

Rank	Team Name	LLM Score					Average
		1	2	3	4	5	
1	Sheffield Causal	1	2	16	52	432	4.813
2	Tredence AICOE	0	2	28	41	432	4.795
3	HSA CORAL	4	1	21	52	425	4.775
4	EMI	2	6	28	49	418	4.739
5	CariMed	14	7	18	57	407	4.662
6	TU Graz Data Team	3	7	25	93	375	4.650
7	LeedsMeng26	8	7	24	83	379	4.614
Total		32	32	160	427	2868	

(a) Spanish Subtask

Rank	Team Name	LLM Score					Average
		1	2	3	4	5	
1	HSA CORAL	3	2	21	33	441	4.814
1	Sheffield Causal	4	2	18	35	441	4.814
3	Tredence AICOE	1	3	24	33	439	4.812
4	Lab Rats	3	9	24	33	431	4.760
5	CariMed	5	6	28	46	415	4.720
6	EMI	6	6	40	26	422	4.704
7	LeedsMeng26	3	11	26	53	407	4.700
8	TU Graz Data Team	8	9	58	55	370	4.662
9	Sarang	1	3	24	33	439	4.540
Total		34	51	263	347	3805	

(b) English Subtask

Table 4: LLM-based ranking of each team’s best submission for the Spanish (a) and English (b) subtasks. Each row reports the team rank, team name, score distribution and the LLM score (avg.) of their best submitted system.

two main techniques: (1) Rephrase-and-Respond (RaR), where the model optimizes the prompt; and (2) Recursive Self-aggregation (RSA), which provides candidate answers and recursively adds subsets to obtain consensus. Stage 1 functions as a causal analyst, stage 2 applies the rephaser (RaR), stage 3 extracts the relevant information, stage 4 aggregates candidates (RSA), and stage 5 validates the candidates before providing the final answer. This approach obtained a score of 4.6620 in Spanish (5th) and 4.720 in English (5th).

Team **HSA CORAL** (Gautam et al., 2026) compares three approaches to address the EQA: (1) encoder-only for token classification using BERT-based models, (2) encoder-decoder for sequence-to-sequence generation using BART-like models, and (3) decoder-only models for generation, enforcing extraction with prompt optimization, few-shot selection, and fine-tuning. They retrieve the most similar cosine-similarity-based C and Q pairs from the training set to those pairs from the test set to include them as shots. They found that fine-tuned generative models outperform encoder-only and encoder-decoder models. Also, 20-shot prompting enables compact models to outperform bigger

models in a zero-shot scenario. Their best system consists of a bilingually fine-tuned GPT-4.1 mini with 20 shots similar to the test sample, achieving a score of 4.7753 in Spanish (3rd) and 4.814 in English (1st).

Team **Tredence AICOE** (Chopra et al., 2026) presented a multilingual financial causality extraction system for English and Spanish based on few-shot prompting, supervised fine-tuning, data augmentation, and ensemble arbitration. They experimented with models such as GPT-5, Gemini 3.0 Pro, Qwen-14B, Gemma-12B, GPT-4.1 mini, GPT-OSS 20B, and GPT-5.1, and explored techniques including joint EN+ES training, LoRA-style efficient fine-tuning, bidirectional translation, synthetic data generation, distilled chain-of-thought reasoning, confidence-based selection, semantic consensus, and voting-based ensembles. Overall, the best results come from multilingual fine-tuning plus selective ensembling, with GPT-5.1-arbitrated voting performing best in English and synthetically augmented GPT-4.1 mini performing best in Spanish. In Spanish, the best score was 4.795 (2nd), achieved by GPT-4.1 mini fine-tuned with synthetic data augmentation. In English, the best result was

obtained with a voting-based ensemble arbitrated by GPT-5.1, which achieved 4.812 (3rd), outperforming all standalone prompting and fine-tuned systems.

Team **EMI** (Attak, 2026) presented a system based on supervised fine-tuning of instruction-following language models, with particular emphasis on prompt repetition as a training strategy to reinforce the relationship between the question format and the expected extractive answer. The authors evaluate both open-weight (Qwen2.5-7B, Qwen2.5-14B) and proprietary models (GPT-4.1-Nano), and show that this simple intervention can improve extraction fidelity, especially for open models, by reducing over-generation and helping the model stay closer to the source span. Their best systems (GPT-4.1-nano) obtained 4.7396 in Spanish (4th) and 4.704 in English (6th).

Team **LeedsMeng26** (Shahrouri et al., 2026) presented a two-stage extractive question answering system for FinCausal 2026. They cast financial causality detection as a QA problem over English and Spanish financial texts, returning a verbatim span from the context rather than generate a free-form answer. Their system works in two steps. First, a model generates an initial candidate span under strict prompting designed to force extractive behavior. Then a second model acts as a verifier and boundary refiner, checking whether the candidate is correct and adjusting its span boundaries when necessary, while still being constrained to output only a contiguous substring from the source text. The paper says this second stage was introduced to fix typical span errors such as truncation, overrun, or occasional paraphrasing. Their final submitted configuration was Qwen-2.5-1.5B-Instruct + Gemini 2.5-flash achieving scores of 4.6143 in Spanish (7th) and 4.7000 in English (7th).

Team **TU Graz Data Team** (Niess and Kern, 2026) introduced SpanDiffusion to approach financial causal question answering as an extractive span prediction problem, but replaced standard start–end classification with a continuous denoising approach in a system called SpanDiffusion. The question and context are encoded with DeBERTa-v3-large plus LoRA, and the target answer is represented as two Gaussian masks marking the start and end of the span. A transformer denoiser trained with flow matching reconstructs these masks from noise, and the final span is obtained by taking the peak of each mask. This design aims to model boundary uncertainty while preserving the extractive nature of the task. In the results, the proposed method proves competitive but not superior to a simpler baseline, namely DeBERTa-v3-large + LoRA + a linear span head. The best SpanDiffusion model reaches 83.0 Exact Match, with notable gains from LoRA and from using flow matching instead of

DDPM. However, the standard span-classification baseline achieves 85.8 Exact Match, outperforming the diffusion model with much lower complexity. Overall, the paper’s contribution is therefore mainly methodological, offering an original alternative to conventional extractive QA rather than a stronger empirical system. Their best systems achieved a score of 4.6501 in Spanish (6th) and 4.662 in English (8th).

Team **Sheffield Causal** (Alqarni et al., 2026) addressed financial causal question answering in English and Spanish as a generative extractive task based on GPT-4.1-mini. The authors combine prompt engineering, retrieval-augmented generation (RAG), and supervised fine-tuning, while enforcing strict verbatim extraction from the input context. Their pipeline has three stages: indexing the training set, retrieving top-k similar examples, and constructing few-shot prompts for either the base or the fine-tuned model. They compare several retrieval strategies, including random example selection, BM25, dense retrieval with text-embedding-3-large and LlamaIndex, a pattern-aware retrieval method that groups questions into CAUSE, EFFECT, or OTHER templates, and a hybrid BM25+dense approach using reciprocal rank fusion. In addition, they evaluate simple, expert, and multilingual prompts, and fine-tune GPT-4.1-mini on bilingual training data formatted as question-context-answer triples. Their system ranked first in both subtasks, reaching a score of 4.813 for Spanish (1st) and 4.814 for English (1st).

5.3.2. English-only Systems

Team **Lab Rats** (Sarda et al., 2026) participation in the competition consists of two main aspects: QLoRA SFT of Qwen3-4B-Instruct on the English dataset, which was adapted to the ChatML instruction format required by the model; and an intra-context TF-IDF retrieval to enrich context for enforcing verbatim span extraction at inference time. The inference pipeline involves four stages: (1) loading the model in 4-bit, (2) intra-context retrieval, comparing each sentence from the given C against the Q , thus filtering the most relevant fragment, (3) constrained prompting, where the model is instructed to extract the span verbatim, and the previously retrieved fragment is also provided, and (4) greedy decoding and deterministic settings aimed at improving extraction accuracy. They achieved a score of 4.760 in English (4th).

Team **Sarang** (Trivedi and Chindukuri, 2026) presented a system based on few-shot prompting and automated prompt optimization for Gemma3-12B. Using DSPy and the MIPROv2 teleprompter, the system optimizes instructions and demonstration examples drawn from the training set, and performs inference locally through Ollama. The paper

compares this setup with RoBERTa and DeBERTa baselines finetuned with other QA datasets and with other few-shot LLM configurations, finding that Gemma3-12B with medium prompt optimization performs best. Their top system obtained an LLM Score of 4.540 in the English shared sub-task (9th), highlighting the effectiveness of prompt optimization for financial causal QA.

5.4. Additional English-only System

Team **YT** utilized a LoRA-finetuned Llama-3.1-8B model using a difficulty-aware training strategy. Their methodology involves a custom labeling scheme that categorizes instances into simple, chain, and abstractive types based on structural and semantic complexity. To ensure strict verbatim extraction, the system employs a three-level span-anchoring mechanism—combining case-insensitive matching, semantic sentence retrieval, and re-prompting—complemented by targeted post-processing rules to truncate redundant clauses. Although the team did not submit official runs for the shared task, their internal evaluations on a partitioned subset of the 2026 dataset demonstrated the effectiveness of combining difficulty-stratified training with post-processing.

6. Evaluation

In FinCausal 2023 (Moreno-Sandoval et al., 2023), the task focused on identifying cause–effect relations linked to events or quantified facts in financial texts. Because it was formulated as an extractive task, system performance was assessed with EM to quantify the proportion of predictions that exactly match the reference span, and token-level precision, token-level recall, and weighted F1 to quantify partial overlap quality for extracted cause and effect spans.

In FinCausal 2025 (Moreno-Sandoval et al., 2025), the extraction task was reformulated as a question-answering task in which questions about causes or effects are posed, and system responses are evaluated with EM and semantic answer similarity metrics. This change accommodates the growing use of generative prompting-based models, many of them based on GPT architectures. For these models, a strict lexical metric such as EM alone is often insufficient, because generative systems can produce answers that are semantically correct but phrased differently from the references. SAS measures semantic similarity between texts rather than exact lexical overlap, making it well suited for abstractive generation tasks. The metric represents texts as vector embeddings using pre-trained models such as BERT (Devlin et al., 2019) or Sentence Transformers (Reimers

and Gurevych, 2019), and computes cosine similarity between them. This allows the evaluation to capture cases where two answers express the same meaning despite differences in wording or structure.

In the FinCausal 2026 edition, system submissions are evaluated through an LLM-as-a-judge framework. For each system prediction, the judge model receives a fixed evaluation rubric and scores the response according to a uniform set of criteria. Specifically, each answer is rated on a five-point Likert scale, whose levels capture different degrees of semantic alignment between the predicted answer and the gold-standard reference. The full FinCausal 2026 rubric is provided in Appendix A. In our setup, the judge model is `openai/gpt-oss-20b` (OpenAI, 2025), a 20-billion-parameter open-weight language model from OpenAI’s GPT-OSS family released in August 2025. We use the model with a medium reasoning configuration and explicitly instruct it in the prompt to generate concise score justifications, thereby improving the transparency and auditability of the evaluation procedure.

6.1. Error Analysis

The following initial analysis is based on a linguistic review of model outputs graded from 1 to 5. The predictions from the best system of each participant were observed, thus enabling the identification of common error patterns found in the lower and middle scoring ranges. Superficially, no errors were found on the scores of 5, but a more thorough analysis should be performed.

