

Conditioning LLMs to Generate Code-Switched Text

Anonymous ACL submission

Abstract

Code-switching (CS) is still a critical challenge in Natural Language Processing (NLP), due to the limited availability of large-scale, diverse CS datasets for robust training and evaluation. Despite recent advances, the true capabilities and limitations of LLMs in handling CS remain underexplored. In this work, we investigate the extent to which LLMs can be used in a framework for CS text generation, focusing on the English-Spanish language pair. Our proposed methodology consists of back-translating natural CS sentences into monolingual English, and using the resulting parallel corpus to fine-tune LLMs to turn monolingual sentences into CS. We thoroughly analyse the models' performance through a study on human preferences, a qualitative error analysis, an evaluation with popular reference-based metrics and LLM-based judgment. Results show that fine-tuning can be a key step to ensure that current LLMs consistently generate fluent code-switched text and that our methodology generates high-quality outputs, expanding research opportunities in CS communication. We find that traditional metrics do not correlate with human judgement when assessing the quality of the generated CS data but LLM-based judgment aligns more closely with human preferences. We release our code and generated dataset under a CC-BY-NC-SA license.¹

1 Introduction

Code-Switching (CS) consists of mixing two or more languages within a single utterance and is a common phenomenon in multilingual settings (Tucker, 2001). Although it is mainly present in spoken interactions, it can also be found in written interactions on-line (Appel and Muysken, 2005; Sarkisov, 2021), where it appears jointly with other features of informal speech. Example 1 shows an utterance where the speaker switches between English and Spanish.

- (1) Why make everybody *sentarse atrás pa' que* everybody has to move *pa' que se salga.*
Why make everybody *sit at the back so that* everybody has to move *so that she may get out.*²

(Poplack, 1980)

Despite the prevalence of code-switching, most research in Natural Language Processing (NLP) assumes monolingualism as a standard for human communication. However, this implicit decision means that state-of-the-art models are not able to properly interpret or generate CS data. Even advances in *multilingual* language modelling (Lin et al., 2022; Chowdhery et al., 2023) have not led to significant improvements, and performance on CS data is still poor compared to performance on monolingual data (Aguilar et al., 2020; Winata et al., 2021). This occurs because there is little CS text available in the multilingual pretraining data. Similarly, there are no parallel datasets available to learn to generate CS in a supervised fashion, as one would expect for tasks such as machine translation. Finally, existing methodology for evaluating automatically generated CS text, which has specific needs different from other text generation tasks, are still not good enough and fail to capture nuances of CS text (Srivastava and Singh, 2021). It is therefore crucial to develop methodologies to enable models to generate natural CS text and simultaneously implement robust evaluation frameworks that can assess how well NLP systems handle CS across multiple tasks. We argue that both of these goals require models that can conditionally generate CS from monolingual text. Consequently, our research focuses on the development of a methodology to fine-tune and evaluate LLMs on the task of CS generation, following three main research questions:

²In all examples of CS featured in this paper, Spanish parts are shown in italics, in both the original instance and its translation.

¹URL to be announced upon acceptance.

RQ1: What are the comparative strengths and limitations of fine-tuned versus non-fine-tuned LLMs in generating fluent and natural code-switched text?

RQ2: How can we leverage LLMs to create high-quality pseudo-parallel data for fine-tuning LLMs in CS text generation?

RQ3: Do automatic metrics for Natural Language Generation (NLG) or LLM judges correlate well with human judgement for the task of CS generation?

Based on these research questions, we propose a novel approach to generate CS from monolingual text using LLMs and apply it to the English-Spanish pair. We create a new parallel English-CS corpus, *EN-CS*, by leveraging natural CS data and using LLMs to perform back-translation from CS into English, resulting in high-quality pseudo-parallel pairs, suitable for training and evaluating models on CS generation (RQ2). We provide a comprehensive comparison of CS generation using LLMs in both zero-shot and fine-tuned settings, and we compare their performance against that of a dedicated machine translation (MT) model (RQ1). Finally, we evaluate our methodology both qualitatively, with a study on human preferences and a manual error analysis, and quantitatively, using automatic NLG metrics and LLM as a judge, which allows us to study the correlation between human and automatic evaluation for this task (RQ3).

2 Related Work

Perspectives in linguistics. CS naturally occurs in communities where two or more languages are in contact, making it a subject of interest to fields like sociolinguistics and psycholinguistics. From a social perspective, it can be affected by speakers’ attitudes towards the languages and the CS phenomenon itself. In this respect, it is related to notions of prestige and identity (Heredia et al., 2025). For example, in bilingual communities where a language is minoritized, CS can be seen as an intrusion of the majority language (Dewaele and Wei, 2014). However, for migrant communities, it may be a way to preserve their mother tongue and as an “emblem of ethnic identity” (Poplack, 1980). Its importance in different social contexts highlights the need to consider CS in NLP research, as it plays a crucial role in linguistic interactions and, consequently, the development of language technologies.

Datasets & benchmarks for CS. Most code-switched data stems from social media, while other popular data sources include recordings and transcriptions (Winata et al., 2023). Shared tasks using such CS data have been organized for the tasks of Language Identification (Solorio et al., 2014; Molina et al., 2016) and Sentiment Analysis (Patwa et al., 2020). Similarly, two benchmarks exist to evaluate model performance on CS text, covering different language pairs and tasks: LINCE (Aguilar et al., 2020), which covers tasks such as Part Of Speech tagging or Sentiment Analysis; and GLUE-CoS (Khanuja et al., 2020), which focuses on NLU tasks for Hindi-English. Unfortunately, GLUECoS cannot be currently used without access to the X API.

CS generation. CS generation has seldom been tackled in previous research. Approaches include linguistically informed techniques to find plausible switching points (Pratapa et al., 2018; Gupta et al., 2020; Gregorius and Okadome, 2022; Hsu et al., 2023), data augmentation (Tarunesh et al., 2021) and, more recently, prompting LLMs for CS generation (Yong et al., 2023). While CS generation is often evaluated by human annotators (Tarunesh et al., 2021; Gregorius and Okadome, 2022), there remains a need for robust automatic evaluation methodologies to assess the naturalness and fluency of the generated texts, with some recent studies already exploring approaches like Judge-LLMs (Kuwanto et al., 2024).

3 Parallel Data Creation

In this work we present a novel approach to generate code-switched text from monolingual sentences. As a first step, we create a synthetic parallel corpus from an initial set of English-Spanish CS sentences from the LINCE benchmark (Aguilar et al., 2020) with their English monolingual equivalents, generated by the Command R model (Cohere For AI, 2024). We exploit the fact that LLMs struggle to generate CS text given a monolingual sentence (c.f. Section 5), but are able to more reliably convert a CS sentence to its corresponding monolingual version, especially when the target language is English. After having created this pseudo-parallel corpus, we use it to fine-tune LLMs on the task of conditional code-switching generation, presented in Section 4.

