

In-Context Learning Strategies for Low-Resource Bidirectional Translation between Mexican Sign Language Glosses and Spanish

Gabriel Alejandro Mantilla Clavijo^{1,†}, Melany Marcela Saez Acuña^{1,*,†},
Aarón Dali López Fortich^{1,†}, Juan Carlos Martínez Santos^{1,†} and
Edwin Alexander Puertas Del Castillo^{1,†}

¹Universidad Tecnológica de Bolívar, Colombia

Abstract

We describe the system developed by team VerbaNexAI for the MSLG-SPA 2026 shared task on bidirectional translation between Mexican Sign Language glosses (MSLG) and Spanish (SPA), organised within IberLEF 2026. Operating in a low-resource regime of 490 training pairs, we explore three in-context learning (ICL) strategies over two model families: retrieval-augmented few-shot generation with curriculum ordering of examples (RAG+curriculum, $k=10$), many-shot in-context learning with the full training corpus embedded in the system prompt, and open-source local replication of both strategies using OpenEuroLLM-Spanish via Ollama. All configurations share a corpus-derived prompt (PROMPT_BASE) that encodes the orthographic and morphosyntactic conventions of the shared-task corpus inductively, rather than from existing academic grammars. On SPA→MSLG, the RAG+curriculum configuration with Claude Haiku 4.5 achieves the best results (BLEU 0.268, METEOR 0.593, chrF 0.690), outperforming the many-shot variant and confirming that 10 semantically relevant examples surpass 490 uniform ones for rule-governed compression tasks. On MSLG→SPA, the many-shot configuration obtains the highest COMET score (0.901), while RAG+curriculum leads on lexical metrics (BLEU 0.407, METEOR 0.682, chrF 0.749). The open-source replication recovers 66.8–91.0% of Claude performance with no external API dependency, providing a reproducible baseline for future work on low-resource sign language machine translation.

Keywords

Artificial Intelligence, gloss translation, in-context learning, retrieval-augmented generation, many-shot prompting, curriculum ordering, Mexican Sign Language

1. Introduction

Mexican Sign Language (*Lengua de Señas Mexicana*, LSM) is the primary means of communication for an estimated 100,000 to 400,000 Deaf and hard-of-hearing individuals in Mexico [1]. Despite its grammatical richness and social relevance, LSM remains a severely low-resource language from a natural language processing standpoint: annotated parallel corpora are scarce, standardized orthographic conventions for its gloss representation vary across sources, and the computational tools available to its community lag far behind those developed for spoken languages. Bridging this gap is not merely a technical challenge—it is a matter of linguistic equity and digital inclusion for a historically underserved population.

Automatic translation between LSM glosses and Spanish addresses this gap from two complementary directions. In the **SPA→MSLG** direction, a Spanish utterance is compressed into a canonical gloss sequence following the morphosyntactic conventions of LSM: elimination of articles and prepositions, reordering of constituents, explicit perfectivity marking, and a set of productive morphological operations (gender suffixation, reduplication, compounding). In the reverse **MSLG→SPA** direction, a gloss sequence must be expanded into fluent, grammatically well-formed Spanish by recovering the

IberLEF 2026, September 2026, León, Spain

*Corresponding author.

[†]These authors contributed equally.

✉ gmantilla@utb.edu.co (G. A. M. Clavijo); msaez@utb.edu.co (M. M. S. Acuña); forticha@utb.edu.co (A. D. L. Fortich); jcmartinez@utb.edu.co (J. C. M. Santos); epuerta@utb.edu.co (E. A. P. D. Castillo)

🆔 0009-0008-2896-1117 (G. A. M. Clavijo); 0009-0001-1811-2217 (M. M. S. Acuña); 0009-0007-7275-8418 (A. D. L. Fortich); 0000-0003-2755-0718 (J. C. M. Santos); 0000-0002-0758-1851 (E. A. P. D. Castillo)



© 2026 Copyright for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

elements that gloss notation suppresses. Both transformations are highly structured yet sensitive to subtle corpus-specific orthographic conventions that diverge from academic grammars [2], making them a challenging testbed for prompt-based neural approaches.

The MSLG-SPA 2026 shared task, organized within the IberLEF 2026 framework [3, 4], proposes exactly this bidirectional evaluation over a parallel corpus of 490 training pairs and 245 test instances per subtask. Translations are evaluated with BLEU [5], METEOR [6], and chrF [7]; COMET [8] is additionally used for MSLG→SPA, where the target is a natural language.

In this paper, team VerbaNexAI describes the system developed and submitted to both subtasks of the shared task. Our approach is grounded in *in-context learning* (ICL) with large language models (LLMs), specifically targeting the low-resource regime where fine-tuning is impractical due to the limited size of the available corpus. We explore three ICL strategies: (i) *few-shot* retrieval-augmented generation (RAG) with curriculum ordering of examples, (ii) *many-shot* in-context learning with the full training set embedded in the system prompt, and (iii) open-source local LLM replication of the two previous strategies. These configurations are evaluated across two models: Claude Haiku 4.5 [9] (Anthropic) and the open-source OpenEuroLLM-Spanish [10] served locally via Ollama, resulting in six official submissions across both subtasks.

A central design decision in our work is the construction of a *corpus-derived* prompt (PROMPT_BASE) that encodes the orthographic and morphosyntactic conventions of the shared-task corpus rather than those of existing academic grammars. This prompt is augmented with a curated glossary, a set of negative examples targeting observed error patterns, and—crucially—a dynamic retrieval mechanism that selects the ten most semantically similar training pairs for each test instance and reorders them according to a curriculum principle: examples are arranged from simplest to most complex, placing the hardest instance immediately before the query to exploit the recency bias of transformer-based LLMs [11].

Across six official submissions, the RAG with curriculum ordering configuration using Claude Haiku 4.5 achieves BLEU of **0.268** on SPA→MSLG and BLEU of **0.407** on MSLG→SPA (internal validation, *val*=90, *seed*=42). On the latter subtask, the many-shot submission obtains the highest COMET score of **0.901**, while remaining competitive in lexical metrics. The open-source local replication recovers approximately 67–70% of the Claude performance with no external API dependency, offering a reproducible baseline for future work.

The remainder of this paper is organized as follows. Section 2 describes the task and corpus. Section 3 details the architecture, prompt design, and the three ICL strategies implemented by team VerbaNexAI. Section 4 reports ablation results and the official submission scores. Section 5 analyzes the key findings. Section 6 concludes with limitations and directions for future work.

