

Are the First Steps of the Traditional NLP Pipeline Solved?

Part-of-Speech Tagging from 97% to 100%: Is It Time for Some Linguistics?

Christopher D. Manning

Departments of Linguistics and Computer Science
Stanford University
353 Serra Mall, Stanford CA 94305-9010
manning@stanford.edu

Abstract. I examine what would be necessary to move part-of-speech tagging performance from its current level of about 97.3% token accuracy (56% sentence accuracy) to close to 100% accuracy. I suggest that

How about:

- ▶ Word segmentation
- ▶ Language classification

How about:

- ▶ Word segmentation
- ▶ Language classification
- ▶ And let's broaden our scope beyond English!

Tokenization

The problem of finding/segmenting tokens (UD):

Input:

If_momma_ain't_happy,_nobody_ain't_happy.

Tokenization:

If_momma_ain't_happy_,_nobody_ain't_happy_.

Multi-word expansions:

If_momma_is_not_happy,_nobody_is_not_happy.

Subword segmentation:

If_mo_##mma_ai_##n_'_t_happy_,_no_##body.ai_##n_'_t_happy_.

Methods

1) **Dr. Dron is his backup.**

2) **s=[.][.][.])} > "]"**\$=\1 \2\3 =g**

3) **biiobiiiiobiobiobiiiiib**

4) **Dr . Dro ##n is his backup .**
 b i b i b b b b

LM for tokenization?

- ▶ Finetuning a language model for this task might be overkill
- ▶ Adapters used before (costly to train)
- ▶ Let's do multi-task learning! add a decoder head for tokenization task and sum loss



<https://github.com/machamp-nlp/machamp/>



Settings

UD 2.10: 122 languages in 20 different scripts

Settings

- ▶ RB: Rule Based
- ▶ ST: Single Task: just tokenization
- ▶ MT: Multi-task: UPOS, morph. tagging, lemmatization, dep. parsing
- ▶ ML+MT: Multi-lingual Multi-task model

In treebank results

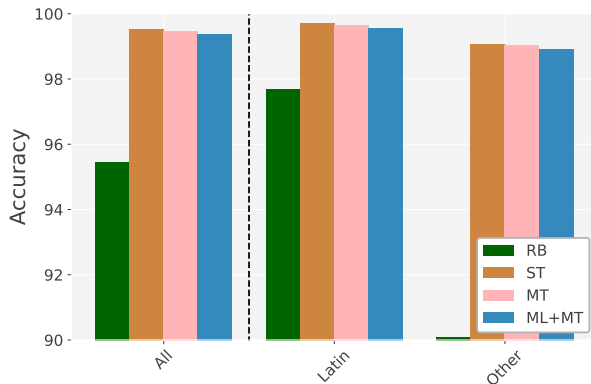
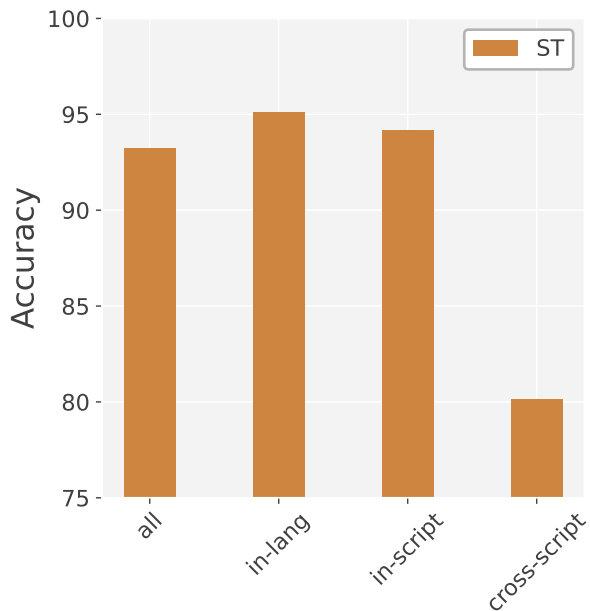


Figure: Results of tokenizers on Latin vs non-Latin languages.
RB=RuleBased, ST=SingleTask, MT=MultiTask,
ML+MT=Multi-Lingual+MultiTask

Cross-treebank results



Analysis

Main takeaways

- ▶ For Latin languages: edge-cases and inconsistencies
- ▶ For other scripts: names, adpositions, compound words
- ▶ More in the paper

