

FrCoLA: a French Corpus of Linguistic Acceptability Judgments

Anonymous ACL submission

Abstract

Large foundation language models and Transformer-based neural language models have exhibited outstanding performance in various downstream tasks. However, there is limited understanding regarding how these models internalize linguistic knowledge, so various linguistic benchmarks have recently been proposed to facilitate syntactic evaluation of language models across languages. This paper introduces FrCoLA (French Corpus of Linguistic Acceptability Judgments), consisting of 25,153 sentences annotated with binary acceptability judgments and categorized into four linguistic phenomena. Specifically, those sentences are manually extracted from an official online resource maintained by a Québec Governments institution, and split into in-domain data splits. Moreover, we also manually extracted 2,675 from a second France-based organization source and created an out-of-domain hold-out split. We then evaluate the linguistic capabilities of three different language models for each of the seven linguistic acceptability judgment benchmarks. The results demonstrated that, for most languages, on average, fine-tuned Transformer-based neural language models are strong baselines on the binary linguistic acceptability classification tasks. However, for the FrCoLA benchmark, on average, a fine-tuned Transformer-based model outperformed other methods tested.

1 Introduction

The introduction of large foundation language models (LLM) and Transformer-based neural language model (Vaswani et al., 2017), such as GPT-3 (Brown et al., 2020) and LLaMA (Touvron et al., 2023), has led to major progress in natural language processing (NLP), substantially increasing the performance of most NLP tasks (Zhang et al., 2023). LLMs and Transformer-based neural language models were initially introduced for English

(Kenton and Toutanova, 2019; Brown et al., 2020), but many other languages were later introduced, such as Norwegian (Kummervold et al., 2021), Russian (Kuratov and Arkhipov, 2019) and French (Martin et al., 2020). NLP research has approached the competencies evaluation of various natural language tasks of LLM with various benchmark corpora such as the English benchmarks GLUE (Wang et al., 2018), SuperGLUE (Wang et al., 2019), and GLGE (Liu et al., 2021) to name a few. These corpora are collections of resources for training, evaluating, and analyzing natural language systems (Gao et al., 2023; Chang et al., 2023). For example, GLUE aims to benchmark an NLP system’s capabilities for natural language understanding (NLU) (Wang et al., 2018). At the same time, GLGE focuses on natural language generation (NLG) tasks such as document summarization (Liu et al., 2021).

Recently, much effort has been put into creating linguistic acceptability resources to assess and benchmark LLM linguistic competency, where recent NLP research formulate linguistic competency as a binary classification acceptability judgments task (Cherniavskii et al., 2022; Proskurina et al., 2023). That is the ability, from a native speaker’s perspective, to distinguish the correct form and naturalness of an acceptable sentence from an unacceptable one (Chomsky, 2014). Recently, similar non-English resources have been proposed to answer this question in typologically diverse languages such as Japanese (Someya et al., 2023), Norwegian (Jentoft and Samuel, 2023), and Chinese (Hu et al., 2023). However, the ability of LLMs to perform linguistic acceptability judgments in French remains understudied.

To this end, we introduce the French Corpus of Linguistic Acceptability Judgments (FrCoLA)¹, a corpus consisting of 25,153 acceptability judgment sentences, making it the second largest linguistic acceptability resources available in the NLP

¹Link removed for double-anonymized anonymity.

literature. Specifically, the sentences were created by French-Canadian linguists and manually extracted from the “*Banque de dépannage linguistique*” (BDL), an official online resource maintained by the Québec Government.

The main contributions of this work are therefore 1) the creation and release of FrCoLA; 2) a set of experiments to assess the performance of Transformer-based models on FrCoLA; 3) a set of experiments using a monolingual and a cross-lingual Transformer-based model on English, Swedish, Italian, Russian, Chinese, Norwegian, Japanese and French, with the potential to open up novel multi-language research perspectives. It is outlined as follows: first, we study the available linguistic acceptability resources corpora and related binary classification neural language model research in [section 2](#). Then, we propose the FrCoLA in [section 3](#), and in [section 4](#) and [section 5](#) we present a set of experiments aimed at testing the performance of Transformer-based binary classifiers on all the linguistic acceptability resources corpora. Finally, in [section 6](#), we conclude and discuss our future work.

2 Related Work

2.1 Linguistic Acceptability Judgments

Linguistic acceptability judgment constitutes a pivotal component of human linguistic competence, underlying individuals’ inherent capacity to distinguish the correct form and naturalness of an acceptable sentence from an unacceptable one, even without formal grammar training. For instance, individuals can inherently distinguish between two sentences and identify the one that is more acceptable or natural-sounding. This assessment is the primary behavioural benchmark employed by generative linguists to investigate the underlying structure of human language (Chomsky, 2014). Through analyzing linguistic acceptability judgments, linguists can learn about the linguistic rules that govern languages and how these rules manifest themselves in native speakers’ speech.

2.2 LLM Evaluation

Historically, evaluation of LLMs and Transformer-based neural language models has been conducted using metrics or benchmark corpora (Chang et al., 2023; Awasthi et al., 2023). The first approach

Acceptable Sentence	Not Acceptable Sentence
The cats annoy Tim.	The cats annoys Tim.

Table 1: Example of a minimal pair (Warstadt and Bowman, 2019).

relies either on task-agnostic metrics, such as perplexity (Jelinek et al., 1977) which measures the quality of the probability distribution of words in a given corpus by a model, or alternatively on task-specific metrics, like the BLEU score that evaluates a model’s performance for machine translation (Papineni et al., 2002). The second approach relies on large corpora designed for NLU or NLG downstream tasks. For example, the GLUE benchmark (Wang et al., 2018) is used to assess a model’s NLU performance on tasks such as semantic similarity, linguistic acceptability judgment and sentiment analysis. In contrast, GLGE (Liu et al., 2021) evaluates language generation tasks such as summarization and question answering.

2.2.1 LLM Linguistic Acceptability Judgments Evaluation

Recently, NLP researchers started using linguistic acceptability judgment tasks to assess the robustness of LLMs’ and Transformer-based neural language models against grammatical errors (Yin et al., 2020; Miaschi et al., 2023) and to probe their grammatical knowledge (Warstadt and Bowman, 2019; Zhang et al., 2021; Choshen et al., 2022; Mikhailov et al., 2022). Two approaches are used to perform this evaluation, namely minimal pairs and binary classification acceptability judgments (Wang et al., 2018; Warstadt and Bowman, 2019; Chang et al., 2023).

In the first approach, a set of minimal pairs of grammatically acceptable and unacceptable sentences, such as the pair illustrated in [Table 1](#), is presented to an LLM that must decide which is grammatically correct. By observing which sentences the LLM assigns a higher correctness probability to, one can assess which grammatical phenomena it is sensitive to (Warstadt and Bowman, 2019) Corpus such as BLIMP in English (Warstadt and Bowman, 2019) and CLiMP in Chinese (Xiang et al., 2021) have been proposed to enable the evaluation of LLM on a wide range of linguistic phenomena.

