



When Being Unseen from mBERT is just the Beginning: Handling New Languages With Multilingual Language Models

Benjamin Muller, Antonios Anastasopoulos, Benoît Sagot, Djamé Seddah

► To cite this version:

Benjamin Muller, Antonios Anastasopoulos, Benoît Sagot, Djamé Seddah. When Being Unseen from mBERT is just the Beginning: Handling New Languages With Multilingual Language Models. NAACL-HLT 2021 - 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Jun 2021, Mexico City, Mexico. hal-03251105

HAL Id: hal-03251105

<https://inria.hal.science/hal-03251105>

Submitted on 3 Feb 2022

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

When Being Unseen from mBERT is just the Beginning: Handling New Languages With Multilingual Language Models

Benjamin Muller^{†*} Antonios Anastasopoulos[‡] Benoît Sagot[†] Djamé Seddah[†]

[†]Inria, Paris, France ^{*}Sorbonne Université, Paris, France

[‡]Department of Computer Science, George Mason University, USA

firstname.lastname@inria.fr antonis@gmu.edu

Abstract

Transfer learning based on pretraining language models on a large amount of raw data has become a new norm to reach state-of-the-art performance in NLP. Still, it remains unclear how this approach should be applied for unseen languages that are not covered by any available large-scale multilingual language model and for which only a small amount of raw data is generally available. In this work, by comparing multilingual and monolingual models, we show that such models behave in multiple ways on unseen languages. Some languages greatly benefit from transfer learning and behave similarly to closely related high resource languages whereas others apparently do not. Focusing on the latter, we show that this failure to transfer is largely related to the impact of the script used to write such languages. We show that transliterating those languages significantly improves the potential of large-scale multilingual language models on downstream tasks. This result provides a promising direction towards making these massively multilingual models useful for a new set of unseen languages.¹

1 Introduction

Language models are now a new standard to build state-of-the-art Natural Language Processing (NLP) systems. In the past year, monolingual language models have been released for more than 20 languages including Arabic, French, German, and Italian (Antoun et al., 2020; Martin et al., 2020; de Vries et al., 2019; Cañete et al., 2020; Kuratov and Arkipov, 2019; Schweter, 2020, inter alia). Additionally, large-scale multilingual models covering more than 100 languages are now available (XLM-R by Conneau et al. (2020) and mBERT by Devlin et al. (2019)). Still, most of the 6500+ spoken languages in the world (Hammarström, 2016) are not covered—remaining unseen—by

those models. Even languages with millions of native speakers like Sorani Kurdish (about 7 million speakers in the Middle East) or Bambara (spoken by around 5 million people in Mali and neighboring countries) are not covered by any available language models at the time of writing.

Even if training multilingual models that cover more languages and language varieties is tempting, the curse of multilinguality (Conneau et al., 2020) makes it an impractical solution, as it would require to train ever larger models. Furthermore, as shown by Wu and Dredze (2020), large-scale multilingual language models are sub-optimal for languages that are under-sampled during pretraining.

In this paper, we analyze task and language adaptation experiments to get usable language model-based representations for under-studied low resource languages. We run experiments on 15 typologically diverse languages on three NLP tasks: part-of-speech (POS) tagging, dependency parsing (DEP) and named-entity recognition (NER).

Our results bring forth a diverse set of behaviors that we classify in three categories reflecting the abilities of pretrained multilingual language models to be used for low-resource languages. We dub those categories Easy, Intermediate and Hard.

Hard languages include both stable and endangered languages, but they predominantly are languages of communities that are majorly underserved by modern NLP. Hence, we direct our attention to these Hard languages. For those languages, we show that the script they are written in can be a critical element in the transfer abilities of pretrained multilingual language models. Translitterating them leads to large gains in performance outperforming non-contextual strong baselines. To sum up, our contributions are the following:

- We propose a new categorization of the low-resource languages that are unseen by available language models: the Hard, the Intermediate and the Easy languages.

¹Code available at <https://github.com/benjamin-mlr/mbert-unseen-languages.git>

- We show that Hard languages can be better addressed by transliterating them into a better-handled script (typically Latin), providing a promising direction towards making multilingual language models useful for a new set of unseen languages.

2 Background and Motivation

As [Joshi et al. \(2020\)](#) vividly illustrate, there is a large divergence in the coverage of languages by NLP technologies. The majority of the 6500+ of the world’s languages are not studied by the NLP community, since most have few or no annotated datasets, making systems’ development challenging.

The development of such models is a matter of high importance for the inclusion of communities, the preservation of endangered languages and more generally to support the rise of tailored NLP ecosystems for such languages ([Schmidt and Wiegand, 2017](#); [Stecklow, 2018](#); [Seddah et al., 2020](#)). In that regard, the advent of the Universal Dependencies project ([Nivre et al., 2016](#)) and the WikiAnn dataset ([Pan et al., 2017](#)) have greatly increased the number of covered languages by providing annotated datasets for more than 90 languages for dependency parsing and 282 languages for NER.

Regarding modeling approaches, the emergence of multilingual representation models, first with static word embeddings ([Ammar et al., 2016](#)) and then with language model-based contextual representations ([Devlin et al., 2019](#); [Conneau et al., 2020](#)) enabled transfer from high to low-resource languages, leading to significant improvements in downstream task performance ([Rahimi et al., 2019](#); [Kondratyuk and Straka, 2019](#)). Furthermore, in their most recent forms, these multilingual models process tokens at the sub-word level ([Kudo and Richardson, 2018](#)). As such, they work in an open vocabulary setting, only constrained by the pretraining character set. This flexibility enables such models to process any language, even those that are not part of their pretraining data.

When it comes to low-resource languages, one direction is to simply train contextualized embedding models on whatever data is available. Another option is to adapt/fine-tune a multilingual pretrained model to the language of interest. We briefly discuss these two options.

Pretraining language models on a small amount of raw data Even though the amount

of pretraining data seems to correlate with downstream task performance (e.g. compare BERT and RoBERTa ([Liu et al., 2020](#))), several attempts have shown that training a model from scratch can be efficient even if the amount of data in that language is limited. Indeed, [Ortiz Suárez et al. \(2020\)](#) showed that pretraining ELMo models ([Peters et al., 2018](#)) on less than 1GB of text leads to state-of-the-art performance while [Martin et al. \(2020\)](#) showed that pretraining a BERT model on as few as 4GB of diverse enough data results in state-of-the-art performance. [Micheli et al. \(2020\)](#) further demonstrated that decent performance was achievable with only 100MB of raw text data.

Adapting large-scale models for low-resource languages Multilingual language models can be used directly on unseen languages, or they can also be adapted using unsupervised methods. For example, [Han and Eisenstein \(2019\)](#) successfully used unsupervised model adaptation of the English BERT model to Early Modern English for sequence labeling. Instead of fine-tuning the whole model, [Pfeiffer et al. \(2020\)](#) recently showed that adapter layers ([Houlsby et al., 2019](#)) can be injected into multilingual language models to provide parameter efficient task and language transfer.