3.1 The LINCE benchmark

We use LINCE as a starting point, a popular benchmark that has been widely used to evaluate CS systems (Aguilar et al., 2020), which is available in 6 language pairs. All sentences in LINCE are tokenized, and each token is annotated with a language tag as well as other categories depending on the task. In our work we focus on the English-Spanish pair and filter all sentences in the data that do not contain CS, similarly discarding all the task-specific annotations. Example 2 shows a random instance from LINCE.

(2) estaba aquí three feet away .
spa spa eng eng eng eng&spa

LINCE comprises around 95,000 train, 20,000 development, and 33,000 test instances for the English-Spanish pair. We deduplicate the instances among splits, and filter and pre-process the instances to ensure that they are suitable for our task by removing links, replacing usernames with the placeholder <user>, and detokenizing all instances with the script provided as part of the Moses toolkit (Koehn et al., 2007). After this preprocessing, we obtain a more natural version of the LINCE data. A preliminary analysis reveals that many sentences in LINCE are monolingual or contain a single word in one language that often correspond to a borrowing, as shown in Example 3. In order to ensure that all of our sentences actually contain CS, we filter sentences that do not have at least two words in each language.

(3) I need a shot of tequila or a glass of scotch to keep me warm right now.

After these pre-processing and filtering steps, we end up with 12,933 train, 2,461 development and 5,353 test instances. The comparison between the original size of LINCE and the final number of sentences selected for our experiments after pre-processing is shown in Table 1.

3.2 EN-CS

The next step in our method requires creating a pseudo-parallel English-CS dataset by translating the natural code-switched instances into monolingual text. As there are no available machine translation systems to convert from English-Spanish CS text to English monolingual text, we instead make use of prompt engineering, using the Command R

	Train	Dev	Test
Original	94,728	19,574	33,361
Pre-processed	12,933	2,461	5,353
EN-CS	10,703	791	1,040

Table 1: Size of original LINCE (EN-ES) compared to the automatically filtered instances and the final set of parallel instances, dubbed EN-CS.

model (Cohere For AI, 2024), one of the strongest publicly available models at the time.

We perform an initial set of experiments to determine the optimal prompt to generate monolingual English versions of the code-switched data. Ideally, we aim for a prompt that generates translations that maintain the meaning of the original sentences, are fluent and natural, whose grammar is correct, and that do not contain any Spanish words or phrases. After extensive testing (see Appendix A), we use the following prompt in a 5-shot setting: *Now convert this code-switched phrase to English. Leave the parts in English as they are, focus on translating the parts in Spanish.* Finally, we filter output instances that contain profanity that was not present in the source texts or irrelevant information, such as “Of course, here’s your translation:”.

In order to create a valid gold standard test set, we perform a manual post-edition of the monolingual test translations for 1,040 instances of the LINCE test set. The post-edition was carried out by three proficient speakers of English and Spanish, who were provided with specific guidelines as shown in Appendix B.

Table 1 shows the final size of the parallel corpus, which we dub EN-CS, after post-processing and post-edition, and Table 2 shows examples of silver and gold instances. The final version of our dataset therefore contains 10,703 train and 791 development instances with automatically translated English sentences matched to their original CS sentences, and 1,040 gold instances with post-edited English translations.

3.2.1 Quality assessment

We evaluate the quality of the automatic translations (train/dev) by measuring two dimensions: the overall *fluency* of the sentences and *adequacy* of the translations in respect to the source texts. Two fluent English-speaking annotators evaluate the same 100 random instances using a 5-point Likert scale (Callison-Burch et al., 2007) and obtain

	Original	English
Silver	you just have to tell me <i>que como te va.</i>	You just have to tell me <i>how it's going.</i>
	osea i know we wanna party <i>pero tampoco no aya asta dallas</i>	like i know we want to party <i>but not all the way to dallas</i>
Gold	<i>hasta venir a plaza se siente</i> like home.	<i>even coming to the square feels</i> like home.
	<i>me siento tan pendejo</i> right now.	<i>i feel so stupid</i> right now.

Table 2: Examples of the *EN-CS* parallel corpus. Left: original code-switched instances, right: generated (silver) or post-edited (gold) English instances.

4.6 (fluency) and 4.5 (adequacy) points on average, which show the quality of the generated translations. A quadratic Cohen’s κ of 0.57 indicates moderate agreement, likely due to lower Likert scores (1, 2, and 3) being rarely selected by the annotators, which is a known problem for κ (Xu and Lorber, 2014; Barnes et al., 2025). In fact, the raw agreement between annotators is substantially higher: 0.71 for fluency and 0.65 for adequacy. See Appendix C for further details.

To estimate the quality of the post-edition process, we compare the post-editions of two annotators on 100 additional random instances. The results show that 75% of the sentences remain unchanged, as they are already adequate. There is a 87.87% similarity between the post-editions of the two annotators, as measured by Levenshtein distance, demonstrating a high degree of consistency and quality in the post-edition process.

4 CS generation experiments

With *EN-CS* as our starting point, we frame CS generation as a machine translation task, with English as the source and CS as the target language, where parts of the source sentence have to be translated to Spanish. In our experiments, we fine-tune four small-sized generative models, namely, Llama3 8B, Llama3 Instruct 8B (Dubey et al., 2024), Mistral 7B and Mistral Instruct 7B (Jiang et al., 2023).

To fine-tune the models, we use the causal language modelling objective, but with appropriate input formats for the base and instruct models. For base models we use templates (Zhu et al., 2024) in the form of “<X>=<Y>”, where <X> and <Y> are placeholders for the input English sentence and generated CS, respectively. At inference, the second code-switched part is left empty for the model to fill. For instruction-tuned models, we provide a system prompt with the instruction, a query by the user in English, and an answer from the assistant with the code-switched target. At inference time, the answer is left blank (See Table 7 in Appendix D

for example prompts).

All models are trained using Quantized Low-Rank Adaptation (QLoRA) (Dettmers et al., 2023) with standard parameters: the model is loaded in 4 bit with NF4 quantization data type and bf16 computational data type. The LoRA rank and scaling factor are set to 16 and the dropout to 0.05. We apply the LoRA update matrices to the attention blocks and do not train bias parameters. Regarding the hyperparameters, we only tune the learning rate ($1e^{-4}$, $5e^{-4}$, $1e^{-3}$ and $5e^{-3}$) and training epoch $\in [1 \dots 10]$, choosing the parameters that give the lowest cross-entropy loss on the development set for each model. We use the transformers package (Wolf et al., 2020) for all training experiments.

