

gaBERT — an Irish Language Model

Anonymous ACL submission

Abstract

The BERT family of neural language models have become highly popular due to their ability to provide sequences of text with rich context-sensitive token encodings which are able to generalise well to many NLP tasks. We introduce, gaBERT, a monolingual BERT model for the Irish language. We compare our gaBERT model to multilingual BERT and monolingual WikiBERT, and we show that gaBERT provides better representations for a downstream parsing task. We also show how different filtering criteria, vocabulary size and the choice of subword tokenisation model affect downstream performance. We release gaBERT and related code to the community.

1 Introduction

The technique of fine-tuning a self-supervised language model has become ubiquitous in Natural Language Processing (NLP) because models trained in this way have advanced evaluation scores on many tasks (Radford et al., 2018; Peters et al., 2018; Devlin et al., 2019). Arguably the most popular architecture is BERT (Devlin et al., 2019) which uses stacks of transformers to predict the identity of a masked token and to predict whether two sequences are contiguous. It has spawned many variants (Liu et al., 2019; Lan et al., 2019) and much analysis (Jawahar et al., 2019; Chi et al., 2020; Rogers et al., 2020). In this paper, we introduce gaBERT, a monolingual model of Irish.

Although Irish is the first official language of the Republic of Ireland, a minority, 1.5% of the

population (CSO, 2016), use it in their everyday lives outside of the education system. As the less dominant language in a bilingual community, the availability of Irish language technology is important since it makes it easier for Irish speakers and learners to use the language in their daily lives. Building on recent progress in data-driven Irish NLP (Lynn et al., 2012, 2015; Walsh et al., 2019), we release gaBERT with the hope that it will contribute to preserving Irish as a living language in the digital age.

While there is evidence to suggest that dedicated monolingual models can be superior to a multilingual model for within-language downstream tasks (de Vries et al., 2019; Virtanen et al., 2019; Farahani et al., 2020), other studies suggest that a multilingual model such as mBERT is a good choice for low-resourced languages (Wu and Dredze, 2020; Rust et al., 2020; Chau et al., 2020). We compare gaBERT to mBERT and to the monolingual Irish WikiBERT, both using Wikipedia as source of training data. We base our comparison on the downstream task of universal dependency (UD) parsing, since we have labelled Irish data in the form of the Irish UD Treebank (Lynn and Foster, 2016; McGuinness et al., 2020). We find that parsing accuracy improves when using gaBERT – by 3.7 and 3.6 LAS points over mBERT and WikiBERT respectively. Continued pretraining of mBERT using the gaBERT training data results in a recovery of 2 LAS points over the off-the-shelf version. The benefit of the gaBERT training data is also shown in a manual analysis which compares the models on their ability to predict a masked token.

We detail our hyperparameter search for our final

model, where we consider the type of text filtering to apply, the vocabulary size and tokenisation model. We release our experiment code through GitHub¹ and our models through the Hugging Face (Wolf et al., 2020) model repository.²

2 Data

We use the following to train gaBERT:

CoNLL17: The Irish data from the CoNLL’17 raw text collection (Ginter et al., 2017) released as part of the 2017 CoNLL Shared Task on UD Parsing (Zeman et al., 2017).

IMT: A collection of Irish texts used in Irish machine translation research (Dowling et al., 2018, 2020), including legal text, general administration and data crawled from public body websites.

NCI: The New Corpus for Ireland (Kilgarriff et al., 2006), which contains a wide range of texts in Irish, including fiction, news reports, informative texts and official documents.

OSCAR: The unshuffled Irish portion of the 2019 OSCAR corpus (Ortiz Suárez et al., 2019), a subset of CommonCrawl.

Paracrawl: The Irish side of the ga-en bitext pair of ParaCrawl v7 (Bañón et al., 2020), which is a collection of parallel corpora crawled from multi-lingual websites.

Wikipedia: Text from Irish Wikipedia, an online encyclopedia.³

The sentence counts in each corpus are listed in Table 1 after tokenisation and segmentation but before filtering described below. See Appendix A for more information on the content of these corpora, including license information. We apply corpus-specific pre-processing, sentence-segmentation and tokenisation described in Appendix B.

3 Experimental Setup

After initial corpus pre-processing, all corpora are merged and we use the WikiBERT pipeline (Pyysalo et al., 2020) to create pretraining data. We experiment with four corpus filtering settings, five vocabulary sizes and three tokenisation models.

3.1 Corpus Filtering

The WikiBERT pipeline contains a number of filters which dictate whether a document should be

Corpus	Num. Sents	Size (MB)
CoNLL17	1.7M	138
IMT	1.4M	124
NCI	1.6M	174
OSCAR	0.8M	89
ParaCrawl	3.1M	380
Wikipedia	0.7M	38
Overall	9.3M	943

Table 1: Sentence counts and plain text file size in megabytes for each corpus after tokenisation and segmentation but before applying sentence filtering.

kept. As we are working with data sources where there may not be clear document boundaries, or where there are no line breaks over a large number of sentences, document-level filtering may be inadequate for such texts. Consequently, we also experiment with using OpusFilter (Aulamo et al., 2020), which filters individual sentences, thereby giving us the flexibility of filtering noisy sentences while not discarding full documents.

For each filter setting below, we train a BERT model on the data which remains after filtering:

- **No-filter:** All collected texts are included in the pre-training data.
- **Document-filter:** The default document-level filtering used in the WikiBERT pipeline.
- **OpusFilter-basic:** We use OpusFilter with basic filtering described in Appendix B.4.
- **OpusFilter-basic-char-lang** We use OpusFilter with basic filtering as well as character-script and language filters described in Appendix B.4.

3.2 Vocabulary Creation

To create a model vocabulary, we experiment with the SentencePiece (Kudo and Richardson, 2018) and WordPiece tokenisers. Using the model with highest median LAS from the filtering experiments, we try vocabulary sizes of 15K, 20K, 30K, 40K and 50K. We then train a WordPiece tokeniser, keeping the vocabulary size that works best for the SentencePiece tokeniser. We also train a BERT model using the union of the two vocabularies.

3.3 BERT Pretraining Parameters

We use the original BERT implementation of Devlin et al. (2019). For the development experiments, we train our BERT model for 500K steps with a sequence length of 128. We use whole word masking

¹GitHub URL will appear here in the published paper.

²Hugging Face URL will appear here in the final paper.

³We use the articles from <https://dumps.wikimedia.org/gawiki/20210520/>

and the default hyperparameters and model architecture of BERT_{BASE} (Devlin et al., 2019).⁴

For the final gaBERT model, we train for 900k steps with sequence length 128 and a further 100k steps with sequence length 512. We train on a TPU-v2-8 with 128GB of memory on Google Compute Engine⁵ and use a batch size of 128.

4 Evaluation Measures

Dependency Parsing The evaluation measure we use to make development decisions is dependency parsing labelled attachment score (LAS). To obtain this measure, we fine-tune a given BERT model in the task of dependency parsing and measure LAS on the development set of the Irish-IDT treebank in version 2.8 of UD. We report the median of five fine-tuning runs with different random initialisation. For the dependency parser, we use a multitask model which uses a graph-based parser with biaffine attention (Dozat and Manning, 2016) as well as additional classifiers for predicting POS tags and morphological features. We use the AllenNLP (Gardner et al., 2018) library to develop our multitask model.

Cloze Test To compile a cloze task test set, 100 strings of Irish text (4–77 words each) containing the pronouns ‘é’ (‘him/it’), ‘í’ (‘her/it’) or ‘iad’ (‘them’) are selected from Irish corpora and online publications. One of these pronouns is masked in each string for the cloze test.⁶

Following Rönqvist et al. (2019), the models are evaluated on their ability to generate the original masked token, and a manual evaluation of the models is performed wherein predictions are classified into the following exclusive categories:

- **Match:** The predicted token fits the context grammatically and semantically. This may occur when the model predicts the original token or another token which also fits the context.
- **Mismatch** The predicted token is a valid Irish word but is unsuitable given the context.
- **Copy** The predicted token is an implausible repetition of another token in the context.
- **Gibberish** The predicted token is not a valid Irish word.