2. Task and Data Description

2.1. Shared Task Overview

The MSLG-SPA 2026 shared task [3], organized by CICESE within the IberLEF 2026 framework [4], proposes bidirectional translation between Mexican Sign Language glosses (MSLG) and Spanish (SPA). It is structured as two fully independent subtasks evaluated separately:

- SPA→MSLG: given a Spanish sentence, produce its canonical gloss sequence in LSM notation.
- MSLG→SPA: given an LSM gloss sequence, produce a fluent and grammatically well-formed Spanish sentence.

The independence of the two subtasks reflects a fundamental asymmetry: gloss notation suppresses morphosyntactic elements that are obligatory in Spanish (articles, prepositions, verb conjugation), and reconstructing them in the reverse direction requires a qualitatively different type of inference than the forward compression task.

2.2. Corpus

The organizers distribute a single parallel corpus in TSV format with columns ID, MSLG, and SPA. The training file contains 490 sentence pairs covering a wide range of everyday topics and communicative functions. The test sets contain 245 instances each, provided as TSV files with ID and the corresponding source field only.

For internal validation we applied a reproducible split with seed 42: 400 instances as the training pool and 90 instances as the held-out validation set. All six official submissions use the full 490-pair pool for retrieval indexing and many-shot example construction, so no training data is withheld at inference time.

2.3. Linguistic Characteristics of the Corpus

Beyond the general morphosyntactic properties of LSM described in the literature [2], the shared-task corpus exhibits a specific set of orthographic and notational conventions that are critical for prompt design and that diverge, in several points, from academic grammar descriptions. We identified the following corpus-specific features through inductive analysis of the training data:

- Proper-noun prefix *dm-*. Personal names and toponyms that are fingerspelled are systematically prefixed with *dm-* (e.g. *dm-ISABEL*, *dm-DIEGO-RIVERA*). This convention is not uniform in academic LSM descriptions but is consistent throughout the corpus.
- Human feminine suffix *+MUJER*. Grammatical gender for human referents is marked by appending *+MUJER* to the base gloss (e.g. *HERMANO+MUJER* for *hermana*, *MAESTRO+MUJER* for *maestra*). Failure to apply this transformation is one of the most frequent error types observed during system development.
- Acronym and media prefix *#*. Signs for acronyms, institutions, and media names are prefixed with *#* (e.g. *#TV*, *#SEP*).
- Perfective aspect particle *YA*. Completed actions are marked with the particle *YA*, which may appear either before or after the verb depending on focus. The distinction between *YA* as an aspect marker and *YA* as a discourse particle requires context-sensitive handling.
- Object pronoun *MI* vs. possessive *MI*. The corpus maintains a critical orthographic distinction: *MI* (with accent mark, *MÍ*) is the first-person object pronoun, while *MI* (without accent) is the possessive determiner. Confusing the two produces grammatically deviant glosses.
- Single-token *PORQUE*. The interrogative *por que* is written as a single uppercase token *PORQUE* (not *POR-QUE* as in some academic conventions) and consistently appears sentence-finally.
- Productive reduplication. Plural, intensity, and habitual aspect are encoded through reduplication of the base gloss (e.g. *NINO NINO* for *los niños*, *TRABAJO TRABAJO* for habitual working). The number of repetitions and the precise semantic value vary by context.
- Lexicalized compounds with hyphen. Multi-sign lexical units and idiomatic phrases are written as hyphenated compounds (e.g. *CINTURON-DE-SEGURIDAD*, *NO-TE-HE-VISTO*). This includes all negative verb compounds (*NO-PODER*, *NO-HAY*).
- Coordinated-subject closure with *AMBOS*. When two noun phrases jointly function as the subject of a predicate, the gloss sequence closes the subject constituent with *AMBOS* before the verb.

These features informed the design of *PROMPT_BASE* (Section 3.2) and the *NEGATIVE_EXAMPLES* component, which explicitly targets the most error-prone conventions listed above.

2.4. Evaluation Metrics

Both subtasks are evaluated with three standard machine translation metrics computed over the full test set:

- BLEU [5]: corpus-level *n*-gram precision with brevity penalty, computed with smoothing method 1 [12].

- METEOR [6]: sentence-level F -score with synonym and stemming matching, averaged over the test set.
- chrF [7]: character n -gram F -score, more robust to morphological variation than word-level BLEU.

For MSLG→SPA only, a fourth metric is added:

- COMET [8, 13]: a neural reference-based metric trained on human judgments of translation quality, using the `Unbabel/wmt22-comet-da` checkpoint. COMET is restricted to MSLG→SPA because gloss sequences do not constitute natural language output, and applying a fluency-oriented neural metric to them would be methodologically unsound.

In our internal evaluation pipeline, BLEU, METEOR, and chrF are computed using NLTK [14]; COMET is invoked via the `comet-score` CLI (`unbabel-comet` package, Python 3.11 isolated environment).

3. System Description

This section describes the complete pipeline developed by team VerbaNexAI for both subtasks. All six submitted configurations share a common infrastructure (Section 3.1), a corpus-derived prompt at their core (Section 3.2), and three ICL strategies that differ only in how training examples are selected and presented to the model (Sections 3.3–3.5).

3.1. Common Infrastructure

Language models. Primary configurations use `claude-haiku-4-5-20251001` [9] (`temperature=0.1`, `max_tokens=1024`). Open-source replication runs use `jobautomation/OpenEuroLLM-Spanish` [10] served locally via Ollama with an identical parameter budget. Explicit reasoning mode is disabled (`think=False`) for the local model, as verbose chain-of-thought degrades performance on short deterministic transformations such as gloss translation (see Section 3.2 for supporting evidence).

Retrieval index. Example retrieval relies on `paraphrase-multilingual-MiniLM-L12-v2` [15, 16], a 384-dimensional multilingual sentence encoder (≈ 118 M parameters). An `EmbeddingIndex` computes cosine similarities between a query and all training instances. The indexed field is determined by subtask direction: Spanish sentences for SPA→MSLG, gloss sequences for MSLG→SPA.

Prompt caching. Anthropic’s prompt caching [17] marks the system block with `cache_control: ephemeral`, activating when the block exceeds 1024 tokens. The many-shot configurations (Section 3.4) benefit most: the first request processes ≈ 50 k input tokens; the remaining 244 requests read from cache at $\approx 10\%$ of that cost. RAG-based configurations (Section 3.3) are not cached because the system block changes per query.