Concurrently, in the second approach, a set of sentences that are either grammatical or ungram-

²<https://vitrinelinguistique.oqlf.gouv.qc.ca/banque-de-depannage-linguistique>

Label	Sentence
0 (Ungrammatical)	Edoardo returned to his last year city
1 (Grammatical)	This woman has impressed me

Table 2: Example sentences from the ItaCoLA dataset (Trotta et al., 2021).

matical, such as the two shown in Table 2, are provided to an LLM which must perform a binary acceptability classification task (Warstadt et al., 2019). Corpora such as CoLA in English (Warstadt et al., 2019) and CoLAC in Chinese (Hu et al., 2023) have been proposed to assess LLMs’ capabilities to discriminate proper grammar from improper in their respective languages. Typically, these datasets comprise sentences collected from syntax textbooks and linguistics journals. However, as of yet no such corpus exists for French.

3 FrCoLA: French Corpus of Linguistic Acceptability Judgments

In this work, we introduce the **French Corpus of Linguistic Acceptability Judgments** (FrCoLA), which will be the first large-scale binary linguistic acceptability judgments task dataset for the French language and the second-largest such corpus in any language. FrCoLA consists of French sentences from the “*Banque de dépannage linguistique*” (BDL), an official online resource from the “*Office québécois de la langue française*”, an official provincial government public organization in Canada. The BDL is a grammatical resource index that offers 2,667 articles divided into eleven categories, such as “Spelling” (*orthographe*), “Grammar” (*grammaire*) and “Syntax” (*syntaxe*). These categories are further subdivided into sub-categories that address either a specific situation in the linguistic literature (e.g. how to use punctuation in a sentence) or a problematic linguistic situation (e.g. the use of borrowed words from English vocabulary). In these articles, the BDL explains various linguistic phenomena that are correct or incorrect and uses examples written by French linguists to illustrate both cases. For example, the “grammar” category includes the “adverbs” (*adverbes*) sub-category that includes an article about the linguistic phenomenon of proper and improper use of the adverb “surrounding” (*alentour*). A snippet of that article is shown in Figure 1. It displays two examples of well-written sentences using the adverb (marked in green) and one example of an



Figure 1: Snipped of the BDL article for the French word “*alentour*”. The text is in French.

erroneous usage (marked in red).

3.1 Data Collection

Sentences in FrCoLa were collected manually from the BDL online resource focusing on French grammar. Specifically, we examined all its 2,667 articles and extracted 25,153 linguistic acceptability judgment sentences. Each sentence was labelled 0 (ungrammatical) or 1 (grammatical) following the BDL green/red colour scheme as illustrated in Figure 1. Furthermore, since the BDL uses a fine-grained category structure to sort various linguistic phenomena, we collected these categories and associated them to labels according to the French linguistic literature (Fagyal et al., 2006; Chesley, 2010; Boivin and Pinsonneault, 2020; Feldhausen and Buchczyk, 2021), and labelled each extracted sentence accordingly. Our linguistic phenomena labels and their BDL-associated categories are listed below, and Table 3 present FrCoLA statistics for each one.

- **Syntax** : agreement violations, corruption of word order, misconstruction of syntactic clauses and phrases, incorrect use of appositions, violations of verb transitivity or argument structure, ellipsis, missing grammatical constituencies or words.
BDL Categories: Editorial (*Rédaction*), Syntax (*Syntaxe*), Punctuation (*Ponctuation*), Typography (*Typographie*), and Proper nouns (*Noms propres*).
- **Morphology** : incorrect derivation or word building, non-existent words.
BDL Categories: Ortograph (*Ortographie*), Grammar (*Grammaire*), Abbreviations and

- 251 symbols (*Abréviations et symboles*), and Pro-
 252 nunciation (*Prononciation*).
- 253 • **Semantic**: incorrect use of negation, violation
 254 of the verb’s semantic argument structure.
 255 **BDL Categories**: Vocabulary (*Vocabulaire*).
 - 256 • **Anglicism**: word and syntactical structure bor-
 257 rowed from English grammar.
 258 **BDL Categories**: Anglicisms (*Anglicisme*).

259 3.2 Analysis of FrCoLA

260 In this section, we compare our FrCoLA dataset to
 261 all related corpora. Table 4 present, for each corpus,
 262 the language, and the number of sentences, percent-
 263 age of acceptable sentences and total vocabulary
 264 in the train, dev and test sets³ as well as for the
 265 entire corpus. The total vocabulary sizes were com-
 266 puted using language-specific SpaCy tokenizers
 267 (Honnibal et al., 2020) that split each sentence into
 268 individual words or punctuation. We can see that
 269 FrCoLA is the second largest corpus behind only
 270 NoCoLA, and is approximately twice the size of
 271 all the other corpora. Moreover, FrCoLA shares a
 272 similar frequency of acceptable sentences to CoLA,
 273 CoLAC and RuCoLA datasets, and like with the
 274 other corpora all splits have a similar frequency
 275 of acceptable sentences. Finally, we can see that
 276 FrCoLA has the third-largest vocabulary size com-
 277 pared to the other datasets.

278 3.3 Out-of-Domain Hold-Out Split

279 Like CoLA, RuCoLA, CoLAC, and JCoLA, our
 280 FrCoLA includes an out-of-domain hold-out data
 281 split in addition to the standard train, dev, and test
 282 dataset splits to assess whether a system trained
 283 with it might suffer from overfitting. The definition
 284 of out-of-domain varies depending on the corpus.
 285 In CoLA, the out-of-domain set includes sources
 286 of varying degrees of domain specificity and time
 287 period compared to those used for the main dataset
 288 (Warstadt and Bowman, 2019). For RuCoLA, they
 289 are sentences generated by an automatic machine
 290 translation system and paraphrase generation mod-
 291 els and annotated by a human annotator (Mikhailov

³It is worth mentioning that for CoLA, RuCoLA and JCoLA, their in-domain test set labels are not publicly available to reduce the risk of overfitting. Thus, like other related work (Trotta et al., 2021; Cherniavskii et al., 2022), we use their out-of-domain dev sets as the test sets. Also, CoLAC does not offer a test set with a label or an out-of-domain dev set. Thus, as per the authors’ recommendation, we have re-sampled the in-domain train and dev set using a 60-10-30% split using the seed 42 to create an in-domain test set.