⁴We use a lower batch size of 32 in order to train on NVIDIA RTX 6000 GPUs with 24 GB RAM.

⁵TPU access was kindly provided to us through the Google Research TPU Research Cloud.

⁶All the masked tokens exist in the vocabularies of the candidate BERT models and are therefore possible predictions.

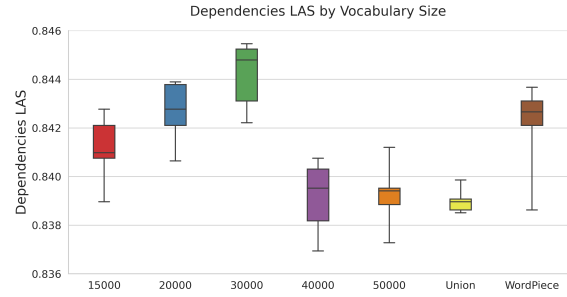


Figure 1: Dependency parsing LAS for each filter type and vocabulary type (five runs each).

Filter	Sentences
No-filter	9.3M
Document-filter	7.9M
OpusFilter-basic	9.0M
OpusFilter-basic-char-lang	7.7M

Table 2: The number of sentences which remain after applying the specific filter.

5 Results

5.1 Development Results

Filter Settings The overall number of sentences which remain after applying each filter are shown in Table 2. The results of training a dependency parser with the gaBERT model produced by each setting are shown in the top half of Fig. 2. Document-Filter has the highest LAS score. As the BERT model requires contiguous text for its next-sentence-prediction task, filtering out full documents may be more appropriate than filtering individual sentences. The two OpusFilter configurations perform marginally worse than the Document-Filter. In the case of OpusFilter-basic-char-lang, perhaps the lower number of sentences overall translates to lower LAS scores. Finally, No-Filter performs in the same range as the two OpusFilter configurations but has the lowest median score, suggesting that some level of filtering is beneficial.

Vocabulary Settings The results of the five runs testing different vocabulary sizes are shown in the bottom half of Fig. 1. A vocabulary size of 30K performs best for the SentencePiece tokeniser, which outperforms the WordPiece tokeniser with the same vocabulary size. The union of the two vocabularies results in 32,314 entries, and does not perform as well as the two vocabularies on their own.

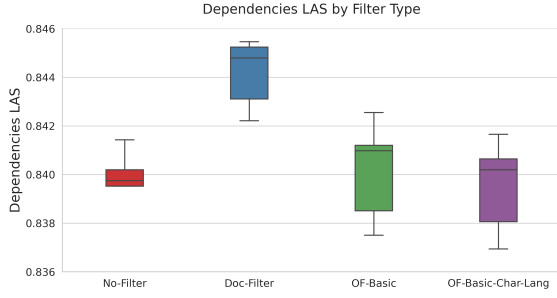


Figure 2: Dependency parsing LAS for each filter type.

Model	LAS		
	UD	Dev	Test
mBERT	2.8	81.8	80.3
WikiBERT	2.8	81.9	80.4
mBERT-cp	2.8	84.3	82.3
gaBERT	2.8	85.6	84.0
Chau et al. (2020)	2.5	-	76.2
gaBERT	2.5	-	77.5

Table 3: LAS in dependency parsing (UD v2.8) for selected models. Median of five fine-tuning runs. Scores are calculated using the official UD evaluation script (*conll18_ud_eval.py*).

5.2 Model Comparison

We compare our final gaBERT model with off-the-shelf mBERT and WikiBERT-ga, as well as an mBERT model obtained with continued pre-training on our corpora.

Dependency Parsing Table 3 shows the results for dependency parsing.⁷ Using mBERT off-the-shelf results in a test set LAS of 80.3. The WikiBERT-ga model performs slightly better than mBERT. By training mBERT for more steps on our corpora, LAS can be improved by 2 points. Our gaBERT model has the highest LAS of 84. The last two rows compare gaBERT, on v2.5 of the treebank, with the system of Chau et al. (2020) who augment the mBERT vocabulary with the 99 most frequent Irish tokens and fine-tune on Irish Wikipedia. Our model outperforms this approach.

Cloze Test Table 4 shows the accuracy of each model with regard to predicting the original masked token. mBERT-cp is the most accurate and gaBERT

⁷A competitive parser, UDPipe, trained on the Irish-IDT Treebank 2.8 without external embeddings achieves 72.59 LAS on the test set.

Model	Original Token Prediction
mBERT	16
WikiBERT	53
mBERT-cp	78
gaBERT	75

Table 4: The number of times the original masked token was predicted (100 test items).

Model	Match	Mism.	Copy	Gib
mBERT	41	42	4	13
WikiBERT	62	31	1	6
mBERT-cp	85	12	1	2
gaBERT	83	14	2	1

Table 5: The number of matches, mismatches, copies and gibberish predicted by each model (100 test items).

is close behind. Table 5 shows the manual evaluation of the tokens generated by each model, accounting for plausible answers deviating from the original token and separately reporting copying of content and production of gibberish. These results echo those of the original masked token prediction evaluation in so far as they rank the models in the same order. Further detail, examples and analysis of the cloze test can be found in Appendix C.

6 Friends of gaBERT

In subsequent experiments, we look at variants of BERT, including RoBERTa (Liu et al., 2019). The multilingual XLM-R_{BASE} (Conneau et al., 2020) clearly outperforms both variants of mBERT but underperforms gaBERT. We tried training a RoBERTa_{BASE} model but could only obtain LAS scores comparable to off-the-shelf mBERT and leave finding suitable hyperparameters to future work. We train an ELECTRA model (Clark et al., 2020), which performs better than both mBERT models and the WikiBERT model but slightly below gaBERT. See Appendix D-F for details.

7 Conclusions

We release gaBERT, a BERT model trained on over 7.9M Irish sentences, combining Irish language text from a variety of sources, and evaluate it in dependency parsing and in a pronoun cloze test task, showing improvements over three baselines, multilingual BERT, WikiBERT-ga and XML-R_{BASE}.

8 Ethical Considerations

No dataset is released with this paper, however most of the corpora are publicly available as described in Appendix A. Furthermore, where an anonymised version of a dataset was available it was used. We release the gaBERT language model based on the BERT_{BASE} (Devlin et al., 2019) autoencoder architecture. We note that an autoregressive architecture may be susceptible to training data extraction, and that larger language models may be more susceptible (Carlini et al., 2021). However, gaBERT is an autoencoder architecture and a smaller language model which may help mitigate this potential vulnerability.

Possible harms of language model pre-trained on web-crawled text have been widely discussed (Bender et al., 2021). Since gaBERT uses CommonCrawl data, there is a risk that the gaBERT model may, for example, produce unsuitable text outputs when used to generate text. To mitigate this possibility we include the following caveat with the released code and model cards:

We note that some data used to pre-train gaBERT was scraped from the web which potentially contains ethically problematic content (bias, hate, adult content, etc.). Consequently, downstream tasks/applications using gaBERT should be thoroughly tested with respect to ethical considerations.

We do not discuss in detail how gaBERT can be used in actual use cases as we expect the use of BERT-style models to be essential knowledge for NLP practitioners up-to-date with current research. There are many downstream tasks which can use gaBERT, including machine translation, educational applications, predictive text, search and games. The authors hope gaBERT will contribute to the ongoing effort to preserve the Irish language as a living language in the technological age. Supporting a low-resourced language like Irish in a bilingual community will make it easier for Irish speakers, and those who wish to be Irish speakers, to use the language in practice.

