podmienky na výpočet skóre podobnosti na filtrovanie hlučných viet. Pre porovnanie s LASER, sme tiež stanovili počiatočné referenčné hodnoty pomocou Bicleaner a Zipporah, dve populárne základné línie, ktoré boli použité v projekte Paracrawl. A dvojité podmienené cross-entropy, ktorá sa ukázala ako najmodernejšia pre úlohu filtrovania vysokozdrojového korpusu (Koehn *et al.*, 2018). Skúmame výkon techník za podobných predbežných podmienok týkajúcich sa filtrovania jazykovej identifikácie a lexikálneho prekrytia. Pozorujeme, že skóre LASER poskytuje jasnú výhodu pre túto úlohu. Nakoniec vykonávame zhromaždenie bodov pochádzajúcich z rôznych metód. Pozorujeme, že keď sú v zmesi zahrnuté skóre LASER, zvýšenie výkonu je relatívne malé. Vo zvyšku tejto časti diskutujeme o nastaveniach pre každú z použitej metódy.

³<http://statmt.org/wmt18/paralelny-korpus-filtering.html>

2.1 Laserové viacjazyčné zastúpenie

Základnou myšlienkou je použiť vzdialenosti medzi dvoma viacjazyčnými reprezentáciami ako pojem paralelizmu medzi dvoma vloženými vetami (Schwenk, 2018). Aby sme to dosiahli, najprv

vycvičíme kódovač, ktorý sa naučí vytvárať viacjazyčné, pevné zastúpenie vety; a potom vypočítajte vzdialenosť medzi dvoma vetami v učnom priestore. Okrem toho používame *maržové* kritérium, ktoré používa prístup k najbližším susedom normalizovať skóre podobnosti vzhľadom na to, že cosine podobnosť nie je globálne konzistentná (Artetxe a Schwenk, 2018a).

Encoder Viacjazyčný kódovač sa skladá z obojsmerného LSTM, a naša veta embeddingssú získané použitím max-pooling cez jeho výstup. Používame jediný kódovač a dekodér v našom systéme, ktoré zdieľajú všetky zúčastnené jazyky. Na tento účel sme trénovali viacjazyčné vety len na poskytnutých paralelných údajoch (podrobnosti pozri v časti 3.2).

Marža Sledujeme definíciu *pomeru* z² (Artetxe a Schwenk, 2018a). Použitím tohto, môže byť skóre podobnosti medzi dvoma vetami (x , y) vypočítané ako

$$\frac{2 \cos(x, y)}{S_y e^{\text{NNfc}(x)} \cos(x, y) + S_x e^{\text{NNfc}(y)} \cos(x, y)}$$

Kde $\text{NNfc}(x)$ označuje k najbližších susedov x v inom jazyku, a analogicky pre $\text{NNfc}(y)$. Všimnite si, že tento zoznam najbližších susedov neobsahuje duplikáty, takže aj v prípade, že daná veta má viac výskytov v korpuse, to by malo (najviac) jeden záznam v zozname.

Okolie Okrem toho sme preskúmali dva spôsoby odberu vzoriek k najbližším susedom. Najprv *globálna* metóda, v ktorej sme využívali okolie, pozostávali z hlučných dát spolu s čistými dátami. Po druhé, *miestna* metóda, v ktorej sme ohodnotili len hlučné dáta pomocou hlučnej štvrte, alebo čisté dáta pomocou čistej štvrte.³

2.2 Iné metódy podobnosti

Zipporah (Xu a Koehn, 2017; Khayrallah *et al.*, 2018), ktorý sa často používa ako východiskové porovnanie, používa jazykový model skóre prekladu slov, s váhami optimalizovanými pre oddelenie čistých a syntetických údajov o hluku. V našom nastavení sme trénovali Zipporah modely pre oba jazykové páry Sinhala-anglický a nepálsky-anglický. Použili sme open source verziu nástroja Zipporah bez úprav. Všetky komponenty modelu

²Preskúmali sme absolútne, vzdialenosť a pomer *maržové* kritériá, ale tie fungovali najlepšie

³Táto posledná časť bola vykonaná len pre tréning súboru

Zipppora (probabilistické prekladové slovníky a jazykové modely) boli školené na poskytnutých čistých údajoch (s výnimkou slovníkov). Jazykové modely boli vycvičené pomocou KenLM (Heafield a kol., 2013) cez čisté paralelné dáta. Nepoužijeme poskytnuté jednojazyčné dáta ako predvolené nastavenie. Na cvičenie sme použili vývojovú sadu z floresových dát.