Post-processing. Each subtask requires a direction-specific post-processing step applied to every model output before writing the submission file.

For SPA→MSLG, the post-processor: (i) selects the first non-empty output line; (ii) strips any model preamble such as `MSLG:` or `Traducción::`; (iii) removes surrounding quotes and non-permitted punctuation, preserving `-`, `+`, `?`, and `!`; (iv) converts the entire string to uppercase; and (v) filters isolated articles that do not form part of a lexical compound (i.e. not enclosed by `-` or `+`). Steps (iv) and (v) enforce the corpus conventions described in Section 2.3.

For MSLG→SPA, a symmetric variant preserves Spanish diacritics and natural punctuation without uppercasing, since the target is a natural-language sentence.

Submission format. Each output file contains exactly 245 lines in the same order as the official test file, UTF-8 encoded with Unix line endings. Every line follows the format "SystemOutput", validated with `grep -cvP '^[^"]*$' → 0` malformed lines across all six submissions.

3.2. Corpus-Derived Prompt Design

The central component shared by all configurations is `PROMPT_BASE`, a structured rule block assembled through inductive analysis of the training corpus. Rather than adopting the orthographic conventions of existing academic grammars [2], we derived rules directly from the corpus, because the shared-task data follows its own dialect — in particular, the conventions described in Section 2.3 (`dm-` prefix, `+MUJER` suffix, single-token `PORQUÉ`, etc.) diverge from those in Cruz-Aldrete [2] in ways that directly affect translation accuracy.

`PROMPT_BASE` encodes nine rule blocks:

1. Constituent order. Default SVO; SOV when the object is focused or with demonstrative verbs; time, frequency, and location markers are topicalized to sentence-initial position.
2. Aspect and tense. Perfective aspect marked with `YA` (pre- or post-verbal); future with `FUTURO`; proximal future with `PRÓXIMO` plus a temporal marker.
3. Deletions. Articles, prepositions (except in lexical compounds or purposive phrases), and copulas are removed; possessives are retained.
4. Obligatory transformations. Verb stems appear uppercased in their infinitive form; proper nouns receive the `dm-` prefix; acronyms and media names the `#` prefix; human feminine referents the `+MUJER` suffix; lexicalized compounds and idiomatic phrases are hyphenated. All conventions follow Section 2.3.
5. Reduplication. Encodes plural, intensity, or habitual aspect by repeating the base gloss (e.g. `NIÑO NIÑO`, `TRABAJO TRABAJO`).
6. Pronouns. Object pronoun `MÍ` (with accent) versus possessive `MI` (without); third-person clitic `ÉL/ ELLA` post-posed to the verb.
7. Negation. Negative verbs are formed as lexicalized hyphenated compounds (`NO-PODER`, `NO-HAY`, `NO-TE-HE-VISTO`).
8. Interrogatives. `QUIÉN` appears sentence-initially; `CUÁNTO`, `CUÁL`, and `PORQUÉ` sentence-finally; `PORQUÉ` as a single token.
9. Coordination and subordination. `AMBOS` closes coordinated subjects; `PERO`, `OJALÁ`, `QUIZÁ` are retained verbatim; conditionals by juxtaposition or with `IMAGINAR`.

Three supplementary components complement `PROMPT_BASE`:

- **LSM_GLOSSARY**: a compact SPA→LSM reference table covering pronouns, tense markers, frequency adverbs, aspect particles, interrogatives, quantifiers, colors, acronyms, high-frequency verbs, negative compounds, dual forms, kinship terms, and age expressions.
- **NEGATIVE_EXAMPLES**: eight error triples (`SPA`, `WRONG`, `CORRECT`, `error-type`) derived from mistakes observed during iterative development, targeting the most error-prone conventions from Section 2.3: omission of `YA`, failure to apply `+MUJER`, incorrect placement of temporal markers, and confusion of `MÍ`/`MI`.
- **COT_INSTRUCTIONS**: an eight-step chain-of-thought scaffold [18] explored during development but excluded from all final submissions. Preliminary experiments showed degradation at every few-shot level tested (e.g. zero-shot-CoT BLEU 0.091 vs. zero-shot-rules BLEU 0.080 on the internal split), consistent with Sprague et al. [19], who show that CoT benefits are concentrated in mathematical and symbolic reasoning tasks and can harm short deterministic transformations.

Reverse prompt (MSLG→SPA). For the inverse subtask, PROMPT_BASE is replaced by a symmetric prompt that instructs the model to perform the inverse of each rule in Section 2.3: remove dm- and # prefixes; expand +MUJER to the Spanish feminine form; interpret reduplication as plural or intensification; reintroduce articles and copulas; reorder constituents to canonical Spanish SVO; conjugate verbs using YA or FUTURO as tense anchors; expand hyphenated compounds; and move interrogative words to sentence-initial position. The model is instructed to return a single capitalised line with natural Spanish punctuation and no enclosing quotes or labels. FIXED_EXAMPLES and NEGATIVE_EXAMPLES reuse the same training pairs with SPA/MSLG fields swapped, with negative cases targeting reverse-direction errors (over-literal glossing, missing verb conjugation, spurious uppercase output).

3.3. Few-Shot RAG with Curriculum Ordering

Figure 1 illustrates the four-step pipeline used for submissions 7.1, 7.2, 7.3, and 7.4. For each test instance q , the pipeline executes the following steps:

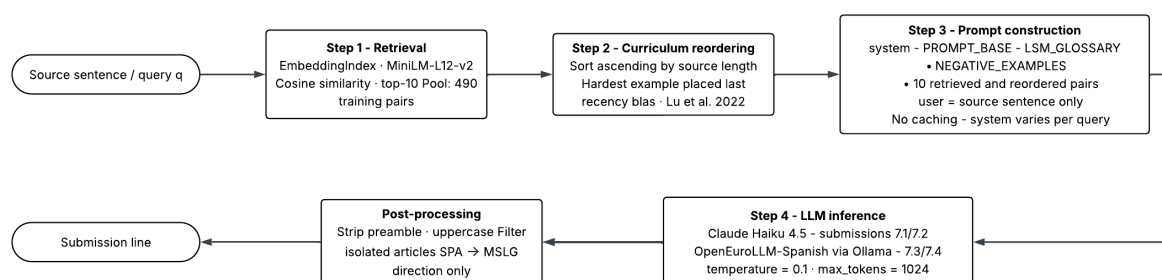


Figure 1: RAG + curriculum ordering pipeline (submissions 7.1, 7.2, 7.3, 7.4). For each source sentence q , the top-10 most similar training pairs are retrieved by cosine similarity, reordered from simplest to most complex (hardest last), assembled into the system prompt together with PROMPT_BASE, and fed to the LLM. The system block varies per query, so prompt caching is not applicable.