Each use case or downstream application may rank the available pre-trained language models differently in terms of suitability. We urge NLP practitioners to compare available models such as those tested in this paper in their application rather than relying on results for a different task.

References

- Mikko Aulamo, Sami Virpioja, and Jörg Tiedemann. 2020. [OpusFilter: A configurable parallel corpus filtering toolbox](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics: System Demonstrations*, pages 150–156, Online. Association for Computational Linguistics.
- Marta Bañón, Pinzhen Chen, Barry Haddow, Kenneth Heafield, Hieu Hoang, Miquel Esplà-Gomis, Mikel L. Forcada, Amir Kamran, Faheem Kirefu, Philipp Koehn, Sergio Ortiz Rojas, Leopoldo Pla Sempere, Gema Ramírez-Sánchez, Elsa Sarrias, Marek Strelec, Brian Thompson, William Waites, Dion Wiggins, and Jaume Zaragoza. 2020. [ParaCrawl: Web-scale acquisition of parallel corpora](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 4555–4567, Online. Association for Computational Linguistics.
- Emily M. Bender, Timnit Gebru, Angelina McMillan-Major, and Shmargaret Shmitchell. 2021. [On the dangers of stochastic parrots: Can language models be too big?](#) In *FAccT '21: Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, pages 610–623.
- Nicholas Carlini, Florian Tramer, Eric Wallace, Matthew Jagielski, Ariel Herbert-Voss, Katherine Lee, Adam Roberts, Tom Brown, Dawn Song, Ulfar Erlingsson, Alina Oprea, and Colin Raffel. 2021. [Extracting training data from large language models](#).
- 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.
- Ethan A. Chi, John Hewitt, and Christopher D. Manning. 2020. [Finding universal grammatical relations in multilingual BERT](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 5564–5577, Online. Association for Computational Linguistics.
- Kevin Clark, Minh-Thang Luong, Quoc V. Le, and Christopher D. Manning. 2020. [ELECTRA: Pre-training text encoders as discriminators rather than generators](#). In *Proceedings of The Eighth International Conference on Learning Representations (ICLR)*.
- 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.

379	CSO. 2016. Census of Population 2016 – Profile 10	434
380	Education, Skills and the Irish Language . Publisher:	435
381	Central Statistics Office.	436
382	Wietse de Vries, Andreas van Cranenburgh, Arianna	437
383	Bisazza, Tommaso Caselli, Gertjan van Noord, and	438
384	Malvina Nissim. 2019. BERTje: A Dutch BERT	439
385	model . ArXiv 1912.09582v1.	440
386	Jacob Devlin, Ming-Wei Chang, Kenton Lee, and	441
387	Kristina Toutanova. 2019. BERT: Pre-training of	442
388	deep bidirectional transformers for language under-	443
389	standing . In <i>Proceedings of the 2019 Conference</i>	444
390	<i>of the North American Chapter of the Association</i>	445
391	<i>for Computational Linguistics: Human Language</i>	446
392	<i>Technologies, Volume 1 (Long and Short Papers)</i> ,	447
393	pages 4171–4186, Minneapolis, Minnesota. Associ-	448
394	ation for Computational Linguistics.	
395	Meghan Dowling, Sheila Castilho, Joss Moorkens,	449
396	Teresa Lynn, and Andy Way. 2020. A human evalua-	450
397	tion of English-Irish statistical and neural machine	451
398	translation . In <i>Proceedings of the 22nd Annual Con-</i>	452
399	<i>ference of the European Association for Machine</i>	453
400	<i>Translation</i> , pages 431–440, Lisboa, Portugal. Euro-	
401	pean Association for Machine Translation.	
402	Meghan Dowling, Teresa Lynn, Alberto Poncelas, and	454
403	Andy Way. 2018. SMT versus NMT: Preliminary	455
404	comparisons for Irish . In <i>Proceedings of the AMTA</i>	456
405	<i>2018 Workshop on Technologies for MT of Low Re-</i>	457
406	<i>source Languages (LoResMT 2018)</i> , pages 12–20,	458
407	Boston, MA. Association for Machine Translation	
408	in the Americas.	
409	Timothy Dozat and Christopher D. Manning. 2016.	459
410	Deep biaffine attention for neural dependency pars-	460
411	ing . <i>CoRR</i> , abs/1611.01734.	461
412	Mehrdad Farahani, Mohammad Gharachorloo,	462
413	Marzieh Farahani, and Mohammad Manthouri.	463
414	2020. ParsBERT: Transformer-based model for Per-	464
415	sian language understanding . ArXiv 2005.12515v1.	465
416	Matt Gardner, Joel Grus, Mark Neumann, Oyvind	466
417	Tafjord, Pradeep Dasigi, Nelson F. Liu, Matthew Pe-	467
418	ters, Michael Schmitz, and Luke Zettlemoyer. 2018.	468
419	AllenNLP: A deep semantic natural language pro-	469
420	cessing platform . In <i>Proceedings of Workshop for</i>	470
421	<i>NLP Open Source Software (NLP-OSS)</i> , pages 1–	
422	6, Melbourne, Australia. Association for Computa-	
423	tional Linguistics.	
424	Filip Ginter, Jan Hajič, Juhani Luotolahti, Milan Straka,	471
425	and Daniel Zeman. 2017. CoNLL 2017 shared task	472
426	- automatically annotated raw texts and word embed-	473
427	dings . LINDAT/CLARIAH-CZ digital library at the	474
428	Institute of Formal and Applied Linguistics (ÚFAL),	475
429	Faculty of Mathematics and Physics, Charles Uni-	476
430	versity.	
431	Ganesh Jawahar, Benoît Sagot, and Djamé Seddah.	477
432	2019. What does BERT learn about the structure	478
433	of language? In <i>Proceedings of the 57th Annual</i>	479
	<i>Meeting of the Association for Computational Lin-</i>	480
	<i>guistics</i> , pages 3651–3657, Florence, Italy. Associa-	481
	tion for Computational Linguistics.	482
	Adam Kilgarriff, Michael Rundell, and Elaine	483
	Uí Dhonnchadha. 2006. Efficient corpus develop-	484
	ment for lexicography: building the New Corpus	485
	for Ireland . <i>Language Resources and Evaluation</i> ,	486
	40:127–152.	487
	Taku Kudo and John Richardson. 2018. SentencePiece:	488
	A simple and language independent subword tok-	489
	enizer and detokenizer for neural text processing . In	
	<i>Proceedings of the 2018 Conference on Empirical</i>	
	<i>Methods in Natural Language Processing: System</i>	
	<i>Demonstrations</i> , pages 66–71, Brussels, Belgium.	
	Association for Computational Linguistics.	
	Zhenzhong Lan, Mingda Chen, Sebastian Goodman,	
	Kevin Gimpel, Piyush Sharma, and Radu Sori-	
	cut. 2019. ALBERT: A lite BERT for self-	
	supervised learning of language representations .	
	ArXiv 1909.11942v6.	
	Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Man-	
	dar Joshi, Danqi Chen, Omer Levy, Mike Lewis,	
	Luke Zettlemoyer, and Veselin Stoyanov. 2019.	
	Roberta: A robustly optimized BERT pretraining ap-	
	proach . ArXiv 1907.11692v1.	
	Teresa Lynn, Ozlem Cetinoglu, Jennifer Foster,	
	Elaine Uí Dhonnchadha, Mark Dras, and Josef van	
	Genabith. 2012. Irish treebanking and parsing: A	
	preliminary evaluation. In <i>Proceedings of the Eight</i>	
	<i>International Conference on Language Resources</i>	
	<i>and Evaluation (LREC-2012)</i> , pages 1939–1946, Is-	
	tanbul, Turkey. European Language Resources Asso-	
	ciation (ELRA).	
	Teresa Lynn and Jennifer Foster. 2016. Universal De-	
	pendencies for Irish. In <i>Proceedings of the Second</i>	
	<i>Celtic Language Technology Workshop</i> , July, pages	
	79–92, Paris.	
	Teresa Lynn, Kevin Scannell, and Eimear Maguire.	
	2015. Minority Language Twitter: Part-of-Speech	
	Tagging and Analysis of Irish Tweets . In <i>Proceed-</i>	
	<i>ings of the Workshop on Noisy User-generated Text</i> ,	
	pages 1–8, Beijing, China. Association for Computa-	
	tional Linguistics.	
	Sarah McGuinness, Jason Phelan, Abigail Walsh, and	
	Teresa Lynn. 2020. Annotating MWEs in the	
	Irish UD treebank . In <i>Proceedings of the Fourth</i>	
	<i>Workshop on Universal Dependencies (UDW 2020)</i> ,	
	pages 126–139, Barcelona, Spain (Online). Associa-	
	tion for Computational Linguistics.	
	Pedro Javier Ortiz Suárez, Benoît Sagot, and Laurent	
	Romary. 2019. Asynchronous pipeline for process-	
	ing huge corpora on medium to low resource infras-	
	tructures . In <i>7th Workshop on the Challenges in the</i>	
	<i>Management of Large Corpora (CMLC-7)</i> , pages 9	
	– 16, Cardiff, United Kingdom. Leibniz-Institut für	
	Deutsche Sprache.	