Early experiments indicated that fine-tuned models usually produce the desired output up to a punctuation mark and then either begin to translate the sentence again or hallucinate more content. We therefore truncate the output up to a punctuation mark where the length is closest to that of the original sentence (Bawden and Yvon, 2023).

As baselines, we include few-shot experiments by directly prompting GPT-4o and Llama3.3 70B Instruct to generate CS text using a 5-shot approach (see Table 7 in the Appendix). We refer to these systems in the experiments as GPT-4o_{fs} and Llama3.3-70B_{fs}, respectively. We also include a strong dedicated MT baseline, developed by fine-tuning the NLLB (team et al., 2022) model (nllb-200-distilled-600M) using *EN-CS*. The model was trained with standard settings.³

Table 3 shows an example of the outputs of the different models, compared to the original code-switched sentence, and the English monolingual sentence that they received as input.

³A batch size of 32, learning rate of 1×10^{-4} , using constant learning rate schedule with 1,000 warmup steps, a gradient clipping threshold of 1.0, and a weight decay of 1×10^{-3} . Training was conducted for 50,000 steps. For evaluation, we selected the checkpoint that achieved the highest BLEU score on the development set.

Model	Generated Output
Original (Gold)	dammm <i>todos se casaron</i> and we still single lol forever alone
English (Source)	damn everyone got married and we're still single lol forever alone
Llama3	damn <i>todos se fueron a casarse y nosotras estamos solitarias</i> lol forever alone
Llama3 Instruct	dammm every1 got married and we're still single lol <i>alonso solit@o</i> foreveerrrr lolololo
Mistral	dammm everyone got married y <i>ni somos esa</i> lol ofrever alone
Mistral Instruct	dude <i>todo el mundo se casó y nosotros hemos quedado solteros</i> lol
Llama3.3-70B _{fs}	damn <i>todo el mundo se casó y nosotros seguimos solteros</i> lol forever alone
GPT-4o _{fs}	damn <i>todos se casaron y nosotros seguimos solteros</i> lol forever alone
NLLB	damn everyone got marry and its still single lol forever alone

Table 3: Example from the test set and the generated outputs of the different models.

5 Qualitative evaluation

As a first step to assess the quality of the outputs produced by the different models, we perform a manual qualitative analysis of the results in two parts: a pairwise tournament-based human evaluation, and an in-depth analysis of the most common errors made by the models and their distribution.

5.1 Preference based evaluation

We perform a tournament-based evaluation that allows us to determine the ranking of models in terms of human preference. A total of 880 instances are matched against each other, corresponding to the outputs of the seven models for 110 English source sentences, as well as the gold standard reference. The evaluation is conducted pairwise, requiring annotators to choose the best out of two sentences or declare a tie. When choosing the best sentence, annotators do not know the original English sentence, nor which model produced what output. This process results in $110 \cdot \binom{8}{2} = 3,080$ comparisons, and was carried out by 14 annotators, with each annotator performing at most 300 random comparisons.

Annotators are provided with a series of criteria to choose between the instances, based on the error analysis described in the next section. They must take into account three main criteria, which must be applied in the following order: a) the presence and naturalness of the CS; b) the content and fluency of the sentences; and c) the orthographical errors of the instances (correct punctuation, presence of typos, etc.). Annotators are furthermore asked to avoid declaring ties, unless completely necessary (e.g., in a case where both sentences are completely monolingual and therefore equally in-

Model	Human		GPT	
	Score	Rank	Score	Rank
Gold Standard	525.0		484.5	
Llama3	423.0		374.0	4
Llama3 Instruct	393.0		367.5	5
NLLB	391.0	4	304.0	8
Mistral	368.0	5	334.5	6
GPT-4o _{fs}	366.5	6	478.0	
Llama3.3-70B _{fs}	307.0	7	428.0	
Mistral Instruct	306.5	8	309.5	7

Table 4: Ranking of models according to human preference (c.f. Section 5.1) and using GPT as a judge (c.f. Section 6.2).

correct) to compel them to develop a preference. The complete annotation guidelines are available in Appendix E. Inter-annotator agreement on a subset of 100 sentence pairs shows substantial agreement ($\kappa = 0.74$).

We calculate a global score for each model, as follows: every time a model is voted, it gets 1 point, and the loser gets 0 points; in case of ties, both models get 0.5 points each. Table 4 shows the global scores, as well as the ranking of human preferences according to said score (second and third columns). We find that the gold standard reference obtains the highest score, as expected, and that fine-tuned Llama3 ranks the highest among the automatic methods. Instruction-tuned models obtain worse scores compared to their base model counterparts, with a similar difference for both Llama3 and Mistral family of models. The NLLB model ranks higher than the Mistral models, but lower than the fine-tuned Llama3 models, showing their potential to rival dedicated models in generation tasks. According to these preferences, fine-tuning

LLMs for CS generation can be critical to ensure better results, since the larger models with few-shot prompting rank only above Mistral Instruct, with GPT-4o_{fs} outranking Llama3.3-10B_{fs}.

5.2 Error analysis

In order to further explore differences between model performance, we analyse the most common errors made by the CS generation models, both quantitatively and qualitatively. We extend the machine translation error typology presented in Popović (2018) to CS generation error analysis. To do so, we randomly select a set of 100 outputs from all models and conduct a detailed examination of the types of errors present in them. This thorough analysis allows us to identify recurring patterns and propose a refined error typology specifically for automatic CS generation. This initial error analysis yields 18 total error categories, which we simplify and group into three main error types: a) CS errors, b) Translation Errors, and c) Format errors. The full error typology, along with detailed descriptions for each error type, is provided in Appendix F, while here we explain the three error categories:

CS Errors: Errors of sentences that are either completely monolingual or switch between languages in an unnatural manner, e.g., by repeating the same word in English and Spanish. In Example 4, Llama3 Instruct preserves the original meaning, but the sentence is fully monolingual.

- (4) **Source** After all these things when we're done.
Output after all these things when we're finished

Translation errors: Critical errors that either change the original meaning of the sentence or introduce mistakes in fluency or grammar, for example, using the wrong tense or word order. Example 5 shows an instance where Mistral Instruct outputs a seemingly natural code-switched sentence, but the phrase “they got hurt” is not adequately translated and the meaning of the sentence is not preserved.