Still, as of today, the availability of monolingual or multilingual language models is limited to approximately 120 languages, leaving many languages without access to valuable NLP technology, although some are spoken by millions of people, including Bambara and Sorani Kurdish, or are an official language of the European Union, like Maltese.

What can be done for unseen languages? Unseen languages strongly vary in the amount of available data, in their script (many languages use non-Latin scripts such as Sorani Kurdish and Mingrelian), and in their morphological or syntactical properties (most largely differ from high-resource Indo-European languages). This makes the design of a methodology to build contextualized models for such languages challenging at best. In this work, by experimenting with 15 typologically diverse unseen languages, (i) we show that there is a diversity of behavior depending on the script, the amount of available data, and the relation to the pretraining languages; (ii) Focusing on the unseen languages that lag in performance compared to their easier-to-handle counterparts, we show that the script plays a

critical role in the transfer abilities of multilingual language models. Transliterating such languages to a script which is used by a related language seen during pretraining can lead to significant improvement in downstream performance.

3 Experimental Setting

We will refer to any languages that are not covered by pretrained language models as “unseen.” We select a small portion of those languages within a large scope of language families and scripts. Our selection is constrained to 15 typologically diverse languages for which we have evaluation data for at least one of our three downstream tasks. Our selection includes low-resource Indo-European and Uralic languages, as well as members of the Bantu, Semitic, and Turkic families. None of these 15 languages are included in the pretraining corpora of mBERT. Information about their scripts, language families, and amount of available raw data can be found in the Appendix in Table 12.

3.1 Raw Data

To perform pretraining and fine-tuning on monolingual data, we use the deduplicated datasets from the OSCAR project (Ortiz Suárez et al., 2019). OSCAR is a corpus extracted from a Common Crawl Web snapshot.² It provides a significant amount of data for all the unseen languages we work with, except for Buryat, Meadow Mari, Erzya and Livvi for which we use Wikipedia dumps and for Narabizi, Naija and Faroese, for which we use data collected by Seddah et al. (2020), Caron et al. (2019) and Biemann et al. (2007) respectively.

3.2 Non-contextual Baselines

For parsing and POS tagging, we use the UDPipe future system (Straka, 2018) as our baseline. This model is a LSTM-based (Hochreiter and Schmidhuber, 1997) recurrent architecture trained with pretrained static word embedding (Mikolov et al., 2013) (hence our non-contextual characterization) along with character-level embeddings. This system was ranked in the very first positions for parsing and tagging in the CoNLL shared task 2018 (Zeman and Hajič, 2018). For NER we use the LSTM-CRF model with character and word level embedding using Qi et al. (2020) implementation.

²<http://commoncrawl.org/>

3.3 Language Models

In all our study, we train our language models using the Transformers library (Wolf et al., 2020).

MLM from scratch The first approach we evaluate is to train a dedicated language model from scratch on the available raw data we have. To do so, we train a language-specific SentencePiece tokenizer (Kudo and Richardson, 2018) before training a Masked-Language Model (MLM) using the RoBERTa (base) architecture and objective functions (Liu et al., 2019). As we work with significantly smaller pretraining sets than in the original setting, we reduce the number of layers to 6 layers in place of the original 12 layers.

Multilingual Language Models We want to assess how large-scale multilingual language models can be used and adapted to languages that are not in their pretraining corpora. We work with the multilingual version of BERT (mBERT) trained on the concatenation of Wikipedia corpora in 104 languages (Devlin et al., 2019). We also ran experiments with the XLM-R base version (Conneau et al., 2020) trained on 100 languages using data extracted from the Web. As the observed behaviors are very similar between both models, we only report results using mBERT. Note that mBERT is highly biased toward Indo-Europeans languages written in the Latin script. More than 77% of the subword vocabulary are in the Latin script while only 1% are in the Georgian script (Ács, 2019).

Adapting Multilingual Language Models to unseen languages with MLM-TUNING Following previous work (Han and Eisenstein, 2019; Karthikeyan et al., 2019; Pfeiffer et al., 2020), we adapt large-scale multilingual models by fine-tuning them with their Mask-Language-Model objective directly on the available raw data in the unseen target language. We refer to this process as MLM-TUNING. We will refer to a MLM-tuned mBERT model as mBERT+MLM.

3.4 Downstream Tasks

We perform experiments on POS tagging, Dependency Parsing (DEP), and Name Entity Recognition (NER). We use annotated data from the Universal Dependency project (Nivre et al., 2016) for POS tagging and parsing, and the WikiAnn dataset (Pan et al., 2017) for NER. For POS tagging and NER, we append a linear classifier layer on top of

Model	UPOS				LAS				NER			
	mBERT	mBERT+MLM	MLM	Baseline	mBERT	mBERT+MLM	MLM	Baseline	mBERT	mBERT+MLM	MLM	Baseline
Faroese	96.3	96.5	91.1	95.4	84.0	86.4	67.6	83.1	52.1	58.3	39.3	44.8
Naija	89.3	89.6	87.1	89.2	71.5	69.2	63.0	68.3	-	-	-	-
Swiss German	76.7	78.7	65.4	75.2	41.2	69.6	30.0	32.2	-	-	-	-
Mingrelian	-	-	-	-	-	-	-	-	53.6	68.4	42.0	48.2

Table 1: **Easy Languages** POS, Parsing and NER scores comparing mBERT, mBERT+MLM and monolingual MLM to strong non-contextual baselines when trained and evaluated on unseen languages. Easy Languages are the ones on which mBERT outperforms strong baselines out-of-the-box. Baselines are LSTM based models from UDPipe-future (Straka, 2018) for parsing and POS tagging and Stanza (Qi et al., 2020) for NER.

the language model. For parsing, following Konratyuk and Straka (2019), we append a Bi-Affine Graph prediction layer (Dozat and Manning, 2017). We refer to the process of fine-tuning a language model in a task-specific way as TASK-TUNING.³

3.5 Dataset Splits

For each task and language, we use the provided training, validation and test dataset split except for the ones that have less than 500 training sentences. In this case, we concatenate the training and test set and perform 8-folds cross-Validation and use the validation set for early stopping. If no validation set is available, we isolate one of the folds for validation and report the test scores as the average of the other folds. This enables us to train on at least 500 sentences in all our experiments (except for Swiss German for which we only have 100 training examples) and reduce the impact of the annotated dataset size on our analysis. Since cross-validation results in training on very limited number of examples, we refer to training in this cross-validation setting as *few-shot* learning.