Bicleaner (Sanchez-Cartagena et al., 2018) používa lexikálny preklad a skóre jazykových modelov a niekoľko plytkých funkcií, ako sú: príslušná dĺžka, zodpovedajúce čísla a interpunkcia. Rovnako ako u Zip-porah, sme použili open source Bicleaner7 toolkit⁶ out-of-the-box. Na výcvik tohto modelu sa použili iba poskytnuté paralelné údaje. Bicleaner používa komponent založený na pravidlách na identifikáciu hlučnejších príkladov v paralelných dátach a školí klasifikátora, aby sa naučil, ako ich oddeliť od ostatných údajov o výcviku. Použitie funkcií jazykového modelu je voliteľné. Používali sme len modely bez skórovacieho komponentu jazykového modelu.⁸

Dual Conditional Cross-Entropy Jeden z

Najlepšie výkonné metódy na túto úlohu boli duálne podmienené cross-entropie **filtrovanie** (Junczys—Dowmunt, 2018), ktoré používa kombináciu dopredu a dozadu modely na výpočet krížovej lingválnej podobnosti skóre. V našich experimentoch sme pre každú jazykovú dvojicu použili poskytnuté čisté školiace údaje na výcvik modelov prekladu neurálneho stroja v oboch smeroch prekladu: zdroj k cieľu a cieľ k zdroju. Vzhľadom na takéto preklad-modelu M, sme silový dekódovať vety páry (x, y) z hlučného paralelného korpusu a získať skóre cross-entropie

$$HM(y|x) = \frac{1}{|y|} \log_{pm}(YTLTY[i, t-i], x) \quad (1)$$

⁶<https://github.com/hainan-xv/zippporah>
⁷<https://github.com/bitextor/bicleaner>
⁸ Zistili sme, že zahrnutie LM ako extra hry viedlo k tomu, že takmer všetky páry viet dostali skóre 0.

Skóre dopredu a dozadu krížovej entropie, $H_F(y|x)$ a $H_B(x|y)$, sú potom v priemere sdodatočným trestom na veľký rozdiel medzi dvoma bodmi $|H_F(y|x) - H_B(x|y)|$.

$$\text{Score}(x, y) = \frac{H_F(y|x) - H_B(x|y)}{|H_F(y|x) - H_B(x|y)|} \quad (2)$$

Vpredu a dozadu sú päťvrstvové enkodérové/dekodérové transformátory trénované pomocou fairseq sparametrami, ktoré sú identické s parametrami používanými v modeli

základných flores⁴⁵. Modely boli vycvičené na čisté paralelné dáta pre 100 epoch. Pre nepálsko-anglickú úlohu sme tiež skúmali pomocou hindsko-anglických dát bez väčších rozdielov vo výsledkoch. Použili sme vývoj flores set na výber modelu, ktorý maximalizuje skóre BLEU.

2.3 Súbor

Aby sme využili silné a slabé stránky rôznych bodovacích systémov, preskúmali sme použitie binárneho klasifikátora na zostavenie súboru. Hoci je banálne získať klady (napr. čisté školiace údaje), ťažba negatívov môže byť skľučujúcou úlohou. Preto používame pozitívne neoznačené (PU) učenie (Mordalea Vert, 2014), ktoré nám umožňuje získať triedičov bez toho, aby sme museli spracovať súbor údajov explicitne pozitívnych a záporných. V tomto nastavení pochádzajú naše pozitívne označenia z čistých paralelných údajov, zatiaľ čo neoznačené údaje pochádzajú z hlučného súboru.

Na dosiahnutie tohto cieľa používame balenie 100 slabých, zaujatých klasifikátorov (t. j. s 2:1 pre neoznačené údaje oproti pozitívnym údajom značenia). Používame podporné vektorové stroje (SVM) s radiálnym jadrom a náhodne podvzorujeme sadu funkcií pre tréning každého základného klasifikátora, čím pomáhame udržať ich rôznorodú a nízku kapacitu.