292 et al., 2022). The JCoLA out-of-domain split com-
 293 prises sentences from the Journal of East Asian
 294 Linguistics, a source with typically more complex
 295 linguistic phenomena than regular text (Someya
 296 et al., 2023). We have used a similar approach to
 297 JCoLA and CoLA; we have used an alternative and
 298 substantially different source to build our out-of-
 299 domain set. The source we picked is the *Académie
 300 française*, a France-based organization founded in
 301 1635 as a “society of scholars” in science and liter-
 302 ature (Académie française, 2024). This organiza-
 303 tion publishes a monthly online journal, *La langue
 304 française: Dire, Ne pas dire*⁴ (“The French lan-
 305 guage: What to say, and not say”), that presents
 306 various articles on the proper and improper use of
 307 the French language, sorted into three categories,
 308 namely “neologisms and anglicisms” (*néologismes
 309 and anglicismes*), “wrongful employment” (*em-
 310 plois fautifs*), and “abusive extensions of meaning”
 311 (*extensions de sens abusives*). We manually ex-
 312 tracted all examples of grammatical and ungram-
 313 matical sentences from the 1,013 articles in the
 314 journal. In total, 2,675 sentences were extracted
 315 and binary labelled.

316 We present in Table 5 the number of sentences,
 317 vocabulary size and percentage of acceptable sen-
 318 tences of all linguistic corpora with an out-of-
 319 domain test set. However, since other corpora
 320 do not distribute their hold-out labels, we could
 321 not compute the percentage of acceptable sen-
 322 tences. We also note that for JCoLA, the out-of-
 323 domain hold-out split was unavailable in their of-
 324 ficial dataset GitHub repository. Once again, we
 325 can see that FrCoLA is the second largest corpus in
 326 terms of number of sentences and vocabulary size,
 327 with nearly as many sentences as RuCoLA. Com-
 328 pared to the main FrCoLA corpus in Table 4, we
 329 can see that the out-of-domain dataset comprises
 330 a much less diverse vocabulary, making it well
 331 distinct from the other splits. Finally, the out-of-
 332 domain hold-out split has a percentage of accept-
 333 able sentences nearly 15% lower than the overall
 334 corpus, making it more robust to highlight overfit-
 335 ting cases in machine learning models.

336 4 Experiments

337 We evaluate three neural-based methods for accept-
 338 ability classification leveraging Transformer-based
 339 architecture.

⁴<https://www.academie-francaise.fr/dire-ne-pas-dire>

Category	# Sen and % Acp	Example
Syntax	5,152 77.56	<i>Dès son arrivée, on s'empessa de lui poser des questions à propos de son voyage.</i> <i>Dès en arrivant, on s'empessa de lui poser des questions à propos de son voyage.</i>
Morphology	10,642 68.18	<i>La journée était calme : point de vent, point de bruit, point de mouvement.</i> <i>La journée était calme : point de vent, poing de bruit, point de mouvement.</i>
Semantic	5,442 72.49	<i>Quand la parade est passée, le vieil homme s'est levé pour aller voir à la fenêtre.</i> <i>Quand la parade est passée, le vieil homme s'est levé debout pour aller voir à la fenêtre.</i>
Anglicism	3,917 58.26	<i>Sauront-ils répondre aux les besoins de l'enfant?</i> <i>Sauront-ils rencontrer les besoins de l'enfant?</i>

Table 3: Number of sentences (# Sen) and the percentage of acceptable sentences (% Acp) per category in FrCoLA (all three splits), and example of a positive and a negative (**bolded** with error underlined) in each category.

	Language	# Sen	Train % Acp	Vocab	# Sen	Dev % Acp	Vocab	# Sen	OODD/Test % Acp	Vocab	# Sen	Total % Acp	Vocab
CoLA (Warstadt et al., 2019)	English	8,551	70.44	5,778	527	69.26	1,375	516	68.60	988	9,594	70.27	6,097
DaLAJ (Volodina et al., 2021)	Swedish	7,682	50.00	6,841	890	50.00	1,799	888	50.00	1,661	9,460	50.00	7,884
ITaCoLA (Trotta et al., 2021)	Italian	7,801	84.39	5,825	946	85.41	1,844	1,888	84.21	1,888	9,722	84.47	6,402
RuCoLA (Mikhailov et al., 2022)	Russian	7,869	74.52	19,057	983	74.57	4,140	1,804	63.69	9,353	10,656	72.69	26,382
CoLAC (Hu et al., 2023)	Chinese	4,134	66.09	3,835	460	66.96	1,024	1,970	67.82	2,636	6,564	66.67	4,759
NoCoLA (Jentoft and Samuel, 2023)	Norwegian	116,195	31.46	32,561	14,289	32.59	8,865	14,383	31.58	8,600	144,867	31.58	37,319
JCoLA (Someya et al., 2023)	Japanese	6,919	83.38	3,730	865	83.93	1,483	684	73.28	896	8,469	82.62	4,146
FrCoLA (This work)	French	15,846	69.49	18,350	1,761	69.51	5,369	7,546	69.49	12,690	25,153	69.49	22,131

Table 4: Comparison of FrCoLA and related corpora for number of sentences (# Sen), percentage of acceptable sentences (% Acp), and total vocabulary (Vocab). “OODD” stands for out-of-domain data split (CoLA, RuCoLA and JCoLA).

	Out-of-Domain Hold-Out		
	# Sen	Vocab	% Acp
CoLA	533	1035	N/A
RuCoLA	2,789	12,211	N/A
CoLAC	931	1,168	N/A
JCoLA	N/A	N/A	N/A
FrCoLA	2,675	1,651	53.91

Table 5: Comparison of FrCoLA with all related corpus with an out-of-domain hold-out set for the number of sentences (# Sen), the vocabulary size (Vocab) and the percentage of acceptable sentences (% Acp).

4.1 Evaluation Metrics

Following Warstadt et al. (2019), performance is measured using the accuracy score (Acc) and Matthews Correlation Coefficient (MCC) (Matthews, 1975). Accuracy on the dev set is used as the target metric for hyperparameter tuning and early stopping. We report the results averaged over ten restarts from different random seeds (i.e. 42, 43, ..., 50, 51).

4.2 Models

4.2.1 Baseline

As a baseline, we evaluate a fine-tuned language-specific pre-trained monolingual and a cross-lingual neural language model using each available linguistic acceptability judgment binary classification benchmark dataset and FrCoLA. We used a different language-specific Transformed-based LLM for each language, which was optimized using different tokenization methods and training corpus. We detail the language-specific models used in our experimentation in Appendix A. We use XLM-RoBERTa-base (Conneau et al., 2020) as our cross-lingual neural language model; the model details are described in Appendix A. For each language, we name these obtained models Monolingual FT and Cross-Lingual FT, respectively.