4 The Three Categories of Unseen Languages

For each unseen language and each task, we experiment with our three modeling approaches: (a) **Training a language model from scratch on the available raw data** and then fine-tuning it on any available annotated data in the target language. (b) **Fine-tuning mBERT with TASK-TUNING** directly on the target language. (c) Finally, **adapting mBERT to the unseen language using MLM-TUNING** before fine-tuning it in a supervised way on the target language. We then compare all these experiments to our non-contextual strong baselines. By doing so, we can assess if language models are

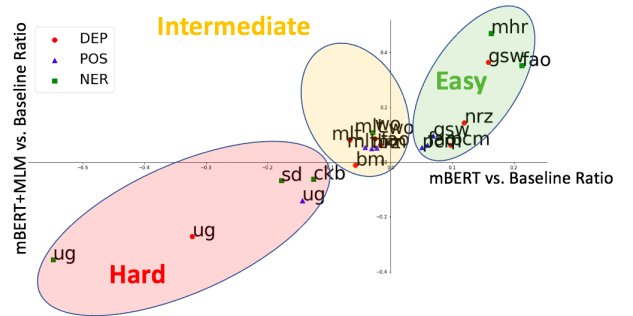


Figure 1: Visualizing our Typology of Unseen Languages. (X,Y) positions are computed for each language and each task as follows: given the score of mBERT denoted s and s_0 the baseline score: $X = \frac{s_{mBERT} - s_0}{s_0}$, $Y = \frac{s_{mBERT+MLM} - s_0}{s_0}$. Easy Languages are the ones on which mBERT performs better than the baseline without MLM-TUNING and Intermediate languages are the ones that require MLM-TUNING. For Hard languages, mBERT under-performs the baselines in all settings.

a practical solution to handle each of these unseen languages.

Interestingly we find a large diversity of behaviors across languages regarding those language model training techniques. As summarized in Figure 1, we observe three clear clusters of languages.

The first cluster, which we dub “Easy”, corresponds to the languages that do not require extra MLM-TUNING for mBERT to achieve good performance. mBERT has the modeling abilities to process such languages without relying on raw data and can outperform strong non-contextual baselines as such. In the second cluster, the “Intermediate” languages require MLM-TUNING. mBERT is not able to beat strong non-contextual baselines using only TASK-TUNING, but MLM-TUNING enables it to do so. Finally, Hard languages are those on which mBERT fails to deliver any decent per-

³Details about optimization can be found in Appendix B

Model	UPOS				LAS				NER			
	mBERT	mBERT+MLM	MLM	Baseline	mBERT	mBERT+MLM	MLM	Baseline	mBERT	mBERT+MLM	MLM	Baseline
Maltese	92.0	96.4	92.05	96.0	74.4	82.1	66.5	79.7	61.2	66.7	62.5	63.1
Narabizi	81.6	84.2	71.3	84.2	56.5	57.8	41.8	52.8	-	-	-	-
Bambara	90.2	92.6	78.1	92.3	71.8	75.4	46.4	76.2	-	-	-	-
Wolof	92.8	95.2	88.4	94.1	73.3	77.9	62.8	77.0	-	-	-	-
Erzya	89.3	91.2	84.4	91.1	61.2	66.6	47.8	65.1	-	-	-	-
Livvi	83.0	85.5	81.1	84.1	36.3	42.3	35.2	40.1	-	-	-	-
Mari	-	-	-	-	-	-	-	-	55.2	57.6	44.0	56.1

Table 2: **Intermediate Languages** POS, Parsing and NER scores comparing mBERT, mBERT+MLM and monolingual MLM to strong non-contextual baselines when trained and evaluated on unseen languages. Intermediate Languages are the ones for which mBERT requires MLM-TUNING to outperform the baselines.

formance even after MLM- and TASK- fine-tuning. mBERT simply does not have the capacity to learn and process such languages.

We emphasize that our categorization of unseen languages is only based on the relative performance of mBERT after fine-tuning compared to strong non-contextual baseline models. We leave for future work the analysis of the absolute performance of the model on such languages (e.g. analysing the impact of the fine-tuning data set size on mBERT’s downstream performance).

In this section, we present our results in detail in each of these language clusters and provide insights into their linguistic properties.

4.1 Easy

Easy languages are the ones on which mBERT delivers high performance out-of-the-box, compared to strong baselines. We classify Faroese, Swiss German, Naija and Mingrelian as easy languages and report performance in Table 1.

We find that those languages match two conditions:

- They are closely related to languages used during MLM pretraining
- These languages use the same script as such closely related languages.

Such languages benefit from multilingual models, as cross-lingual transfer is easy to achieve and hence quite effective.

More details about those languages can be found in Appendix C.

4.2 Intermediate

The second type of languages (which we dub “Intermediate”) are generally harder to process for pretrained MLMs out-of-the-box. In particular, pretrained multilingual language models are typically outperformed by a non-contextual strong baselines. Still, MLM-TUNING has an important impact and

leads to usable state-of-the-art models.

A good example of such an intermediate language is Maltese, a member of the Semitic language but using the Latin script. Maltese has not been seen by mBERT during pretraining. Other Semitic languages though, namely Arabic and Hebrew, have been included in the pretraining languages. As seen in Table 2, the non-contextual baseline outperforms mBERT. Additionally, a monolingual MLM trained on only 50K sentences matches mBERT performance for both NER and POS tagging. However, the best results are reached with MLM-TUNING: the proper use of monolingual data and the advantage of similarity to other pretraining languages render Maltese a *tackle-able* language as shown by the performance gain over our strong non-contextual baselines.

Our Maltese dependency parsing results are in line with those of Chau et al. (2020), who also showed that MLM-TUNING leads to significant improvements. They also additionally showed that a small vocabulary transformation allowed fine-tuning to be even more effective and gain 0.8 LAS points more. We further discuss the vocabulary adaptation technique of Chau et al. (2020) in section 6.

We consider Narabizi (Seddah et al., 2020), an Arabic dialect spoken in North-Africa written in the Latin script and code-mixed with French, to fall in the same Intermediate category, because it follows the same pattern. For both POS tagging and parsing, the multilingual models outperform the monolingual NarabiziBERT. In addition, MLM-TUNING leads to significant improvements over the non-language-tuned mBERT baseline, also outperforming the non-contextual dependency parsing baseline.

We also categorize Bambara, a Niger-Congo Bantu language spoken in Mali and surrounding countries, as Intermediate, relying mostly on the

Model	UPOS				LAS				NER			
	mBERT	mBERT+MLM	MLM	Baseline	mBERT	mBERT+MLM	MLM	Baseline	mBERT	mBERT+MLM	MLM	Baseline
Uyghur	77.0	88.4	87.4	90.0	45.5	48.9	57.3	67.9	24.3	34.6	41.4	53.8
Sindhi	-	-	-	-	-	-	-	-	42.3	47.9	45.2	51.4
Sorani Kurdish	-	-	-	-	-	-	-	-	70.4	75.6	80.6	80.5

Table 3: **Hard Languages** POS, Parsing and NER scores comparing mBERT, mBERT+MLM and monolingual MLM to strong non-contextual baselines when trained and evaluated on unseen languages. Hard Languages are the ones for which mBERT fails to reach decent performance even after MLM-TUNING.

POS tagging results which follow similar patterns as Maltese and Narabizi. We note that the BambaraBERT that we trained achieves notably poor performance compared to the non-contextual baseline, a fact we attribute to the extremely low amount of available data (1000 sentences only). We also note that the non-contextual baseline is the best performing model for dependency parsing, which could also potentially classify Bambara as a “Hard” language instead.

Our results in Wolof follow the same pattern. The non-contextual baseline achieves a 77.0 in LAS outperforming mBERT. However, MLM-TUNING achieves the highest score of 77.9.