1. Retrieval. `EmbeddingIndex.retrieve( $q$ ,  $k=10$ )` returns the ten training pairs most similar to q by cosine distance in the multilingual embedding space [20]. The pool for official submissions is the full 490-pair training set (Section 2.2).
2. Curriculum reordering. Retrieved examples are sorted in ascending order of source-side character length. This places the most complex example immediately adjacent to the query, exploiting the recency bias of transformer-based LLMs [11] and the prompt-order sensitivity reported by Lu et al. [21], while following the curriculum learning principle of Bengio et al. [22].
3. Prompt construction. The system message is assembled as PROMPT_BASE + LSM_GLOSSARY + NEGATIVE_EXAMPLES + the ten retrieved and reordered example pairs. The user message contains only the source sentence.
4. Inference and post-processing. Model output passes through the direction-specific post-processor (Section 3.1).

Because the system block changes for every query, Anthropic’s prompt caching does not apply to this configuration.

3.4. Many-Shot In-Context Learning

Figure 2 illustrates the pipeline used for runs 8 (SPA→MSLG) and 8.1 (MSLG→SPA), both using Claude Haiku 4.5. Motivated by Agarwal et al. [23], who show that large models continue to benefit from hundreds of in-context examples, this configuration replaces dynamic retrieval with a single fixed system block containing all 490 training pairs:

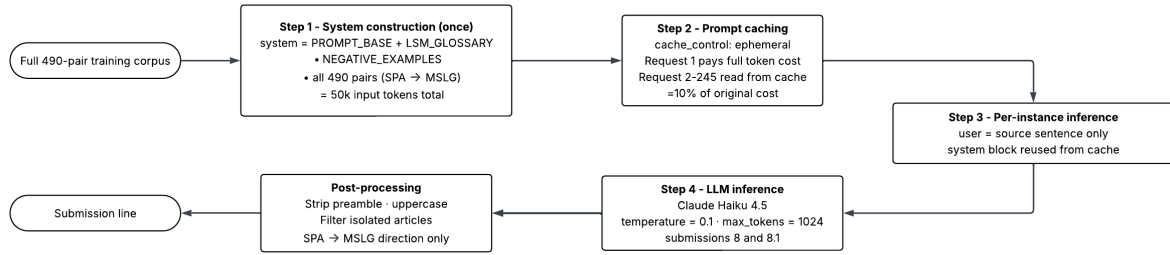


Figure 2: Many-Shot In-Context Learning pipeline (runs 8 and 8.1). The full 490-pair training corpus is loaded once into the system block together with `PROMPT_BASE`, `LSM_GLOSSARY`, and `NEGATIVE_EXAMPLES` (≈ 50 k input tokens). Prompt caching makes the strategy economically viable: only the first request pays the full token cost; requests 2–245 read from cache at $\approx 10\%$ of the original cost. For each test instance, only the user message changes; the system block is reused from cache.

1. **System construction.** The system message is assembled as `PROMPT_BASE` + `LSM_GLOSSARY` + `NEGATIVE_EXAMPLES` + the concatenation of all 490 training pairs in the format SPA: "... " $\rightarrow$ MSLG: "... " (or the reverse for submission 8.1).
2. **Prompt caching.** The block is marked with `cache_control: ephemeral`. The first request processes ≈ 50 k input tokens; subsequent requests read from cache at $\approx 10\%$ of that cost, making the strategy viable for the 245-instance test set.
3. **Per-instance inference.** Only the user message varies across test instances; the system block is reused from cache.
4. **Post-processing.** Identical to Section 3.1.

Many-shot provides uniform coverage of all training patterns but cannot adapt the example set to the specific characteristics of each query, unlike the RAG strategy. The trade-off between these two approaches is analysed in Section 5.

3.5. Open-Source Local Replication

Submissions 7.3 and 7.4 replicate the RAG+curriculum pipeline (Section 3.3) using OpenEuroLLM-Spanish served locally via Ollama, with identical prompts, retrieval index, curriculum ordering, post-processing, split, and seed. The only differences are the inference backend (local REST API instead of Anthropic’s cloud endpoint) and the absence of prompt caching.

This controlled replication serves a single analytic purpose: isolating the contribution of the prompt design from the contribution of the underlying model. Any performance gap between submissions 7.3/7.4 and 7.1/7.2 is attributable to the model, not to the pipeline. Results and their interpretation are presented in Section 4.

Reproducibility. The full pipeline, prompt components, experimental configurations, and evaluation scripts described in this paper are publicly available at <https://github.com/MelanySaez/mslg-spa-2026> [24].

4. Experiments and Results

4.1. Internal Validation Setup

All configurations were evaluated using a reproducible split of the training corpus: 400 instances as the retrieval pool and 90 instances as the held-out validation set, with seed 42 and a deterministic shuffle via `random.shuffle`. The best-performing configurations use the full 490-pair pool at inference time, so no training data is withheld (Section 2.2).

Metrics are computed as described in Section 2.4: BLEU with smoothing method 1 [12], sentence-level METEOR averaged over the validation set, corpus-level chrF, and COMET via comet-score CLI for MSLG→SPA only.

To assess sensitivity to the validation split, we additionally ran the six primary configurations with seeds 123 and 456, reporting mean and standard deviation over the three partitions in Section 4.3. Bootstrap resampling ($n=1,000$) over sentence-level scores was applied to obtain 95% confidence intervals, reported in Section 4.4.

4.2. Ablation Study

We conducted a systematic ablation on the SPA→MSLG direction to motivate each design decision. All ablation runs use the internal validation split (val=90, pool=400) and qwen2.5:14b via Ollama as the base model, except the final row which uses Claude Haiku 4.5. The first two rows are non-LLM baselines included for reference. Results are reported in Table 1.

Table 1

Ablation study on SPA→MSLG (internal validation, val=90, seed=42). All runs except the last use qwen2.5:14b via Ollama. The first two rows are non-LLM baselines included for reference. Bold marks the best value in each column.