Score 1: Structural failure. At this level, models show major structural problems. While they often identify causal markers (like “due to” or “thanks to”), they extract the wrong information or focus on unrelated events nearby. Another common issue is empty referencing, where the model only extracts the connector (e.g., “Due to the above”) without the actual explanation. Additionally, models at this stage sometimes reverse the cause and effect or skip essential steps in a causal chain. These models are also prone to neglecting or hallucinating quantitative financial data; a failure severely penalized by the LLM judge, which assigns the minimum score to responses lacking numerical precision.

Score 2: Incomplete and inaccurate. Errors here involve missing information or technical mistakes. Models often cut the answer too short, leaving out the specific details that provide the actual argument. They also tend to confuse similar financial terms or acronyms (like DVA instead of CVA). We also observed a metric bias: LLM judges sometimes give higher scores to fluent-sounding answers, even if the content is partially incorrect.

This behaviour is typical of metrics based in neural models (Kovacs et al., 2024; Freitag et al., 2024).

Score 3: Partial and selective. These responses overlap with the correct answer but are incomplete. Models often pick only the first factor in a list and ignore the others, providing a one-sided view of the event. They also frequently leave out important time-related details (like specific dates or years). While the core reason is usually captured, these models tend to ignore secondary details or consequences, resulting in an over-simplified version of the reference.

Score 4: Precise with minor noise. A score of 4 represents an accurate answer that includes unnecessary noise. The most common error is marker dragging, where the model includes extra words like connectors or introductory verbs (e.g., “resulting in” or “allowing”) that were not part of the target answer. In other cases, the model adds extra context that is true but goes beyond the exact boundaries set by the experts.

The errors progress from total confusion at score 1 to minor noise at score 4. The main challenges for these models are navigating complex financial texts with many variables and staying within the strict boundaries of the required answer. Finally, the metric bias seen in scores 2 and 3 suggests that LLM judges often favor smooth, fluent writing over literal, word-for-word accuracy. Therefore, the differences between scores of 2 and 3 are often blurry, while the scores of 1 and 4 seem reliable.

7. Conclusions

This FinCausal 2026 edition had higher participation than previous editions in terms of effective system-description submissions. Consequently, it enabled us to gain a clearer overview of the current state of causal EQA in the financial domain. Some insights can be drawn from the results obtained by participants.

Bilingual fine-tuning performs consistently better than monolingual fine-tuning, as shown by direct comparisons in the HSA CORAL and Tredence AICOE submissions. Additionally, as highlighted by HSA CORAL, smaller fine-tuned models can outperform larger zero-shot models.

Another clear trend was the use of retrieval-based techniques and semantic matching. Four teams used similarity measures at different stages of their pipelines for different purposes. HSA CORAL applied vector similarity at inference time to retrieve few-shot examples similar to the test C and Q for in-context learning. Sheffield Causal followed a similar approach to HSA CORAL but conducted a broader evaluation of retrieval methods, including hybrid BM25 and dense retrieval. In contrast, Lab Rats performed intra-context filtering by compar-

ing each sentence in the provided C with the Q to isolate the most relevant fragment before generation. Finally, YT integrated semantic similarity at two points: for the initial difficulty-based classification of the dataset and for a span-anchoring mechanism during post-processing.

On a different note, complex pipelines that rely on zero-shot settings still appear insufficient to consistently enforce verbatim extractive answers, as seen in CariMed’s VERSA system and Sarang’s prompt-optimization strategy. Compared with similarly complex pipelines that include fine-tuning, such as Tredence AICOE’s, these results suggest that fine-tuning remains essential for stronger performance.

Regarding the main model choice, GPT-4.1 mini was used by several teams and proved to be a strong option despite its size. Sheffield Causal used it to obtain first place in both subtasks, tying with HSA CORAL in English. HSA CORAL also used GPT-4.1 mini in its best-performing system, and Tredence AICOE used it for their best Spanish run. A smaller variant, GPT-4.1 nano, was used by EMI. These results suggest an advantage for proprietary models in this edition. Given the compact size of the mini and nano variants, they also offer an attractive cost-performance ratio. On the open-source side, Qwen2.5 appeared in the systems of LeedsMeng26, while Qwen3 appeared in Lab Rats’. Gemma 3 was used by Sarang, while TU Graz Data Team used DeBERTa-v3 together with a span diffusion approach; YT used Llama-3.1-8B.

To conclude, despite the difficulty of the task, all participants achieved strong scores. This highlights both the quality of the participants’ work and the consistency and coherence of the dataset in both languages. In addition, the shift to an LLM-as-a-judge metric has proved useful in light of the error analysis. However, a deeper future analysis of its decisions and of its alignment with human judgment is still needed. For future FinCausal editions, some important changes should be considered to maintain participant interest; for instance, including an additional subtask formulated as a pure QA task with abstractive answers, while preserving the current EQA approach.

Acknowledgments

We would like to thank FinCausal 2026 participants for their outstanding contributions to this shared task.

This work is framed under the Spanish National Project GRESEL (PID2023-151280OB-C21). It was also partially funded by grant PTA2023-023812-I (awarded to Yanco Amor Torterolo Orta) through MICIU/AEI/10.13039/501100011033 and the European Social Fund Plus (ESF+).

8. Bibliographical References

- Aali Abdullah Alqarni, Mark Stevenson, and Arif Dwi Laksito. 2026. A Comparative Study of RAG Approaches and Fine-Tuning for Causal QA in Financial Text. In *Proceedings of the 7th Financial Narrative Processing Workshop (FNP 2026) at LREC-COLING 2026*, Palma de Mallorca, Spain. European Language Resources Association (ELRA). To appear.
- Sanae Attak. 2026. Improving Verbatim Financial Causality Extraction with Supervised Fine-Tuning and Prompt Repetition. In *Proceedings of the 7th Financial Narrative Processing Workshop (FNP 2026) at LREC-COLING 2026*, Palma de Mallorca, Spain. European Language Resources Association (ELRA). To appear.
- Ankush Chopra, Shubham Sharma, and Ashmani Kumar. 2026. Extracting Financial Causality: A Multilingual Approach with SLM Fine-Tuning and LLM-Arbitrated Ensembles. In *Proceedings of the 7th Financial Narrative Processing Workshop (FNP 2026) at LREC-COLING 2026*, Palma de Mallorca, Spain. European Language Resources Association (ELRA). To appear.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pages 4171–4186.
- Markus Freitag, Nitika Mathur, Daniel Deutsch, Chikui Lo, Eleftherios Avramidis, Ricardo Rei, Brian Thompson, Frederic Blain, Tom Kocmi, Jiayi Wang, David Ifeoluwa Adelani, Marianna Buchichio, Chrysoula Zerva, and Alon Lavie. 2024. [Are LLMs breaking MT metrics? results of the WMT24 metrics shared task](#). In *Proceedings of the Ninth Conference on Machine Translation*, pages 47–81, Miami, Florida, USA. Association for Computational Linguistics.
- Akash Kumar Gautam, Serhii Hamotskyi, and Christian Hånig. 2026. Causal Connections: Leveraging Multilingual Fine-Tuning for Financial QA@FinCausal 2026. In *Proceedings of the 7th Financial Narrative Processing Workshop (FNP 2026) at LREC-COLING 2026*, Palma de Mallorca, Spain. European Language Resources Association (ELRA). To appear.
- Aldan Jay, Rafael Berlanda, Yoelvis Moreno, and Vincent Santamarta. 2026. VERSA: Verbatim Extraction via Rephrasing and Self-Aggregation for Financial Causality. In *Proceedings of the 7th Financial Narrative Processing Workshop (FNP 2026) at LREC-COLING 2026*, Palma de Mallorca, Spain. European Language Resources Association (ELRA). To appear.
- Geza Kovacs, Daniel Deutsch, and Markus Freitag. 2024. [Mitigating metric bias in minimum Bayes risk decoding](#). In *Proceedings of the Ninth Conference on Machine Translation*, pages 1063–1094, Miami, Florida, USA. Association for Computational Linguistics.
- Dominique Mariko, Hanna Abi Akl, Estelle Labidurie, Stephane Durfort, Hugues de Mazancourt, and Mahmoud El-Haj. 2021. The Financial Document Causality Detection Shared Task (FinCausal 2021). In *Proceedings of the 3rd Financial Narrative Processing Workshop*, pages 58–60, Lancaster, United Kingdom. Association for Computational Linguistics.
- Dominique Mariko, Hanna Abi-Akl, Kim Trottier, and Mahmoud El-Haj. 2022. The Financial Causality Extraction Shared Task (FinCausal 2022). In *Proceedings of the 4th Financial Narrative Processing Workshop @LREC2022*, pages 105–107, Marseille, France. European Language Resources Association.
- Antonio Moreno-Sandoval, Blanca Carbajo-Coronado, Jordi Porta-Zamorano, Yanco-Amor Torterolo-Orta, and Doaa Samy. 2025. [The Financial Document Causality Detection Shared Task \(FinCausal 2025\)](#). In *Proceedings of the Joint Workshop of the 9th Financial Technology and Natural Language Processing (FinNLP), the 6th Financial Narrative Processing (FNP), and the 1st Workshop on Large Language Models for Finance and Legal (LLMFinLegal)*, pages 214–221, Abu Dhabi, UAE. Association for Computational Linguistics.
- Antonio Moreno Sandoval, Ana Gisbert, and Elena Montoro. 2020. FinT-esp: A corpus of financial reports in Spanish. In Miguel Fuster-Márquez, Carmen Gregori-Signes, and José Santaemilia Ruiz, editors, *Multiperspectives in analysis and corpus design*, pages 89–102. Comares, Granada.
- Antonio Moreno-Sandoval, Jordi Porta-Zamorano, Blanca Carbajo-Coronado, Doaa Samy, Dominique Mariko, and Mahmoud El-Haj. 2023. The Financial Document Causality Detection Shared Task (FinCausal 2023). In *2023 IEEE International Conference on Big Data (BigData)*, pages 2855–2860.
- Georg Niess and Roman Kern. 2026. SpanDiffusion: Flow Matching over Continuous Span

Masks for Financial Causal Question Answering. In *Proceedings of the 7th Financial Narrative Processing Workshop (FNP 2026) at LREC-COLING 2026*, Palma de Mallorca, Spain. European Language Resources Association (ELRA). To appear.

OpenAI. 2025. [gpt-oss-120b](#) [gpt-oss-20b model card](#).

Nils Reimers and Iryna Gurevych. 2019. Sentence-BERT: Sentence Embeddings using Siamese BERT-Networks. In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing*.

Julian Risch, Timo Möller, Julian Gutsch, and Malte Pietsch. 2021. [Semantic Answer Similarity for Evaluating Question Answering Models](#). In *Proceedings of the 3rd Workshop on Machine Reading for Question Answering*, pages 149–157, Punta Cana, Dominican Republic. Association for Computational Linguistics.

Bavya Sarda, Pulkit Chatwal, and Sonal Dabral. 2026. QRAFT: QLoRA Retrieval-Augmented Fine-Tuning for Causal Span Extraction in Financial Documents. In *Proceedings of the 7th Financial Narrative Processing Workshop (FNP 2026) at LREC-COLING 2026*, Palma de Mallorca, Spain. European Language Resources Association (ELRA). To appear.

Zaid Shahrouri, Ayomide Ivienagbor, Idrees Asad, Rijul Shrestha, Yasemin Bal, and Zahaab Nadeem. 2026. LeedsMEng26: Qwen + Gemini for FinCausal 2026 Causality Detection in Financial Narrative Texts. In *Proceedings of the 7th Financial Narrative Processing Workshop (FNP 2026) at LREC-COLING 2026*, Palma de Mallorca, Spain. European Language Resources Association (ELRA). To appear.