Absolvovali sme dve iterácie tréningu tohto súboru. V prvej iterácii sme použili pôvodné pozitívne neoznačené údaje popísané vyššie. Na druhú iteráciu sme použili poučeného klasifikátora na preznačenie údajov o výcviku. Preskúmali sme niekoľko prístupov k opätovnému označeniu (napr. stanovenie prahovej hodnoty, ktorá maximalizuje skóre $F1$). Zistili sme však, že stanovenie triednej hranice na zachovanie pôvodného pomeru pozitív k neoznačeniu najlepšie fungovalo. Pozorovali sme tiež, že výkon sa po dvoch iteráciách zhoršil.

3 Nastavenie experimentu

Experimentovali sme s rôznymi metódami pomocou nastavenia, ktoré presne odráža oficiálne skóre zdieľanej úlohy. Všetky metódy sú školené na poskytnutých čistých paralelných údajoch (pozri tabuľku 1). Nepoužili sme dané jednojazyčné údaje. Na vývojové účely sme použili dodávané flores dev set. Na vyhodnotenie sme vyškolili systémy strojového prekladu na vybraných podmnožinách (1M, 5M) hlučných paralelných tréningových dát pomocou fairseq s predvolenou konfiguráciou flores tréningového parametra. Nahlasujeme Sacrebleu skóre na flores DevTest set. Vybrali sme si náš

⁴⁵<https://github.com/facebookresearch/Flores#Vlak-a-základný-transformer-model>

hlavný systém na základe najlepších skóre na DevTest sade pre 1M stav.

	si-sk	ne-sk	hi-en
Rozsudky	646k	573k	1.5
Anglické slová	3.7	3.7 MIL.	20.7

Tabuľka 1: Dostupné bitexty na tréovanie filtračných priblížení.

3.1 Predbežné spracovanie

Použili sme sadu filtračných techník podobných tým, ktoré sa používajú v LASER (Artetxea Schwenk, 2018a) a priradili sme skóre —1 na hlučné vety založené na nesprávnom jazyku buď na zdroji alebo na cieľovej strane, alebo s prekrývaním aspoň 60 % medzi zdrojom a cieľovými tokenmi. Na filtrovanie jazykových id sme použili fastText10. Vzhľadom k tomu, LASER počíta skóre podobnosti pre pár viet pomocou týchto filtračných techník, experimentovali sme pridaním týchto k ostatným modelom, ktoré sme použili pre túto zdieľanú úlohu.

3.2 Školenie laserového Encoderu

Pre naše experimenty a oficiálne podanie sme tréovali viacjazyčný kódovač viet s použitím povolených zdrojov v tabuľke 1. Vyškoliť sme jediný encoder využívajúci všetky paralelné dáta pre Sinhala-anglický, nepálsky-anglický a hindský jazyk. Keďže Hindi a Nepali zdieľajú rovnaké písmo, zhromaždili sme ich korpuse do jedného paralelného korpusu. Aby sme zohľadnili rozdiel vo veľkosti paralelných tréningových údajov, preverili sme bitexty Sinhala-anglického a Nepálu/Hindi-anglického v pomere 5:3. To malo za následok zhruba 3,2 M výcvikové vety pre každý smer jazyka, tj. Sinhala a kombinovať Nepali-Hindi.

¹⁰⁹<https://fasttext.cc/docs/en/jazyková-identifikácia.html>

Modely boli vycvičené s použitím rovnakého nastavenia ako verejný kód LASER, ktorý zahŕňa normalizáciu textov a tokenizáciu s Mojžišovými nástrojmi (späť do anglického režimu). Najprv sa naučíme spoločnú 50k BPE slovnú zásobu na skrátené školiace dáta pomocou fastBPE.⁶ Encoder vidí Sinhala, Nepali, hindčina a anglické vety na vstupe, bez akýchkoľvek informácií o aktuálnom jazyku. Tento vstup je vždy preložený do angličtiny.⁷ Experimentovali sme s rôznymi technikami na pridanie šumu do anglických vstupných viet podobných tomu, čo sa používa v neurálnom strojovom preklade bez dozoru, napr. (Artetxe et al., 2018; Lample a kol., 2018), ale to-

⁶ <https://github.com/glample/fastBPE>

¹²To znamená, že musíme tréovať anglický autoenkodér. Nezdalo sa, že by to bolelo, pretože ten istý encoder ovláda aj ďalšie tri jazyky.

nezlepšilo výsledky.

Enkodér je päťvrstvový BLSTM s 512 rozmerovými vrstvami. Dekodér LSTM má jednu skrytú vrstvu veľkosti 2048, tréovaný s optimalizátorom Adam. Pre vývoj vypočítame chybu podobnosti na zoskupenie flores dev sady pre Sinhala-anglický a nepálsky-anglický. Naše modely boli tréované pre sedem epoch asi 2,5 hodiny na 8 GPU Nvidia.