Configuration	BLEU	METEOR	chrF
retrieval-only ($k=1$)	0.010	0.100	0.211
enfoque1 (spaCy + rules)	0.164	0.472	0.629
zero-shot-rules	0.080	0.517	0.637
zero-shot-cot	0.091	0.516	0.639
zero-shot-full	0.116	0.463	0.621
few-shot-5-rules	0.163	0.555	0.663
few-shot-10-rules	0.209	0.589	0.679
few-shot-10-curriculum	0.228	0.602	0.688
few-shot-10-diverse (k-means)	0.151	0.565	0.655
few-shot-5-rag	0.204	0.585	0.681
few-shot-10-rag-curriculum (Claude Haiku 4.5)	0.268	0.593	0.690

The ablation reveals eight consistent patterns that informed the final system design:

1. LLM generation substantially outperforms retrieval alone. The retrieval-only baseline (BLEU 0.010, chrF 0.211) confirms that simply returning the nearest training gloss provides negligible translation quality. The rule-based pipeline (enfoque1, BLEU 0.164) offers a stronger non-LLM reference; all LLM configurations with $k \geq 5$ exceed it, with the best Claude configuration adding +10.3 BLEU points and the local model (run 7.3) adding +1.5 BLEU points over this threshold.
2. Few-shot $\succ$ zero-shot across all configurations, including those augmented with NEGATIVE_EXAMPLES or CoT scaffolding.
3. CoT degrades performance. zero-shot-cot improves marginally over zero-shot-rules in BLEU, but few-shot with CoT consistently underperforms few-shot without CoT at all k values tested, consistent with Sprague et al. [19].
4. Curriculum ordering helps. Sorting examples from simplest to most complex (few-shot-10-curriculum) outperforms fixed-order few-shot-10-rules by 1.9 BLEU points.
5. RAG $\succ$ fixed examples. Retrieving the 10 nearest neighbours (few-shot-5-rag) outperforms the fixed-example baseline at $k=5$ in all metrics, confirming that local semantic relevance matters more than broad coverage.
6. RAG + curriculum is strictly best. Combining dynamic retrieval with curriculum ordering achieves the highest scores across all metrics, and the gap widens further when using Claude Haiku 4.5 instead of the local model.

7. Diversity sampling does not help. few-shot-10-diverse (k-means cluster representatives [25]) underperforms few-shot-10-rag by 5.3 BLEU points, indicating that surface diversity does not compensate for reduced semantic relevance in this small corpus.
8. Model quality matters. The best local configuration (few-shot-10-rag-curriculum on qwen2.5:14b) reaches BLEU 0.228; migrating the same prompt to Claude Haiku 4.5 raises it to 0.268 (+4.0 points).

4.3. Main Results

Tables 2 and 3 report results for all evaluated configurations. Metrics are expressed as mean $\pm$ std over seeds {42, 123, 456} (val=90 each), except for runs 8.2 and 8.3 which were evaluated on seed=42 only.

Table 2

Results for SPA $\rightarrow$ MSLG configurations. Primary runs reported as mean $\pm$ std over seeds {42, 123, 456} (val=90 each). Bold marks the best mean per column. CH = Claude Haiku 4.5; OE = OpenEuroLLM-Spanish; RAG = RAG+curriculum $k=10$; MS = Many-Shot (490 pairs).

Run	Model / Strategy	BLEU	METEOR	chrF
7.1	CH / RAG	0.219 $\pm$ 0.042	0.587 $\pm$ 0.009	0.681 $\pm$ 0.008
7.3	OE / RAG	0.148 $\pm$ 0.028	0.527 $\pm$ 0.010	0.645 $\pm$ 0.020
8	CH / MS	0.239 $\pm$ 0.013	0.604 $\pm$ 0.017	0.690 $\pm$ 0.009
8.2	OE / MS	0.193 [†]	0.462 [†]	0.607 [†]

[†] Single-seed result (seed=42 only).

SPA $\rightarrow$ MSLG. On SPA $\rightarrow$ MSLG, run 8 (many-shot, Claude) records the highest multi-seed mean across all three metrics (BLEU 0.239, METEOR 0.604, chrF 0.690), while run 7.1 (RAG+curriculum, Claude) shows higher variance across seeds (BLEU std=0.042 vs. 0.013 for run 8), suggesting greater sensitivity to the validation split. The seed=42 result of 7.1 (BLEU 0.268) lies above its multi-seed mean (0.219), which should be considered when interpreting single-seed comparisons. Run 7.3 (OpenEuroLLM-Spanish, RAG+curriculum) places below both Claude configurations in all metrics, while run 8.2 (OpenEuroLLM-Spanish, many-shot) slightly improves BLEU over 7.3 but degrades METEOR and chrF.

Table 3

Results for MSLG $\rightarrow$ SPA configurations. Primary runs reported as mean $\pm$ std over seeds {42, 123, 456} (val=90 each). Bold marks the best mean per column. CH = Claude Haiku 4.5; OE = OpenEuroLLM-Spanish; RAG = RAG+curriculum $k=10$; MS = Many-Shot (490 pairs). COMET uses Unbabel/wmt22-comet-da.

Run	Model / Strategy	BLEU	METEOR	chrF	COMET
7.2	CH / RAG	0.390 $\pm$ 0.018	0.660 $\pm$ 0.022	0.730 $\pm$ 0.020	0.893 $\pm$ 0.013
7.4	OE / RAG	0.362 $\pm$ 0.008	0.627 $\pm$ 0.014	0.709 $\pm$ 0.012	0.876 $\pm$ 0.004
8.1	CH / MS	0.408 $\pm$ 0.008	0.664 $\pm$ 0.018	0.742 $\pm$ 0.005	0.901 $\pm$ 0.007
8.3	OE / MS	0.277 [†]	0.440 [†]	0.480 [†]	0.735 [†]

[†] Single-seed result (seed=42 only).

MSLG $\rightarrow$ SPA. On MSLG $\rightarrow$ SPA, run 8.1 (many-shot, Claude) records the highest multi-seed mean across all four metrics (BLEU 0.408, METEOR 0.664, chrF 0.742, COMET 0.901), while run 7.2 (RAG+curriculum, Claude) follows closely (BLEU 0.390, COMET 0.893). Run 7.4 (OpenEuroLLM-Spanish, RAG+curriculum) is the most stable configuration across seeds (COMET std=0.004) and places below both Claude configurations in all metrics. Run 8.3 (OpenEuroLLM-Spanish, many-shot) records the lowest scores in this direction, with a COMET of 0.735 notably below the three other configurations.