The importance of script We now turn our focus to Uralic languages. Finnish, Estonian, and Hungarian are high-resource representatives of this language family that are typically included in multilingual LMs, also having task-tuning data available in large quantities. However, for several smaller Uralic languages, task-tuning data are generally very scarce.

We report in Table 2 the performance for two low-resource Uralic languages, namely Livvi and Erzya using 8-fold cross-validation, with each run only using around 700 training instances. Note the striking difference between the parsing performance (LAS) of mBERT on Livvi, written with the Latin script, and on Erzya that uses the Cyrillic script. This suggests that the script could be playing a critical role when transferring to those languages. We explore this hypothesis in detail in section 5.2.

4.3 Hard

The last category of the hard unseen language is perhaps the most interesting one, as these languages are very hard to process. mBERT is outperformed by non-contextual baselines as well as by monolingual language models trained from scratch on the available raw data. At the same time, MLM-TUNING on the available raw data has a minimal

impact on performance.

Uyghur, a Turkic language with about 10-15 million speakers in central Asia, is a prime example of a hard language for current models. In our experiments, outlined in Table 3, the non-contextual baseline outperforms all contextual variants, both monolingual and multilingual, in all the tasks with up to 20 points difference compared to mBERT for parsing. Additionally, the monolingual UyghurBERT trained on only 105K sentences outperforms mBERT even after MLM-TUNING.

We attribute this discrepancy to script differences: Uyghur uses the Perso-Arabic script, when the other Turkic languages that were part of mBERT pretraining use either the Latin (e.g. Turkish) or the Cyrillic script (e.g. Kazakh).

Sorani Kurdish (also known as Central Kurdish) is a similarly hard language, mainly spoken in Iraqi Kurdistan by around 8 million speakers, which uses the Sorani alphabet, a variant of the Arabic script. We can solely evaluate on the NER task, where the non-contextual baseline and the monolingual SoraniBERT perform similarly around 80.5 F1-score outperforming significantly mBERT which only reaches 70.4 in F1-score. MLM-TUNING on 380K sentences of Sorani texts improves mBERT performance to 75.6 F1-score, but it is still lagging behind the baseline. Our results in Sindhi follow the same pattern. The non-contextual baseline achieves a 51.4 F1-score outperforming with a large margin our language models (a monolingual SindhiBERT achieves an F1-score of 45.2, and mBERT is worse at 42.3).

5 Tackling Hard Languages with Multilingual Language Models

Our intermediate Uralic language results provide initial supporting evidence for our argument on the importance of having pretrained LMs on languages with similar scripts, even for generally high-resource language families. Our hypothesis is that the script is a key element for language models to

Model	POS	LAS	NER	Model	NER
Uyghur (Arabic→Latin)			Sorani (Arabic→Latin)		
UyghurBERT	87.4→86.2	57.3→54.6	41.4→41.7	SoraniBERT	80.6→78.9
mBERT	77.0→87.9	45.7→65.0	24.3→35.7	mBERT	70.5→77.8
mBERT+MLM	77.3→ 89.8	48.9→ 66.8	34.7→ 55.2	mBERT+MLM	75.6→ 82.7
Buryat (Cyrillic→Latin)			Meadow Mari (Cyrillic→Latin)		
BuryatBERT	75.8→75.8	31.4→31.4	–	MariBERT	44.0→45.5
mBERT	83.9→81.6	50.3→45.8	–	mBERT	55.2→58.2
mBERT+MLM	86.5 →84.6	52.9 →51.9	–	mBERT+MLM	57.6→ 65.9
Erzya (Cyrillic→Latin)			Mingrelian (Georgian→Latin)		
ErzyaBERT	84.4→84.5	47.8→47.8	–	MingrelianBERT	42.0→42.2
mBERT	89.3→88.2	61.2→58.3	–	mBERT	53.6→41.8
mBERT+MLM	91.2 →90.5	66.6 →65.5	–	mBERT+MLM	68.4 →62.6

Table 4: Transliterating low-resource languages into the Latin script leads to significant improvements in languages like Uyghur, Sorani, and Meadow Mari. For languages like Erzya and Buryat transliteration, does not significantly influence results, while it does not help for Mingrelian. In all cases, mBERT+MLM is the best approach.

correctly process unseen languages.

To test this hypothesis, we assess the ability of mBERT to process an unseen language after transliterating it to another script present in the pretraining data. We experiment on six languages belonging to four language families: Erzya, Buryat and Meadow Mari (Uralic), Sorani Kurdish (Iranian, Indo-European), Uyghur (Turkic) and Mingrelian (Kartvelian). We apply the following transliterations:

- Erzya/Buryat/Mari: Cyrillic → Latin Script
- Uyghur: Arabic Script → Latin Script
- Sorani: Arabic Script → Latin Script
- Mingrelian: Georgian Script → Latin Script

5.1 Linguistically-motivated transliteration

The strategy we used to transliterate the above-listed language is specific to the purpose of our experiments. Indeed, our goal is for the model to take advantage of the information it has learned during training on a related language written in the Latin script. The goal of our transliteration is therefore to transcribe each character in the source script, which we assume corresponds to a phoneme, into the most frequent (sometimes only) way this phoneme is rendered in the closest related language written in the Latin script, hereafter the target language. This process is not a transliteration strictly speaking, and it needs not be reversible. It is not a phonetization either, but rather a way to render the source language in a way that maximizes the similarity between the transliterated source language and the target language.

We have manually developed transliteration

scripts for Uyghur and Sorani Kurdish, using respectively Turkish and Kurmanji Kurdish as target languages, only Turkish being one of the languages used to train mBERT. Note however that Turkish and Kurmanji Kurdish share a number of conventions for rendering phonemes in the Latin script (for instance, /ʃ/, rendered in English by “sh”, is rendered in both languages by “ş”; as a result, the Arabic letter “ش”, used in both languages, is rendered as “ş” by both our transliteration scripts). As for Erzya, Buryat and Mari, we used the readily available transliteration package *transliterate*,⁴ which performs a standard transliteration.⁵ We used the Russian transliteration module, as it covers the Cyrillic script. Finally, for our control experiments on Mingrelian, we used the Georgian transliteration module from the same package.

5.2 Transfer via Transliteration

We train mBERT with MLM-TUNING and TASK-TUNING as well as monolingual BERT model trained from scratch on the transliterated data. We also run controlled experiments on high-resource languages written in the Latin script on which mBERT was pretrained on, namely Arabic, Japanese and Russian (reported in Table 5).

Our results with and without transliteration are listed in Table 4. Transliteration for Sorani and Uyghur has a noticeable positive impact. For instance, transliterating Uyghur to Latin leads to an improvement of 16 points in parsing and 20 points

⁴<https://pypi.org/project/transliterate/>

⁵In future work, we intend to develop dedicated transliteration scripts using the strategy described above, and to compare the results obtained with it with those described here.

in NER. For one of the low-resource Uralic languages, Meadow Mari, we observe an 8 F1-score points improvement on NER, while for other Uralic languages like Erzya the effect of transliteration is very minor. The only case where transliteration to the Latin script leads to a drop in performance for mBERT and mBERT+MLM is Mingrelian.