- (5) **Source** I wasn't happy because they got hurt.
Output no estuve happy porque me dieron mal

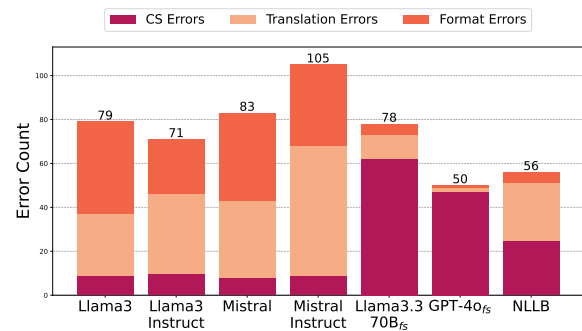


Figure 1: Error distribution by model, obtained by counting the number of instances that present errors of each type.

Format errors: Errors in form that do not make the sentences unintelligible nor change their meaning, such as repetitions of a word or phrase or incorrect punctuation. Example 6, by the model Llama3, accurately preserves the original meaning and introduces CS, but removes the username and adds a smiley face.

- (6) **Source** <user> old mexican remedies
Output old school *remedios mexicanos* :)

We classify 700 additional instances (100 instances per model, obtained from the same source sentences) into these kind of errors, and show the results in Figure 1.

GPT-4o_{fs} makes the fewest errors overall (50), closely followed by NLLB (56). However, 90% of GPT’s errors (45) and 45% on NLLB’s (25) are CS related, indicating that while these systems preserve the meaning of sentences and generate few formatting errors, they often produce entirely monolingual outputs, which is a critical error. In comparison, CS-related mistakes are the least common in fine-tuned LLMs, accounting for less than 15% of the overall error count. This analysis shows that fine-tuned LLMs have effectively learned to switch between languages naturally, though they may still be prone to other less critical types of errors.

Among the fine-tuned LLMs, Llama models present 19 fewer errors on average than the Mistral family. Base models of both families struggle mainly with format errors, which make up 50.68% of their errors on average, whereas instruction-tuned models present more meaning-related issues, 53, 45%. This suggests that the linguistic knowledge of the models degrades when tuned on instruc-

Model	BLEU	BERTScore	chrF
Llama3 8B	<u>34.49</u>	81.64	<u>53.17</u>
Llama3 8B Instruct	33.42	81.77	52.01
Mistral	31.65	80.93	50.56
Mistral Instruct	25.98	78.66	44.58
Llama3.3-70B _{fs}	22.41	79.77	44.57
GPT-4o _{fs}	32.25	<u>83.09</u>	50.48
NLLB	35.56	84.11	54.74
Identity	33.34	82.31	45.51

Table 5: Results of reference-based metrics the *EN-CS* test set. Best results in bold, second best results underlined.

tions, a phenomenon that has been observed on other related areas (Fu et al., 2024).

6 Automatic Evaluation

Previous research highlights the challenges of automatically evaluating code-switching (CS) generation, with many existing metrics showing low correlation with human judgments (Srivastava and Singh, 2021; Kuwanto et al., 2024). On the other hand, recent studies show that using LLMs as evaluators or judges can offer a promising alternative to evaluate generation tasks, as they show a higher alignment with human ratings (Chiang and Lee, 2023; Wang et al., 2023). In this section, we present the results of an automatic evaluation with reference-based NLG metrics (BLEU, BERTScore, chrF) and GPT-4o as a judge, as well as their correlation with human preferences.

6.1 Reference-based metrics

We report the results of BLEU (Papineni et al., 2002), BERTScore (Zhang et al., 2020), and chrF (Popović, 2015), implemented with the *evaluate* library. All three are task-agnostic quality metrics that give results between 0-1, based on character-level F-score, n-gram precision and semantic similarity using contextual embeddings⁴ respectively. We compute the metrics for all systems, and include an Identity system that simply returns the provided input as the output.

The results of the evaluation can be seen in Table 5. The best model is NLLB, with the highest scores for the three metrics. It is closely followed by GPT-4o_{fs}, with the second highest BERTScore, and fine-tuned Llama3, with the second highest BLEU and chrF. They are closely followed by Llama3

⁴BERTscore has been calculated using the embeddings from the model *Bert Base Multilingual Cased*.

Instruct and Mistral. The Identity system scores nearly as well as the top-performing models. The strong results from the Identity baseline and few-shot models, which often produce monolingual outputs, suggest that reference-based metrics assign high scores to models that match only the English part of the reference. This reflects the nature of the task and dataset, and highlights the limitations and artifacts of using reference-based metrics to evaluate code-switched generation.

6.2 GPT as a judge

As a complementary automatic assessment of the outputs of our models, we have implemented a zero-shot pair-wise evaluation using GPT-4o as a judge, mimicking the settings of the human evaluation. Details about the implementation are included in Appendix G. Results are shown in the third and fourth columns of Table 4. GPT shows a strong preference for few-shot models, whereas these models are ranked third- and second-to-last by humans. Based on the error analysis (c.f Section 5.2), few-shot models tend to make many CS-related mistakes, although they are the most fluent. Thus, this disagreement may be caused by the humans adhering to the guidelines and taking the presence of CS as the main criterion, whereas GPT is making decisions based on the style and fluency of the answers. Regarding fine-tuned LLMs, they all share the same relative ranking in both evaluations. Finally, NLLB is ranked last by GPT, while it is the fourth best model overall as ranked by humans. Further research is needed to explain this behaviour, but there may be some stylistic features of the NLLB model’s outputs that are affecting GPT’s preferences.

6.3 Correlation With Human Evaluation

The reference-based metrics used in Section 6.1 are known to have weak correlations with human judgment in NLG tasks (Sai et al., 2022), whereas JudgeLLM-based evaluations seem promising (Chiang and Lee, 2023; Wang et al., 2023). In this section, we compare reference-based metrics and GPT scores with the preference-based scores obtained in Section 5.1.

We calculate Pearson’s (ρ) correlation coefficient at instance-level, using the 700 instances employed for the error classification and human evaluation (the output of 7 models for 100 source sentences).⁵

⁵We do not consider the reference CS sentences when

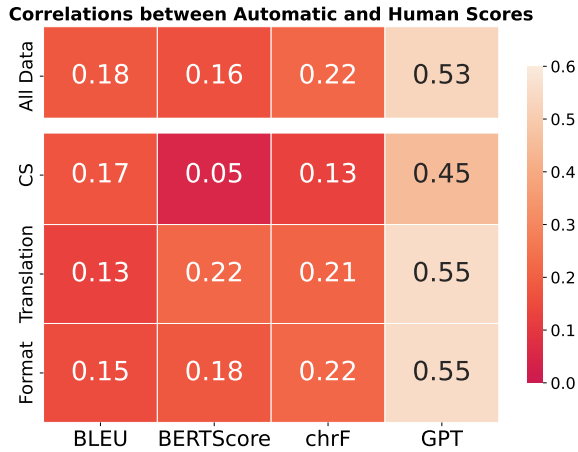


Figure 2: Heatmap of the correlations between human scores and reference-based metrics and scores given by GPT, calculated using the Pearson Correlation Coefficient. The correlations are calculated for all instances, as well as for different subsets of instances, according to the type of errors they exhibit.