490	Matthew Peters, Mark Neumann, Mohit Iyyer, Matt	<i>System Demonstrations</i> , pages 38–45, Online. Association for Computational Linguistics.	547
491	Gardner, Christopher Clark, Kenton Lee, and Luke		548
492	Zettlemoyer. 2018. Deep contextualized word representations .		
493	In <i>Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long Papers)</i> , pages	Shijie Wu and Mark Dredze. 2020. Are all languages created equal in multilingual BERT? In <i>Proceedings of the 5th Workshop on Representation Learning for NLP</i> , pages 120–130, Online. Association for Computational Linguistics.	549
494	2227–2237, New Orleans, Louisiana, USA. Association for Computational Linguistics.		550
495			551
496			552
497			553
498			
499	Sampo Pyysalo, Jenna Kanerva, Antti Virtanen, and	Daniel Zeman, Joakim Nivre, Mitchell Abrams,	554
500	Filip Ginter. 2020. WikiBERT models: deep transfer learning for many languages . ArXiv	Elia Ackermann, Noëmi Aepli, Hamid Aghaei,	555
501	2006.01538v1.	Željko Agić, Amir Ahmadi, Lars Ahrenberg,	556
502		Chika Kennedy Ajede, Gabrielë Aleksandravičiūtė,	557
503		Ika Alfina, Lene Antonsen, Katya Aplonova, An-	558
504	Alec Radford, Karthik Narasimhan, Tim Salimans,	gelina Aquino, Carolina Aragon, Maria Jesus Aran-	559
505	and Ilya Sutskever. 2018. Improving language understanding by generative pre-training . OpenAI	zabe, Hórunn Arnardóttir, Gashaw Arutie, Jes-	560
506	Preprint.	sica Naraiswari Arwidarasti, Masayuki Asahara,	561
507		Luma Ateyah, Furkan Atmaca, Mohammed Attia,	562
508	Anna Rogers, Olga Kovaleva, and Anna Rumshisky.	Aitziber Atutxa, Liesbeth Augustinus, Elena Bad-	563
509	2020. A primer in BERTology: What we know about how BERT works . <i>Transactions of the Association for Computational Linguistics</i> , 8:842–866.	maeva, Keerthana Balasubramani, Miguel Balle-	564
510		steros, Esha Banerjee, Sebastian Bank, Verginica	565
511		Barbu Mititelu, Victoria Basmov, Colin Batche-	566
512	Samuel Rönnqvist, Jenna Kanerva, Tapio Salakoski,	lor, John Bauer, Seyyit Talha Bedir, Kepa Ben-	567
513	and Filip Ginter. 2019. Is multilingual BERT fluent in language generation? In <i>Proceedings of the First NLP Workshop on Deep Learning for Natural Language Processing</i> , pages 29–36, Turku, Finland. Linköping University Electronic Press.	goetxea, Gözde Berk, Yevgeni Berzak, Irshad Ah-	568
514		mad Bhat, Riyaz Ahmad Bhat, Erica Biagetti, Eck-	569
515		hard Bick, Agnė Bielinskienė, Kristín Bjarnadóttir,	570
516		Rogier Blokland, Victoria Bobicev, Loïc Boizou,	571
517	Phillip Rust, Jonas Pfeiffer, Ivan Vulić, Sebastian	Emanuel Borges Völker, Carl Börstell, Cristina	572
518	Ruder, and Iryna Gurevych. 2020. How good is your tokenizer? on the monolingual performance of multilingual language models . ArXiv 2012.15613v2.	Bosco, Gosse Bouma, Sam Bowman, Adriane Boyd,	573
519		Kristina Brokaitė, Aljoscha Burchardt, Marie Can-	574
520		dito, Bernard Caron, Gauthier Caron, Tatiana Cav-	575
521	Milan Straka and Jana Straková. 2017. Tokenizing, POS tagging, lemmatizing and parsing UD 2.0 with UDPipe . In <i>Proceedings of the CoNLL 2017 Shared Task: Multilingual Parsing from Raw Text to Universal Dependencies</i> , pages 88–99, Vancouver, Canada. Association for Computational Linguistics.	alcanti, Gülşen Cebiroğlu Eryiğit, Flavio Massimil-	576
522		iano Cecchini, Giuseppe G. A. Celano, Slavomír Čé-	577
523		plö, Savas Cetin, Özlem Çetinoğlu, Fabricio Chalub,	578
524		Ethan Chi, Yongseok Cho, Jinho Choi, Jayeol	579
525		Chun, Alessandra T. Cignarella, Silvie Cinková, Au-	580
526		rémie Collomb, Çağrı Çöltekin, Miriam Connor, Ma-	581
527	Antti Virtanen, Jenna Kanerva, Rami Ilo, Jouni Luoma,	rine Courtin, Elizabeth Davidson, Marie-Catherine	582
528	Juhani Luotolahti, Tapio Salakoski, Filip Ginter, and	de Marneffe, Valeria de Paiva, Mehmet Oguz	583
529	Sampo Pyysalo. 2019. Multilingual is not enough: BERT for Finnish . ArXiv 1912.07076v1.	Derin, Elvis de Souza, Arantza Diaz de Ilar-	584
530		raza, Carly Dickerson, Arawinda Dinakaramani,	585
531	Abigail Walsh, Teresa Lynn, and Jennifer Foster. 2019.	Bamba Dione, Peter Dirix, Kaja Dobrovoljc, Tim-	586
532	Ilfhocail: A lexicon of Irish MWEs . In <i>Proceedings of the Joint Workshop on Multiword Expressions and WordNet (MWE-WN 2019)</i> , pages 162–168, Florence, Italy. Association for Computational Linguistics.	othy Dozat, Kira Drokanova, Puneet Dwivedi,	587
533		Hanne Eckhoff, Marhaba Eli, Ali Elkahky, Binyam	588
534		Ephrem, Olga Erina, Tomaž Erjavec, Aline Eti-	589
535		enne, Wograinne Evelyn, Sidney Facundes, Richárd	590
536		Farkas, Marília Fernanda, Hector Fernandez Al-	591
537	Thomas Wolf, Lysandre Debut, Victor Sanh, Julien	calde, Jennifer Foster, Cláudia Freitas, Kazunori	592
538	Chaumond, Clement Delangue, Anthony Moi, Pier-	Fujita, Katarína Gajdošová, Daniel Galbraith, Mar-	593
539	ric Cistac, Tim Rault, Remi Louf, Morgan Funtow-	cos Garcia, Moa Gärdenfors, Sebastian Garza,	594
540	icz, Joe Davison, Sam Shleifer, Patrick von Platen,	Fabício Ferraz Gerardi, Kim Gerdes, Filip Ginter,	595
541	Clara Ma, Yacine Jernite, Julien Plu, Canwen Xu,	Iakes Goenaga, Koldo Gojenola, Memduh	596
542	Teven Le Scao, Sylvain Gugger, Mariama Drame,	Gökırmak, Yoav Goldberg, Xavier Gómez Guino-	597
543	Quentin Lhoest, and Alexander Rush. 2020. Transformers: State-of-the-art natural language processing . In <i>Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing</i> :	vart, Berta González Saavedra, Bernadeta Griciūtė,	598
544		Matias Grioni, Loïc Grobol, Normunds Grūzītis,	599
545		Bruno Guillaume, Céline Guillot-Barbance, Tunga	600
546		Güngör, Nizar Habash, Hinrik Hafsteinsson, Jan	601
		Hajič, Jan Hajič jr., Mika Hämäläinen, Linh	602
		Hà Mỹ, Na-Rae Han, Muhammad Yudistira Han-	603
		ifmuti, Sam Hardwick, Kim Harris, Dag Haug,	604
		Johannes Heinecke, Oliver Hellwig, Felix Hen-	605
		nig, Barbora Hladká, Jaroslava Hlaváčová, Florinel	606
		Hociung, Petter Hohle, Eva Huber, Jena Hwang,	607
		Takumi Ikeda, Anton Karl Ingason, Radu Ion,	608