4.2.2 State-Of-The-Art

The state-of-the-art approach to binary linguistic acceptability judgments is the topological data analysis (TDA) proposed by Cherniavskii et al. (2022). This approach extracts the attention maps of a fine-tuned Transformers-based neural language model to use as linguistic features to train a binary logistic regression. The authors report that this approach significantly outperformed previous approaches, in-

creasing by up to 0.24 Matthew’s correlation coefficient score on linguistic acceptability for three languages (English, Italian and Swedish). In our case, we use the attention maps from the monolingual models we fine-tuned as baselines. For each language, we name this model LA-TDA.

4.3 Training Settings

Each language model is fine-tuned using the train and dev split and evaluated using the test or out-of-domain valid split (OODD) (CoLA, RuCoLA and JCoLA) following the standard procedure under the HuggingFace library (Wolf et al., 2020). Each model is fine-tuned for four epochs and uses the AdamW optimizer (Loshchilov and Hutter, 2018), with a learning rate of $3e-5$ and a weights decay of $1e-2$. Since the corpora are unbalanced, we use a weighted balanced loss based on the training percentage of acceptable sentences. We use a batch size of 32 and the other default hyperparameters. For each language model, we use the default tokenizer with a maximum sequence length of 64 tokens without lowercasing during tokenization.

5 Results and Discussion

Table 6 presents the accuracy and the MCC of our three models for each benchmark dataset on the dev and test set, with **bolded** value indicating the best score per benchmark. The table reports the average and one standard deviation over the ten restarts. We observe that, for most languages, on average LA-TDA outperforms other fine-tuned methods, but not on all metrics and with a smaller margin than reported by Cherniavskii et al. (2022). The two exceptions to this are CoLA and our FrCoLA. FrCoLA’s performance is always slightly better when using the fine-tuned Monolingual FT model. Considering that LA-TDA is computed asymptotically in quadratic time (Cherniavskii et al., 2022), the performance gains seem marginal compared to the added computational expense. These results show that fine-tuned Transformer-based neural language models are strong baselines for the binary linguistic acceptability classification tasks.

We present in Table 7 the accuracy and the MCC of our three models trained using FrCoLA over the dataset’s four categories. The table reports the average and one standard deviation over the ten restarts. We can see that the category “anglicism” has the lowest performance. For the two approaches us-

ing monolingual LLM (i.e. Monolingual FT and LA-TDA), we hypothesize that this situation is due to occurrences of anglicism in the LLM training dataset. Indeed, using word and syntactical structure borrowed from English grammar is more common over web-based (Laviosa, 2010; Planchon and Stockemer, 2019; Solano, 2021; Šukalić et al., 2022) and even official educational text (Simon et al., 2021). Thus, fine-tuning the pre-trained LLM model can be more challenging, considering that the “anglicism” category contains the least examples. For the cross-lingual approach, since the LLM has learned word representation over English during training, we hypothesize that sentences using English words or syntax are considered more probable for the model; thus, it is more challenging for the classifier to classify these examples correctly.

Finally, we present in Table 8 the accuracy and the MCC of our three models trained using FrCoLA but evaluated using our out-of-domain hold-out set. The table reports the average and one standard deviation over the ten restarts. We can see that, once again, the Monolingual FT model outperforms the LA-TDA model. However, all three models show significant performance drops, of nearly 22% in accuracy and nearly 50% for the MCC. It shows that the fine-tuned models have overfitted over the train and dev dataset.

	Out-of-Domain Hold-Out Acc (%)	MCC
Monolingual FT	62.69 ± 1.13	0.286 ± 0.020
Cross-Lingual FT	55.99 ± 4.36	0.107 ± 0.088
LA-TDA	61.36 ± 0.90	0.090 ± 0.019

Table 8: Acceptability binary classification result on the FrCoLA out-of-domain hold-out set. The best score per benchmark is **bolded**.

6 Conclusion and Future Works

This article introduced FrCoLA, the French Corpus of Linguistic Acceptability Judgments, a dataset comprising 15,846 sentences annotated with binary acceptability judgments. It is the first such corpus in French, and the second-biggest one in any language. The sentences it comprises were manually extracted from an official online linguistic resource maintained by a Québec Government institution, and an additional out-of-domain dataset was compiled from an equivalent French institution. We then evaluated the linguistic performances of three fine-tuned models on FrCoLA and

Model	Dev		Test/OOD	
	Acc (%)	MCC	Acc (%)	MCC
CoLA				
Monolingual FT	83.61 ± 2.56	0.639 ± 0.030	80.89 ± 1.15	0.544 ± 0.025
Cross-Lingual FT	82.24 ± 1.35	0.575 ± 0.033	77.25 ± 2.42	0.452 ± 0.041
LA-TDA	84.91 ± 1.24	0.633 ± 0.031	80.70 ± 1.38	0.532 ± 0.034
DaLAJ				
Monolingual FT	69.12 ± 1.53	0.411 ± 0.029	72.33 ± 1.40	0.467 ± 0.025
Cross-Lingual FT	55.18 ± 5.90	0.131 ± 0.144	55.21 ± 5.89	0.124 ± 0.137
LA-TDA	70.08 ± 1.24	0.411 ± 0.024	73.54 ± 1.05	0.475 ± 0.020
ITACoLA				
Monolingual FT	83.29 ± 3.71	0.420 ± 0.051	83.45 ± 3.34	0.446 ± 0.050
Cross-Lingual FT	79.97 ± 6.22	0.105 ± 0.121	79.12 ± 5.99	0.117 ± 0.124
LA-TDA	87.51 ± 0.88	0.423 ± 0.050	86.59 ± 0.93	0.422 ± 0.054
RuCoLA				
Monolingual FT	74.49 ± 2.56	0.352 ± 0.027	66.81 ± 3.56	0.379 ± 0.030
Cross-Lingual FT	71.84 ± 3.00	0.276 ± 0.038	56.81 ± 3.18	0.189 ± 0.026
LA-TDA	77.56 ± 0.61	0.337 ± 0.022	71.09 ± 0.92	0.382 ± 0.018
CoLAC				
Monolingual FT	75.93 ± 1.35	0.444 ± 0.027	77.78 ± 1.43	0.482 ± 0.023
Cross-Lingual FT	73.37 ± 2.72	0.337 ± 0.022	71.09 ± 0.92	0.382 ± 0.018
LA-TDA	77.33 ± 1.79	0.469 ± 0.044	79.01 ± 0.86	0.502 ± 0.023
NoCoLA				
Monolingual FT	77.90 ± 0.96	0.560 ± 0.009	77.90 ± 0.98	0.560 ± 0.009
Cross-Lingual FT	73.92 ± 1.40	0.504 ± 0.017	73.79 ± 1.37	0.505 ± 0.015
LA-TDA	81.58 ± 0.29	0.582 ± 0.007	82.01 ± 0.31	0.589 ± 0.009
JCoLA				
Monolingual FT	81.34 ± 4.48	0.039 ± 0.062	73.17 ± 0.61	0.067 ± 0.111
Cross-Lingual FT	72.64 ± 8.11	0.262 ± 0.058	72.86 ± 4.61	0.328 ± 0.059
LA-TDA	83.49 ± 0.68	0.252 ± 0.051	75.30 ± 1.25	0.230 ± 0.070
FrCoLA				
Monolingual FT	84.51 ± 0.78	0.619 ± 0.02	82.92 ± 0.61	0.578 ± 0.015
Cross-Lingual FT	70.67 ± 15.13	0.243 ± 0.263	69.91 ± 14.61	0.222 ± 0.240
LA-TDA	84.00 ± 0.48	0.606 ± 0.013	82.79 ± 0.45	0.574 ± 0.012

Table 6: Acceptability binary classification results and MCC by language. The best score per benchmark is **bolded**. “OODD” stands for out-of-domain data split (CoLA, RuCoLA and JCoLA).