Open challenges in language identification

- ▶ Many tools/benchmarks available
- ▶ When to use which?:

Open challenges in language identification

- ▶ Many tools/benchmarks available
- ▶ When to use which?:
 - ▶ # languages
 - ▶ input size
 - ▶ # training instances per language
 - ▶ scripts
 - ▶ language families
 - ▶ domains

Open challenges in language identification

- ▶ Many tools/benchmarks available
- ▶ When to use which?:
 - ▶ # languages
 - ▶ input size
 - ▶ # training instances per language
 - ▶ scripts
 - ▶ language families
 - ▶ domains
- ▶ (Language: ISO 639-3 code)

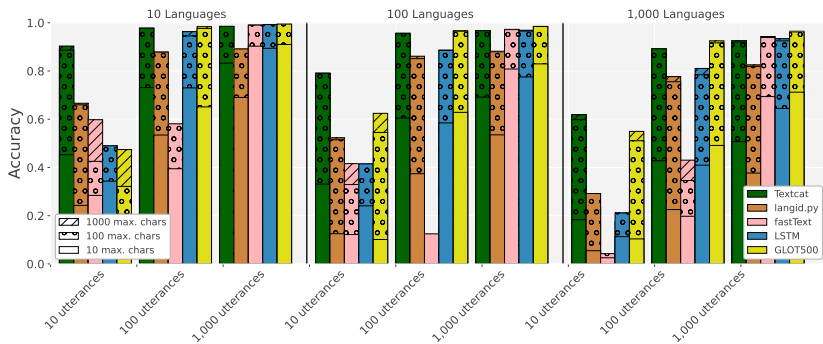
Data

Dataset	langs	scripts	fams	domains
OpenLID	139	25	16	literature, news, wiki, social, grammar, subtitles, spoken
UDHR	397	38	61	rights
LTI LangID	2,110	47	139	wiki, political, religious, grammar
TwitUser	59	20	13	social
MassiveSumm	77	24	13	news
UD2.12	54	11	17	medical, news, academic, wiki, legal, nonfiction, learner-essays, fiction, social, grammar-examples, reviews, religious, spoken
Total	2176/ 7850	51/ 163	145/ 298	

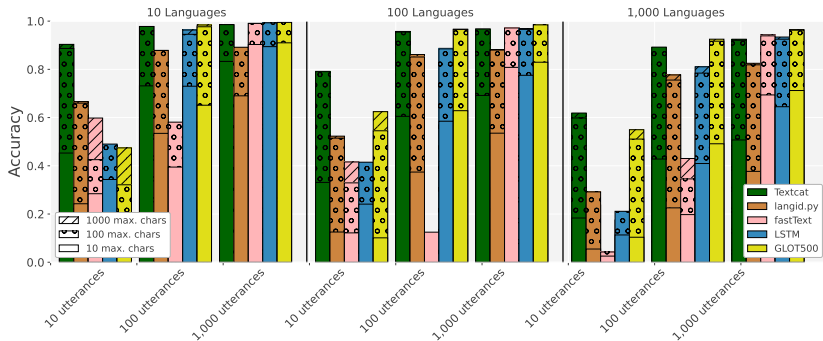
Models

- ▶ Heuristics: `textcat`
- ▶ Naive Bayes: `langid.py`
- ▶ Embeddings: `FastText`
- ▶ Neural: `BiLSTM`
- ▶ CLM: `Glott500`