Each data point corresponds to the CS output of one particular model for an English source sentence, and we compute the correlation using two values: the score obtained by the model for this instance in the human preference-based evaluation of Section 5.1, and the score it attains if we apply the same strategy using the values of the reference-based metrics to determine the winner, or, in case of JudgeLLM, the scores given by GPT.

The correlation coefficients are shown in Figure 2. The top part of the figure shows the correlation using all the instances, whereas the bottom part only considers those instances that showed some type of error, according to the error analysis described in Section 5.2.

If we consider all the instances, the maximum ρ correlation value with reference-based metrics is 0.22, which indicates a low alignment with human scores. The metric with highest correlation is chrF, in line with previous research (Popović, 2015). GPT-4o_{fs} shows a ρ of 0.53, which is stronger than reference-based metrics, but still too weak to be regarded as a reliable measure for assessing CS generation.

If we instead consider instances with errors from the human evaluation, there is again a higher correlation between human scores and GPT’s judgments, with a margin of at least 0.28 points. Instances with CS errors show the lowest overall correlation. This likely derives from the fact that hu-

calculating the correlations.

man evaluators never prefer an instance without CS as instructed in the guidelines, but reference-based and JudgeLLM-based metrics are not sensitive to these nuances, and may assign high scores to instances regardless of whether they contain CS or not.

All in all, these results confirm that several of the most commonly used reference-based metrics for NLG have a weak correlation with human judgments when evaluating CS generation. This underscores the need to research more specialized evaluation methods designed specifically to capture the nuances of this task that correlate with human judgement.

7 Conclusion

In this work, we have presented a methodology to leverage LLMs in the generation of code-switched text from monolingual instances, specifically for the English-Spanish language pair.

Our framework consists of back-translating natural code-switched instances (EN-ES) into monolingual English sentences, and using the resulting parallel corpus, dubbed *EN-CS*, to fine-tune autoregressive models to translate monolingual sentences into CS. This has the advantage of ensuring that the target sentences contain completely natural CS, which has the potential to improve the naturalness of CS generation.

We experiment with fine-tuning base and instruction-tuned LLMs on our dataset using LoRA. For baselines, we include few-shot LLMs (Llama3.3-70B and GPT-4o) and a pretrained NLLB translation system that we also finetune using our dataset. The results indicate that fine-tuned LLMs show higher ranking in a human preference-based evaluation and fewer critical errors than the other baselines, performing better even than proprietary models such as GPT-4o.

We also perform a meta-evaluation of reference-based NLG metrics commonly used for CS evaluation, as well as an LLM judge (GPT-4o). Our analyses show low correlation between human and reference-based evaluations, while the LLM judge achieves moderate correlations. However, particularly in cases with CS errors, no metric is adequate for assessing CS generation. We therefore advocate for more research in specialized evaluation methods.

Limitations

Our research focuses on testing the capabilities of LLMs for CS generation, a field of interest in the research of many applications, yet still in need of more research. While our findings highlight promising potential, we also identify key areas for refinement and improvement, as well as promising lines for future research in this domain.

We only perform an in-domain evaluation where the train, validation and test sets had the same origin. Additionally, we would like to test the efficiency of our models in an out-of-domain setting, since one of the use-cases of a CS generation model is to create parallel corpora to evaluate the abilities of models to perform different tasks when there is CS.

We want to acknowledge the fact that our approach is dependent on having an initial set of code-switched sentences, which may not be available for all pairs of languages, especially in a low-resource scenario. We believe that it would be interesting to explore the possibility of a cross-lingual approach using our methodology, with English and/or Spanish as pivot languages, that could be useful for transfer knowledge into other less-resourced language pairs.

Finally, as we have pointed out, we are aware of the problems of the automatic metrics that we have used to evaluate the outputs of our models, that do not capture the nuances of our task. In the future, we would like to investigate how to improve this evaluation by designing new methods to automatically evaluate CS generation, focusing on a more linguistic approach able to capture the linguistic and social intricacies of CS.

References

- Gustavo Aguilar, Sudipta Kar, and Tamar Solorio. 2020. [LinCE: A centralized benchmark for linguistic code-switching evaluation](#). In *Proceedings of the Twelfth Language Resources and Evaluation Conference*, pages 1803–1813, Marseille, France. European Language Resources Association.
- Rene Appel and Pieter C. Muysken. 2005. *Language Contact and Bilingualism*. Amsterdam University Press.
- Jeremy Barnes, Naiara Perez, Alba Bonet-Jover, and Begoña Altuna. 2025. [Summarization metrics for spanish and basque: Do automatic scores and llm-judges correlate with humans?](#)
- Rachel Bawden and François Yvon. 2023. [Investigating the translation performance of a large multilingual language model: the case of BLOOM](#). In *Proceedings of the 24th Annual Conference of the European Association for Machine Translation*, pages 157–170, Tampere, Finland. European Association for Machine Translation.
- Chris Callison-Burch, Cameron Fordyce, Philipp Koehn, Christof Monz, and Josh Schroeder. 2007. [\(meta-\) evaluation of machine translation](#). In *Proceedings of the Second Workshop on Statistical Machine Translation*, pages 136–158, Prague, Czech Republic. Association for Computational Linguistics.
- Cheng-Han Chiang and Hung-yi Lee. 2023. [Can large language models be an alternative to human evaluations?](#) In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 15607–15631, Toronto, Canada. Association for Computational Linguistics.
- Aakanksha Chowdhery, Sharan Narang, Jacob Devlin, Maarten Bosma, Gaurav Mishra, Adam Roberts, Paul Barham, Hyung Won Chung, Charles Sutton, Sebastian Gehrmann, Parker Schuh, Kensen Shi, Sasha Tsvyashchenko, Joshua Maynez, Abhishek Rao, Parker Barnes, Yi Tay, Noam Shazeer, Vinodkumar Prabhakaran, Emily Reif, Nan Du, Ben Hutchinson, Reiner Pope, James Bradbury, Jacob Austin, Michael Isard, Guy Gur-Ari, Pengcheng Yin, Toju Duke, Anselm Levskaya, Sanjay Ghemawat, Sunipa Dev, Henryk Michalewski, Xavier Garcia, Vedant Misra, Kevin Robinson, Liam Fedus, Denny Zhou, Daphne Ippolito, David Luan, Hyeontaek Lim, Barret Zoph, Alexander Spiridonov, Ryan Sepassi, David Dohan, Shivani Agrawal, Mark Omernick, Andrew M. Dai, Thanumalayan Sankaranarayanan Pillai, Marie Pellat, Aitor Lewkowycz, Erica Moreira, Rewon Child, Oleksandr Polozov, Katherine Lee, Zongwei Zhou, Xuezhi Wang, Brennan Saeta, Mark Diaz, Orhan Firat, Michele Catasta, Jason Wei, Kathy Meier-Hellstern, Douglas Eck, Jeff Dean, Slav Petrov, and Noah Fiedel. 2023. [Palm: Scaling language modeling with pathways](#). *Journal of Machine Learning Research*, 24(240):1–113.
- Cohere For AI. 2024. [c4ai-command-r-v01 \(revision 8089a08\)](#).
- Tim Dettmers, Artidoro Pagnoni, Ari Holtzman, and Luke Zettlemoyer. 2023. [Qlora: Efficient finetuning of quantized llms](#).
- Jean-Marc Dewaele and Li Wei. 2014. [Attitudes towards code-switching among adult mono- and multilingual language users](#). *Journal of Multilingual and Multicultural Development*, 35(3):235–251.
- Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Amy Yang, Angela Fan, et al. 2024. The llama 3 herd of models. *arXiv preprint arXiv:2407.21783*.