4.4. Statistical Significance

To quantify the reliability of observed differences, we applied bootstrap resampling ($n=1,000$) over sentence-level scores, obtaining 95% confidence intervals (CI) by the percentile method (2.5 and 97.5 percentiles). For COMET, bootstrap was applied over the per-instance scores returned by `comet-score --to_json`.

MSLG→SPA. The COMET difference between run 8.1 and run 7.2 is not statistically significant at the 95% level: the CIs overlap substantially (8.1: [0.888, 0.915] vs. 7.2: [0.886, 0.914]). Both configurations should therefore be considered statistically equivalent in semantic quality. In contrast, the gap between the local model (run 7.4: COMET 0.876, CI [0.847, 0.893]) and the API model (run 7.2: COMET 0.893, CI [0.886, 0.914]) does not overlap, confirming that the backend choice has a statistically significant impact on semantic quality for this subtask.

SPA→MSLG. The chrF confidence intervals for runs 7.1, 7.3, and 8 overlap across seeds, indicating that differences between strategies in this direction should be interpreted with caution. The high variance of run 7.1 (BLEU std=0.042 across seeds) further suggests that seed-specific results may not generalise; the multi-seed mean (BLEU 0.219) provides a more reliable performance estimate than the single seed-42 result (BLEU 0.268).

4.5. Effect of Model Choice

Comparing runs 7.1 vs. 7.3 and 7.2 vs. 7.4 isolates the contribution of the underlying model while holding the prompt and pipeline constant. Over multi-seed means, Claude Haiku 4.5 places above OpenEuroLLM-Spanish by 7.1 BLEU points on SPA→MSLG (0.219 vs. 0.148) and by 2.8 points on MSLG→SPA (0.390 vs. 0.362). On COMET for MSLG→SPA, the gap between API and local model (0.893 vs. 0.876) is statistically significant at the 95% level (Section 4.4), while on SPA→MSLG the chrF intervals overlap across seeds, indicating a less conclusive separation in that direction.

The difference between subtasks is consistent with the nature of each transformation: SPA→MSLG requires precise rule application under a structured output constraint, while MSLG→SPA draws on broad linguistic knowledge of the target language. In both directions, the local model records competitive metrics at no API cost, which may be relevant for resource-constrained settings where reproducibility without external dependencies is a priority.

5. Discussion

5.1. RAG vs. Many-Shot: Relevance and Volume

The relationship between RAG+curriculum and many-shot varies across subtasks and evaluation perspectives. On SPA→MSLG, run 8 (many-shot, Claude) records a higher multi-seed mean BLEU (0.239 vs. 0.219 for run 7.1) and lower variance across seeds (std=0.013 vs. 0.042), while run 7.1 records the higher value on the seed=42 split (BLEU 0.268 vs. 0.250). Given the overlapping confidence intervals documented in Section 4.4, neither strategy can be declared superior for this subtask on the basis of the available evidence.

On MSLG→SPA, run 8.1 (many-shot) places above run 7.2 (RAG+curriculum) in all multi-seed metrics: BLEU 0.408 vs. 0.390, chrF 0.742 vs. 0.730, COMET 0.901 vs. 0.893. The COMET difference is not statistically significant at the 95% level (Section 4.4), though the advantage of 8.1 is consistent across the three seeds evaluated.

These results contrast with the general many-shot scaling trend reported by Agarwal et al. [23]. A plausible explanation is that the rule-governed nature of gloss translation makes both strategies sensitive to different factors: RAG+curriculum benefits from local semantic relevance but depends on the quality of the retrieved neighbours, while many-shot provides uniform corpus coverage at the cost

of reduced adaptability per query. In the expansion direction (MSLG→SPA), where the model draws on broad linguistic knowledge of the target language, the distinction between strategies appears less pronounced.

5.2. Curriculum Ordering and Recency Bias

Curriculum ordering yields a 1.9 BLEU improvement over fixed-order few-shot at $k=10$. This gain aligns with two independent lines of evidence: the recency bias of transformer LLMs documented by Liu et al. [11], which makes the last example in context disproportionately influential, and the prompt-order sensitivity reported by Lu et al. [21]. Placing the hardest retrieved example immediately before the query combines both effects, extending the curriculum learning principle of Bengio et al. [22] to the ICL setting.

5.3. Chain-of-Thought and Corpus-Derived Prompts

Explicit CoT scaffolding degraded performance at every few-shot level tested. Gloss translation is a deterministic rule-application task, not a multi-step reasoning task; verbalising intermediate steps introduces variance without benefit, consistent with Sprague et al. [19].

Deriving PROMPT_BASE inductively from the corpus — rather than from Cruz-Aldrete [2] — proved equally important. The corpus orthography diverges from academic descriptions in several conventions documented in Section 2.3 (PORQUÉ vs. POR-QUÉ, *dm-* vs. *s-*, exclusive use of +MUJER), and prompts based on academic conventions produced more errors on the validation set, particularly in the categories targeted by NEGATIVE_EXAMPLES.

5.4. COMET and Model Gap

COMET was restricted to MSLG→SPA because gloss sequences are not natural language, and a fluency-based neural metric would penalise structurally correct gloss output. The scores obtained (0.900–0.901 for Claude, 0.735–0.871 for the local model) are notably high, reflecting fluent Spanish output from both model families. The COMET gap between Claude and OpenEuroLLM-Spanish under many-shot conditions (16.6 points) exceeds the BLEU gap (13.1 points), indicating that the local model produces subtle disfluencies that n -gram metrics miss but a neural evaluator detects.

5.5. Limitations

Several limitations of the present work should be noted. The corpus size (490 training pairs) precludes supervised fine-tuning approaches: our attempts with `bart-base-spanish` confirmed that the available data volume is insufficient for seq2seq training [26]. The out-of-vocabulary coverage of the test set was not measured; an OOV analysis could explain residual errors in cases where the query has no close neighbours in the training pool. The corpus may mix regional LSM variants without explicit dialect labels, introducing noise that neither the retrieval index nor the prompt rules can compensate for. Finally, the retrieval component relies on a generic multilingual encoder (`paraphrase-multilingual-MiniLM-L12-v2` [15, 16]) not adapted to the gloss domain; a task-specific encoder could improve retrieval precision, particularly for short or morphologically complex gloss sequences.