Size

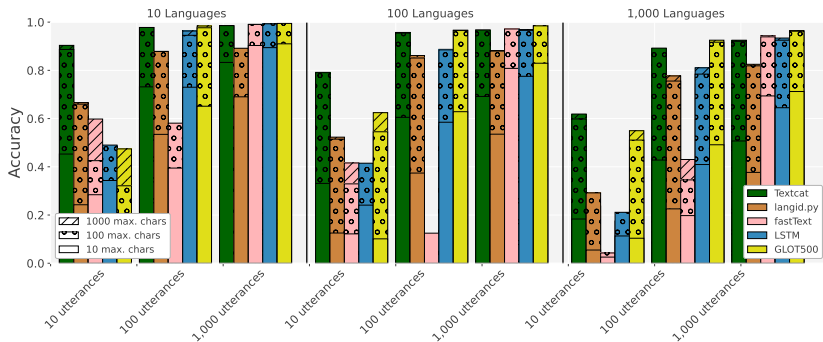


Size



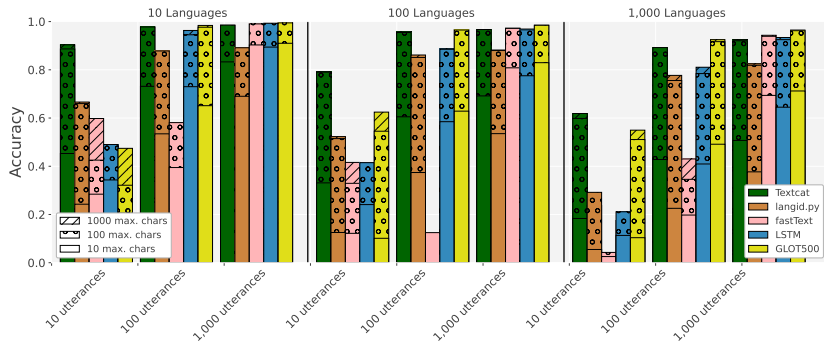
- ▶ 100 characters is enough.
- ▶ Number of utterances more influential, but 100 already quite good

Size: languages



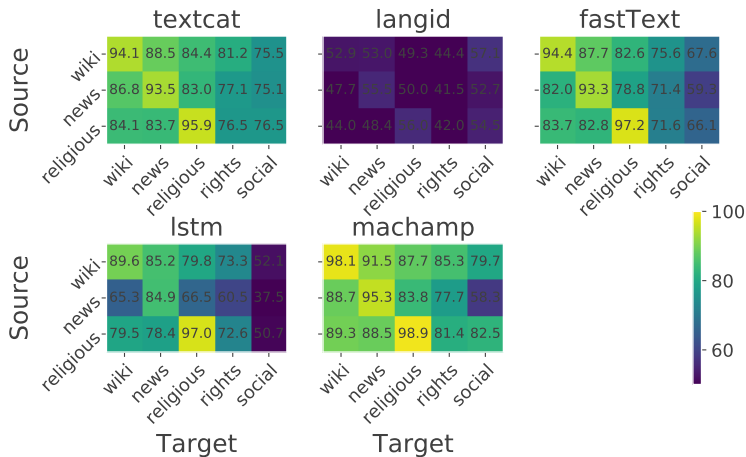
- # of languages is not very influential when there are enough utterances (i.e. 100)

Size: models

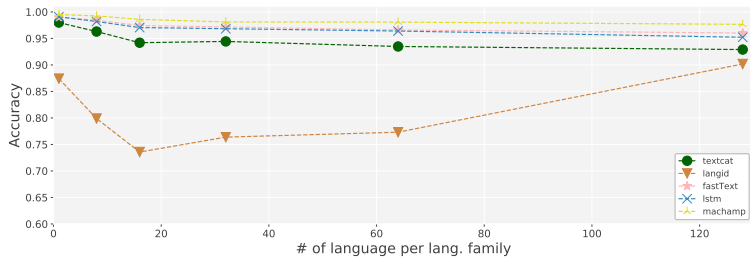


- ▶ Glot500 most robust
- ▶ Character n-gram overlap still impressively good

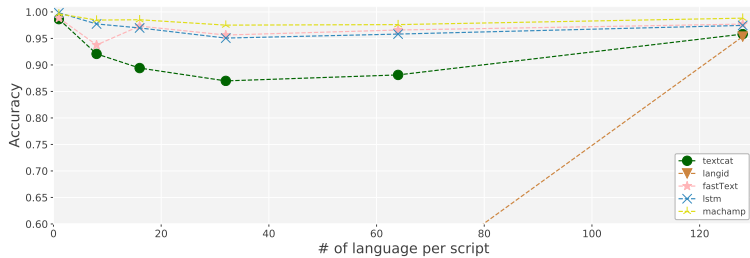
Domains



Language families



Scripts



Other things:

Robustness:

- ▶ Normalization of social media data
- ▶ Multi-task learning (encoder models)
- ▶ Multi-lingual learning
- ▶ Language varieties

Useful (perhaps?)

Obtain information about ISO 639-3 codes: https://bitbucket.org/robvanderger/lang_dist/src/master/