We interpret our results as follows. When running MLM-TUNING and TASK-TUNING, mBERT associates the target unseen language to a set of similar languages seen during pretraining based on the script. In consequence, mBERT is not able to associate a language to its related language if they are not written in the same script. For instance, transliterating Uyghur enables mBERT to match it to Turkish, a language which accounts for a sizeable portion of mBERT pretraining. In the case of Mingrelian, transliteration has the opposite effect: transliterating Mingrelian in the Latin script is harming the performance as mBERT is not able to associate it to Georgian which is seen during pretraining and uses the Georgian script.

This is further supported by our experiments on high resource languages (cf. Table 5). When transliterating pretrained languages such as Arabic, Russian or Japanese, mBERT is not able to compete with the performance reached when using the script seen during pretraining. Transliterating the Arabic script and the Cyrillic script to Latin does not automatically improve mBERT performance as it does for Sorani, Uyghur and Meadow Mari. For instance, transliterating Arabic to the Latin script leads to a drop in performance of 1.5, 4.1 and 6.9 points for POS tagging, parsing and NER respectively.⁶

Our findings are generally in line with previous work. Transliteration to English specifically (Lin et al., 2016; Durrani et al., 2014) and named entity transliteration (Kundu et al., 2018; Grundkiewicz and Heafield, 2018) has been proven useful for cross-lingual transfer in tasks like NER, entity linking (Rijhwani et al., 2019), morphological inflection (Murikinati et al., 2020), and Machine Translation (Amrhein and Sennrich, 2020).

The transliteration approach provides a viable path for rendering large pretrained models like mBERT useful for all languages of the world. Indeed, as reported in Table 4, transliterating both Uyghur and Sorani leads to matching or outper-

Model	Original Script → Latin Script		
	POS	LAS	NER
Arabic	96.4 → 94.9	82.9 → 78.8	87.8 → 80.9
Russian	98.1 → 96.0	88.4 → 84.5	88.1 → 86.0
Japanese	97.4 → 95.7	88.5 → 86.9	61.5 → 55.6

Table 5: mBERT TASK-TUNED on high resource languages for POS tagging, parsing and NER. We compare fine-tuning done on data written the original language script with fine-tuning done on Latin transliteration. In all cases, transliteration degrades downstream performance.

forming the performance of non-contextual strong baselines and deliver usable models (e.g. +12.5 POS accuracy in Uyghur).

6 Discussion and Conclusion

Pretraining ever larger language models is a research direction that is currently receiving a lot of attention and resources from the NLP research community (Raffel et al., 2019; Brown et al., 2020). Still, a large majority of human languages are under-resourced making the development of monolingual language models very challenging in those settings. Another path is to build large scale multilingual language models.⁷ However, such an approach faces the inherent zipfian structure of human languages, making the training of a single model to cover all languages an unfeasible solution (Conneau et al., 2020). Reusing large scale pretrained language models for new unseen languages seems to be a more promising and reasonable solution from a cost-efficiency and environmental perspective (Strubell et al., 2019).

Recently, Pfeiffer et al. (2020) proposed to use adapter layers (Houlsby et al., 2019) to build parameter efficient multilingual language models for unseen languages. However, this solution brings no significant improvement in the supervised setting, compared to a more simple Masked-Language Model finetuning. Furthermore, developing a language agnostic adaptation method is an unreasonable wish with regard to the large typological diversity of human languages.

On the other hand, the promising vocabulary adaptation technique of Chau et al. (2020) which leads to good dependency parsing results on unseen languages when combined with task-tuning has

⁶Details and complete results on these controlled experiments can be found in Appendix E.

⁷Even though we explore a different research direction, recent advances in small scale and domain specific language models suggest such models could also have an important impact for those languages (Micheli et al., 2020).

so far been tested only on Latin script languages (Singlish and Maltese). We expect that it will be orthogonal to our transliteration approach, but we leave for future work the study of its applicability and efficacy on more languages and tasks.

In this context, we bring empirical evidence to assess the efficiency of language models pretraining and adaptation methods on 15 low-resource and typologically diverse unseen languages. Our results show that the “Hard” languages are currently out-of-the-scope of any currently available language models and are therefore left outside of the current NLP progress. By focusing on those, we find that this challenge is mostly due to the script. Transliterating them to a script that is used by a related higher resource language on which the language model has been pretrained on leads to large improvements in downstream performance. Our results shed some new light on the importance of the script in multilingual pretrained models. While previous work suggests that multilingual language models could transfer efficiently across scripts in zero-shot settings (Pires et al., 2019; Karthikeyan et al., 2019), our results show that such cross-script transfer is possible only if the model has seen related languages in the same script during pretraining.

Our work paves the way for a better understanding of the mechanics at play in cross-language transfer learning in low-resource scenarios. We strongly believe that our method can contribute to bootstrapping NLP resources and tools for low-resource languages, thereby favoring the emergence of NLP ecosystems for languages currently under-served by the NLP community.

Acknowledgments

The Inria authors were partly funded by two French Research National agency projects, namely projects PARSITI (ANR-16-CE33-0021) and SoSweet (ANR-15-CE38-0011), as well as by Benoit Sagot’s chair in the PRAIRIE institute as part of the “Investissements d’avenir” programme under the reference ANR-19-P3IA-0001. Antonios Anastasopoulos is generously supported by NSF Award 2040926 and is also thankful to Graham Neubig for very insightful initial discussions on this research direction.