Model	Category			
	Syntax	Morphology	Semantic	Anglicism
Test Accuracy (%)				
Monolingual FT	88.59 ± 0.60	81.76 ± 0.74	85.82 ± 0.40	74.36 ± 1.40
Cross-Lingual FT	83.31 ± 4.31	74.93 ± 4.70	79.84 ± 4.88	63.79 ± 4.66
LA-TDA	88.40 ± 0.23	81.49 ± 0.51	85.39 ± 0.53	74.18 ± 1.44
Test MCC				
Monolingual FT	0.654 ± 0.018	0.563 ± 0.017	0.620 ± 0.011	0.506 ± 0.028
Cross-Lingual FT	0.403 ± 0.279	0.327 ± 0.226	0.378 ± 0.261	0.223 ± 0.156
LA-TDA	0.649 ± 0.009	0.555 ± 0.013	0.609 ± 0.014	0.405 ± 0.026

Table 7: Acceptability binary classification results and MCC for FrCoLA per category. The best score is **bolded**.

equivalent datasets in other languages. Our results demonstrated that Transformer-based neural language models achieve high results on the binary classification task and are strong baselines. When fine-tuned on FrCoLA, a Transformer-based neural language model even outperforms the state-of-the-art TDA method proposed by Cherniavskii et al. (2022).

In future works, we plan to extend our dataset linguistic phenomena granularity and generate the complementary grammatical or ungrammatical sentence of each sentence in the dataset to create the first French minimal pair benchmark dataset. Moreover, we would also like to qualitatively explore the linguistic phenomena errors generated by the French fine-tuned Transformer-base model.

Limitations

All the sentences included in FrCoLA have been extracted from an official linguistic source on theoretical syntax. Therefore, those sentences are guaranteed to be theoretically meaningful, making FrCoLA a challenging dataset. However, the categories extracted automatically from the official source are skewed. Indeed, as shown in Table 3, nearly 42% of the dataset comprises morphological linguistic phenomena.

Ethical Considerations

FrCoLA may serve as training data for binary linguistic acceptability judgment classifiers (Batra et al., 2021), which may benefit the quality of generated texts. We acknowledge that such text generation progress could lead to misusing LLMs for malicious purposes, such as disinformation or harmful text generation and online harassment (Weidinger

et al., 2021; Bender et al., 2021). Nevertheless, our corpus can be used to train adversarial defence against such misuse and to train artificial text detection models (Lewis and White, 2023; Kumar et al., 2023).

Acknowledgements

Removed for double-anonymized.

References

- Académie française. 2024. *L’institution et l’organisation (The Institution and the Organization)*. Accessed: 2024-02-10.
- Raghav Awasthi, Shreya Mishra, Dwarikanath Mahapatra, Ashish Khanna, Kamal Maheshwari, Jacek Cywinski, Frank Papay, and Piyush Mathur. 2023. Humanely: Human Evaluation of Llm Yield, Using a Novel Web-Based Evaluation Tool. *medRxiv*, pages 2023–12.
- Soumya Batra, Shashank Jain, Peyman Heidari, Ankit Arun, Catharine Youngs, Xintong Li, Pinar Donmez, Shawn Mei, Shiunzu Kuo, Vikas Bhardwaj, Anuj Kumar, and Michael White. 2021. *Building adaptive acceptability classifiers for neural NLG*. In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pages 682–697, Online and Punta Cana, Dominican Republic. 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 *Proceedings of the ACM conference on fairness, accountability, and transparency*, pages 610–623.
- Marie-Claude Boivin and Reine Pinsonneault. 2020. La catégorisation des erreurs linguistiques: une grille de codage fondée sur la grammaire moderne. *Le français aujourd’hui*, 2(209):89–116.