Avinash Trivedi and Mallikarjuna Chindukuri. 2026. Financial Causal QA via Instruction and Prompt Tuning of Gemma3-12B. In *Proceedings of the 7th Financial Narrative Processing Workshop (FNP 2026) at LREC-COLING 2026*, Palma de Mallorca, Spain. European Language Resources Association (ELRA). To appear.

A. FinCausal 2026 Rubric

You are an expert evaluator. Your task
 → is to rate how good an
 → STUDENT_ANSWER is
 for a given QUESTION, a CONTEXT and a
 → REFERENCE_ANSWER according to
 the rubric below.

RUBRIC (score from 1 to 5):

- 5: Excellent quality: The prediction
 → is semantically identical or fully
 → equivalent to the gold standard
 → response. Minor formal variations
 → (e.g., punctuation, casing) are
 → tolerated. The content perfectly
 → addresses the causal question,
 → showing no signs of omission or
 → irrelevant inclusion. The answer is
 → deemed fully appropriate and
 → reliable.
- 4: Good quality: The prediction
 → matches the gold-standard answer in
 → full but included small additional
 → content from the context. These
 → additions did not compromise the
 → semantic correctness of the answer
 → but extended it slightly. Therefore,
 → the prediction could be seen to be
 → complete, relevant and informative,
 → although wordy.
- 3: Medium quality: The prediction
 → contains the central idea or a
 → correct causal link, but is either
 → incomplete or diluted with unrelated
 → information. It may capture the
 → start of a causal phrase but miss
 → important qualifiers or follow-up
 → clauses. Alternatively, the
 → prediction might include correct
 → content but extend unnecessarily
 → beyond the relevant span.
- 2: Low quality: The predictions
 → demonstrates only a superficial
 → connection to the question or to the
 → correct answer. Typically, large
 → portions of the expected content are
 → omitted, and irrelevant elements may
 → have been added. The response might
 → contain a partial clue or
 → topic-related phrase, but failed to
 → provide a clear, informative, or
 → accurate answer. This category
 → captures both underinformative and
 → noisy outputs.
- 1: Very poor quality: Predictions in
 → this category failed entirely to
 → answer the question. These responses
 → were irrelevant, incorrect, or
 → confusing, and often exhibited no
 → meaningful overlap with the
 → gold-standard answer. Even if some
 → surface text matched the source, the
 → essential causal content was absent.

First, briefly explain your reasoning.
 Then assign a single integer score from
 → 1 to 5.

Return your response as pure JSON with
 → this exact schema:

```
{{
```

```
"score": <integer 1-5>,  
"reasoning": "<short explanation>"  
}}
```

```
CONTEXT:  
{context}
```

```
QUESTION:  
{question}
```

```
REFERENCE_ANSWER:  
{reference}
```

```
STUDENT_ANSWER:  
{answer}
```

B. Language Resource References

Blanca Carbajo-Coronado, Antonio Moreno-Sandoval, Yanco Amor Torterolo Orta, and Paula Gozalo. 2025. [The Financial Document Causality Detection Shared Task \(FinCausal 2025\): Dataset](#).

Mahmoud El-Haj, Steven Young, and Paul Rayson. 2019. [Annual reports key sections corpora 2003 to 2017](#).

Antonio Moreno-Sandoval, Blanca Carbajo-Coronado, and Jordi Porta. 2023. [The financial document causality detection shared task \(FinCausal 2023\): Dataset](#).

Antonio Moreno-Sandoval, Yanco Amor Torterolo Orta, Maria Alexia Stanescu, and Melina Chatzi. 2026. [The Financial Document Causality Detection Shared Task \(FinCausal 2026\): Dataset](#).

Sheffield NLP at FinCausal 2026: A Comparative Study of RAG Approaches and Fine-Tuning for Causal Q&A in Financial Texts

Aali Alqarni, Mark Stevenson and Arif Laksito

School of Computer Science, University of Sheffield
Sheffield, United Kingdom
{aalqarni1, mark.stevenson, alaksito1}@sheffield.ac.uk

Abstract

This paper describes our approach to the FinCausal 2026 shared task, which addresses causal question answering from financial documents in English and Spanish. We investigated the effectiveness of fine-tuned generative models combined with Retrieval-Augmented Generation (RAG). Our approach compares five retrieval strategies across base and fine-tuned GPT models (GPT-4.1-mini). RAG-based few-shot selection performed better than random sampling, particularly for the base model. In the FinCausal 2026 official run, this approach was ranked first in both the English and Spanish sub-tasks, obtaining LLM scores of 4.8140 and 4.8131 out of 5, respectively.

Keywords: Question and Answering (Q&A), Causality, Large Language Model (LLM), Generative Pre-trained Transformer (GPT), Retrieval-Augmented Generation (RAG)

1. Introduction

Financial documents, including earnings reports and financial analysis articles, contain rich causal relationships that drive market movements and trends. Understanding causality in these documents is fundamental for risk assessment, investment decision making, and market analysis (Cor-mack et al., 2009). Manual extraction of these relations is time-consuming and labour intensive. Automated methods offer a solution but face challenges due to the complexity of financial language and the diversity of causal expressions (Li et al., 2024; Cao et al., 2022).

To advance research in this area, the FinCausal 2026 shared task (Moreno-Sandoval et al., 2026) builds on the 2025 edition, which introduced a question and answering (Q&A) framework for extracting causal relationships from financial texts (Moreno Sandoval et al., 2025). The core objective of the 2026 edition remains the identification of events that explain financial outcomes such as revenue changes, market movements, or corporate decisions, while also introducing more rigorous benchmarks. The dataset has been expanded with over 500 complex causal cases, including multi-element causal chains, and questions have been rephrased to reduce reliance on sentence-level similarity. The task has also shifted to LLM-as-a-judge evaluation rather than the Exact Match (EM) and Semantic Answer Similarity (SAS) measures used previously (Risch et al., 2021).

1.1. Task Description

Objective: Given a natural language question, Q , and corresponding financial text context, C , systems must extract a verbatim span, A , from C that

captures the underlying causal relationship.

Example:

Q : “What explains their strong performance in health and safety?”

C : “As a result of **these very high standards and relentless focus**, we have a strong performance in health and safety” (A indicated by **bold font**).

2. Related Work

The FinCausal shared tasks have been organised as a benchmark for evaluation methods that detect causal relationships in financial texts. Early editions (2020-22) formulated causality detection as either causal sentence classification or cause–effect span identification, with most systems relying on extractive models such as BERT and BiLSTM-CRFs (Mariko et al., 2020, 2022). The 2023 edition extended the task to a multilingual setting by requiring systems to identify causal relationships in financial texts written in English and Spanish. This edition marked a notable transition to the use of generative large language models (LLMs) such as GPT (Moreno-Sandoval et al., 2023). The 2025 edition reframed the task as a Q&A problem, requiring systems to extract answer spans from financial contexts given a natural language question.

LLM-based causal extraction. Recent work in cause–effect span detection tasks has explored GPT-based approaches for causal extraction from financial texts. Shukla et al. (2023) combined Retrieval-Augmented Generation (RAG) with few-shot prompting using GPT-4, while the LTRC II-ITH team explored chain-of-thought (CoT) prompting to enhance reasoning performance (Moreno-Sandoval et al., 2023). These approaches suggested that in-context learning can effectively iden-

ID	Context	Question
English		
29	At the start of the year, we also reduced the size of our transport fleet by 10% in light of network changes. At this point, we decided against further reducing the size of the fleet due to the future demand from new contracts. As a result, the business carried significant excess costs of under-utilised trucks during the current financial year.	What accounts for the business carrying significant excess costs derived from under-utilized trucks during the current financial year?
Spanish		
76	Por el contrario, la Comunidad de Estados Independientes registró un descenso de la producción del 28%. Esta reducción fue menos acusada que la de la demanda interna debido a las exportaciones.	¿Qué efecto tuvieron las exportaciones?

Table 1: Example entries from the English and Spanish datasets. Answers are highlighted in **bold** within the context.

tify cause–effect relationships without fine-tuning. **LLM-based Q&A.** Niess et al. (2025) compared the performance of extractive models, such as BERT, with generative LLMs, including Llama, in zero-shot and few-shot settings for causal Q&A. Their findings indicated that instruction-tuned LLMs without task-specific fine-tuning were competitive. However, a fine-tuned Llama model trained on a multilingual dataset significantly outperformed both approaches.

Beyond stand-alone model comparisons, recent work has explored the integration of LLMs within retrieval-augmented Q&A frameworks to enhance factual grounding and reduce hallucinations. For instance, Yang et al. (2026) proposed Structured-Semantic RAG (SSRAG), a hybrid architecture that incorporates LLM-based query augmentation, prompt-guided source routing, and unified graph–vector retrieval to improve answer faithfulness across open-domain Q&A benchmarks. In the medical domain, Aljohani and Alsanoosy (2026) proposed a modular hybrid RAG framework that integrates sparse retrieval (BM25) with dense biomedical retrieval (MedCPT; Jin et al. 2023) to enhance medical Q&A. Their approach demonstrated substantial improvements in retrieval recall, precision, and generation faithfulness across PubMedQA, MedMCQA, and MedQA-US benchmarks (Pal et al., 2022; Jin et al., 2019, 2021). Despite these advances, limited work has systematically examined hybrid retrieval approaches for domain-specific causal Q&A in financial texts, especially in multilingual settings. This gap highlights the need for retrieval strategies that are specifically designed for causal financial Q&A tasks.

3. FinCausal 2026 Dataset

The FinCausal 2026 dataset (Moreno-Sandoval et al., 2026) contains English text from UK financial reports dated 2017 and Spanish text from Spanish financial reports dated 2014 to 2018. Entries

consist of an abstractive question asking about a cause or effect, a context passage, and a verbatim answer extracted from the context.

The dataset includes different types of causality. Some have explicit causal markers such as ‘due to’ and ‘because’, while others require reasoning from the context to identify implicit causal connections. It includes complex cause-and-effect relationships, such as causal chains of three or more elements where multiple events are linked sequentially (e.g., *A causes B, which leads to C*). Questions can be classified into cause-seeking (e.g., ‘What caused the...?’), effect-seeking (e.g., ‘What was the impact of...?’), and others that do not follow a specific causal pattern (e.g., ‘What did these considerations entail?’).

Dataset Description. Context passages vary in length from a single sentence to multiple sentences, and answer spans range from a short phrase to multiple clauses. Table 1 shows representative examples from the English and Spanish datasets.

Dataset Statistics. The dataset comprises 2,000 labelled training instances per language, with 500 blind test instances for English and 503 for Spanish. We applied an 80:20 split to the training data. This resulted in 1,600 instances per language for fine-tuning and retrieval indexing, and 400 per language for development (including ablation studies and prompt engineering).

4. Methodology

4.1. Model Selection

A generative approach was adopted following previous work which demonstrated that fine-tuned generative models notably outperform traditional extractive Q&A systems (Moreno Sandoval et al., 2025). Figure 1 shows the overall system architecture.