609	Elena Irimia, Olájidé Ishola, Tomáš Jelínek, Anders	672
610	Johannsen, Hildur Jónsdóttir, Fredrik Jørgensen,	673
611	Markus Juutinen, Sarveswaran K, Hüner Kaşıkara,	674
612	Andre Kaasen, Nadezhda Kabaeva, Sylvain Ka-	675
613	hane, Hiroshi Kanayama, Jenna Kanerva, Boris	676
614	Katz, Tolga Kayadelen, Jessica Kenney, Václava	677
615	Kettnerová, Jesse Kirchner, Elena Klementieva,	678
616	Arne Köhn, Abdullatif Köksal, Kamil Kopacewicz,	679
617	Timo Korkiakangas, Natalia Kotsyba, Jolanta Ko-	680
618	valevskaitė, Simon Krek, Parameswari Krishna-	681
619	murthy, Sookyoung Kwak, Veronika Laippala, Lu-	682
620	cia Lam, Lorenzo Lambertino, Tatiana Lando,	683
621	Septina Dian Larasati, Alexei Lavrentiev, John Lee,	684
622	Phùng Lê Hồng, Alessandro Lenci, Saran Lert-	685
623	pradit, Herman Leung, Maria Levina, Cheuk Ying	686
624	Li, Josie Li, Keying Li, Yuan Li, KyungTae Lim,	687
625	Krister Lindén, Nikola Ljubešić, Olga Loginova,	688
626	Andry Luthfi, Mikko Luukko, Olga Lyashevskaya,	689
627	Teresa Lynn, Vivien Macketanz, Aibek Makazhanov,	690
628	Michael Mandl, Christopher Manning, Ruli Manu-	691
629	runung, Cătălina Măranduc, David Mareček, Katrin	692
630	Marheinecke, Héctor Martínez Alonso, André Mar-	693
631	tins, Jan Mašek, Hiroshi Matsuda, Yuji Matsumoto,	694
632	Ryan McDonald, Sarah McGuinness, Gustavo Men-	695
633	donça, Niko Miekka, Karina Mischenkova, Mar-	696
634	garita Misirpashayeva, Anna Missilä, Cătălin Mi-	697
635	titelu, Maria Mitrofan, Yusuke Miyao, AmirHossein	698
636	Mojiri Foroushani, Amirsaeid Moloodi, Simonetta	699
637	Montemagni, Amir More, Laura Moreno Romero,	700
638	Keiko Sophie Mori, Shinsuke Mori, Tomohiko	701
639	Morioka, Shigeki Moro, Bjartur Mortensen, Bohdan	702
640	Moskalevskiy, Kadri Muischnek, Robert Munro,	703
641	Yugo Murawaki, Kaili Müürisep, Pinkey Nain-	704
642	wani, Mariam Nakhilė, Juan Ignacio Navarro Horñi-	705
643	acek, Anna Nedoluzhko, Gunta Nešpore-Bērzkalne,	706
644	Lùng Nguyễn Thị, Huyền Nguyễn Thị Minh,	
645	Yoshihiro Nikaido, Vitaly Nikolaev, Rattima Nitis-	
646	aroj, Alireza Nourian, Hanna Nurmi, Stina Ojala,	
647	Atul Kr. Ojha, Adédayò Olùòkun, Mai Omura,	
648	Emeka Onwuegbuzia, Petya Osenova, Robert	
649	Östling, Lilja Övrelid, Şaziye Betül Özates, Arzucan	
650	Özgür, Balkız Öztürk Başaran, Niko Partanen, Elena	
651	Pascual, Marco Passarotti, Agnieszka Patejuk, Guil-	
652	herme Paulino-Passos, Angelika Peljak-Łapińska,	
653	Siyao Peng, Cenel-Augusto Perez, Natalia Perkova,	
654	Guy Perrier, Slav Petrov, Daria Petrova, Jason Phe-	
655	lan, Jussi Piitulainen, Tommi A Pirinen, Emily Pitler,	
656	Barbara Plank, Thierry Poibeau, Larisa Ponomareva,	
657	Martin Popel, Lauma Pretkalniņa, Sophie Prévost,	
658	Prokopis Prokopidis, Adam Przepiórkowski, Tiina	
659	Puolakainen, Sampo Pyysalo, Peng Qi, Andriela	
660	Rääbis, Alexandre Rademaker, Taraka Rama, Lo-	
661	ganathan Ramasamy, Carlos Ramisch, Fam Rashel,	
662	Mohammad Sadegh Rasooli, Vinit Ravishankar,	
663	Livy Real, Petru Rebeja, Siva Reddy, Georg Rehm,	
664	Ivan Riabov, Michael Rießler, Erika Rimkutė,	
665	Larissa Rinaldi, Laura Rituma, Luisa Rocha, Eiríkur	
666	Rögnvaldsson, Mykhailo Romanenko, Rudolf Rosa,	
667	Valentin Roşca, Davide Rovati, Olga Rudina, Jack	
668	Rueter, Kristján Rúnarsson, Shoval Sadde, Pegah	
669	Safari, Benoît Sagot, Aleksí Sahala, Shadi Saleh,	
670	Alessio Salomoni, Tanja Samardžić, Stephanie Sam-	
671	son, Manuela Sanguinetti, Dage Särög, Baiba Saulīte,	
	Yanin Sawanakunanon, Kevin Scannell, Salvatore	707
	Scarlata, Nathan Schneider, Sebastian Schuster,	708
	Djamé Seddah, Wolfgang Seeker, Mojgan Seraji,	709
	Mo Shen, Atsuko Shimada, Hiroyuki Shirasu, Muh	710
	Shohibussirri, Dmitry Sichinava, Einar Freyr Sig-	711
	urðsson, Aline Silveira, Natalia Silveira, Maria Simi,	712
	Radu Simionescu, Katalin Simkó, Mária Šimková,	713
	Kiril Simov, Maria Skachodubova, Aaron Smith, Is-	714
	abela Soares-Bastos, Carolyn Spadine, Steinhór Ste-	715
	ingrímsson, Antonio Stella, Milan Straka, Emmett	716
	Strickland, Jana Strnadová, Alane Suhr, Yogi Les-	717
	mana Sulestio, Umut Sulubacak, Shingo Suzuki,	718
	Zsolt Szántó, Dima Taji, Yuta Takahashi, Fabio Tam-	719
	burini, Mary Ann C. Tan, Takaaki Tanaka, Sam-	720
	son Tella, Isabelle Tellier, Guillaume Thomas, Li-	721
	isi Torga, Marsida Toska, Trond Trosterud, Anna	722
	Trukhina, Reut Tsarfaty, Utku Türk, Francis Ty-	723
	ers, Sumire Uematsu, Roman Untilov, Zdeňka Ure-	724
	šová, Larraitz Uria, Hans Uszkoreit, Andrius Utkā,	725
	Sowmya Vajjala, Daniel van Niekerk, Gertjan van	726
	Noord, Viktor Varga, Eric Villemonte de la Clerg-	727
	erie, Veronika Vincze, Aya Wakasa, Joel C. Wallen-	728
	berg, Lars Wallin, Abigail Walsh, Jing Xian Wang,	729
	Jonathan North Washington, Maximilian Wendt,	730
	Paul Widmer, Seyi Williams, Mats Wirén, Chris-	731
	tian Wittern, Tsegay Woldemariam, Tak-sum Wong,	732
	Alina Wróblewska, Mary Yako, Kayo Yamashita,	733
	Naoki Yamazaki, Chunxiao Yan, Koichi Yasuoka,	
	Marat M. Yavrumyan, Zhuoran Yu, Zdeněk Žabokrt-	
	ský, Shorouq Zahra, Amir Zeldes, Hanzhi Zhu, and	
	Anna Zhuravleva. 2020. Universal dependencies 2.7 .	
	LINDAT/CLARIAH-CZ digital library at the Insti-	
	tute of Formal and Applied Linguistics (ÚFAL), Fac-	
	ulty of Mathematics and Physics, Charles Univer-	
	sity.	
	Daniel Zeman, Martin Popel, Milan Straka, Jan Ha-	707
	jič, Joakim Nivre, Filip Ginter, Juhani Luotolahti,	708
	Sampo Pyysalo, Slav Petrov, Martin Potthast, Fran-	709
	cis Tyers, Elena Badmaeva, Memduh Gokirmak,	710
	Anna Nedoluzhko, Silvie Cinková, Jan Hajič jr.,	711
	Jaroslava Hlaváčová, Václava Kettnerová, Zdeňka	712
	Urešová, Jenna Kanerva, Stina Ojala, Anna Mis-	713
	silä, Christopher D. Manning, Sebastian Schuster,	714
	Siva Reddy, Dima Taji, Nizar Habash, Herman Le-	715
	ung, Marie-Catherine de Marneffe, Manuela San-	716
	guinetti, Maria Simi, Hiroshi Kanayama, Valeria	717
	de Paiva, Kira Droganova, Héctor Martínez Alonso,	718
	Çağrı Çöltekin, Umut Sulubacak, Hans Uszkoreit,	719
	Vivien Macketanz, Aljoscha Burchardt, Kim Harris,	720
	Katrin Marheinecke, Georg Rehm, Tolga Kayadelen,	721
	Mohammed Attia, Ali Elkahky, Zhuoran Yu, Emily	722
	Pitler, Saran Lertpradit, Michael Mandl, Jesse Kirchner,	723
	Hector Fernandez Alcalde, Jana Strnadová,	724
	Esha Banerjee, Ruli Manurung, Antonio Stella, At-	725
	suko Shimada, Sookyoung Kwak, Gustavo Men-	726
	donça, Tatiana Lando, Rattima Nitisaroj, and Josie	727
	Li. 2017. CoNLL 2017 shared task: Multilingual	728
	parsing from raw text to Universal Dependencies . In	729
	<i>Proceedings of the CoNLL 2017 Shared Task: Mul-</i>	730
	<i>tilingual Parsing from Raw Text to Universal Depen-</i>	731
	<i>dependencies</i> , pages 1–19, Vancouver, Canada. Associa-	732
	tion for Computational Linguistics.	733