855	Ananya B. Sai, Akash Kumar Mohankumar, and	Genta Winata, Alham Fikri Aji, Zheng Xin Yong, and	913
856	Mitesh M. Khapra. 2022. A survey of evaluation	Thamar Solorio. 2023. The decades progress on code-	914
857	metrics used for nlg systems . <i>ACM Comput. Surv.</i> ,	switching research in NLP: A systematic survey on	915
858	55(2).	trends and challenges . In <i>Findings of the Association</i>	916
859	E. Sarkisov. 2021. Interlingual interference as a linguistic	for Computational Linguistics: ACL 2023 , pages	917
860	and cultural characteristic of the current online	2936–2978, Toronto, Canada. Association for Com-	918
861	communication. <i>Russian Journal of Bilingualism</i>	putational Linguistics.	919
862	<i>Studies</i> , 3:16–21.		
863	Thamar Solorio, Elizabeth Blair, Suraj Mahar-	Genta Indra Winata, Samuel Cahyawijaya, Zihan Liu,	920
864	jan, Steven Bethard, Mona Diab, Mahmoud	Zhaojiang Lin, Andrea Madotto, and Pascale Fung.	921
865	Ghoneim, Abdelati Hawwari, Fahad AlGhamdi, Julia	2021. Are multilingual models effective in code-	922
866	Hirschberg, Alison Chang, and Pascale Fung. 2014.	switching? In <i>Proceedings of the Fifth Workshop</i>	923
867	Overview for the first shared task on language iden-	on Computational Approaches to Linguistic Code-	924
868	tification in code-switched data . In <i>Proceedings of</i>	Switching , pages 142–153, Online. Association for	925
869	<i>the First Workshop on Computational Approaches to</i>	Computational Linguistics.	926
870	<i>Code Switching</i> , pages 62–72, Doha, Qatar. Associa-		
871	tion for Computational Linguistics.		
872	Vivek Srivastava and Mayank Singh. 2021. Challenges	Thomas Wolf, Lysandre Debut, Victor Sanh, Julien	927
873	and limitations with the metrics measuring the com-	Chaumond, Clement Delangue, Anthony Moi, Pier-	928
874	plexity of code-mixed text . In <i>Proceedings of the</i>	ric Cistac, Tim Rault, Rémi Louf, Morgan Funtowicz,	929
875	<i>Fifth Workshop on Computational Approaches to Lin-</i>	Joe Davison, Sam Shleifer, Patrick von Platen, Clara	930
876	<i>guistic Code-Switching</i> , pages 6–14, Online. Associa-	Ma, Yacine Jernite, Julien Plu, Canwen Xu, Teven Le	931
877	tion for Computational Linguistics.	Scao, Sylvain Gugger, Mariama Drame, Quentin	932
		Lhoest, and Alexander M. Rush. 2020. Transform-	933
878	Ishan Tarunesh, Syamantak Kumar, and Preethi Jyothi.	ers: State-of-the-art natural language processing . In	934
879	2021. From machine translation to code-switching:	<i>Proceedings of the 2020 Conference on Empirical</i>	935
880	Generating high-quality code-switched text . In <i>Pro-</i>	<i>Methods in Natural Language Processing: System</i>	936
881	<i>ceedings of the 59th Annual Meeting of the Associa-</i>	<i>Demonstrations</i> , pages 38–45, Online. Association	937
882	<i>tion for Computational Linguistics and the 11th Inter-</i>	for Computational Linguistics.	938
883	<i>national Joint Conference on Natural Language Pro-</i>		
884	<i>cessing (Volume 1: Long Papers)</i> , pages 3154–3169,	Shu Xu and Michael F. Lorber. 2014. Interrater agree-	939
885	Online. Association for Computational Linguistics.	ment statistics with skewed data: evaluation of alter-	940
		natives to cohen’s kappa . <i>Journal of consulting and</i>	941
886	Nllb team, Marta Ruiz Costa-jussà, James Cross, Onur	<i>clinical psychology</i> , 82(6):1219—1227.	942
887	cCelebi, Maha Elbayad, Kenneth Heafield, Kevin		
888	Heffernan, Elahe Kalbassi, Janice Lam, Daniel Licht,	Zheng Xin Yong, Ruochen Zhang, Jessica Forde, Skyler	943
889	Jean Maillard, Anna Sun, Skyler Wang, Guillaume	Wang, Arjun Subramonian, Holy Lovenia, Samuel	944
890	Wenzek, Alison Youngblood, Bapi Akula, Loïc Bar-	Cahyawijaya, Genta Winata, Lintang Sutawika, Jan	945
891	rault, Gabriel Mejia Gonzalez, Prangthip Hansanti,	Christian Blaise Cruz, Yin Lin Tan, Long Phan,	946
892	John Hoffman, Semarley Jarrett, Kaushik Ram	Long Phan, Rowena Garcia, Thamar Solorio, and	947
893	Sadagopan, Dirk Rowe, Shannon L. Spruit, C. Tran,	Alham Fikri Aji. 2023. Prompting multilingual large	948
894	Pierre Yves Andrews, Necip Fazil Ayan, Shruti	language models to generate code-mixed texts: The	949
895	Bhosale, Sergey Edunov, Angela Fan, Cynthia Gao,	case of south East Asian languages . In <i>Proceedings</i>	950
896	Vedanuj Goswami, Francisco Guzmán, Philipp	<i>of the 6th Workshop on Computational Approaches to</i>	951
897	Koehn, Alexandre Mourachko, Christophe Rop-	<i>Linguistic Code-Switching</i> , pages 43–63, Singapore.	952
898	pers, Safiyyah Saleem, Holger Schwenk, and Jeff	Association for Computational Linguistics.	953
899	Wang. 2022. No language left behind: Scal-		
900	ing human-centered machine translation . <i>ArXiv</i> ,	Tianyi Zhang, Varsha Kishore, Felix Wu, Kilian Q.	954
901	abs/2207.04672.	Weinberger, and Yoav Artzi. 2020. Bertscore: Evalu-	955
902	G Richard Tucker. 2001. A global perspective on bilin-	ating text generation with bert .	956
903	gualism and bilingual education. <i>Georgetown Uni-</i>		
904	<i>versity Round table on Languages and Linguistics</i>	Wenhao Zhu, Hongyi Liu, Qingxiu Dong, Jingjing Xu,	957
905	<i>1999</i> .	Shujian Huang, Lingpeng Kong, Jiajun Chen, and	958
		Lei Li. 2024. Multilingual machine translation with	959
906	Jiaan Wang, Yunlong Liang, Fandong Meng, Zengkui	large language models: Empirical results and anal-	960
907	Sun, Haoxiang Shi, Zhixu Li, Jinan Xu, Jianfeng Qu,	ysis . In <i>Findings of the Association for Computa-</i>	961
908	and Jie Zhou. 2023. Is ChatGPT a good NLG evalua-	<i>tional Linguistics: NAACL 2024</i> , pages 2765–2781,	962
909	tor? a preliminary study . In <i>Proceedings of the 4th</i>	Mexico City, Mexico. Association for Computational	963
910	<i>New Frontiers in Summarization Workshop</i> , pages	Linguistics.	964
911	1–11, Singapore. Association for Computational Lin-		
912	guistics.	A Prompt-tuning for CS-EN translation	965
		For CS→EN translation of the LINCE benchmark,	966
		we test the prompts in Table 6, combined with 0-,	967
		1- and 5-shot strategies. The prompts include the	968