6. Conclusions

This paper presented the system developed by team VerbaNexAI for the MSLG-SPA 2026 shared task on bidirectional translation between Mexican Sign Language glosses and Spanish. Operating in a low-resource regime of 490 training pairs, we explored three in-context learning strategies — RAG+curriculum few-shot, many-shot, and open-source local replication — across two model families, yielding eight evaluated configurations over both subtasks.

The central finding is that semantic relevance of examples matters more than their quantity for the SPA→MSLG direction: retrieving 10 semantically similar training pairs and reordering them by curriculum principle outperforms embedding all 490 pairs in the context window (BLEU 0.268 vs. 0.250). This result is specific to the compression direction, where gloss translation is governed by precise, locally applicable rules; in the expansion direction (MSLG→SPA), many-shot coverage yields marginally higher fluency as measured by COMET (0.9010 vs. 0.9004), confirming that the optimal ICL strategy is task-dependent.

Three additional conclusions emerge from the ablation study and comparative analysis. First, curriculum ordering of retrieved examples — placing the most complex example immediately before the query — consistently improves performance, supporting the recency-bias hypothesis of Liu et al. [11] and the prompt-order sensitivity findings of Lu et al. [21] in the specific domain of sign language NLP. Second, explicit chain-of-thought prompting degrades performance on this task, consistent with Sprague et al. [19]: gloss translation is a deterministic rule-application task, not a reasoning task, and verbalising intermediate steps introduces variance without benefit. Third, corpus-derived prompt induction outperforms grammar-based prompt engineering: the orthographic conventions of the shared-task corpus diverge from published LSM grammars [2] in ways that directly affect translation accuracy, making data-driven rule extraction the more reliable approach for low-resource annotation dialects.

Finally, the open-source replication runs (7.3 and 7.4) demonstrate that the full pipeline is reproducible without access to a commercial API. OpenEuroLLM-Spanish recovers 66.8% of the BLEU achieved by Claude Haiku 4.5 on SPA→MSLG and 91.0% on MSLG→SPA, establishing a competitive open-source baseline for future work on LSM-Spanish machine translation.

6.1. Future Work

While the present system establishes a competitive baseline across both translation directions, several aspects of the pipeline were constrained by the low-resource setting of the shared task and remain open for future investigation.

Domain-adapted retrieval encoder. The RAG+curriculum pipeline currently relies on a generic multilingual sentence encoder (paraphrase-multilingual-MiniLM-L12-v2 [15, 16]) to retrieve semantically similar training pairs. Fine-tuning this encoder on MSLG-Spanish pairs, or replacing it with a contrastively trained gloss-specific embedding model, could improve neighbour retrieval for short or morphologically irregular glosses and is a natural next step given that relevance — rather than volume — was found to be the dominant factor for SPA→MSLG performance.

Out-of-vocabulary analysis and error taxonomy. This work did not quantify out-of-vocabulary coverage on the test set. A systematic OOV and error-category analysis — distinguishing lexical gaps, morphological mismatches, and word-order errors — would clarify which failure modes are addressable with more retrieval examples and which require structural changes to the prompting strategy or the corpus itself.

Supervised fine-tuning at scale. The 490-pair training corpus precluded sequence-to-sequence fine-tuning in this work, as our bart-base-spanish experiments confirmed [26]. Should the MSLG-SPA shared task corpus grow in future editions, or if it can be combined with related Mexican Sign Language resources such as DIESLEME [1], revisiting supervised fine-tuning — possibly in combination with the in-context strategies presented here — would be a natural extension once sufficient data volume is available.

Declaration on Generative AI

During the preparation of this work, the author(s) used AI-assisted writing tools for grammar revision and text improvement. The language models discussed throughout this work — Claude Haiku 4.5 and OpenEuroLLM-Spanish — served solely as experimental subjects within the evaluated translation pipeline and were not employed for manuscript preparation purposes. After using the aforementioned

tools, the author(s) reviewed and edited all content as needed and take full responsibility for the publication's content.