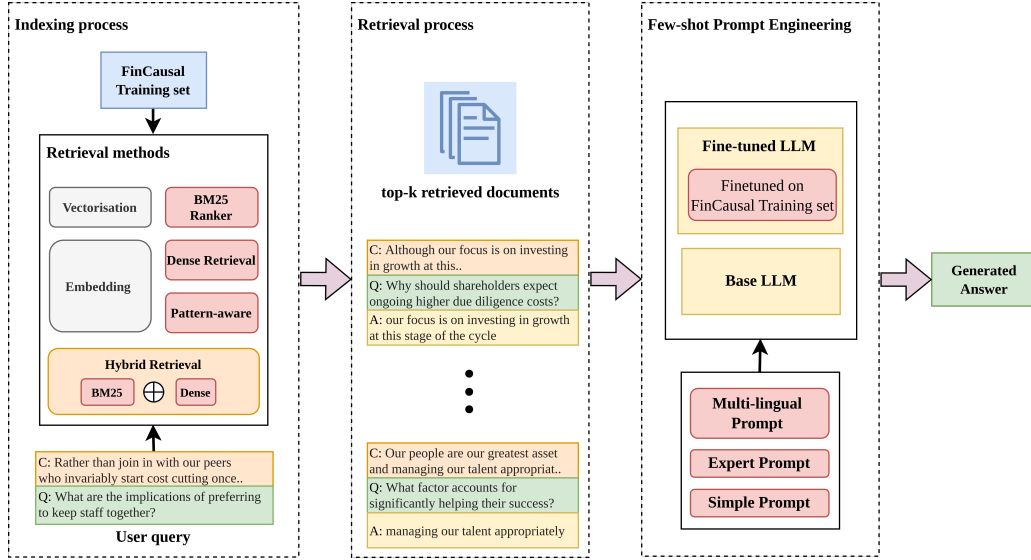


Figure 1: Overview of system architecture. The pipeline consists of three stages: (1) **Indexing**, where the training set is indexed using multiple retrieval methods (Random, RAG-BM25, RAG-Dense, RAG-Pattern, and RAG-Hybrid); (2) **Retrieval**, where the top- k most relevant training examples are retrieved for each test query; and (3) **Few-shot Prompt Engineering**, where the retrieved examples are combined with the test query and passed to either the base or fine-tuned LLM with different prompt strategies to generate the answer.

4.2. Prompt Engineering

Prompt engineering was employed to guide the GPT model to answer causal questions from the provided passages. The model was configured with a temperature of 0.1 and top-p of 1 to ensure deterministic outputs across multiple runs, and a maximum token limit of 512. We began with zero-shot prompting and progressively refined our approach using few-shot examples from the training set. To meet the task’s extractive requirements, we implemented strict instructional constraints: the model was explicitly directed to answer the given question by extracting the answer verbatim from the context C and it was strictly prohibited from paraphrasing.

Three prompting strategies were explored, as illustrated in Figure 1: (1) a *simple prompt*, which provides minimal instructions to extract the answer from the context; (2) an *expert prompt*, which assigns the model a domain-expert role (e.g., “You are a financial causal analysis expert”) and includes detailed extraction rules; and (3) a *multilingual prompt*, which extends the expert prompt with language-aware instructions to handle both English and Spanish inputs (e.g., “The passage and question may be in Spanish.”).¹

Two strategies were explored for selecting few-shot examples:

1) Random sampling. We sampled k examples

randomly from the training set. The model was provided with k question-context-answer triplets as examples before being asked to answer the target question. We experimented with $k \in \{5, 10\}$.

2) RAG-based approaches. The most relevant training examples for each test instance were retrieved using the RAG approach (Lewis et al., 2020). Multiple retrieval strategies were explored: (1) BM25 (Robertson and Zaragoza, 2009), a sparse lexical retrieval method widely used for relevance ranking in information retrieval (RAG-BM25); (2) a dense semantic retrieval approach that encodes texts into dense vector representations and retrieves the most semantically similar training examples based on vector similarity (RAG-Dense) (Karpukhin et al., 2020); (3) a pattern-aware variant (RAG-Pattern), in which each question is first classified into a causal template (CAUSE, EFFECT, or OTHER), and retrieval is restricted to examples of the same type; and (4) a hybrid approach that combines RAG-BM25 and RAG-Dense using Reciprocal Rank Fusion (RRF) (Cormack et al., 2009) to merge the ranked lists into a single unified ranking (RAG-Hybrid).

For RAG-BM25, every question-context pair in the training dataset was indexed using a BM25 retriever with PyStemmer-based language-specific stemming for both English and Spanish.² For RAG-Dense, embeddings were generated for the same training examples using OpenAI’s `text-`

¹Prompt templates: <https://github.com/aaalgarni/fincausal-2026/blob/main/prompts/README.md>

²<https://pypi.org/project/PyStemmer/>

embedding-3-large³ model and indexed using LlamaIndex’s in-memory vector store (Liu, 2022). For RAG-Pattern, each question was first classified into the causal class using regular expression patterns, and retrieval was restricted to training examples of the same template class. At inference time, each test question–context pair was embedded using the same encoder, and the top- k most similar training examples were retrieved and provided as few-shot examples to the base and the fine-tuned GPT-4.1-mini models.

4.3. Fine-tuning

Following the approach of the top-performing teams in FinCausal 2025 (Moreno Sandoval et al., 2025), GPT-4.1-mini was fine-tuned on the provided training data using OpenAI’s fine-tuning API. The training set was formatted as question-context-answer in JSONL format, where each entry contained the question, context, and expected extractive answer. The training process ran for 3 epochs with a learning rate multiplier of 2 and a batch size of 3. This configuration was specifically chosen to accelerate convergence and to avoid overfitting, given the limited size of the dataset.⁴

5. Experimental Setup

Experimental Stages. Experiments were conducted in three different stages. First, various prompt templates and few-shot configurations ($k \in \{0, 5, 10\}$) were evaluated on the development set using the base GPT-4.1-mini model. Second, GPT-4.1-mini was fine-tuned on the 3,200 training instances and its performance was benchmarked against the base version across all retrieval strategies. Third, ablation studies were conducted on the five retrieval methods described in Section 4, varying the number of few-shot examples. In the hybrid configuration, RRF was used with equal weighting between the lexical (RAG-BM25) and semantic (RAG-Dense) retrieval examples.

Final Submission. For the final submission, the retrieval indexes were rebuilt using all 2,000 training instances per language, and predictions were generated on the blind test sets using the best-performing configuration from the development experiments.

Evaluation Metrics. EM and SAS (Risch et al., 2021) are reported for the development set. EM measures whether the predicted answer exactly matches the gold answer string. SAS measures the semantic similarity between the pre-

dicted and gold answers using a cross-encoder model. For SAS, all-MiniLM-L6-v2 is used for English and paraphrase-multilingual-MiniLM-L12-v2 for Spanish. The official competition ranking uses LLM-as-a-judge scoring on the blind test set, which rates system outputs on a 1–5 adequacy scale.

6. Results and Discussion

For English, the simple prompt performed best, while the multilingual prompt produced the best results for Spanish. Table 2 presents our results across all configurations for both English and Spanish sub-tasks. Results showed a large performance gap between the base and fine-tuned models. Fine-tuning GPT-4.1-mini on bilingual (EN+ES) data increased English EM from .3533 to .8875 and Spanish EM from .0250 to .8625. This improvement was consistent across all retrieval configurations, suggesting that fine-tuning helps the model follow the strict verbatim extraction requirement rather than relying only on in-context examples.

For the fine-tuned model, all retrieval strategies produced comparable results. English EM ranged from .8650 to .8875, and Spanish EM from .8425 to .8625. The difference between the best configuration RAG-Dense ($k=5$) and the worst was below 2 percentage points in both languages. In contrast, the base model benefited more from retrieval. For example, RAG-Hybrid ($k=10$) achieved .5950 EM in English compared to .3533 in the zero-shot setting. This suggests that retrieval approaches play a larger role when the model is not fine-tuned, whereas fine-tuning reduces the additional gains from complex retrieval strategies. RAG-Pattern showed higher EM for the base model in English (RAG-Pattern $k=10$: .7550 vs RAG-Dense $k=10$: .5875), but underperformed on SAS and showed no consistent gains for the fine-tuned model, suggesting that fine-tuning implicitly captures causal directionality.

RAG benefits (for Spanish). The base model showed lower zero-shot performance on Spanish (EM = .0250) compared to English (EM = .3533), likely due to a higher prevalence of English financial corpora in the pre-training data. However, RAG strategies narrowed this gap. Specifically, RAG-Dense ($k=10$) increased the Spanish EM score to .5075, a substantial relative improvement over the zero-shot baseline. This suggests that RAG few-shot examples are particularly beneficial in lower-resource scenarios.

Official Blind Test Results. On the official blind test set, evaluated using LLM-as-a-judge, the fine-tuned model with RAG-Dense ($k=5$) achieved the highest English score (4.8140), while RAG-Dense ($k=10$) obtained the best Spanish score (4.8131).

³<https://developers.openai.com/api/docs/models/text-embedding-3-large>

⁴<https://developers.openai.com/api/docs/guides/model-optimization>

Model	Mode	k	English			Spanish		
			Validation		Blind	Validation		Blind
			EM	SAS	LLM Score	EM	SAS	LLM Score
GPT-4.1-mini	Zero-shot	0	.3533	.8605	4.4600	.0250	.8876	4.3002
	Random	5	.5450	.9083	4.4620	.3450	.9282	4.3857
		10	.5850	.9354	4.4760	.3925	.9383	4.4076
	RAG-BM25	5	.5675	.9113	4.4920	.4400	.9393	4.4513
		10	.5800	.9135	4.5280	.4850	.9395	4.5189
	RAG-Dense	5	.5550	.9136	4.5780	.4525	.9347	4.4712
		10	.5875	.9151	4.5840	.5075	.9408	4.5030
	RAG-Pattern	5	.7300	.8593	4.6220	.4075	.9406	4.4135
		10	.7550	.8622	4.6500	.4175	.9432	4.4334
	RAG-Hybrid	10	.5950	.9152	4.5940	.4575	.9372	4.4712
Finetuned GPT-4.1-mini (EN+ES)	Zero-shot	0	.8800	.9755	4.7440	.8550	.9734	4.7853
	Random	5	.8800	.9763	4.7940	.8600	.9735	4.7952
		10	.8675	.9776	4.7860	.8575	.9732	4.8012
	RAG-BM25	5	.8775	.9741	4.7980	.8550	.9740	4.7932
		10	.8825	.9770	4.7920	.8475	.9734	4.8012
	RAG-Dense	5	.8875	.9790	4.8140	.8625	.9720	4.8032
		10	.8825	.9759	4.7940	.8425	.9698	4.8131
	RAG-Pattern	5	.8650	.8796	4.7380	.8500	.9780	4.7714
		10	.8700	.8801	4.7080	.8550	.9780	4.7773
	RAG-Hybrid	10	.8850	.9773	4.7980	.8600	.9724	4.7932

Table 2: Results across English and Spanish sub-tasks. EM (Exact Match) and SAS (Semantic Answer Similarity) are computed on the 400-instance validation set per language. LLM denotes the official blind test score evaluated by an LLM-as-judge on a 1–5 scale. Bold values indicate the best result per metric within each model group.

Model	Mode	k	English	Spanish
Finetuned GPT-4.1-mini	RAG-Dense	(5, 10)	4.8140	4.8131
Our Ranking			1 st (tie)	1 st

Table 3: Official rankings based on LLM scores (out of 5) on the blind test set for the FinCausal 2026 shared task. The system used 5-shot prompting for English and 10-shot prompting for Spanish.

These configurations were ranked first among all participating systems for both sub-tasks, as illustrated in Table 3. A small variance was observed across fine-tuned settings (4.7440–4.8140 for English; 4.7853–4.8131 for Spanish), indicating stable performance across different RAG approaches. These results demonstrate that fine-tuned GPT-4.1-mini combined with RAG-Dense is the most effective approach for causal Q&A in financial texts.

Error Analysis. Error analysis was conducted on the development set, as gold answers for the official test set were not released to participants. We analysed the mismatched predictions from the best-performing fine-tuned configuration RAG-Dense ($k=5$) and identified two main types of errors: (1) *span boundary errors*, where the model extracted a slightly longer or shorter span than the gold answer, typically by including or omitting a leading

clause; and (2) *causal direction confusion*, where the model confused the direction of causality when multiple cause–effect relations were present in the context. For instance:

Q: “Why did Legendary increase its stake in Virtual Stock from 6.8% to 7.2%?”