4 Výsledky

Z výsledkov v tabuľke 2 pozorujeme niekoľko trendov: (i) skóre pre 5M stav je vo všeobecnosti nižšie ako pri 1 M stave. Zdá sa, že táto podmienka sa zhoršuje použitím jazykového id a prekrývania filtrovania. (ii) LASER vykazuje trvale dobrý výkon. Miestna štvrt' funguje lepšie ako tá globálna. V tomto nastavení je LASER v priemere 0,71 BLEU nad najlepším systémom non-LASER. Tieto medzery sú vyššie pri 1M stave (0,94 BLEU). (iii) Najlepšia konfigurácia súboru poskytuje malé vylepšenia v porovnaní s najlepšou konfiguráciou LASER. Pre Sinhala-anglickú najlepšiu konfiguráciu zahŕňa všetky ostatné bodovacie metódy (ALL). Pre nepali-anglický je najlepšou konfiguráciou súbor LASER skóre. (iv) Dvojitá krížová entropia vykazuje zmiešané výsledky. V prípade Sinhala-anglického jazyka funguje iba vtedy, keď je povolené filtrovanie id jazyka, čo je v súlade s predchádzajúcimi pozorovaniami (Junczys-Dowmunt, 2018). Pre Nepali – angličtinu, poskytuje skóre hlboko pod ostatnými bodovacími metódami. Všimnite si, že sme nevykonali prieskum architektúry.

Metóda	česk		si-sk	
	1M	5M	1M	5M
Zipporah				
základňa	5.03	2.09	4.86	4.53
+ VEKO	5.30	1.53	5.53	3.16
+ Prekrytie	5.35	1.34	5.18	3.14
Duálny X-Ent.				
základňa	2.83	1.88	0.33	4.63 ⁺
+ VEKO	2.19	0.82	6.42	3.68
+ Prekrytie	2.23	0.91	6.65	4.31
Bicleaner				
základňa	5.91	2.54 ⁺	6.20	4.25
+ VEKO	5.88	2.09	6.36	3.95
+ Prekrytie	6.12 ⁺	2.14	6.66 ⁺	3.26
LASER				
miestne	7.37	3.15	7.49	5.01
globálne	6.98	2.98	7.27	4.76
Súbor				
VŠETKY	6.17	2.53	7.64	5.12
Laserový glob. + Lokálny	7.49	2.76	7.27	5.08

Tabuľka 2: Sacrebleu skóre na flores DevTest set. Tučným písmom zvýrazňujeme najlepšie skóre pre každú podmienku. Kurzívou* zvýrazníme bežca hore. Tiež signalizujeme najlepšiu non-LASER metódu s +.

4.1 Diskusia

Jedna prirodzená otázka, ktorú treba preskúmať, je, ako by metóda LASER mala prospech, ak by mala prístup k dodatočným údajom. Aby sme to preskúmali, použili sme súbor LASER open-source toolkit, ktorý poskytuje vyškolený encoder-pokrývajúci 93 jazykov, ale nezahŕňa nepali. V tabuľke 4 pozorujeme, že vopred vyškolený model LASER prekonáva miestny model LASERo 0.4

Podanie Pre oficiálne podanie sme použili súbor ALL pre úlohu Sinhala-anglický a LASER *globálny* + *lokálny* súbor pre nepálsko-anglickú úlohu. Tiež sme predložili LASER *lokálny* ako kontrastný systém. Ako môžeme vidieť v tabuľke 3, výsledky hlavných a kontrastných podaní sú veľmi blízke. V jednom prípade, kontrastný roztok (jedno LASER) model prináša lepšie výsledky ako súbor. Tieto výsledky umiestnili naše 1M príspevky 1.3 a 1.4 bodov BLEU nad bežcov pre úlohy nepáľčiny – angličtiny a Sinhaly-anglického. Ako už bolo spomenuté, naše systémy dosahujú horšie výsledky pri 5M stave. Takisto sme poznamenali, že čísla v tabuľke 2 sa mierne líšia od čísel vykázaných v (Koehn a kol., 2019). Tento rozdiel pripisujeme efektu tréningu v 4 (náš) GPU vs. 1 (ich).