535	Tom Brown, Benjamin Mann, Nick Ryder, Melanie	Philip Gage. 1994. A New Algorithm for Data Com-	590
536	Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind	pression. <i>C Users Journal</i> , 12(2):23–38.	591
537	Neelakantan, Pranav Shyam, Girish Sastry, Amanda		
538	Askill, et al. 2020. Language Models Are Few-Shot	Leo Gao, Jonathan Tow, Baber Abbasi, Stella Biderman,	592
539	Learners. <i>Advances in neural information processing</i>	Sid Black, Anthony DiPofi, Charles Foster, Laurence	593
540	<i>systems</i> , 33:1877–1901.	Golding, Jeffrey Hsu, Alain Le Noac’h, Haonan Li,	594
		Kyle McDonell, Niklas Muennighoff, Chris Ociepa,	595
541	Yupeng Chang, Xu Wang, Jindong Wang, Yuan Wu,	Jason Phang, Laria Reynolds, Hailey Schoelkopf,	596
542	Linyi Yang, Kaijie Zhu, Hao Chen, Xiaoyuan Yi,	Aviya Skowron, Lintang Sutawika, Eric Tang, An-	597
543	Cunxiang Wang, Yidong Wang, et al. 2023. A Survey	ish Thite, Ben Wang, Kevin Wang, and Andy Zou.	598
544	on Evaluation of Large Language Models. <i>ACM</i>	2023. A Framework for Few-Shot Language Model	599
545	<i>Transactions on Intelligent Systems and Technology</i> .	Evaluation .	600
546	Wanxiang Che, Zhenghua Li, and Ting Liu. 2010. LTP:	Matthew Honnibal, Ines Montani, Sofie Van Lan-	601
547	A Chinese Language Technology Platform. In <i>Inter-</i>	degheem, and Adriane Boyd. 2020. SpaCy: Industrial-	602
548	<i>national Conference on Computational Linguistics</i> ,	strength Natural Language Processing in Python .	603
549	page 13. Citeseer.		
550	Daniil Cherniavskii, Eduard Tulchinskii, Vladislav	Hai Hu, Ziyin Zhang, Weifang Huang, Jackie Yan-Ki	604
551	Mikhailov, Irina Proskurina, Laida Kushnareva, Eka-	Lai, Aini Li, Yina Ma, Jiahui Huang, Peng Zhang,	605
552	terina Artemova, Serguei Barannikov, Irina Pio-	and Rui Wang. 2023. Revisiting Acceptability Judge-	606
553	ntkovskaya, Dmitri Piontkovski, and Evgeny Bur-	ments. <i>arXiv:2305.14091</i> .	607
554	naev. 2022. Acceptability judgements via examining		
555	the topology of attention maps . In <i>Findings of the</i>	Fred Jelinek, Robert L Mercer, Lalit R Bahl, and	608
556	<i>Association for Computational Linguistics: EMNLP</i> ,	James K Baker. 1977. Perplexity—a Measure of the	609
557	pages 88–107, Abu Dhabi, United Arab Emirates.	Difficulty of Speech Recognition Tasks. <i>The Journal</i>	610
558	Association for Computational Linguistics.	<i>of the Acoustical Society of America</i> , 62(S1):S63–	611
		S63.	612
559	Paula Chesley. 2010. Lexical Borrowings in French:	Matias Jentoft and David Samuel. 2023. NoCoLA: The	613
560	Anglicisms as a Separate Phenomenon. <i>Journal of</i>	Norwegian Corpus of Linguistic Acceptability. In	614
561	<i>French Language Studies</i> , 20(3):231–251.	<i>Proceedings of the Nordic Conference on Computa-</i>	615
		<i>tional Linguistics</i> , pages 610–617.	616
562	Noam Chomsky. 2014. <i>Aspects of the Theory of Syntax</i> .		
563	11. MIT press.	Jacob Devlin Ming-Wei Chang Kenton and Lee Kristina	617
		Toutanova. 2019. BERT: Pre-training of Deep Bidi-	618
564	Leshem Choshen, Guy Hacohen, Daphna Weinshall,	rectional Transformers for Language Understanding.	619
565	and Omri Abend. 2022. The grammar-learning trajec-	In <i>Proceedings of NAACL-HLT</i> , pages 4171–4186.	620
566	tories of neural language models . In <i>Proceedings of</i>		
567	<i>the Annual Meeting of the Association for Computa-</i>	Taku Kudo. 2005. MeCab: Yet Another Part-Of-Speech	621
568	<i>tional Linguistics</i> , pages 8281–8297, Dublin, Ireland.	and Morphological Analyzer. http://mecab.sourceforge.net/ .	622
569	Association for Computational Linguistics.		623
570	Alexis Conneau, Kartikay Khandelwal, Naman Goyal,	Taku Kudo and John Richardson. 2018. SentencePiece:	624
571	Vishrav Chaudhary, Guillaume Wenzek, Francisco	A Simple and Language Independent Subword Tok-	625
572	Guzmán, Edouard Grave, Myle Ott, Luke Zettle-	enizer and Detokenizer for Neural Text Processing.	626
573	moyer, and Veselin Stoyanov. 2020. Unsupervised	In <i>Proceedings of the Conference on Empirical Meth-</i>	627
574	cross-lingual representation learning at scale . In <i>Pro-</i>	<i>ods in Natural Language Processing: System Demon-</i>	628
575	<i>ceedings of the Annual Meeting of the Association for</i>	<i>strations</i> , pages 66–71.	629
576	<i>Computational Linguistics</i> , pages 8440–8451, Online.		
577	Association for Computational Linguistics.	Sachin Kumar, Vidhisha Balachandran, Lucille Njoo,	630
		Antonios Anastasopoulos, and Yulia Tsvetkov. 2023.	631
578	Yiming Cui, Wanxiang Che, Ting Liu, Bing Qin, and	Mitigating societal harms in large language mod-	632
579	Ziqing Yang. 2021. Pre-training With Whole Word	els . In <i>Proceedings of the Conference on Empirical</i>	633
580	Masking for Chinese Bert. <i>IEEE/ACM Transac-</i>	<i>Methods in Natural Language Processing: Tutorial</i>	634
581	<i>tions on Audio, Speech, and Language Processing</i> ,	<i>Abstracts</i> , pages 26–33, Singapore. Association for	635
582	29:3504–3514.	Computational Linguistics.	636
583	Zsuzsanna Fagyal, Douglas Kibbee, and Frederic Jenk-	Per Kummervold, Freddy Wetjen, and Javier	637
584	ins. 2006. <i>French: A Linguistic Introduction</i> . Cam-	De La Rosa. 2022. The Norwegian Colossal Corpus:	638
585	bridge University Press.	A Text Corpus for Training Large Norwegian	639
		Language Models. In <i>Proceedings of the Language</i>	640
586	Ingo Feldhausen and Sebastian Buchczyk. 2021. Re-	<i>Resources and Evaluation Conference</i> , pages	641
587	visiting Subjunctive Obviation in French: A Formal	3852–3860.	642
588	Acceptability Judgment Study. <i>Glossa: a journal of</i>		
589	<i>general linguistics</i> . 6 (1): 59.		