C: “Subsequent to this investment and also in July 2017, Legendary increased its stake in Virtual Stock from 6.8% (subsequent to **the dilution due to Notion’s investment**) to 7.2% increasing the carrying value of its investment in Virtual Stock to £4.3m.” (*A* indicated by **bold font**)

A (predicted): *increasing the carrying value of its investment in Virtual Stock to £4.3m.*

These errors highlight the difficulty of distinguishing between causal explanations and subsequent financial outcomes in complex financial narratives.

7. Conclusion and Future Work

This paper presented a systematic comparison of RAG approaches combined with fine-tuning for causal Q&A in financial texts in English and Spanish. Fine-tuning GPT-4.1-mini on bilingual data improved performance across both languages. RAG-Dense further improved performance, achieving the best results in both sub-tasks. RAG-Pattern improved base model performance but showed no consistent gains after fine-tuning, suggesting that fine-tuning reduced reliance on causal pattern matching. Our system ranked first in both sub-tasks. Future work will explore more advanced RAG approaches and extend the approach to other domains and languages.

Code Availability

Code to reproduce experimental results reported in this paper is publicly available⁵.

Bibliographical References

- Bushra Aljohani and Tawfeeq Alsanoosy. 2026. [Enhancing medical question answering with llms via a hybrid retrieval-augmented generation framework](#). *Information*, 17(2):133.
- Lang Cao, Shihua Zhang, and Juxing Chen. 2022. [Cbcp: A method of causality extraction from unstructured financial text](#). In *Proceedings of the 2021 5th International Conference on Natural Language Processing and Information Retrieval*, NLPPIR '21, page 135–140, New York, NY, USA. Association for Computing Machinery.
- Gordon V. Cormack, Charles L A Clarke, and Stefan Buettcher. 2009. [Reciprocal rank fusion outperforms condorcet and individual rank learning methods](#). In *Proceedings of the 32nd International ACM SIGIR Conference on Research and Development in Information Retrieval*, SIGIR '09, page 758–759, New York, NY, USA. Association for Computing Machinery.
- Qiao Jin, Won Kim, Qingyu Chen, Donald C Comeau, Lana Yeganova, W John Wilbur, and Zhiyong Lu. 2023. Medcpt: Contrastive pre-trained transformers with large-scale pubmed search logs for zero-shot biomedical information retrieval. *Bioinformatics*, 39(11):btad651.
- Vladimir Karpukhin, Barlas Oguz, Sewon Min, Patrick Lewis, Ledell Wu, Sergey Edunov, Danqi Chen, and Wen-tau Yih. 2020. [Dense passage retrieval for open-domain question answering](#). In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 6769–6781, Online. Association for Computational Linguistics.
- Patrick Lewis, Ethan Perez, Aleksandra Piktus, Fabio Petroni, Vladimir Karpukhin, Naman Goyal, Heinrich Küttler, Mike Lewis, Wen-tau Yih, Tim Rocktäschel, Sebastian Riedel, and Douwe Kiela. 2020. Retrieval-augmented generation for knowledge-intensive nlp tasks. In *Proceedings of the 34th International Conference on Neural Information Processing Systems*, NIPS '20, Red Hook, NY, USA. Curran Associates Inc.
- Ying Li, Xiaosha Xue, Zhipeng Liu, Peibo Duan, and Bin Zhang. 2024. [Implicit-causality-exploration-enabled graph neural network for stock prediction](#). *Information*, 15:743.
- Jerry Liu. 2022. [LlamaIndex](#).
- Dominique Mariko, Hanna Abi-Akl, Kim Trottier, and Mahmoud El-Haj. 2022. [The financial causality extraction shared task \(FinCausal 2022\)](#). In *Proceedings of the 4th Financial Narrative Processing Workshop @LREC2022*, pages 105–107, Marseille, France. European Language Resources Association.
- Dominique Mariko, Hanna Abi Akl, Estelle Labidurie, Stéphane Durfort, Hugues de Mazancourt, and Mahmoud El-Haj. 2020. [Financial document causality detection shared task \(fincausal 2020\)](#).
- Antonio Moreno Sandoval, Blanca Carbajo Coronado, Jordi Porta Zamorano, Yanco Amor Torterolo Orta, and Doaa Samy. 2025. [The financial document causality detection shared task \(FinCausal 2025\)](#). In *Proceedings of the Joint Workshop of the 9th Financial Technology and Natural Language Processing (FinNLP), the 6th Financial Narrative Processing (FNP), and the 1st Workshop on Large Language Models for Finance and Legal (LLMFinLegal)*, pages 214–221, Abu Dhabi, UAE. Association for Computational Linguistics.
- Antonio Moreno-Sandoval, Jordi Porta, Yanco Torterolo, Alexia Stanescu, Melina Chatzi, and Sofía Roseti. 2026. The financial document causality detection shared task (fincausal 2026). In *Proceedings of the 7th Financial Narrative Processing Workshop (FNP 2026) at LREC 2026*, Palma de Mallorca, Spain. ELRA. To appear.
- Antonio Moreno-Sandoval, Jordi Porta-Zamorano, Blanca Carbajo-Coronado, Doaa Samy, Dominique Mariko, and Mahmoud El-Haj. 2023. [The](#)

⁵<https://github.com/aaalqarni/fincausal-2026>

- financial document causality detection shared task (fincausal 2023). In *2023 IEEE International Conference on Big Data (BigData)*, pages 2855–2860.
- Georg Niess, Houssam Razouk, Stasa Mandic, and Roman Kern. 2025. [Addressing hallucination in causal Q&A: The efficacy of fine-tuning over prompting in LLMs](#). In *Proceedings of the Joint Workshop of the 9th Financial Technology and Natural Language Processing (FinNLP), the 6th Financial Narrative Processing (FNP), and the 1st Workshop on Large Language Models for Finance and Legal (LLMFinLegal)*, pages 253–258, Abu Dhabi, UAE. Association for Computational Linguistics.
- Julian Risch, Timo Möller, Julian Gutsch, and Malte Pietsch. 2021. [Semantic answer similarity for evaluating question answering models](#). In *Proceedings of the 3rd Workshop on Machine Reading for Question Answering*, pages 149–157, Punta Cana, Dominican Republic. Association for Computational Linguistics.
- Stephen Robertson and Hugo Zaragoza. 2009. [The probabilistic relevance framework: Bm25 and beyond](#). *Found. Trends Inf. Retr.*, 3(4):333–389.
- Neelesh K Shukla, Raghu Katikeri, Msp Raja, Gowtham Sivam, Shlok Yadav, Amit Vaid, and Shreenivas Prabhakararao. 2023. Investigating large language models for financial causality detection in multilingual setup. In *2023 IEEE International Conference on Big Data (BigData)*, pages 2866–2871. IEEE.
- Tianyi Yang, Nashrah Haque, Vaishnave Jonnalagadda, Yuya Jeremy Ong, Zhehui Chen, Yanzhao Wu, Lei Yu, Divyesh Jadav, and Wenqi Wei. 2026. [Augmenting question answering with a hybrid rag approach](#).
- Chatzi, Melina. 2026. *The Financial Document Causality Detection Shared Task (FinCausal 2026): Dataset*. e-cienciaDatos.
- Ankit Pal, Logesh Kumar Umapathi, and Malaikanan Sankarasubbu. 2022. [Medmcqa : A large-scale multi-subject multi-choice dataset for medical domain question answering](#).

Language Resource References

- Di Jin, Eileen Pan, Nassim Oufattole, Wei-Hung Weng, Hanyi Fang, and Peter Szolovits. 2021. What disease does this patient have? a large-scale open domain question answering dataset from medical exams. *Applied Sciences*, 11(14):6421.
- Qiao Jin, Bhuwan Dhingra, Zhengping Liu, William W. Cohen, and Xinghua Lu. 2019. [PubMedqa: A dataset for biomedical research question answering](#).
- Moreno-Sandoval, Antonio and Torterolo Orta, Yanco Amor and Stanescu, Maria Alexia and

Causal Connections: Leveraging Multilingual Fine-Tuning for Financial QA@FinCausal 2026

Akash Kumar Gautam, Serhii Hamotskyi, Christian Hänig

Anhalt University of Applied Sciences

{akash-kumar.gautam, serhii.hamotskyi, christian.haenig}@hs-anhalt.de

Abstract

This paper describes team HSA_CORAL's submission to the FinCausal 2026 shared task on extracting cause–effect relations from financial narratives via extractive question answering in English and Spanish. We compare three modeling families: (i) encoder-only token tagging with multilingual BERT, (ii) encoder–decoder generation with multilingual BART, and (iii) decoder-only LLMs (Llama 3.1 and GPT variants) using prompt refinement, few-shot demonstrations, and supervised fine-tuning. Across settings, prompting and few-shot examples yield competitive performance, while supervised fine-tuning provides the largest gains. Our best system, GPT-4.1 Mini fine-tuned on combined English and Spanish training data, achieves a tied highest score on the English subtask (score 4.8140) and ranks third on Spanish (score 4.7753) under the shared task's LLM-as-a-judge metric. Overall, the results highlight the value of task-specific adaptation and multilingual fine-tuning for cross-lingual transfer in financial causality QA.

Keywords: Large Language Models, Financial NLP, Financial Question Answering.

1. Introduction

Understanding cause–effect relationships in financial texts is essential for informed decision-making. They reveal drivers of stock prices, economic changes, market behavior, and regulatory decisions across different countries. For analysts and investors, identifying causal connections provides valuable insights into financial risks, potential investments, and strategic planning (Gopalakrishnan et al., 2023).

The FinCausal shared task has steadily advanced how we detect causality in finance. Earlier editions focused on identifying causal phrases directly in text. Later versions introduced more complex challenges, such as recognizing implicit causality and multi-step reasoning (Mariko et al., 2022, 2020; Zavitsanos et al., 2023). The most recent task required generative models to answer open-ended questions about causes and effects, using Exact Match (EM) and Semantic Answer Similarity (SAS) (Risch et al., 2021) as evaluation metrics (Moreno-Sandoval et al., 2025).

The 2026 edition introduces several new elements (Moreno-Sandoval et al., 2026). The dataset includes fragments with new complex causal relationships, rephrased questions that demand more sophisticated reasoning, and randomly divided train-test splits. A novel challenge is the use of an LLM-as-a-judge for evaluation, which rates responses on a 1–5 adequacy scale. The task pushes models beyond simple text matching toward understanding both explicit and implicit causal relationships in detailed financial narratives.

We explored three approaches to extractive question-answering: (i) token classification with encoder-only models like BERT, (ii) sequence-to-sequence generation with encoder-decoder models like BART, and (iii) few-shot prompting with decoder-only large language models.

Across our experiments, we find that fine-tuned generative models consistently outperform encoder-only and encoder-decoder architectures on the extractive QA task. Notably, GPT-4.1 Mini fine-tuned on a multilingual corpus achieves the best overall performance, securing the top position on the English subtask and third place on Spanish. Adding 20 few-shot examples proves critical, enabling even a smaller model to surpass its larger counterparts in zero-shot settings. These results demonstrate that multilingual fine-tuning and few-shot learning together provide a reliable approach for causal relation extraction in financial documents.