References

- Judit Ács. 2019. Exploring BERT’s vocabulary. [Http://juditacs.github.io/2019/02/19/bert-tokenization-stats.html](http://juditacs.github.io/2019/02/19/bert-tokenization-stats.html).
- Waleed Ammar, George Mulcaire, Yulia Tsvetkov, Guillaume Lample, Chris Dyer, and Noah A Smith. 2016. Massively multilingual word embeddings. arXiv:1602.01925.
- Chantal Amrhein and Rico Sennrich. 2020. [On Romanization for model transfer between scripts in neural machine translation](#). In *Findings of the Association for Computational Linguistics: EMNLP 2020*, pages 2461–2469, Online. Association for Computational Linguistics.
- Wissam Antoun, Fady Baly, and Hazem Hajj. 2020. [AraBERT: Transformer-based model for Arabic language understanding](#). In *Proceedings of the 4th Workshop on Open-Source Arabic Corpora and Processing Tools, with a Shared Task on Offensive Language Detection*, pages 9–15, Marseille, France. European Language Resource Association.
- Chris Biemann, Gerhard Heyer, Uwe Quasthoff, and Matthias Richter. 2007. The Leipzig Corpora collection-monolingual corpora of standard size. *Proceedings of Corpus Linguistic*, 2007.
- Tom B Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, et al. 2020. Language models are few-shot learners. arXiv:2005.14165.
- Bernard Caron, Marine Courtin, Kim Gerdes, and Sylvain Kahane. 2019. A surface-syntactic UD treebank for Naija. In *Proceedings of the 18th International Workshop on Treebanks and Linguistic Theories (TLT, SyntaxFest 2019)*, pages 13–24.
- José Cañete, Gabriel Chaperon, Rodrigo Fuentes, and Jorge Pérez. 2020. [Spanish pre-trained BERT model and evaluation data](#). In *PML4DC at ICLR 2020*.
- Ethan C. Chau, Lucy H. Lin, and Noah A. Smith. 2020. [Parsing with multilingual BERT, a small corpus, and a small treebank](#). In *Findings of the Association for Computational Linguistics: EMNLP 2020*, pages 1324–1334, Online. Association for Computational Linguistics.
- Alexis Conneau, Kartikay Khandelwal, Naman Goyal, Vishrav Chaudhary, Guillaume Wenzek, Francisco Guzmán, Edouard Grave, Myle Ott, Luke Zettlemoyer, and Veselin Stoyanov. 2020. [Unsupervised cross-lingual representation learning at scale](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 8440–8451, Online. Association for Computational Linguistics.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. [BERT: Pre-training of deep bidirectional transformers for language understanding](#). In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pages 4171–4186, Minneapolis, Minnesota. Association for Computational Linguistics.
- Timothy Dozat and Christopher D. Manning. 2017. [Deep biaffine attention for neural dependency parsing](#). In *5th International Conference on Learning Representations, ICLR 2017, Toulon, France, April 24–26, 2017, Conference Track Proceedings*. OpenReview.net.
- Nadir Durrani, Hassan Sajjad, Hieu Hoang, and Philipp Koehn. 2014. [Integrating an unsupervised transliteration model into statistical machine translation](#). In *Proceedings of the 14th Conference of the European Chapter of the Association for Computational Linguistics, volume 2: Short Papers*, pages 148–153, Gothenburg, Sweden. Association for Computational Linguistics.
- Roman Grundkiewicz and Kenneth Heafield. 2018. [Neural machine translation techniques for named entity transliteration](#). In *Proceedings of the Seventh Named Entities Workshop*, pages 89–94, Melbourne, Australia. Association for Computational Linguistics.
- Harald Hammarström. 2016. Linguistic diversity and language evolution. *Journal of Language Evolution*, 1(1):19–29.
- Xiaochuang Han and Jacob Eisenstein. 2019. [Unsupervised domain adaptation of contextualized embeddings for sequence labeling](#). In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)*, pages 4238–4248, Hong Kong, China. Association for Computational Linguistics.
- Sepp Hochreiter and Jürgen Schmidhuber. 1997. Long short-term memory. *Neural computation*, 9(8):1735–1780.
- Neil Houlsby, Andrei Giurgiu, Stanislaw Jastrzebski, Bruna Morrone, Quentin De Laroussilhe, Andrea Gesmundo, Mona Attariyan, and Sylvain Gelly. 2019. Parameter-efficient transfer learning for NLP. arXiv:1902.00751.
- Pratik Joshi, Sebastin Santy, Amar Budhiraja, Kalika Bali, and Monojit Choudhury. 2020. [The state and fate of linguistic diversity and inclusion in the NLP world](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 6282–6293, Online. Association for Computational Linguistics.
- K Karthikeyan, Zihan Wang, Stephen Mayhew, and Dan Roth. 2019. Cross-lingual ability of multilingual BERT: An empirical study. In *International Conference on Learning Representations*.

- Diederik P. Kingma and Jimmy Ba. 2015. [Adam: A method for stochastic optimization](#). In *3rd International Conference on Learning Representations, ICLR 2015, San Diego, CA, USA, May 7-9, 2015, Conference Track Proceedings*.
- Dan Kondratyuk and Milan Straka. 2019. [75 languages, 1 model: Parsing universal dependencies universally](#). In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)*, pages 2779–2795, Hong Kong, China. Association for Computational Linguistics.
- Taku Kudo and John Richardson. 2018. [SentencePiece: A simple and language independent subword tokenizer and detokenizer for neural text processing](#). In *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing: System Demonstrations*, pages 66–71, Brussels, Belgium. Association for Computational Linguistics.
- Soumyadeep Kundu, Sayantan Paul, and Santanu Pal. 2018. [A deep learning based approach to transliteration](#). In *Proceedings of the Seventh Named Entities Workshop*, pages 79–83, Melbourne, Australia. Association for Computational Linguistics.
- Yuri Kuratov and Mikhail Arhipov. 2019. [Adaptation of deep bidirectional multilingual transformers for Russian language](#). arXiv:1905.07213.
- Ying Lin, Xiaoman Pan, Aliya Deri, Heng Ji, and Kevin Knight. 2016. [Leveraging entity linking and related language projection to improve name transliteration](#). In *Proceedings of the Sixth Named Entity Workshop*, pages 1–10, Berlin, Germany. Association for Computational Linguistics.
- Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Mandar Joshi, Danqi Chen, Omer Levy, Mike Lewis, Luke Zettlemoyer, and Veselin Stoyanov. 2019. RoBERTa: A robustly optimized BERT pretraining approach. arXiv:1907.11692.
- Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Mandar Joshi, Danqi Chen, Omer Levy, Mike Lewis, Luke Zettlemoyer, and Veselin Stoyanov. 2020. [Ro{bert}a: A robustly optimized {bert} pretraining approach](#).
- Louis Martin, Benjamin Muller, Pedro Javier Ortiz Suárez, Yoann Dupont, Laurent Romary, Éric de la Clergerie, Djamé Seddah, and Benoît Sagot. 2020. [CamemBERT: a tasty French language model](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 7203–7219, Online. Association for Computational Linguistics.
- Vincent Micheli, Martin d’Hoffschmidt, and François Fleuret. 2020. [On the importance of pre-training data volume for compact language models](#). In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 7853–7858, Online. Association for Computational Linguistics.
- Tomas Mikolov, Ilya Sutskever, Kai Chen, Greg S Corrado, and Jeff Dean. 2013. Distributed representations of words and phrases and their compositionality. In *Advances in neural information processing systems*, pages 3111–3119.
- Nikitha Murikinati, Antonios Anastasopoulos, and Graham Neubig. 2020. [Transliteration for cross-lingual morphological inflection](#). In *Proceedings of the 17th SIGMORPHON Workshop on Computational Research in Phonetics, Phonology, and Morphology*, pages 189–197, Online. Association for Computational Linguistics.
- Joakim Nivre, Marie-Catherine De Marneffe, Filip Ginter, Yoav Goldberg, Jan Hajic, Christopher D Manning, Ryan McDonald, Slav Petrov, Sampo Pyysalo, Natalia Silveira, et al. 2016. Universal dependencies v1: A multilingual treebank collection. In *Proceedings of the Tenth International Conference on Language Resources and Evaluation (LREC’16)*, pages 1659–1666.
- Pedro Javier Ortiz Suárez, Laurent Romary, and Benoît Sagot. 2020. [A monolingual approach to contextualized word embeddings for mid-resource languages](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 1703–1714, Online. Association for Computational Linguistics.
- Pedro Javier Ortiz Suárez, Benoît Sagot, and Laurent Romary. 2019. [Asynchronous pipelines for processing huge corpora on medium to low resource infrastructures](#). In *Proceedings of the Workshop on Challenges in the Management of Large Corpora (CMLC-7) 2019, Cardiff, 22nd July 2019*, pages 9 – 16, Mannheim. Leibniz-Institut für Deutsche Sprache.
- Xiaoman Pan, Boliang Zhang, Jonathan May, Joel Nothman, Kevin Knight, and Heng Ji. 2017. [Cross-lingual name tagging and linking for 282 languages](#). In *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 1946–1958, Vancouver, Canada. Association for Computational Linguistics.
- Matthew E. Peters, Mark Neumann, Mohit Iyyer, Matt Gardner, Christopher Clark, Kenton Lee, and Luke Zettlemoyer. 2018. Deep contextualized word representations. In *Proc. of NAACL*.
- Jonas Pfeiffer, Ivan Vulić, Iryna Gurevych, and Sebastian Ruder. 2020. [MAD-X: An Adapter-Based Framework for Multi-Task Cross-Lingual Transfer](#). In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 7654–7673, Online. Association for Computational Linguistics.