A Data Licenses

This Appendix provides specific details of the licence for each of the datasets used in the experiments.

A.1 CoNLL17

The Irish annotated CoNLL17 corpus can be found here: <http://hdl.handle.net/11234/1-1989> (Ginter et al., 2017).

The automatically generated annotations on the raw text data are available under the CC BY-SA-NC 4.0 licence. Wikipedia texts are available under the CC BY-SA 3.0 licence. Texts from Common Crawl are subject to Common Crawl Terms of Use, the full details of which can be found here: <https://commoncrawl.org/terms-of-use/full/>.

A.2 IMT

The Irish Machine Translation datasets contains text from the following sources:

- Text crawled from the Citizen’s Information website, contains Irish Public Sector Data licensed under a Creative Commons Attribution 4.0 International (CC BY 4.0) licence: <https://www.citizensinformation.ie/ga/>. 778 779
- Text crawled from Comhairle na Gaelscolaíochta website: <https://www.comhairle.org/gaeilge/>. 780 781
- Text crawled from the FÁS website (<http://www.fas.ie/>), accessed in 2017. The website has since been dissolved. 782 783
- Text crawled from the Galway County Council website: <http://www.galway.ie/ga/>. 784 785
- Text crawled from <https://www.gov.ie/ga/>, the central portal for government services and information. 786 787
- Text crawled from articles on the Irish Times website. 788 789
- Text crawled from the Kerry County Council website: <https://ciarrai.ie/>. 790 791
- Text crawled from the Oideas Gael website: <http://www.oideas-gael.com/ga/>. 792 793
- Text crawled from articles generated by Teagasc, available under PSI licence. 794 795
- Text generated by Conradh na Gaeilge, shared with us for research purposes. 796 797
- The Irish text from a parallel English–Irish corpus of legal texts from the Department of Justice. This dataset is available for reuse on the ELRC-SHARE repository under a PSI license: <https://elrc-share.eu> 798 799
- Text from the Directorate-General for Translation (DGT), available for download from the European Commission website. Reuse of the texts are subject to Terms of Use, found on the website: <https://ec.europa.eu/jrc/en/language-technologies/dgt-translation-memory>. 800 801 802 803
- Text reports and notices generated by Dublin City Council, shared with us for research purposes. 804 805
- Text uploaded to ELRC-share via the National Relay Station, shared with us for research purposes. 806 807
- Text reports and reference files generated by the Language Commissioner, available on ELRC-share under PSI license: <https://elrc-share.eu/>. 808 809
- Text generated by the magazine Nós, shared with us for research purposes. 810 811
- Irish texts available for download on OPUS, under various licenses: <https://opus.nlpl.eu/> 812 813 814
- Text generated from in-house translation provided by the then titled Department of Culture, Heritage and Gaeltacht (DCHG), provided for research purposes. The anonymised dataset is available on ELRC-share, under a CC-BY 4.0 license: <https://elrc-share.eu/>. 815 816 817
- Text reports created by Údarás na Gaeilge, uploaded to ELRC-share available under PSI license: <https://elrc-share.eu/>. 818 819
- Text generated by the University Times, shared with us for research purposes. 820 821 822

A.3 NCI

The corpus is compiled and owned by Foras na Gaeilge and is provided to us for research purposes.

A.4 OSCAR

The unshuffled version of the Irish part of the OSCAR corpus was provided to us by the authors for research purposes.

A.5 ParaCrawl

Text from ParaCrawl v7, available here: <https://www.paracrawl.eu/v7>. The texts themselves are not owned by ParaCrawl, the actual packaging of these parallel data are under the Creative Commons CC0 licence ("no rights reserved").