2. Related Work

Causal Information Extraction Early approaches to causal relationship detection in financial texts relied heavily on rule-based systems and traditional machine learning methods such as Support Vector Machines (SVMs) and decision trees (Ghosh and Naskar, 2022; Baranes et al., 2019). While these models demonstrated some success in identifying patterns in financial reports and news articles, they required extensive feature engineering and often struggled to capture the complexity and temporal dynamics of causal relations in domain-specific documents.

The introduction of BERT (Devlin et al., 2019) and its multilingual variants marked a significant shift in the field (Yang et al., 2019; Wan and Li, 2022). These pre-trained language models enabled more nuanced understanding of contextual information with minimal feature extraction, substantially improving performance on tasks involving document-level comprehension (Zhang and Jankowski, 2022).

Subsequent work has focused on fine-tuning these models on domain-specific financial datasets, achieving state-of-the-art results in tasks such as sentiment analysis and event extraction (Mariko et al., 2020). Fine-tuned pre-trained language models have consistently outperformed traditional machine learning approaches, particularly when working with large-scale financial reports (Jin et al., 2023; Huang et al., 2023; Sarmah et al., 2023). Liu et al. (2023) propose an implicit cause-effect interaction framework to improve event causality extraction.

Pretrained Language Models More recently, proprietary large language models such as GPT-4 have been explored for question-answering tasks in the financial domain (Zhang et al., 2023; Kalpakchi and Boye, 2023). These models have demonstrated strong few-shot learning capabilities, enabling them to generalize from limited examples without task-specific fine-tuning (Xiao et al., 2022; Guo et al., 2023). This line of work underscores the growing potential of generative models for specialized NLP tasks, including the extraction of causal relationships in finance (Nayak et al., 2022; Kim et al., 2023).

3. Methodology

We compare three approaches to extractive QA for FinCausal 2026: (1) token classification using BERT-based models, (2) extractive QA using encoder-decoder models, and (3) generative models. These approaches are compared across a variety of pretrained large language models in few-shot settings using a multilingual dataset.

3.1. Encoder-Based Extractive QA

This approach utilizes text embedding models such as BERT for token classification, following a similar methodology to Yoon et al. (2022). The process begins by tokenizing both the context and the question, which are then concatenated with a special [SEP] token. During training, each sample is annotated using an IO tagging scheme, where

the answer span is mapped to its first occurrence within the passage.

We compute the cross-entropy loss between the predicted token classes and the ground truth labels derived from the training data. To refine the loss calculation, a loss mask is applied to restrict the computation exclusively to tokens originating from the passage. This masking strategy excludes tokens from the question and special symbols (e.g., [SEP], padding). This focuses learning on the passage tokens.

3.2. Encoder-Decoder Extractive QA

Our second approach treats extractive QA as a sequence-to-sequence generation task using BART (Lewis et al., 2020). The input consists of the question and context concatenated with a delimiter, and the model is trained to generate the answer token-by-token.

We fine-tune BART by minimizing the negative log-likelihood of the target answer sequence. During inference, beam search is employed to decode the most likely answer. This formulation is effective for extractive tasks as it leverages BART’s pre-trained language understanding while flexibly handling answer spans. We evaluate multilingual BART variants on both the English and Spanish subtasks.

3.3. Decoder-Based Extractive QA

The flexibility of large language models makes them well-suited for extractive QA tasks. However, it is important to ensure that models follow instructions and avoid hallucinations beyond the provided context (Xu et al., 2024). To address this, we implement a multi-step strategy combining prompt optimization, few-shot selection, and fine-tuning.

First, we perform prompt optimization through several iterative refinements on a small subset of the training dataset. The final version of the prompt used is described in Appendix A. This process helps identify instruction formats and phrasings that consistently yield extractive, context-grounded answers. We incorporate few-shot examples within the prompt, selecting relevant QA demonstrations using cosine similarity. Specifically, given a test context and question, we retrieve the most similar QA pairs from the training set based on multilingual embedding similarity¹, and include them as examples in the prompt to guide the model’s output.

¹<https://huggingface.co/sentence-transformers/paraphrase-multilingual-MiniLM-L12-v2>

Model	Fine-tuned	Corpus	English	Spanish
BERT Base Multilingual	yes	Multilingual (en+es)	3.9800	3.9810
BART Facebook Base	yes	Multilingual (en+es)	4.1200	4.0300
Llama-3.1 8B	yes	Multilingual (en+es)	4.0200	3.9100
GPT-3.5 turbo	no	-	4.7040	4.7060
GPT-4.1 mini	yes	Monolingual (en)	4.7560	4.7141
GPT-4.1 mini	yes	Monolingual (es)	4.7210	4.7674
GPT-4.1 mini	yes	Multilingual (en+es)	4.8140	4.7753
GPT-5.2	no	-	4.7600	4.7350

Table 1: Evaluation results on blinded English and Spanish test sets, with scores provided by an external LLM judge. The *Corpus* column indicates whether models were fine-tuned on monolingual or multilingual data. Bold entries denote the highest score in each language column. Decoder-only models were tested with varying numbers of few-shot examples; results, obtained with 20 examples, are reported here. Best results using our approach are presented in **bold**.

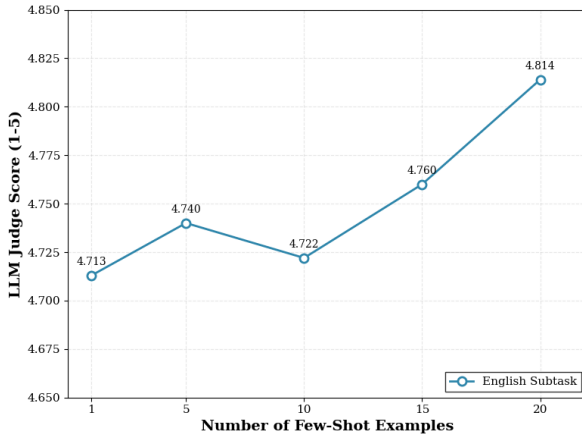


Figure 1: External LLM-judge score (English subtask) as a function of the number of few-shot demonstrations in the prompt for the best-performing decoder-only configuration.

As a further enhancement, we fine-tune the model on up to 2,000 samples, experimenting with three configurations: training on English data only, Spanish data only, and a combined bilingual dataset. Fine-tuning reinforces task-specific behavior and significantly reduces hallucinations, leading to more reliable extractive answers.

This hybrid approach—combining prompt engineering, dynamic few-shot selection, and targeted fine-tuning—allows us to leverage the generative strength of LLMs while maintaining the extractive precision required for the task.

4. Experiments

4.1. Dataset

The dataset comprises financial narratives in English and Spanish designed for an extractive question answering task focused on causal relationships [Moreno-Sandoval et al. \(2026\)](#). Given a context and a question, the task requires identifying a text span that expresses a causal relation within the financial narrative ([Moreno-Sandoval et al., 2025](#)).

Questions are formulated abstractly, targeting either the cause or the effect described in the text, with causes defined as agents or facts that can be extracted verbatim from the provided context.

For complex causal structures—such as causal chains or non-linear relationships—up to two questions are included per context to ensure comprehensive coverage. The English dataset is sourced from financial annual reports from 2017, drawn from the UCREL corpus and the English portion of the 2018 FinT-esp corpus, while the Spanish dataset is extracted from a corpus of Spanish financial annual reports spanning 2014 to 2018. The training set for both subtasks consists of 2,000 samples, with test sets of 500 samples for English and 503 samples for Spanish, respectively.

4.2. Model Selection

BERT² was used as an encoder model, multilingual BART³ for encoder-decoder model and we evaluated Llama-3.1 and several GPT⁴ variants for

²<https://huggingface.co/google-bert/bert-base-multilingual-cased>

³<https://huggingface.co/facebook/bart-base>

⁴<https://developers.openai.com/api/docs/models>

decoder-only setups. For fine-tuning of Llama3.1 we used Low-Rank Adaptation (LoRA) to speed the process (Hu et al., 2022).

4.3. Results

Table 1 reports the best-performing configuration for each model family. Scores correspond to the shared task’s external LLM-as-a-judge evaluation on the blind test set, rated on a 1–5 adequacy scale for both English and Spanish.

Encoder and Encoder-Decoder Setups. Extractive models such as BERT perform reasonably well at identifying exact spans corresponding to causal relationships, particularly when the answer is explicitly stated in the context. However, BART consistently outperforms BERT across both languages, benefiting from its sequence-to-sequence formulation which allows for more flexible and accurate span generation. Interestingly, the multilingual variant of BERT fine-tuned on combined English and Spanish data yields better performance on the blinded test set than its monolingual counterparts trained on individual languages. For brevity, we report only the best-performing configurations.

Few-Shot Learning with Decoder-Only Models

In our decoder-only experiments, we evaluated Llama 3.1 and several variants of GPT. For fine-tuning Llama 3.1, we employed Low-Rank Adaptation (LoRA) to reduce computational overhead while maintaining task performance. Our findings indicate that while well-structured prompts with clear instructions are effective at reducing hallucinations, the inclusion of few-shot examples substantially improves the quality and precision of the generated answer spans. Notably, adding examples from the training set resulted in substantial improvements for GPT-4.1 Mini compared to zero-shot settings—even enabling it to outperform a much larger model (GPT-5.2⁵) on both Spanish and English subtasks.

Figure 1 plots the LLM-as-a-judge scores against the increasing number of few-shot examples included in the prompt. The trend shows a clear improvement in output quality as more examples are added, with scores increasing correspondingly on the leaderboard. However, increasing the number of examples beyond a certain point did not yield further gains; in some cases, it even led to the generation of hallucinated content.

Given the context window constraints of each model, we systematically varied the number of few-

shot examples. The best results across all decoder-only models were achieved with 20 few-shot examples, which we adopt as our final configuration.

Fine-Tuning Our experiments reveal a clear advantage of fine-tuned generative models over both encoder-only and encoder-decoder architectures. Multilingual fine-tuning consistently outperforms monolingual fine-tuning on both English and Spanish subtasks, demonstrating strong cross-lingual transfer for causal relation extraction in financial narratives. Among all configurations, GPT-4.1 Mini with multilingual fine-tuning achieved the best results across both languages.

Notably, this fine-tuned model outperforms a much larger model, GPT-5.2, in zero-shot settings. This suggests that the test dataset’s contexts and questions contain lexical characteristics and causal relationships (for both Spanish and English) that can be reliably identified only when models have access to task-specific annotation guidelines through fine-tuning. The observation points to an interesting conclusion: fine-tuned pre-trained language models of moderate size may offer better performance for specialized tasks than much larger models used in inference-only mode, highlighting the importance of fine-tuning over parameter scaling for domain-specific applications.

Error Analysis and Limitations While our models achieved strong overall results, error analysis reveals persistent errors in handling nested causal structures and contexts containing multiple potential causes. In these challenging cases, models often selected the wrong causal pair, with errors more frequent in the Spanish subtask—suggesting that cross-lingual generalization remains challenging for complex causal structures.

The lack of detailed annotation guidelines for identifying relevant spans further compounds this issue, as prompt optimization alone cannot fully resolve such ambiguities. Looking ahead, alignment techniques such as Direct Preference Optimization (Rafailov et al., 2023) and Reinforcement Learning from Human Feedback (Bai et al., 2022) offer promising ways for learning the implicit patterns underlying human annotations for structurally complex examples.

5. Conclusion

This work demonstrates that generative models outperform encoder and encoder-decoder architectures on causal question-answering tasks for

⁵OpenAI currently supports supervised fine-tuning for GPT models up to GPT-4.1-2025-04-14.

financial narrative documents, provided that hallucinations are mitigated through careful prompt engineering and few-shot examples. Multilingual fine-tuning with a large number of training samples from both Spanish and English, significantly boosts performance, highlighting the effectiveness of cross-lingual transfer in this domain. Future work may explore LLM-based evaluation as a means to further enhance output quality, as well as more sophisticated strategies for handling nested and ambiguous causal structures.