643	Per E Kummervold, Javier De la Rosa, Freddy Wetjen,	Pedro Javier Ortiz Suárez, Benoît Sagot, and Laurent	696
644	and Svein Arne Brygfeld. 2021. Operationalizing	Romary. 2019. Asynchronous pipelines for process-	697
645	a national digital library: The case for a Norwegian	ing huge corpora on medium to low resource infras-	698
646	transformer model . In <i>Proceedings of the Nordic</i>	tructures . In <i>Proceedings of the Workshop on Chal-</i>	699
647	<i>Conference on Computational Linguistics</i> , pages 20–	<i>lenges in the Management of Large Corpora</i> , pages	700
648	29, Reykjavik, Iceland (Online). Linköping Univer-	9 – 16, Mannheim. Leibniz-Institut für Deutsche	701
649	sity Electronic Press, Sweden.	<i>Sprache</i> .	702
650	Yuri Kuratov and Mikhail Arkhipov. 2019. Adaptation	Kishore Papineni, Salim Roukos, Todd Ward, and Wei-	703
651	of Deep Bidirectional Multilingual Transformers for	Jing Zhu. 2002. BLEU: a Method for Automatic	704
652	Russian Language. <i>arXiv:1905.07213</i> .	Evaluation of Machine Translation. In <i>Proceedings</i>	705
653	Sara Laviosa. 2010. Corpus-Based Translation Studies	<i>of the annual meeting of the Association for Compu-</i>	706
654	15 Years On: Theory, Findings, Applications.	<i>tational Linguistics</i> , pages 311–318.	707
655	Ashley Lewis and Michael White. 2023. Mitigating	Cecile Planchon and Daniel Stockemer. 2019. Angli-	708
656	harms of LLMs via knowledge distillation for a vir-	cisms, French Equivalents, and Language Attitudes	709
657	tual museum tour guide . In <i>Proceedings of the Work-</i>	Among Quebec Undergraduates. <i>British Journal of</i>	710
658	<i>shop on Taming Large Language Models: Controlla-</i>	<i>Canadian Studies</i> , 32(12):93–118.	711
659	<i>bility in the era of Interactive Assistants!</i> , pages 31–	Irina Proskurina, Ekaterina Artemova, and Irina Pio-	712
660	45, Prague, Czech Republic. Association for Compu-	ntkovskaya. 2023. Can BERT eat RuCoLA? Topo-	713
661	tational Linguistics.	logical Data Analysis to Explain. In <i>Proceedings of</i>	714
662	Dayiheng Liu, Yu Yan, Yeyun Gong, Weizhen Qi, Hang	<i>the Workshop on Slavic Natural Language Process-</i>	715
663	Zhang, Jian Jiao, Weizhu Chen, Jie Fu, Linjun Shou,	<i>ing</i> , pages 123–137.	716
664	Ming Gong, et al. 2021. GLGE: A New General Lan-	Stefan Schweter. 2020. Italian BERT and ELECTRA	717
665	guage Generation Evaluation Benchmark. In <i>Find-</i>	Models .	718
666	<i>ings of the Association for Computational Linguistics:</i>	Ramon Marti Solano. 2021. Anglicisms and Corpus	719
667	<i>ACL-IJCNLP</i> , pages 408–420.	Linguistics: Corpus-Aided Research into the Infl-	720
668	Ilya Loshchilov and Frank Hutter. 2018. Decoupled	uence of English on European Languages. Introduc-	721
669	Weight Decay Regularization. In <i>International Con-</i>	tion.	722
670	<i>ference on Learning Representations</i> .	Taiga Someya, Yushi Sugimoto, and Yohei Oseki. 2023.	723
671	Martin Malmsten, Love Börjeson, and Chris Haffenden.	JCoLA: Japanese Corpus of Linguistic Acceptability.	724
672	2020. Playing with Words at the National Library of	<i>arXiv:2309.12676</i> .	725
673	Sweden – Making a Swedish BERT .	Delaludina Šukalić, Edina Rizvić-Eminović, and Adnan	726
674	Louis Martin, Benjamin Muller, Pedro Javier Ortiz	Bujak. 2022. A Corpus-Based Study of Anglicisms	727
675	Suárez, Yoann Dupont, Laurent Romary, Éric Ville-	Across Different Text Types of Online News. <i>Journal</i>	728
676	monte de la Clergerie, Djamel Seddah, and Benoît	<i>of French Language Studies</i> , 20(3):231–251.	729
677	Sagot. 2020. CamemBERT: a Tasty French Lan-	Masatoshi Suzuki and Ryo Takahashi. 2019.	730
678	guage Model. In <i>Proceedings of the Annual Meeting</i>	BERT base Japanese (IPA dictionary).	731
679	<i>of the Association for Computational Linguistics</i> .	https://huggingface.co/tohoku-nlp/	732
680	Brian W Matthews. 1975. Comparison of the Predicted	bert-base-japanese . Accessed: 2024-02-10.	733
681	and Observed Secondary Structure of T4 Phage	Jörg Tiedemann. 2016. OPUS-Parallel Corpora for Ev-	734
682	Lysozyme. <i>Biochimica et Biophysica Acta (BBA)-</i>	eryone. <i>Baltic Journal of Modern Computing</i> , 4(2).	735
683	<i>Protein Structure</i> , 405(2):442–451.	Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier	736
684	Alessio Miaschi, Dominique Brunato, Felice	Martinet, Marie-Anne Lachaux, Timothée Lacroix,	737
685	Dell’Orletta, and Giulia Venturi. 2023. On	Baptiste Rozière, Naman Goyal, Eric Hambro, Faisal	738
686	Robustness and Sensitivity of a Neural Language	Azhar, et al. 2023. Llama: Open and Efficient Found-	739
687	Model: A Case Study on Italian L1 Learner Errors .	ation Language Models. <i>arXiv:2302.13971</i> .	740
688	<i>IEEE/ACM Transactions on Audio, Speech, and</i>	Daniela Trotta, Raffaele Guarasci, Elisa Leonardelli,	741
689	<i>Language Processing</i> , 31:426–438.	and Sara Tonelli. 2021. Monolingual and Cross-	742
690	Vladislav Mikhailov, Tatiana Shamardina, Max	Lingual Acceptability Judgments With the Italian	743
691	Ryabinin, Alena Pestova, Ivan Smurov, and Ekate-	Cola Corpus. In <i>Findings of the Association for Com-</i>	744
692	rina Artemova. 2022. RuCoLA: Russian Corpus of	<i>putational Linguistics: EMNLP</i> , pages 2929–2940.	745
693	Linguistic Acceptability. In <i>Proceedings of the Con-</i>	Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob	746
694	<i>ference on Empirical Methods in Natural Language</i>	Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz	747
695	<i>Processing</i> , pages 5207–5227.	Kaiser, and Illia Polosukhin. 2017. Attention Is All	748
		You Need. <i>Advances in neural information process-</i>	749
		<i>ing systems</i> , 30.	750