instructions explained in different ways, including more or less information.

For the few-shot strategies, the prompt includes the following template at the beginning, alongside a set of manually selected examples that are representative of some phenomena we want to cover in our prompt:

Here are {n} examples of a code-switched text that has been converted to {lang}:

Testing the different prompts, we are able to choose the one whose outputs are closest to our needs, taking into consideration the trade-off between including too little and too much level of specificity in the instructions to the models.

Regarding the few-shot strategies, we find out that giving some examples to the models results in outputs that are more aligned with the expected output, which is logical, since this allows the models to more faithfully replicate the examples provided. The more examples given, the more the model is able to comply to leaving the punctuation marks as they are and not standardizing the spelling, but also it tends to add more colloquial terms and alternate spellings.

B Post-edition Guidelines

The original sentence should contain **CS** and be **translatable**. The main reasons to **remove** an instance altogether are:

- If the sentence is very clearly monolingual and the CS has been detected incorrectly (eg, the case of interlingual homographs such as *has*).
- When the sentence is bilingual for metalinguistic reasons, because it makes the translation tricky and hard to understand, and in most cases it's not even CS.
- The part that is in the other language is a named entity, such as a title, a name, ...
- If the code-switched part is not translatable or very hard to translate, probably because it's a borrowing. Ambiguous and a little bit up to the annotator.
- If the tweet is saying the same thing in both languages (making it monolingual doesn't make sense).

- Some instances are tweets that are part of a conversation or thread and taken out of context are very hard to understand/intelligible.
- Some tweets are not translatable because of wordplay that doesn't transfer to monolingual speech.

The result should be a **monolingual** sentence that has roughly the **same meaning** as the original sentence. The main reasons to edit a translation are:

- If the meaning changes or the model has hallucinated extra information that wasn't present in the original sentence.
 - If there are still some words in the Spanish.
 - Attempts to translate named entities.
 - Remove "meta comments" from the model about the task.
- It is not necessary to correct things like:
- Punctuation marks.
 - Different spellings of the same word.
 - Words or phrases that the model has changed for synonyms.

C Inter-annotator agreement

Figure 3 shows the distribution of the scores given by to annotators on the fluency and adequacy of the instances translated by Command R. Although the inter-annotator agreement shows moderate agreement ($\kappa = 0.57$), the distributions between the annotators are very similar to each other.

D Fine-tuning / Few-shot prompting

In Table 7, we can see the prompt used for a) fine-tuning of base models; b) fine-tuning of instruction-tuned models; and c) 5-shot prompting. For both fine-tuning settings, at inference time the second part of the prompt that contains the target CS sentence is left blank for the model to complete.

E Pairwise Annotation Guidelines

The main objective of this task is to evaluate a pair of sentences that **should contain code-switching between English and Spanish**. It should be noted that models have been trained with texts extracted from social media and informal conversations, so

Convert this code-switched phrase to English.
Convert this code-switched phrase to English without correcting the original spelling, focus on translating the parts in Spanish.
Convert this code-switched phrase to English. Leave the parts in Spanish as they are, focus on translating the parts in Spanish.
Convert this code-switched phrase to English. Directly output the translation and don't correct the original spelling, focus on translating the parts in Spanish.

Table 6: Different prompts that have been used to convert the code-switched instances into English, with different levels of specificity. Final prompt in bold.

Base model
<i>I want to not work and make money. = quiero no trabajar and make money</i>
Instruction-tuned model
system prompt: "You are a bilingual speaker of English and Spanish. Translate the following English sentence into code-switched text between both languages:"
user: "I want to not work and make money."
assistant: "quiero no trabajar and make money"
Few-shot prompting
system prompt: "You are a bilingual speaker of English and Spanish. Translate the following English sentence into code-switched text between both languages. Do not add any comments or explanations:"
user: Source example n
assistant: Target example n
user: "I want to not work and make money."
assistant: ...

Table 7: Examples and format of prompts used for fine-tuning base and instruction-tuned models and for few-shot prompting

the outputs of the models are expected to present traits of informality, such as common typos, that at first should not be considered errors, because they are within the expected behaviour of the models. The criteria to choose between both sentences is to be applied in the following order:

1. Code-switching

- 1.1. **Presence of code-switching:** For a sentence to be a suitable candidate it must have tokens in both languages. A completely monolingual sentence will always be wrong.
- 1.2. **Naturalness of the code-switching:** A switch between both languages can be unnatural. There are different linguistic constraints. For example, a switch is only

possible at a point in a sentence where it does not violate the syntactic rules of either language.