- Telmo Pires, Eva Schlinger, and Dan Garrette. 2019. How multilingual is multilingual BERT? In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, pages 4996–5001.
- Peng Qi, Yuhao Zhang, Yuhui Zhang, Jason Bolton, and Christopher D. Manning. 2020. *Stanza: A python natural language processing toolkit for many human languages*. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics: System Demonstrations*, pages 101–108, Online. Association for Computational Linguistics.
- Colin Raffel, Noam Shazeer, Adam Roberts, Katherine Lee, Sharan Narang, Michael Matena, Yanqi Zhou, Wei Li, and Peter J Liu. 2019. Exploring the limits of transfer learning with a unified text-to-text transformer. arXiv:1910.10683.
- Afshin Rahimi, Yuan Li, and Trevor Cohn. 2019. *Masively multilingual transfer for NER*. In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, pages 151–164, Florence, Italy. Association for Computational Linguistics.
- Shruti Rijhwani, Jiateng Xie, Graham Neubig, and Jaime Carbonell. 2019. Zero-shot neural transfer for cross-lingual entity linking. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 33, pages 6924–6931.
- Anna Schmidt and Michael Wiegand. 2017. *A survey on hate speech detection using natural language processing*. In *Proceedings of the Fifth International Workshop on Natural Language Processing for Social Media*, pages 1–10, Valencia, Spain. Association for Computational Linguistics.
- Stefan Schweter. 2020. *BERTurk - BERT models for Turkish*.
- Djamé Seddah, Farah Essaidi, Amal Fethi, Matthieu Futral, Benjamin Muller, Pedro Javier Ortiz Suárez, Benoît Sagot, and Abhishek Srivastava. 2020. *Building a user-generated content North-African Arabizi treebank: Tackling hell*. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 1139–1150, Online. Association for Computational Linguistics.
- Steve Stecklow. 2018. Why Facebook is losing the war on hate speech in Myanmar, Reuters. <https://www.reuters.com/investigates/special-report/myanmar-facebook-hate/>.
- Milan Straka. 2018. UDPipe 2.0 prototype at CoNLL 2018 UD shared task. In *Proceedings of the CoNLL 2018 Shared Task: Multilingual Parsing from Raw Text to Universal Dependencies*, pages 197–207.
- Emma Strubell, Ananya Ganesh, and Andrew McCalum. 2019. *Energy and policy considerations for deep learning in NLP*. In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, pages 3645–3650, Florence, Italy. Association for Computational Linguistics.
- Wietse de Vries, Andreas van Cranenburgh, Arianna Bisazza, Tommaso Caselli, Gertjan van Noord, and Malvina Nissim. 2019. Bertje: A dutch BERT model. arXiv:1912.09582.
- Thomas Wolf, Lysandre Debut, Victor Sanh, Julien Chaumond, Clement Delangue, Anthony Moi, Pierric Cistac, Tim Rault, Remi Louf, Morgan Funtowicz, Joe Davison, Sam Shleifer, Patrick von Platen, Clara Ma, Yacine Jernite, Julien Plu, Canwen Xu, Teven Le Scao, Sylvain Gugger, Mariama Drame, Quentin Lhoest, and Alexander Rush. 2020. *Transformers: State-of-the-art natural language processing*. In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing: System Demonstrations*, pages 38–45, Online. Association for Computational Linguistics.
- Shijie Wu and Mark Dredze. 2020. *Are all languages created equal in multilingual BERT?* In *Proceedings of the 5th Workshop on Representation Learning for NLP*, pages 120–130, Online. Association for Computational Linguistics.
- Daniel Zeman and Jan Hajič, editors. 2018. *Proceedings of the CoNLL 2018 Shared Task: Multilingual Parsing from Raw Text to Universal Dependencies*. Association for Computational Linguistics, Brussels, Belgium.

A Languages

We list the 15 typologically diverse unseen languages we experiment with in Table 12 with information on their language family, script, origin and number of sentences available along with the categories we classified them in.

Data Sources We base our experiments on data originated from two sources. The Universal Dependency project (Nivre et al., 2016) downloadable here <https://lindat.mff.cuni.cz/repository/xmlui/handle/11234/1-2988> and the WikiNER dataset (Pan et al., 2017). We also use of the CoNLL-2003 shared task NER English dataset <https://www.clips.uantwerpen.be/conll2003/>

B Reproducibility

Infrastructure Our experiments were ran on a shared cluster on the equivalent of 15 Nvidia Tesla T4 GPUs.⁸

Optimization For all pretraining and fine-tuning runs, we use the Adam optimizer (Kingma and Ba, 2015). For fine-tuning, following Devlin et al. (2019), we only back-propagate through the first

⁸<https://www.nvidia.com/en-sg/data-center/tesla-t4/>

<i>Params.</i>	Parsing	NER	POS	Bounds
batch size	32	16	16	[1,256]
learning rate	5e-5	3.5e-5	5e-5	[1e-6,1e-3]
epochs (best)	15	6	6	[1,+ inf]
#grid	60	60	180	-
Run-time	32	24	75	-

Table 6: Fine-tuning best hyper-parameters for each task as selected on the validation set with bounds. #grid: number of grid search trial. Run-time is reported in average for training and evaluation. Run-time indicated in minutes.

Parameter	Value
batch size	64
learning rate	5e-5
optimizer	Adam
warmup	linear
warmup steps	10% total
epochs (best of)	10

Table 7: Unsupervised fine-tuning hyper-parameters

token of each word. We select the hyperparameters that minimize the loss on the validation set. The reported results are the average score of 5 runs with different random seeds computed on the test splits. We report the hyperparameters in Table 6-7.

C Easy Languages

We describe here in more details languages that we classify as Easy in section 4.1.

In practice, one can obtain very high performance even in zero-shot settings for such languages, by performing task-tuning on related languages.