Acknowledgments

This work has been completed as part of project CORAL (Constrained Retrieval-Augmented Language Models), funded by the German Federal Ministry of Research, Technology, and Space (BMFTR) under grant number 16IS24077C.

6. References

- Yuntao Bai, Andy Jones, Kamal Ndousse, Amanda Askell, Anna Chen, Nova DasSarma, Dawn Drain, Stanislav Fort, Deep Ganguli, Tom Henighan, et al. 2022. Training a helpful and harmless assistant with reinforcement learning from human feedback. *arXiv preprint arXiv:2204.05862*.
- Amos Baranes, Rimona Palas, et al. 2019. Earning movement prediction using machine learning-support vector machines (svm). *Journal of Management Information and Decision Sciences*, 22(2):36–53.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. Bert: Pre-training of deep bidirectional transformers for language understanding. In *Proceedings of the 2019 conference of the North American chapter of the association for computational linguistics: human language technologies, volume 1 (long and short papers)*, pages 4171–4186.
- Sohom Ghosh and Sudip Kumar Naskar. 2022. Lipi at fincausal 2022: Mining causes and effects from financial texts. In *Proceedings of the 4th financial narrative processing workshop@LREC2022*, pages 121–123.
- Seethalakshmi Gopalakrishnan, Victor Zitian Chen, Wenwen Dou, Gus Hahn-Powell, Sreekar Neddunuri, and Wlodek Zadrozny. 2023. Text to causal knowledge graph: A framework to synthesize knowledge from unstructured business texts into causal graphs. *Information*, 14(7):367.
- Biyang Guo, Xin Zhang, Ziyuan Wang, Minqi Jiang, Jinran Nie, Yuxuan Ding, Jianwei Yue, and Yupeng Wu. 2023. How close is chatgpt to human experts? comparison corpus, evaluation, and detection. *arXiv preprint arXiv:2301.07597*.
- Edward J Hu, Yelong Shen, Phillip Wallis, Zeyuan Allen-Zhu, Yuanzhi Li, Shean Wang, Liang Wang, Weizhu Chen, et al. 2022. Lora: Low-rank adaptation of large language models. *ICLR*, 1(2):3.
- Allen H Huang, Hui Wang, and Yi Yang. 2023. Finbert: A large language model for extracting information from financial text. *Contemporary Accounting Research*, 40(2):806–841.
- Yiqiao Jin, Xiting Wang, Yaru Hao, Yizhou Sun, and Xing Xie. 2023. Prototypical fine-tuning: Towards robust performance under varying data sizes. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 37, pages 12968–12976.
- Dmytro Kalpakchi and Johan Boye. 2023. Quasi: a synthetic question-answering dataset in swedish using gpt-3 and zero-shot learning. In *Proceedings of the 24th Nordic Conference on Computational Linguistics (NoDaLiDa)*, pages 477–491.
- Yuheun Kim, Lu Guo, Bei Yu, and Yingya Li. 2023. Can chatgpt understand causal language in science claims? In *Proceedings of the 13th Workshop on Computational Approaches to Subjectivity, Sentiment, & Social Media Analysis*, pages 379–389.
- Mike Lewis, Yinhan Liu, Naman Goyal, Marjan Ghazvininejad, Abdelrahman Mohamed, Omer Levy, Veselin Stoyanov, and Luke Zettlemoyer. 2020. 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*, pages 7871–7880.
- Jintao Liu, Zequn Zhang, Kaiwen Wei, Zhi Guo, Xian Sun, Li Jin, and Xiaoyu Li. 2023. Event causality extraction via implicit cause-effect interactions. In *Proceedings of the 2023 conference on empirical methods in natural language processing*, pages 6792–6804.
- Dominique Mariko, Hanna Abi-Akl, Estelle Labidurie, Stephane Durfort, Hugues De Mazancourt, and Mahmoud El-Haj. 2020. The financial document causality detection shared task (fincausal 2020). In *Proceedings of the 1st Joint Workshop on Financial Narrative Processing and MultiLing Financial Summarisation*, pages 23–32.
- Dominique Mariko, Hanna Abi-Akl, Kim Trottier, and Mahmoud El-Haj. 2022. The financial causality

- extraction shared task (fincausal 2022). In *Proceedings of the 4th Financial Narrative Processing Workshop@ LREC2022*, pages 105–107.
- Antonio Moreno-Sandoval, Blanca Carbajo-Coronado, Jordi Porta Zamorano, Yanco Amor Torterolo Orta, and Doaa Samy. 2025. The financial document causality detection shared task (fincausal 2025). In *Proceedings of the Joint Workshop of the 9th Financial Technology and Natural Language Processing (FinNLP), the 6th Financial Narrative Processing (FNP), and the 1st Workshop on Large Language Models for Finance and Legal (LLMFinLegal)*, pages 214–221.
- Antonio Moreno-Sandoval, Jordi Porta, Yanco Torterolo, Alexia Stanescu, Melina Chatzi, and Sofía Roseti. 2026. The Financial Document Causality Detection Shared Task (FinCausal 2026). In *Proceedings of the 7th Financial Narrative Processing Workshop (FNP 2026) at LREC 2026*, Palma de Mallorca, Spain. ELRA.
- Tapas Nayak, Soumya Sharma, Yash Butala, Koustuv Dasgupta, Pawan Goyal, and Niloy Ganguly. 2022. A generative approach for financial causality extraction. In *Companion Proceedings of the Web Conference 2022*, pages 576–578.
- Rafael Rafailov, Archit Sharma, Eric Mitchell, Christopher D Manning, Stefano Ermon, and Chelsea Finn. 2023. Direct preference optimization: Your language model is secretly a reward model. *Advances in neural information processing systems*, 36:53728–53741.
- Julian Risch, Timo Möller, Julian Gutsch, and Malte Pietsch. 2021. Semantic answer similarity for evaluating question answering models. In *Proceedings of the 3rd Workshop on Machine Reading for Question Answering*, pages 149–157.
- Bhaskarjit Sarmah, Dhagash Mehta, Stefano Pasquali, and Tianjie Zhu. 2023. Towards reducing hallucination in extracting information from financial reports using large language models. In *Proceedings of the Third International Conference on AI-ML Systems*, pages 1–5.
- Chang-Xuan Wan and Bo Li. 2022. Financial causal sentence recognition based on bert-cnn text classification. *The Journal of Supercomputing*, 78(5):6503–6527.
- Jiayu Xiao, Liang Li, Chaofei Wang, Zheng-Jun Zha, and Qingming Huang. 2022. Few shot generative model adaption via relaxed spatial structural alignment. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 11204–11213.
- Ziwei Xu, Sanjay Jain, and Mohan Kankanhalli. 2024. Hallucination is inevitable: An innate limitation of large language models. *arXiv preprint arXiv:2401.11817*.
- Wei Yang, Yuqing Xie, Aileen Lin, Xingyu Li, Luchen Tan, Kun Xiong, Ming Li, and Jimmy Lin. 2019. End-to-end open-domain question answering with bertserini. In *Proceedings of the 2019 conference of the North American chapter of the association for computational linguistics (Demonstrations)*, pages 72–77.
- Wonjin Yoon, Richard Jackson, Aron Lagerberg, and Jaewoo Kang. 2022. Sequence tagging for biomedical extractive question answering. *Bioinformatics*, 38(15):3794–3801.
- Elias Zavitsanos, Aris Kosmopoulos, George Giannakopoulos, Marina Litvak, Blanca Carbajo-Coronado, Antonio Moreno-Sandoval, and Mo El-Haj. 2023. The financial narrative summarisation shared task (fns 2023). In *2023 IEEE International Conference on Big Data, BigData 2023*, pages 2890–2896. Institute of Electrical and Electronics Engineers.
- Le Zhang, Yihong Wu, Fengran Mo, Jian-Yun Nie, and Aishwarya Agrawal. 2023. Moqagpt: Zero-shot multi-modal open-domain question answering with large language model. In *Findings of the Association for Computational Linguistics: EMNLP 2023*, pages 1195–1210.
- Ning Zhang and Maciej Jankowski. 2022. Hierarchical bert for medical document understanding. *arXiv preprint arXiv:2204.09600*.

7. Language Resource References

- Antonio Moreno-Sandoval, Yanco Amor Torterolo Orta, Maria Alexia Stanescu, and Melina Chatzi. 2026. [The Financial Document Causality Detection Shared Task \(FinCausal 2026\): Dataset](#).

A. Prompt

Figure 2 describes the prompt used by decoder only language models for generating the final response as part of the extractive question-answering task.

Task Description

Given a financial context and a question, extract an exact answer from the context about cause or effect that addresses the question. The answer will be either the cause or the effect to a specific event mentioned in the context.

Instructions

1. **Read the context carefully:**
 - Understand the events and relationship described in the context.
2. **Understand the question:**
 - Determine if the question is asking for cause or effect.
 - Identify the specific events or statement the question refers to.
3. **Extract the answer verbatim:**
 - Locate the exact sentence or phrase in the context that answers the question.
 - **Do not paraphrase, summarise, or add any external information. The answer must be copied word-for-word from the context.**
4. **Provide only the answer:**
 - **Do not include any instructions, explanations or formatting.**
 - **Output only the extracted answer and nothing else.**
5. **Answer should be in the language in which question is asked and the context is mentioned:**
 - **Understand the context very well.**

Examples

```
{ { formatted_examples } }
```

Your Task

Context:

```
{ { context } }
```

Question:

```
{ { question } }
```

Answer:

[Provide only the exact answer from the context]

Remember

- **Output only the answer. Do not include any additional text. Do not include “Answer” in your output.**
- **The answer must exactly match a portion of the context.**
- **Do not add instructions, explanations, or any extra information.**

Figure 2: Prompt used for extractive QA by decoder models used in the FinCausal 2026 shared task.

VERSA: Verbatim Extraction via Rephrasing and Self-Aggregation for Financial Causality

Aldan Jay , Rafael Berlanga , Yoelvis Moreno, Vicent Santamarta

Escola de Doctorat, Universitat Jaume I, Castellón de la Plana, Spain

Universitat Jaume I, Castellón de la Plana, Spain

{jay, berlanga, alcayde, santamav}@uji.es

Abstract

Financial causality detection, the task of identifying and extracting verbatim causal spans from financial narratives, remains a challenging problem in Natural Language Processing (NLP). Large Language Models (LLMs), while powerful reasoners, frequently paraphrase source text or produce imprecise span boundaries when used in zero-shot extraction settings, leading to poor Exact Match scores. In this paper, we present VERSA, our system for the FinCausal 2026 Shared Task, a multi-agent pipeline that integrates two complementary inference strategies: Rephrase-and-Respond (RaR) and Recursive Self-Aggregation (RSA). The pipeline decomposes the extraction task into five sequential stages, each handled by a specialised agent: (1) causal structure analysis, (2) question reformulation via RaR, (3) diverse candidate population generation, (4) iterative refinement through RSA, and (5) verbatim validation with word-boundary alignment. We evaluate our approach on both the English and Spanish subsets of the FinCausal 2026 dataset. An ablation study demonstrates the individual and combined contributions of RaR and RSA, showing that the full pipeline substantially outperforms a zero-shot baseline in Exact Match and token-level F1.