2. Content and fluency

- 2.1. **Content:** Sentences must have meaning as a whole, they have to be understandable, without extra content disconnected from the rest of the message or abrupt interruptions.
- 2.2. **Agreement:** Sentences must have the right gender and number agreement.
- 2.3. **Conjugation:** Verbs have to be correctly conjugated.

3. Form: Additional errors that can be used in case none of the above are applicable.

- 3.1. **Repetitions** of the same word or phrase.
- 3.2. **Misspelled words / uncommon typos**
- 3.3. **Wrong punctuation marks**
- 3.4. **Extra characters**

Ties are only contemplated in two situations:

- Two sentences that are **equally wrong**, that is to say, they are both either completely monolingual or unintelligible.
- Two sentences that are **exactly the same** and thus no criteria can be used to break the tie.

In case no criteria is applicable to a pair, we ask the annotators to choose their preferred sentence, using their own judgement or additional criteria they might observe in the specific pair of sentences.

F Error Typology

1. CS errors

- 1.1. **No CS** - the sentence is entirely monolingual.

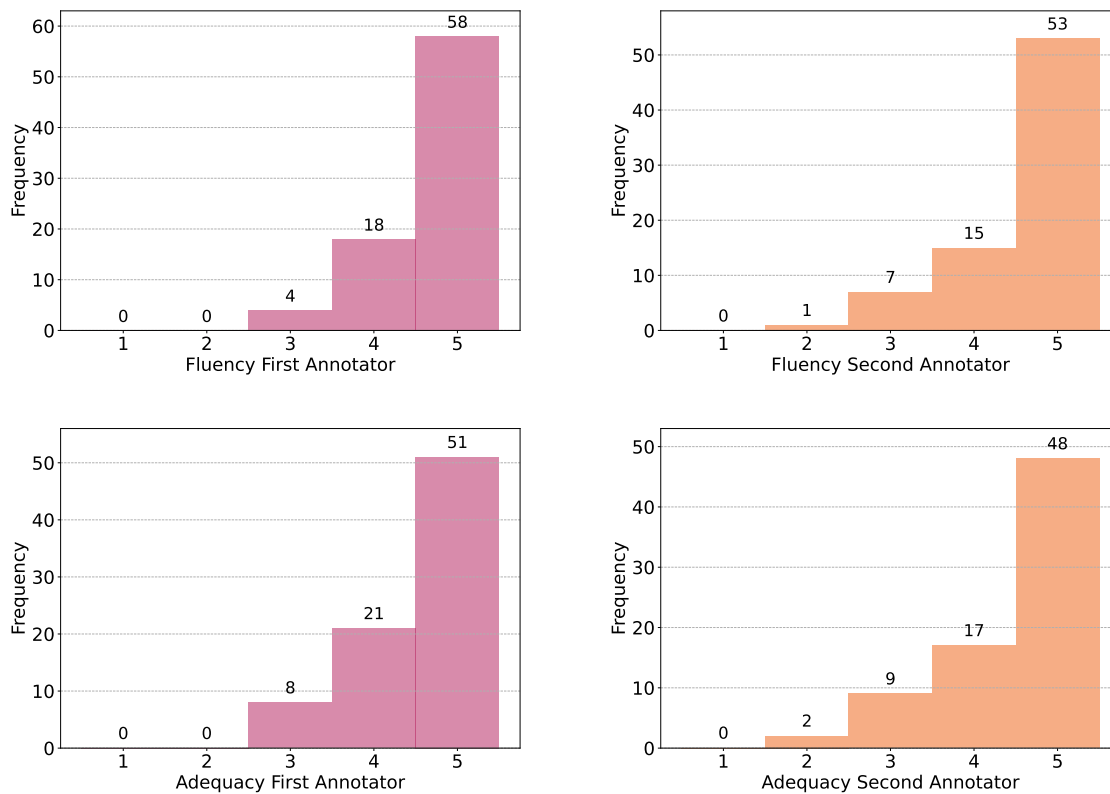


Figure 3: Distribution of Adequacy and Fluency scores per annotator.

- 1.2. **Unnatural CS** - the sentence contains unnatural CS, either due to unnatural switching points, or unnatural register.
- 1.3. **Repetition in both languages** - the sentence contains the same information repeated in both languages, rather than CS.

2. Translation errors

- 2.1. **Made-up words** - the words in the output look like English or Spanish but do not actually exist.
- 2.2. **Wrong translation** - the translation of a word or phrase is incorrect.
- 2.3. **Wrong conjugation** - a verb is translated with the right lexeme but a seemingly made-up conjugation.
- 2.4. **Wrong agreement** - there is a mistake in agreement in gender or number.
- 2.5. **Wrong meaning** - a word or phrase has been translated into a sense that does not fit into the context.
- 2.6. **Wrong order** - the words are right but they are written in the wrong order.
- 2.7. **Wrong tense** - the verbal tense is not consistent through the sentence.

- 2.8. **Unintelligible** - it is not possible to understand the sentence in English nor in Spanish.
- 2.9. **Instruction misunderstanding** - the task has been misunderstood, e.g., the model makes a "comment" about the content of the output or explains a word.

3. Format errors

- 3.1. **Extra words** - the sentence contains seemingly random extra words that do not affect its meaning.
- 3.2. **Extra characters** - the sentence contains more non-word characters than the original, e.g., '???' instead of '??'.
- 3.3. **Hallucinations** - the sentence contains new words or phrases not derived from the original text.
- 3.4. **Start over** - the sentence is finalized, but the model begins a second translation of the same sentence.
- 3.5. **Duplications** - some words or phrases of the sentence are duplicated.

developer: "You are a helpful bilingual system that knows how to code-switch between English and Spanish and how to distinguish natural sentences. Your only job is to judge sentences and output a verdict A, B or T."

user:"Which one of the next two automatically generated sentences with code-switching is more natural? The most important criterion is that the sentences must have code-switching to even be considered eligible. Two sentences can be tied if they are equally wrong.

A: {s1}

B: {s2}

Answer(A/B/T):"

Table 8: Prompt used for GPT to act as a judge.

G Implementation of GPT as judge

For the implementation of GPT as judge, the developer and user prompts in Table 8 have been used to prompt GPT-4o. To calculate the scores, we first check the answers the directly contain the desired format, "A", "B or "T", which are the most common. For the rest of the outputs that did not follow this format, it is possible to extract the labels using simple regular expressions. Then, the scores are calculated just like the human score: every time a model is voted, it gets 1 point, and the loser gets 0 points; in case of ties, both models get 0.5 points each.