A.6 Wikipedia

The texts used are available under a CC BY-SA 3.0 licence and/or a GNU Free Documentation License.

B Corpus Pre-processing

This appendix provides specific details on corpus pre-processing, and the OpusFilter filters used.

CoNLL17 The CoNLL17 corpus is already tokenised, as it is provided in CoNLL-U format, which we convert to one-sentence-per-line tokenised plain text.

IMT, OSCAR and ParaCrawl The text files from the IMT, OSCAR and ParaCrawl contain raw sentences requiring tokenisation. We describe the tokenisation process for these corpora in Appendix B.1.

Wikipedia For the Wikipedia articles, the Irish Wikipedia dump is downloaded and the WikiExtractor tool⁸ is then used to extract plain text. Article headers are included in the extracted text files. Once the articles have been converted to plain text, they are tokenised using the tokeniser described in Appendix B.1.

NCI As many of the NCI segments marked up with $\langle s \rangle$ tags contain multiple sentences, we further split these segments with heuristics described in Appendix B.3.

B.1 Tokenisation and Segmentation

Raw texts from the IMT, OSCAR, ParaCrawl and Wikipedia corpora are tokenised and segmented with UDPipe (Straka and Straková, 2017) trained on a combination of the Irish-IDT and English-EWT corpora from version 2.7 of the Universal

Dependencies (UD) treebanks (Zeman et al., 2020). We include the English-EWT treebank in the training data to expose the tokeniser to more incidences of punctuation symbols which are prevalent in our pre-training data. This also comes with the benefit of supporting the tokenisation of code-mixed data. We upsample the Irish-IDT treebank by ten times to offset the larger English-EWT treebank size. This tokeniser is applied to all corpora apart from the NCI, which is already tokenised by Kilgarriff et al. (2006), and the CoNLL17 corpus as this corpus is already tokenised in CoNLL-U format.

B.2 NCI

Foras na Gaeilge provided us with a .vert file⁹ containing 33,088,532 tokens in 3,485 documents. We extract the raw text from the first tab-separated column and carry out the following conversions (number of events):

- Replace $\&\text{quot};$ with a neutral double quote (4408).
- Replace the standard xml/html entities quot, lt, gt and amp tokenised into three tokens, e.g. $\&\text{lt}\text{quot}\text{lt};$, with the appropriate characters (128).
- Replace the numeric html entities 38, 60, 147, 148, 205, 218, 225, 233, 237, 243 and 250, again spanning three tokens, e.g. $\&\text{lt}\#38\text{lt};$, with the appropriate Unicode characters (3679).
- Repeat from the start until the text does not change.

We do not modify the seven occurrences of $\backslash x\backslash x13$ as it is not clear from their contexts how they should be replaced. After pre-processing and treating all whitespace as token separators, e.g. in the NCI token "go leor", we obtain 33,472,496 tokens from the NCI.

B.3 Sentence Boundary Detection

Many of the NCI segments marked up with $\langle s \rangle$ tags contain multiple sentences. We treat each segment boundary as a sentence boundary and further split segments into sentences recursively, finding the best split point according to the following heuristics, splitting the segment into two halves and applying the same procedure to each half until no suitable split point is found.

⁸<https://github.com/attardi/wikiextractor>

⁹MD5 7be5c0e9bc473fb83af13541b1cd8d20

912	• Reject if the left half contains no letters and is	these are most likely to be part of enumera-	956
913	short. This covers cases where the left half is	tions. An exception is “Airteagal” followed	957
914	only a decimal number.	by a token ending with a full-stop, a num-	958
915	• Reject if the right half has no letters and is	ber, a full-stop, another number and another	959
916	short or is an ellipsis.	full-stop. Here, we implemented a preference	960
917	• Reject if the right half’s first letter, skipping	for splitting after the first separated full-stop,	961
918	enumerations, is lowercase.	assuming the last number is part of an enumer-	962
919	• Reject if the left half only contains a Roman	ation.	963
920	number (in addition to the full-stop).	• Prefer a split point balancing the lengths of	964
921	• Reject if inside round, square, curly or angle	the halves in characters.	965
922	brackets and the brackets not far away from		
923	the candidate split point.	B.4 OpusFilter Filters	966
924	• If sentence-ending punctuation is followed by	For OpusFilter-basic , we include the following	967
925	two quote tokens we also consider splitting	filters:	968
926	between the quotes and prefer this split point	• LengthFilter : Filter sentences contain-	969
927	if not rejected by above rules.	ing more than 512 words.	970
928	• If sentence-ending punctuation is followed by	• LongWordFilter : Filter sentences con-	971
929	a closing bracket we also consider splitting	taining words longer than 40 characters.	972
930	after the closing bracket and prefer this split	• HTMLTagFilter : Filter sentences contain-	973
931	point if not rejected by above rules.	ing HTML tags.	974
932	• If a question mark is followed by more ques-	• PunctuationFilter : Filter sentences	975
933	tion marks we also consider splitting after the	which are over 60% punctuation.	976
934	end of the sequence of question marks and	• DigitsFilter : Filter sentences which are	977
935	prefer this split point if not rejected by above	over 60% numeric symbols.	978
936	rules.	For OpusFilter-basic-char-lang , we use the	979
937	• If a exclamation mark is followed by more	same filters as in OpusFilter-basic but include the	980
938	exclamation marks we also consider splitting	following character script and language ID filters:	981
939	after the end of the sequence of exclamations	• CharacterScoreFilter : Filter sen-	982
940	marks and prefer this split point if not rejected	tences which are below a ratio r of Latin char-	983
941	by above rules.	acters, where $r \in \{1.0\}$.	984
942	• If a full-stop is the first full-stop in the overall	• LanguageIDFilter : Filter sentences	985
943	segment, the preceding token is “1”, there	where the language ID tools have a lower con-	986
944	are more tokens before this “1” and the token	confidence score than c , where $c \in \{0.8\}$.	987
945	directly before “1” is not a comma or semi-		
946	colon we assume that this is an enumeration	C Cloze Test Examples	988
947	following a heading and prefer splitting before	C.1 Prediction Classification	989
948	the “1”.	Table 6 provides one example per classification cat-	990
949	• We do not insert new sentence boundaries at	egory of masked token predictions generated by the	991
950	a full-stop after “DR”, “Prof” and “nDr”, and,	language models during our cloze test evaluation.	992
951	if followed by a decimal number, after “No”,	In the <i>match</i> example in Table 6, the original	993
952	“Vol” and “Iml”.	meaning (“What are those radical roots?”) differs	994
953	• Splitting after a full-stop following decimal	to the meaning of the resulting string (“What about	995
954	numbers in all other cases is dispreferred, giv-	those radical roots?”) in which the masked token	996
955	ing the largest penalty to small numbers as	is replaced by the predicted by mBERT-cp. How-	997
		ever, the latter construction is grammatically and	998
		semantically acceptable.	999

Context Cue	Masked Word	Model	Prediction	Classification
<i>Céard [MASK] na préamhacha raidiciúla sin?</i> (‘What [MASK] those radical roots?’)	<i>iad</i> (‘them’)	mBERT-cp	<i>faoi</i> (‘about’)	match
<i>Agus seo [MASK] an fhadhb mhór leis an bhfógra seo.</i> (‘And this [MASK] the big problem with this advert.’)	<i>í</i> (‘it’)	WikiBERT	<i>thaitin</i> (‘liked’)	mismatch
<i>Cheannaigh Seán leabhar agus léigh sé [MASK].</i> (‘Seán bought a book and he read [MASK].’)	<i>é</i> (‘it’)	gaBERT	<i>leabhar</i> (‘a book’)	copy
<i>Ní h[MASK] sin aidhm an chláir.</i> (‘[MASK] is not the aim of the programme.’)	<i>##é</i> (‘it’)	mBERT	- (minus sign)	gibberish

Table 6: Examples of cloze test predictions and classifications.