Metóda	česk		sisk	
	1M	5M	1M	5M
Main – Ensemble	6.8	2.8	6.4	4.0
To je Constr. — LASER	6.9	2.5	6.2	3.8
Najlepší (iný)	5.5	3.4	5.0	4.4

Tabuľka 3: Oficiálne výsledky hlavných a sekundárnych podaní na sadetestov flores vyhodnotených s konfiguráciou NMT. Pre porovnanie, zahrňujeme najlepšie skóre iným systémom.

5 Závery a budúca práca

V tomto dokumente popisujeme naše podanie do úlohy filtrovania nízkozdrojového paralelného korpusu WMT. Používame viacjazyčné vety z LASER na filtrovanie hlučných viet. Pozorujeme, že LASER môže dosiahnuť lepšie výsledky ako základné línieso širokým rozpätím. Zahrnutie skóre z iných techník a vytvorenie súboru prináša ďalšie zisky. Naš hlavný príspevok k zdieľanej úlohe je založený na tom najlepšom z konfigurácie súboru a naše kontrastné podanie je založené na najlepšej konfigurácii LASER. Naše systémy plnia najlepšie 1M podmienky pre nepálsko-anglické a sinhalsko-

BLEU. Pre nepálsko-anglický stav je situácia opačná: Laser *lokálne* poskytuje oveľa lepšie výsledky. Avšak výsledky predtrénovaného LASERu sú len o niečo horšie ako výsledky Bicleaner (6.12), ktorý je najlepšou metódou, ktorá nepatrí do systému LASER. To naznačuje, že LASER môže dobre fungovať v scenároch s nulovým záberom (t.j. nepálsky-anglický), ale funguje ešte lepšie, keď má dodatočný dohľad nad jazykmi, na ktorých sa testuje.

Metóda	česk		sisk	
	1M	5M	1M	5M
Vopred vyškolený	6.06	1.49	7.82	5.56
Laserová lokalita	7.37	3.15	7.49	5.01

Tabuľka 4: Porovnanie výsledkov na flores DevTest set s použitím obmedzených a vopred vyškolených verzií LASER.

anglické úlohy. Analyzujeme výkonnosť vopred vycvičenej verzie LASER a pozorujeme, že dokáže vykonávať filtračnú úlohu dobre aj v scenároch s nulovými zdrojmi, čo je veľmi sľubné.

V budúcnosti chceme túto techniku vyhodnotiť pre scenáre s vysokým zdrojom a sledovať, či sa rovnaké výsledky prenású do tohto stavu. Okrem toho plánujeme preskúmať, ako veľkosť školiacich dát (paralelné, jednojazyčné) ovplyvňuje úlohu nízkozdrojového filtrovania viet.

Referencie

- Mikel Artetxe, Gorka Labaka, Eneko Agirre a Kyunghyun Cho. 2018. [Nekontrolovaný neurálny strojový preklad](#). Na medzinárodnej konferencii o zastúpeniach vzdelávania (ICLR).
- Mikel Artetxe a Holger Schwenk. 2018a. [Maržová paralelná baníctvo korpusu s viacjazyčným vložením-trestu](#). arXiv preprint arXiv:1811.01136.
- Mikel Artetxe a Holger Schwenk. 2018b. [Massively Multilingual Sentence Embeddings for Zero-Shot Cross-Lingual Transfer and Beyond](#). arXiv preprint arXiv:1812.10464.
- Francisco Guzman, Peng-Jen Chen, Myle Ott, Juan Pino, Guillaume Lample, Philipp Koehn, Vishrav Chaudhary a Marc’Aurelio Ranzato. 2019. [Dva nové hodnotiace súbory údajov pre nízkozdrojový strojový preklad: Nepali-angličtina a sinhala-english](#). arXiv preprint arXiv:1902.01382.
- Kenneth Heafield, Ivan Pouzyrevsky, Jonathan H. Clark a Philipp Koehn. 2013. [Scalable modifikovaný model jazyka Kneser-Ney](#). V pokračovaní 51. výročného zasadnutia Asociácie pre výpočtovú lingvistiku, strany 690–696, Sofia, Bulharsko.
- Marcin Junczys-Dowmunt. 2018. [Dvojité podmienené cross-entropia filtrácia hlučných paralelných korpusov](#). In *Proceedings of the Third Conference on Machine Translation, zväzok 2: Dokumenty o zdieľanej úlohe*, strany 901–908, Belgicko, Brusel. Asociácia pre počítačovú lingvistiku.
- Huda Khayrallah a Philipp Koehn. 2018. [O vplyve rôznych](#)