751	Elena Volodina, Yousuf Ali Mohammed, and Julia Klezl.	Fan Yin, Quanyu Long, Tao Meng, and Kai-Wei Chang.	808
752	2021. Dalaj—a dataset for linguistic acceptability	2020. On the robustness of language encoders	809
753	judgments for swedish. In <i>Proceedings of the Work-</i>	against grammatical errors . In <i>Proceedings of the</i>	810
754	<i>shop on NLP for Computer Assisted Language Learn-</i>	<i>Annual Meeting of the Association for Computational</i>	811
755	<i>ing</i> , pages 28–37.	<i>Linguistics</i> , pages 3386–3403, Online. Association	812
		for Computational Linguistics.	813
756	Alex Wang, Yada Pruksachatkun, Nikita Nangia, Aman-	Shengyu Zhang, Linfeng Dong, Xiaoya Li, Sen Zhang,	814
757	preet Singh, Julian Michael, Felix Hill, Omer Levy,	Xiaofei Sun, Shuhe Wang, Jiwei Li, Runyi Hu,	815
758	and Samuel Bowman. 2019. Superglue: A Stickier	Tianwei Zhang, Fei Wu, et al. 2023. Instruction	816
759	Benchmark for General-Purpose Language Under-	Tuning for Large Language Models: A Survey.	817
760	standing Systems. <i>Advances in neural information</i>	<i>arXiv:2308.10792</i> .	818
761	<i>processing systems</i> , 32.		
762	Alex Wang, Amanpreet Singh, Julian Michael, Felix	Yian Zhang, Alex Warstadt, Xiaocheng Li, and	819
763	Hill, Omer Levy, and Samuel Bowman. 2018. GLUE:	Samuel R. Bowman. 2021. When do you need bil-	820
764	A multi-task benchmark and analysis platform for nat-	lions of words of pretraining data? In <i>Proceedings</i>	821
765	ural language understanding . In <i>Proceedings of the</i>	<i>of the Annual Meeting of the Association for Com-</i>	822
766	<i>EMNLP Workshop BlackboxNLP: Analyzing and In-</i>	<i>putational Linguistics and the International Joint</i>	823
767	<i>terpreting Neural Networks for NLP</i> , pages 353–355,	<i>Conference on Natural Language Processing</i> , pages	824
768	Brussels, Belgium. Association for Computational	1112–1125, Online. Association for Computational	825
769	Linguistics.	Linguistics.	826
770	Alex Warstadt and Samuel R Bowman. 2019. Linguis-	Yukun Zhu, Ryan Kiros, Richard Zemel, Ruslan	827
771	tic Analysis of Pretrained Sentence Encoders With	Salakhutdinov, Raquel Urtasun, Antonio Torralba,	828
772	Acceptability Judgments. <i>arXiv:1901.03438</i> .	and Sanja Fidler. 2015. Aligning Books and Movies:	829
773	Alex Warstadt, Amanpreet Singh, and Samuel R Bow-	Towards Story-like Visual Explanations by Watching	830
774	man. 2019. Neural Network Acceptability Judg-	Movies and Reading Books. In <i>Proceedings of the</i>	831
775	ments. <i>Transactions of the Association for Com-</i>	<i>IEEE international conference on computer vision</i> ,	832
776	<i>putational Linguistics</i> , 7:625–641.	pages 19–27.	833
777	Laura Weidinger, John Mellor, Maribeth Rauh, Conor	Dmitry Zmitrovich, Alexander Abramov, Andrey	834
778	Griffin, Jonathan Uesato, Po-Sen Huang, Myra	Kalmykov, Maria Tikhonova, Ekaterina Taktasheva,	835
779	Cheng, Mia Glaese, Borja Balle, Atoosa Kasirzadeh,	Danil Astafurov, Mark Baushenko, Artem Sne-	836
780	et al. 2021. Ethical and Social Risks of Harm From	girev, Tatiana Shavrina, Sergey Markov, Vladislav	837
781	Language Models. <i>arXiv:2112.04359</i> .	Mikhailov, and Alena Fenogenova. 2023. A Fam-	838
		ily of Pretrained Transformer Language Models for	839
		Russian .	840
782	Guillaume Wenzek, Marie-Anne Lachaux, Alexis Con-	Simona Şimon, Claudia E Stoian, Anca Dejjica-Carţiş,	841
783	neau, Vishrav Chaudhary, Francisco Guzmán, Ar-	and Andrea Kriston. 2021. The use of anglicisms	842
784	mand Joulin, and Édouard Grave. 2020. CCNet: Ex-	in the field of education: A comparative analysis	843
785	tracting High Quality Monolingual Datasets from	of Romanian, German, and French. <i>Sage Open</i> ,	844
786	Web Crawl Data. In <i>Proceedings of the Language</i>	11(4):21582440211053241.	845
787	<i>Resources and Evaluation Conference</i> , pages 4003–		
788	4012.		
789	Wikimedia Foundation. 2024. Wikimedia downloads .	A Pre-Trained Monolingual and	846
790	Thomas Wolf, Lysandre Debut, Victor Sanh, Julien	Cross-Lingual Neural Language	847
791	Chaumond, Clement Delangue, Anthony Moi, Pier-	Models Details	848
792	ric Cistac, Tim Rault, Remi Louf, Morgan Funtowicz,	We selected the models based on their performance	849
793	Joe Davison, Sam Shleifer, Patrick von Platen, Clara	and number of downloads on the HuggingFace	850
794	Ma, Yacine Jernite, Julien Plu, Canwen Xu, Teven	Hub ⁵ .	851
795	Le Scao, Sylvain Gugger, Mariama Drame, Quentin		
796	Lhoest, and Alexander Rush. 2020. Transformers:		
797	State-of-the-art natural language processing . In <i>Pro-</i>		
798	<i>ceedings of the Conference on Empirical Methods in</i>		
799	<i>Natural Language Processing: System Demonstra-</i>		
800	<i>tions</i> , pages 38–45, Online. Association for Compu-		
801	tational Linguistics.		
802	Beilei Xiang, Changbing Yang, Yu Li, Alex Warstadt,		
803	and Katharina Kann. 2021. CLiMP: A Benchmark		
804	for Chinese Language Model Evaluation. In <i>Proceed-</i>		
805	<i>ings of the Conference of the European Chapter of</i>		
806	<i>the Association for Computational Linguistics: Main</i>		
807	<i>Volume</i> , pages 2784–2790.		

⁵<https://huggingface.co/>

Dataset	Transformer Model	# Layers	# Att Hds	Hid St Dim	Tokenization	Training Dataset
CoLA	bert-base-cased (Kenton and Toutanova, 2019)	12	12	768	WordPiece (Kenton and Toutanova, 2019)	BooksCorpus (Zhu et al., 2015) and English Wikipedia (Wikimedia Foundation, 2024)
DaLAJ	bert-base-swedish-cased (Malmsten et al., 2020)	12	12	768	SentencePiece (Kudo and Richardson, 2018)	National Library of Sweden Corpus (Malmsten et al., 2020)
ITACoLA	bert-base-italian-cased (Schweter, 2020)	12	12	768	Not specified	OPUS Corpus (Tiedemann, 2016)
RuCoLA	ruBert-base (Zmitrovich et al., 2023)	12	12	768	BPE (Gage, 1994)	Various Russian corpus (i.e. Russian Wikipedia and Russian news) (Zmitrovich et al., 2023)
CoLAC	bert-base-chinese (Cui et al., 2021)	12	12	768	LTP (Che et al., 2010) and WordPiece	Chinese Wikipedia
NoCoLA	nb-bert-base (Kummervold et al., 2021)	12	12	768	WordPiece	Norwegian Colossal Corpus (Kummervold et al., 2022)
JCoLA	bert-base-japanese (Suzuki and Takahashi, 2019)	12	12	768	MeCab (Kudo, 2005) and WordPiece	Japanese Wikipedia
FrCoLA	camembert-base (Martin et al., 2020)	12	12	768	SentencePiece	OSCAR (Ortiz Suárez et al., 2019)
Cross-lingual	xlm-roberta-base (Conneau et al., 2020)	12	12	768	SentencePiece	CommonCrawl (Wenzek et al., 2020)

Table 9: Details of the pre-trained transformer models used for each linguistic acceptability corpus. For each model, it presents the number (#) of layers and attention heads (Att Hds) along with hidden states dimension (Hid St Dim), tokenization and training dataset.