Model	UPOS	LAS	NER
<i>Zero-Shot</i>			
(1) FaroeseBERT	66.4	35.8	-
(2) mBERT	79.4	67.5	-
(3) mBERT +MLM	83.4	67.8	-
<i>Few-Shot (CV with around 500 instances)</i>			
(4) Baseline	95.36	83.02	44.8
(5) FaroeseBERT	91.12	67.66	39.3
(6) mBERT	96.31	84.02	52.1
(7) mBERT +MLM	96.52	86.41	58.3

Table 8: Faroese is an “easy” unseen language: a multilingual model (+ language-specific MLM) easily outperforms all baselines. Zero-shot performance, after task-tuning only on related languages (Danish, Norwegian, Swedish) is also high.

Perhaps the best example of such an “easy” set-

ting is Faroese. mBERT has been trained on several languages of the north Germanic genus of the Indo-European language family, all of which use the Latin script. As a result, the multilingual mBERT model performs much better than the monolingual FaroeseBERT model that we trained on the available Faroese text (cf rows 1–2 and 5–6 in Table 8). Fine-tuning mBERT on the Faroese text is even more effective (rows 3 and 6 in Table 8), leading to further improvements, reaching more than 96.5% POS-tagging accuracy, 86% LAS for dependency parsing, and 58% NER F1 in the few-shot setting, surpassing the non-contextual baseline. In fact, even in zero-shot conditions, where we task-tune only on related languages (Danish, Norwegian, and Swedish), the model achieves remarkable performance of over 83% POS-tagging accuracy and 67.8% LAS dependency parsing.

Model	UPOS	LAS
<i>Zero-Shot</i>		
(1) SwissGermanBERT	64.7	30.0
(2) mBERT	62.7	41.2
(3) mBERT +MLM	87.9	69.6
<i>Few-Shot (CV with around 100 instances)</i>		
(4) Baseline	75.22	32.18
(5) SwissGermanBERT	65.42	30.0
(6) mBERT	76.66	41.2
(7) mBERT +MLM	78.68	69.6

Table 9: Swiss German is an “easy” unseen language: a multilingual model (+ language-specific MLM) outperforms all baselines in both zero-shot (task-tuning on the related High German) and few-shot settings.

Swiss German is another example of a language for which one can easily adapt a multilingual model and obtain good performance even in zero-shot settings. As in Faroese, simple MLM fine-tuning of the mBERT model with 200K sentences leads to an improvement of more than 25 points in both POS tagging and dependency parsing (Table 9) in zero-shot settings, with similar improvement trends in the few-shot setting.

The potential of similar-language pretraining along with script similarity is also showcased in the case of Naija (also known as Nigerian English or Nigerian Pidgin), an English creole spoken by millions in Nigeria. As Table 10 shows, with results after language- and task-tuning on 6K training examples, the multilingual approach surpasses the monolingual baseline.

On a side note, we can rely on [Han and Eisen-](#)

Model	UPOS	LAS
NaijaBERT	87.1	63.02
mBERT	89.3	71.6
mBERT +MLM	89.6	69.2

Table 10: Performance on Naija, an English creole, is very high, so we also classify it as an “easy” unseen language.

stein (2019) to also classify Early Modern English as an easy language. Similarly, the work of Chau et al. (2020) allows us to also classify Singlish (Singaporean English) as an easy language. In both cases, these languages are technically unseen by mBERT, but the fact that they are variants of English allows them to be easily handled by mBERT.

D Additional Uralic languages experiments

Following a similar procedure as in the Appendix C, we start with mBERT, perform task-tuning on Finnish and Estonian (both of which use the Latin script) and then do zero-shot experiments on Livvi, and Komi, all low-resource Uralic languages (results on the top part of Table 11). We also report results on the Finnish treebanks after task-tuning, for better comparison. The difference in performance on Livvi (which uses the Latin script) and the other languages that use the Cyrillic script is striking.

Although they are not easy enough to be tackled in a zero-shot setting, we show that the low-resource Uralic languages fall in the “Intermediate” category, since mBERT has been trained on similar languages: a small amount of annotated data are enough to improve over mBERT using task-tuning.

For both Livvi and Erzya, the multilingual model along with MLM-TUNING achieves the best performance, outperforming the non-contextual baseline by more than 1.5 point for parsing and POS tagging.

E Controlled experiment: Transliterating High-Resource Languages

To have a broader view on the effect of transliteration when using mBERT (section 5.2), we study the impact of transliteration to the Latin script on high resource languages seen during mBERT pre-training such as Arabic, Japanese and Russian. We compare the performance of mBERT fine-tuned

Language	UPOS	LAS
<i>Task-tuned – Latin script</i>		
Finnish (FTB)	93.1	77.5
Finnish (TDT)	95.0	78.9
Finnish (PUD)	96.8	83.5
<i>Zero-Shot Experiments</i>		
<i>Latin script</i>		
Livvi	72.3	40.3
<i>Cyrillic script</i>		
Erzya	51.5	18.6
<i>Few-Shot Experiments (CV)</i>		
<i>Livvi – Latin script</i>		
Baseline	84.1	40.1
mBERT	83.0	36.3
mBERT +MLM	85.5	42.3
<i>Erzya – Cyrillic script</i>		
Baseline	91.1	65.1
mBERT	89.3	61.2
mBERT +MLM	91.2	66.6

Table 11: The script matters for the efficacy of cross-lingual transfer. The zero-shot performance on Livvi, which is written in the same script as the task-tuning languages (Finnish, Estonian), is almost twice as good as the performance on the Uralic languages that use the Cyrillic script.

and evaluated on the original script with mBERT fine-tuned and evaluated on the transliterated text. As reported in Table 5, transliterating those languages to the Latin script leads to large drop in performance for all the three tasks.

Language (iso)	Script	Family	#sents	source	Category
Faroese (fao)	Latin	North Germanic	297K	(Biemann et al., 2007)	Easy
Mingrelian (xmf)	Georg.	Kartvelian	29K	Wikipedia	Easy
Naija (pcm)	Latin	English Pidgin	237K	(Caron et al., 2019)	Easy
Swiss German (gsw)	Latin	West Germanic	250K	OSCAR	Easy
Bambara (bm)	Latin	Niger-Congo	1K	OSCAR	Intermediate
Wolof (wo)	Latin	Niger-Congo	10K	OSCAR	Intermediate
Narabizi (nrz)	Latin	Semitic*	87K	(Seddah et al., 2020)	Intermediate
Maltese (mlt)	Latin	Semitic	50K	OSCAR	Intermediate
Buryat (bxu)	Cyrillic	Mongolic	7K	Wikipedia	Intermediate
Mari (mhr)	Cyrillic	Uralic	58K	Wikipedia	Intermediate
Erzya (myv)	Cyrillic	Uralic	20K	Wikipedia	Intermediate
Livvi (olo)	Latin	Uralic	9.4K	Wikipedia	Intermediate
Uyghur (ug)	Arabic	Turkic	105K	OSCAR	Hard
Sindhi (sd)	Arabic	Indo-Aryan	375K	OSCAR	Hard
Sorani (ckb)	Arabic	Indo-Iranian	380K	OSCAR	Hard

Table 12: Unseen Languages used for our experiments. #sents indicates the number of sentences used for training from scratch Monolingual Language Models as well as for MLM-TUNING mBERT

*code-mixed with French