- typov hluku na neurálny strojový preklad. In *Proceedings of the 2nd Workshop on Neural Machine Translation and Generation*, strany 74 – 83, Melbourne, Austrália. Asociácia pre počítačovú lingvistiku.
- Huda Khayrallah, Hainan Xu a Philipp Koehn. 2018. Paralelné korpusové systémy JHU pre WMT 2018. In *Proceedings of the Third Conference on Machine Translation*, zväzok 2: *Dokumenty o zdieľanej úlohe*, strany 909 – 912, Belgicko, Brusel. Asociácia pre počítačovú lingvistiku.
- Philipp Koehn, Francisco Guzman, Vishrav Chaudhary a Juan M. Pino. 2019. Zistenia zdieľanej úlohy WMT 2019 o paralelnom korpusovom filtrovaní pre nízke zdroje. In *Proceedings of the Fourth Conference on Machine Translation*, zväzok 2: *Common Task Papers*, Florencia, Taliansko. Asociácia pre počítačovú lingvistiku.
- Philipp Koehn, Hieu Hoang, Alexandra Birch, Chris Callison-Burch, Marcello Federico, Nicola Bertoldi, Brooke Cowan, Wade Shen, Christine Moran, Richard Zens, Ondrej Bojar, Chris Dyer, Alexandra Constantin, a Evan Herbst. 2007. Mojžiš: Súbor nástrojov s otvoreným zdrojovým kódom na preklad štatistických strojov. Čo
Výročné zasadnutie Asociácie pre počítačovú lingvistiku (ACL), demo zasadnutie.
- Philipp Koehn, Huda Khayrallah, Kenneth Heafield a Mikel L. Forcada. 2018. Výsledky zdieľanej úlohy WMT 2018 o paralelnom korpusovom filtrovaní. In *Proceedings of the Third Conference on Machine Translation*, zväzok 2: *Common Task Papers*, s. 726739, Belgicko, Brusel. Asociácia pre počítačovú lingvistiku.
- Guillaume Lample, Myle Ott, Alexis Conneau, Ludovic Denoyer, a Marc'Aurelio Ranzato. 2018. Frázy založené na & neurálny bez dozoru strojový preklad. *Empirical Methods in Natural Language Processing (EMNLP)*, strany 5039 – 5049, Belgicko, Brusel. Asociácia pre počítačovú lingvistiku.
- Fantine Mordelet a J-P Vert. 2014. Bagging svm učiť sa z pozitívnych a neoznačených príkladov. *Listy o rozpoznávaní vzorov*, 37: 201 – 209.
- Myle Ott, Sergej Edunov, Alexej Baevski, Angela Fan, Sam Gross, Nathan Ng, David Grangier a Michael Auli. 2019. fairseq: Rýchly, rozšíriteľný súbor nástrojov pre sekvenčné modelovanie. V *konaní o NAACL-HLT 2019: Ukážky*.
- Matt Post. 2018. Výzva na zrozumiteľnosť pri oznamovaní výsledkov bleu. In *Proceedings of the Third Conference on Machine Translation (WMT)*, zväzok 1: *Research Papers*, zväzok 1804.08771, strany 186 – 191, Belgicko, Brusel. Asociácia pre počítačovú lingvistiku.
- Victor M. Sanchez-Cartagena, Marta Banon, Sergio Ortiz-Rojas a Gema Ramirez. 2018. Prompts itové podanie do WMT 2018 paralelného korpusu zdieľanej úlohy. V *konaní o tretej konferencii o strojovom preklade: Dokumenty o zdieľanej úlohe*, strany 955 – 962, Belgicko, Brusel. Asociácia pre počítačovú lingvistiku.
- Holger Schwenk. 2018. Filtrovanie a ťažba paralelných údajov v spoločnom viacjazyčnom priestore. V *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Krátko dokumenty)*, strany 228234, Austrália, Melbourne. Asociácia pre počítačovú lingvistiku.
- Hainan Xu a Philipp Koehn. 2017. Zipporah: Rýchly a škálovateľný systém čistenia dát pre hlučnú webplazil sa paralelné korpusy. In *Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing*, pages 2945 – 2950, Dánsko, Copenhagen. Asociácia pre počítačovú lingvistiku.