Model	Short	Medium	Long
mBERT	20.69%	55.56%	41.67%
wikibert	51.72%	58.33%	74.29%
mBERT-cp	75.86%	83.33%	94.29%
gaBERT	79.31%	83.33%	85.71%
gaELECTRA	79.31%	77.78%	88.57%

Table 7: Accuracy of language models segmented by length of context cue where short: 4–10 tokens, medium: 11–20 tokens, and long: 21–77 tokens.

In the *mismatch* example in Table 6, the predicted token is a valid Irish word, however the resulting generated text is nonsensical.

Though technically grammatical, the predicted token in the *copy* example in Table 6 results in a string with an unnatural repetition of a noun phrase where a pronoun would be highly preferable (‘Seán bought a book and he read a book.’).

In the *gibberish* example in Table 6, the predicted token does not form a valid Irish word and the resulting sentence is ungrammatical and meaningless.

C.2 Effect of Length of Context on Accuracy of Prediction

In order to observe the effect that the amount of context provided has on the accuracy of the model, Table 7 shows the proportion of matches achieved by each language model when the results are segmented by the length of the context cues.

All the models tested are least accurate when tested on the group of short context cues. All except mBERT achieved the highest accuracy on the group of long sentences.

C.3 Easy and Difficult Context Cues

A context cue may be considered easy or difficult based on:

- Whether the tokens occur frequently in the training data
- The number of context clues
- The distance of the context clues from the masked token

Two Irish language context cues which vary in terms of difficulty are exemplified below.

Bean, agus í cromtha thar thralaí bia agus [MASK] ag ithe a sáithe.

‘A woman, bent over a food trolley while eating her fill.’

We can consider the above sentence to be easy for the task of token prediction due to the following context clues:

- ‘Bean’ is a frequent feminine singular noun.
- ‘í’ is a repetition of the feminine singular pronoun to be predicted.
- The lack of lenition on ‘sáithe’ further indicates that the noun it refers to may not be masculine.

These clues indicate that the missing pronoun will be feminine and singular.

Seo béile aoibhinn fuirist nach dtógann ach timpeall leathuair a chloig chun [MASK] a ullmhú.
‘This is an easy, delicious meal that only takes about half an hour to prepare.’

None of the language models tested predicted a plausible token for the above sentence. This example is more challenging as the only context clue

is the feminine singular noun ‘béile’ which is 11 tokens in distance from the masked token.

D gaELECTRA Model

In addition to the gaBERT model of the main paper, we release gaELECTRA, an ELECTRA model (Clark et al., 2020) trained on the same data as gaBERT. ELECTRA replaces the MLM pre-training objective of BERT with a binary classification task discriminating between authentic tokens and alternative tokens generated by a smaller model for higher training efficiency. We use the default settings of the “Base” configuration of the official implementation¹⁰ and train on a TPU-v3-8. As with BERT, we train for 1M steps and evaluate every 100k steps. However, we train on more data per step as the batch size is increased from 128 to 256 and a sequence length of 512 is used throughout.

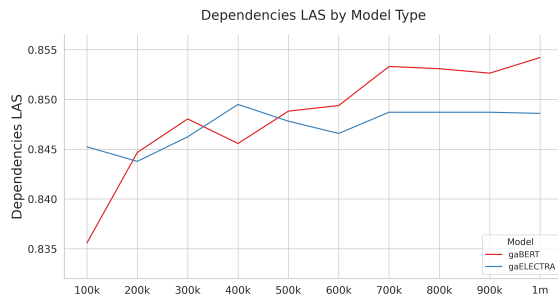


Figure 3: Dependency parsing LAS for each model type. Every 100k steps, we show the median of five LAS scores obtained from fine-tuning the respective model five times with different initialisation.

Figure 3 shows the development LAS of gaELECTRA and gaBERT for each checkpoint. The best gaBERT checkpoint is reached at step 1 million, which may indicate that there are still gains to be made from training for more steps. The highest median LAS for gaELECTRA is reached at step 400k. It is worth noting that although the two models are compared at the same number of steps, the different pretraining hyperparameters mean they are not trained on the same number of tokens per step.

We also compare the results of the gaELECTRA model to the other models in Tables 8 and 9. gaELECTRA performs slightly below gaBERT but better than both mBERT models and the WikiBERT model.

¹⁰<https://github.com/google-research/electra>

In terms of the Cloze test experiments: First, for the original masked token prediction (Table 4), gaELECTRA predicted the correct token 75 times, which is the same number as gaBERT and is slightly below mBERT with continued pretraining, which has a score of 78. Second, for the manual evaluation of the tokens generated by each model (Table 5), gaELECTRA predicted 82 matches, 8 mismatches, 1 copy, and 9 gibberish tokens; compared to 83, 14, 2 and 1 predicted by gaBERT, respectively.

E XLM-R Baseline

We add another off-the-shelf baseline by fine-tuning XLM-R_{BASE}, which is a multilingual RoBERTa model introduced by Conneau et al. (2020), in the task of multitask dependency parsing and POS and morphological features tagging. This model performs better than both variants of mBERT as well as the WikiBERT model but underperforms our two monolingual models, gaBERT and gaELECTRA.

F Full Model Results

This section examines the results produced by each of our models in more detail and also presents the scores of the additional models we examine, namely XLM-R_{BASE} and gaELECTRA.¹¹ Tables 8 and 9 list the accuracies for predicting universal part of speech (UPOS), treebank-specific part of speech (XPOS) and morphological features, as well as the unlabelled and labelled attachment score (UAS and LAS, respectively) for all models discussed in this paper.

For the multilingual models, mBERT performs worse than XLM-R_{BASE}, which is a strong multilingual baseline. The monolingual WikiBERT model performs slightly better than mBERT in terms of LAS but is worse than XLM-R_{BASE}. The continued pretraining of mBERT on our data enables us to close the gap between mBERT and XLM-R_{BASE}. gaBERT is still the strongest model for all metrics in terms of test set scores. gaELECTRA performs slightly below that of gaBERT but better than XLM-R_{BASE}. It should be noted that each row selects the model based on median LAS, therefore, all other metrics are those that this selected model achieved.

¹¹We tried training a RoBERTa_{BASE} model on our data but could not obtain satisfactory LAS scores (a fine-tuned model achieved a dev LAS of 81.8, which is comparable to mBERT) and leave finding suitable hyperparameters for this architecture to future work.

Model	UD	UPOS	XPOS	FEATS	UAS	LAS
mbert-os	2.8	95.7	94.7	89.2	86.9	81.8
xlmr-base-os	2.8	96.4	95.1	90.6	88.3	84.0
wikibert-os	2.8	95.9	94.9	89.4	86.8	81.9
mbert-cp	2.8	97.2	95.8	92.3	88.1	84.3
gabert	2.8	97.1	96.2	93.1	89.2	85.6
gaelectra	2.8	97.3	96.1	92.8	89.1	85.3

Table 8: Full model results on development data. For model name abbreviations, see test result table.

Model	UD	UPOS	XPOS	FEATS	UAS	LAS
mbert-os	2.8	95.4	94.3	88.6	86.2	80.3
xlmr-base-os	2.8	96.1	95.1	90.0	87.7	82.5
wikibert-os	2.8	95.7	94.4	88.3	85.9	80.4
mbert-cp	2.8	96.7	95.5	91.7	87.1	82.3
gabert	2.8	97.0	95.7	91.8	88.4	84.0
gaelectra	2.8	96.9	95.5	91.5	87.6	83.1

Table 9: Full model results on test data (os = fine-tuned off-the-shelf model, cp = continued pre-training before fine-tuning).