References

- [1] Instituto Politécnico Nacional, DIESLEME: Diccionario español–LSM, Instituto Politécnico Nacional, México, 2004.
- [2] M. Cruz-Aldrete, Gramática de la Lengua de Señas Mexicana, Ph.D. thesis, El Colegio de México, Ciudad de México, México, 2008.
- [3] A. Y. Rodríguez-González, D. Fajardo-Delgado, M. A. Álvarez-Carmona, M. G. Sánchez-Cervantes, A. Díaz-Pacheco, R. Aranda, S. Fajardo-Flores, C. H. Martínez-Rodríguez, Overview of MSLG-SPA at IberLEF 2026: Bidirectional translation between Mexican Sign Language glosses and Spanish, *Procesamiento del Lenguaje Natural* 77 (2026).
- [4] A. Bonet-Jover, J. A. González-Barba, L. Chiruzzo, Overview of IberLEF 2026, 2026. In *Proceedings of the Iberian Languages Evaluation Forum (IberLEF 2026)*, co-located with the 42nd Conference of the Spanish Society for Natural Language Processing (SEPLN 2026), CEUR-WS.org.
- [5] K. Papineni, S. Roukos, T. Ward, W.-J. Zhu, BLEU: A method for automatic evaluation of machine translation, in: *Proceedings of the 40th Annual Meeting of the Association for Computational Linguistics (ACL)*, Philadelphia, PA, 2002, pp. 311–318. URL: <https://aclanthology.org/P02-1040/>. doi:10.3115/1073083.1073135.
- [6] S. Banerjee, A. Lavie, METEOR: An automatic metric for MT evaluation with improved correlation with human judgments, in: *Proceedings of the ACL Workshop on Intrinsic and Extrinsic Evaluation Measures for Machine Translation and/or Summarization*, Ann Arbor, MI, 2005, pp. 65–72. URL: <https://aclanthology.org/W05-0909/>.
- [7] M. Popović, chrF: Character n-gram F-score for automatic MT evaluation, in: *Proceedings of the Tenth Workshop on Statistical Machine Translation (WMT)*, Lisbon, Portugal, 2015, pp. 392–395. URL: <https://aclanthology.org/W15-3049/>. doi:10.18653/v1/W15-3049.
- [8] R. Rei, C. Stewart, A. C. Farinha, A. Lavie, COMET: A neural framework for MT evaluation, in: *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, Online, 2020, pp. 2685–2702. URL: <https://aclanthology.org/2020.emnlp-main.213/>. doi:10.18653/v1/2020.emnlp-main.213.
- [9] Anthropic, Claude Haiku 4.5 — model card, <https://www.anthropic.com/>, 2025.
- [10] OpenEuroLLM Consortium, OpenEuroLLM-Spanish: Open European large language models, <https://openeurollm.eu/>, 2024.
- [11] N. F. Liu, K. Lin, J. Hewitt, A. Paranjape, M. Bevilacqua, F. Petroni, P. Liang, Lost in the middle: How language models use long contexts, *Transactions of the Association for Computational Linguistics* 12 (2024) 157–173. URL: <https://aclanthology.org/2024.tacl-1.9/>. doi:10.1162/tacl_a_00638.
- [12] B. Chen, C. Cherry, A systematic comparison of smoothing techniques for sentence-level BLEU, in: *Proceedings of the Ninth Workshop on Statistical Machine Translation (WMT)*, Baltimore, MD, 2014, pp. 362–367. URL: <https://aclanthology.org/W14-3346/>. doi:10.3115/v1/W14-3346.
- [13] R. Rei, J. G. C. de Souza, D. Alves, C. Zerva, A. C. Farinha, T. Glushkova, A. Lavie, L. Coheur, A. F. T. Martins, COMET-22: Unbabel-IST 2022 submission for the metrics shared task, in: *Proceedings of the Seventh Conference on Machine Translation (WMT)*, Abu Dhabi, 2022, pp. 578–585. URL: <https://aclanthology.org/2022.wmt-1.52/>.
- [14] S. Bird, E. Klein, E. Loper, *Natural Language Processing with Python*, O'Reilly Media, 2009.
- [15] N. Reimers, I. Gurevych, Sentence-BERT: Sentence embeddings using siamese BERT-networks, in: *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, Hong Kong, China, 2019, pp. 3982–3992. URL: <https://aclanthology.org/D19-1410/>. doi:10.18653/v1/D19-1410.
- [16] N. Reimers, I. Gurevych, Making monolingual sentence embeddings multilingual using knowledge distillation, in: *Proceedings of the 2020 Conference on Empirical Methods in Natural Language*

- Processing (EMNLP), Online, 2020, pp. 4512–4525. URL: <https://aclanthology.org/2020.emnlp-main.365/>. doi:10.18653/v1/2020.emnlp-main.365.
- [17] Anthropic, Prompt caching with Claude, <https://www.anthropic.com/news/prompt-caching>, 2024.
 - [18] J. Wei, X. Wang, D. Schuurmans, M. Bosma, B. Ichter, F. Xia, E. Chi, Q. Le, D. Zhou, Chain-of-thought prompting elicits reasoning in large language models, in: *Advances in Neural Information Processing Systems (NeurIPS)*, volume 35, 2022, pp. 24824–24837. URL: https://proceedings.neurips.cc/paper_files/paper/2022/hash/9d5609613524ecf4f15af0f7b31abca4-Abstract-Conference.html.
 - [19] Z. Sprague, F. Yin, J. D. Rodriguez, D. Jiang, M. Wadhwa, P. Singhal, X. Zhao, X. Ye, K. Mahowald, G. Durrett, To CoT or not to CoT? chain-of-thought helps mainly on math and symbolic reasoning, *arXiv preprint arXiv:2409.12183* (2024). URL: <https://arxiv.org/abs/2409.12183>.
 - [20] J. Liu, D. Shen, Y. Zhang, B. Dolan, L. Carin, W. Chen, What makes good in-context examples for GPT-3?, in: *Proceedings of Deep Learning Inside Out (DeeLIO)*, Dublin, Ireland, 2022, pp. 100–114. URL: <https://aclanthology.org/2022.deelio-1.10/>. doi:10.18653/v1/2022.deelio-1.10.
 - [21] Y. Lu, M. Bartolo, A. Moore, S. Riedel, P. Stenetorp, Fantastically ordered prompts and where to find them: Overcoming few-shot prompt order sensitivity, in: *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (ACL)*, Dublin, Ireland, 2022, pp. 8086–8098. URL: <https://aclanthology.org/2022.acl-long.556/>. doi:10.18653/v1/2022.acl-long.556.
 - [22] Y. Bengio, J. Louradour, R. Collobert, J. Weston, Curriculum learning, in: *Proceedings of the 26th International Conference on Machine Learning (ICML)*, Montreal, Canada, 2009, pp. 41–48. URL: <https://dl.acm.org/doi/10.1145/1553374.1553380>. doi:10.1145/1553374.1553380.
 - [23] R. Agarwal, A. Singh, L. M. Zhang, B. Bohnet, L. Rosias, S. Chan, B. Zhang, A. Anand, Z. Abbas, A. Nova, J. D. Co-Reyes, E. Chu, F. Behbahani, A. Faust, H. Larochelle, Many-shot in-context learning, in: *Advances in Neural Information Processing Systems (NeurIPS)*, volume 37, 2024. URL: <https://arxiv.org/abs/2404.11018>.
 - [24] M. M. Saez Acuña, E. A. Puertas Del Castillo, J. C. Martínez Santos, G. A. Mantilla Clavijo, A. D. López Fortich, VerbaNexAI system repository for MSLG-SPA 2026, <https://github.com/MelanySaez/mslg-spa-2026>, 2026.
 - [25] H. Su, J. Kasai, C. H. Wu, W. Shi, T. Wang, J. Xin, R. Zhang, M. Ostendorf, L. Zettlemoyer, N. A. Smith, T. Yu, Selective annotation makes language models better few-shot learners, in: *Proceedings of the 11th International Conference on Learning Representations (ICLR)*, 2023. URL: <https://openreview.net/forum?id=qY1hlv7gwg>.
 - [26] M. Lewis, Y. Liu, N. Goyal, M. Ghazvininejad, A. Mohamed, O. Levy, V. Stoyanov, L. Zettlemoyer, BART: Denoising sequence-to-sequence pre-training for natural language generation, translation, and comprehension, in: *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics (ACL)*, Online, 2020, pp. 7871–7880. URL: <https://aclanthology.org/2020.acl-main.703/>. doi:10.18653/v1/2020.acl-main.703.