Keywords: financial causality, extractive question answering, multi-agent systems, large language models, recursive self-aggregation

1. Introduction

The automatic extraction of causal relationships from financial documents is a long-standing challenge in Financial NLP. Causal reasoning is central to financial analysis: understanding why revenue declined, what drove a loss, or which factors contributed to market movements requires precise identification of cause–effect spans within narrative text. The FinCausal shared task series (Mariko et al., 2020, 2022; Moreno-Sandoval et al., 2023; Moreno Sandoval et al., 2025) has established a rigorous evaluation framework for this problem, requiring systems to return *verbatim* text spans that correctly identify the cause or effect queried in a given question.

Modern Large Language Models (LLMs) have demonstrated strong performance in a wide range of NLP tasks (Brown et al., 2020), including question answering and information extraction. However, when applied to extractive tasks requiring exact span matching, LLMs exhibit systematic weaknesses. In zero-shot settings, they tend to *paraphrase* rather than *extract*, they *hallucinate* causal connectives (e.g., prepending “due to” to the extracted span), and they produce *inconsistent span boundaries*—for instance, omitting an initial determiner or including trailing punctuation. These behaviours, while often semantically harmless, are severely penalised by Exact Match (EM) metrics.

The 2026 edition of FinCausal (Moreno-Sandoval et al., 2026) introduces additional complexity: over 500 new fragments containing

multi-element causal chains, abstractive question rephrasing in approximately 10% of the dataset, and a new LLM-as-a-judge evaluation metric scored on a 1–5 adequacy scale (Zheng et al., 2024). These changes demand systems capable of deep reasoning beyond surface-level lexical matching.

To address these challenges, we propose **VERSA** (*Verbatim Extraction via Rephrasing and Self-Aggregation*), a multi-agent pipeline that decomposes the extraction task into five specialised stages. Our approach integrates two key techniques from recent research:

1. **Rephrase-and-Respond (RaR)** (Deng et al., 2023): a prompting strategy in which the model reformulates the input question to align it with its own internal reasoning frame, thereby reducing ambiguity prior to extraction.
2. **Recursive Self-Aggregation (RSA)** (Venktraman et al., 2025): an iterative refinement mechanism that maintains a population of candidate answers and recursively aggregates subsets to converge towards a robust consensus, rather than relying on a single inference pass or simple majority voting.

VERSA operates on both the English and Spanish portions of the FinCausal 2026 dataset. The remainder of this paper is organised as follows: Section 2 reviews related work; Section 3 describes our methodology in detail; Section 4 presents the experimental setup; Section 5 reports results and

an ablation study; and Section 6 offers concluding remarks.

2. Related Work

2.1. Financial Causality Detection

The FinCausal shared task series, initiated at COLING 2020 (Mariko et al., 2020), formulated financial causality detection as an extractive question answering (QA) problem. Subsequent editions (Mariko et al., 2022; Moreno-Sandoval et al., 2023; Moreno Sandoval et al., 2025) refined the annotation scheme and expanded coverage to Spanish and multi-element causal chains, culminating in the 2026 edition (Moreno-Sandoval et al., 2026) evaluated in this work. Given a financial text passage and a causal question, systems must return the exact text span containing the queried cause or effect. Previous editions have seen strong participation from systems based on fine-tuned Transformer encoders such as BERT (Devlin et al., 2019) and XLM-RoBERTa (Conneau et al., 2020), often coupled with Conditional Random Fields (CRFs) for token-level sequence labelling (Lafferty et al., 2001). While effective, these approaches require task-specific fine-tuning and struggle with complex, multi-hop causal chains where the answer spans multiple clauses.

2.2. Prompting Strategies for Extraction

The advent of instruction-tuned LLMs has enabled zero-shot and few-shot extractive QA without task-specific training (Brown et al., 2020). However, LLMs frequently reformulate rather than extract, a fundamental tension between their generative nature and the extractive requirement of tasks like FinCausal. Chain-of-Thought prompting (Wei et al., 2022) and Self-Consistency (Wang et al., 2023) have improved reasoning quality, but neither directly addresses the verbatim constraint. The Rephrase-and-Respond (RaR) framework (Deng et al., 2023) proposes that LLMs reformulate ambiguous questions in their own terms before answering, effectively aligning the query with the model’s internal processing frame. We adapt this insight specifically for extractive QA, using the rephrased question to generate explicit extraction hints.

2.3. Aggregation-Based Inference

Recursive Self-Aggregation (RSA) (Venkatraman et al., 2025) extends the self-consistency paradigm by maintaining a population of N candidate solutions and iteratively recombining randomly sampled subsets of size K over T iterations. Unlike majority voting, which selects the most frequent answer, RSA synthesises new candidates from subsets,

enabling the correction of partial errors and convergence towards more complete answers. This approach has demonstrated improvements in mathematical reasoning and code generation, but has not, to our knowledge, been applied to extractive information extraction tasks.

3. Methodology

We propose a sequential multi-agent pipeline comprising five specialised agents, each responsible for a distinct stage of the extraction process. Figure 1 illustrates the overall workflow.

3.1. Stage 1: Causal Structure Analysis

The first agent analyses the input pair (c, q) —where c denotes the financial context and q the causal question—to produce a structured analysis $\mathcal{A}$ that guides subsequent stages. This analysis comprises three components:

Trigger detection. The agent identifies explicit causal markers in the context using pattern matching over language-specific lexicons. For English, these include expressions such as “due to,” “as a result of,” “driven by,” and “consequently.” For Spanish: “debido a,” “como consecuencia de,” “gracias a,” and “por lo que.” Each detected trigger is classified as *causal* (indicating a cause), *resultative* (indicating an effect), or *temporal-causal* (e.g., “following,” “tras”).

Directionality inference. The agent determines whether the question seeks the *cause* or the *effect* of the described relationship, a distinction critical for selecting the correct span when both are present in the context.

Complexity assessment. The number and arrangement of causal triggers are used to classify the relationship as *simple* (single cause–effect pair), *chain* (sequential cascade), or *multiple* (several concurrent causes or effects).

3.2. Stage 2: Question Reformulation (RaR)

Following Deng et al. (2023), who demonstrated that LLMs perform better when questions are restated in terms aligned with the model’s internal reasoning, we apply a Rephrase-and-Respond step. Rather than forwarding the original question q directly to the extraction stage, the Rephraser agent generates a reformulated question q' and a set of extraction hints H .

The reformulation incorporates the causal analysis $\mathcal{A}$: it makes the target direction explicit (e.g., transforming “Why did X happen?” into “Identify the cause of X,”) names specific entities from the context, and specifies the expected syntactic form

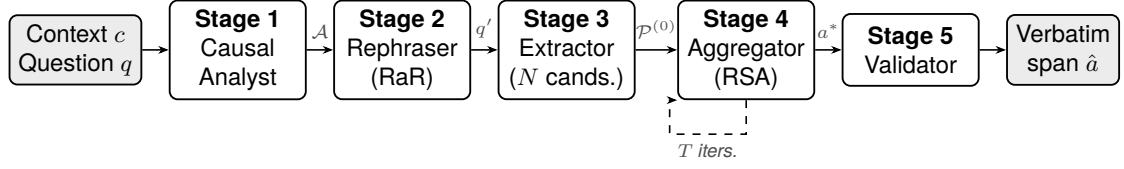


Figure 1: Overview of the proposed multi-agent pipeline. Stages 1–5 are executed sequentially. Stage 4 (Aggregator) performs T recursive iterations over the candidate population. Notation: $\mathcal{A}$ = causal analysis metadata; q' = rephrased question; $\mathcal{P}^{(0)}$ = initial candidate population; a^* = selected answer; $\hat{a}$ = final verbatim span.

of the answer (e.g., “a noun phrase following the marker ‘due to.’”) This step reduces the ambiguity inherent in abstractive questions—particularly relevant for the 10% of FinCausal 2026 questions that have been deliberately rephrased away from the source text.

3.3. Stage 3: Candidate Population Generation

Rather than producing a single extraction, the Extractor agent generates a *population* $\mathcal{P}^{(0)} = \{a_1, a_2, \dots, a_N\}$ of N candidate spans. This design choice is motivated by the observation that individual LLM inferences are stochastic: varying the decoding temperature or prompt formulation yields different, partially overlapping spans whose union often contains the correct answer.

The population is composed using two complementary methods:

- **LLM-based extraction:** The majority of candidates ($\lceil 3N/4 \rceil$) are obtained by querying the language model with the reformulated question q' and hints H at varying temperature values $\tau \in \{0.3, 0.5, 0.7, 1.0\}$. Each inference is independent, promoting diversity.
- **Machine Reading Comprehension (MRC):** The remaining candidates ($\lfloor N/4 \rfloor$) are produced by a pre-trained multilingual extractive QA model (XLM-RoBERTa fine-tuned on SQuAD 2.0; Rajpurkar et al., 2018; Conneau et al., 2020), providing a non-generative anchor that is inherently verbatim.

Exact duplicate spans are removed to ensure diversity. In our experiments, we set $N = 8$.

3.4. Stage 4: Recursive Self-Aggregation (RSA)

The core refinement mechanism of our pipeline applies RSA (Venkatraman et al., 2025) to the candidate population. At each iteration $t \in \{1, \dots, T\}$, we construct a new population $\mathcal{P}^{(t)}$ from $\mathcal{P}^{(t-1)}$ as follows:

1. For each position $i \in \{1, \dots, N\}$, sample a subset $S_i \subset \mathcal{P}^{(t-1)}$ of size K uniformly at random without replacement.
2. Present the K candidates in S_i , together with the context c and the original question q , to the Aggregator agent.
3. The agent compares the candidates, reasons about their respective merits, and synthesises an improved candidate $a_i^{(t)}$ that must appear verbatim in c .
4. Set $\mathcal{P}^{(t)} = \{a_1^{(t)}, \dots, a_N^{(t)}\}$.

Algorithm 1 formalises this procedure. Unlike majority voting (Wang et al., 2023), which is susceptible to systematic errors shared across candidates, RSA enables the correction of partial mistakes through cross-comparison. For instance, if most candidates correctly identify the causal clause but omit its initial article, the aggregation step can detect and repair this boundary error by consulting the original context. In our experiments, we set $K = 3$ and $T = 4$.

After T iterations, the final answer a^* is selected from $\mathcal{P}^{(T)}$ by majority vote over the converged population. Ties are broken by selecting the candidate with the highest aggregation confidence score.

Algorithm 1 Recursive Self-Aggregation (RSA)

Require: Population $\mathcal{P}^{(0)}$, context c , question q , subset size K , iterations T

Ensure: Final answer a^*

```

1: for  $t = 1$  to  $T$  do
2:    $\mathcal{P}^{(t)} \leftarrow \emptyset$ 
3:   for  $i = 1$  to  $|\mathcal{P}^{(t-1)}|$  do
4:      $S_i \leftarrow \text{RandomSample}(\mathcal{P}^{(t-1)}, K)$ 
5:      $a_i^{(t)} \leftarrow \text{Aggregate}(S_i, c, q)$ 
6:      $\mathcal{P}^{(t)} \leftarrow \mathcal{P}^{(t)} \cup \{a_i^{(t)}\}$ 
7:   end for
8: end for
9:  $a^* \leftarrow \text{MajorityVote}(\mathcal{P}^{(T)})$ 
10: return  $a^*$ 

```

3.5. Stage 5: Verbatim Validation

The final agent enforces the strict extractive constraint of the task. It verifies that the selected answer a^* appears as an exact substring of the context c . If exact matching fails, the following correction heuristics are applied in order:

1. **Normalisation**: whitespace collapsing and Unicode normalisation (NFC).
2. **Fuzzy matching**: the closest substring in c is identified using edit distance (Levenshtein, 1966), accepting matches below a threshold δ .
3. **Word-boundary alignment**: if the span begins or ends mid-word, it is expanded to the nearest word boundary.
4. **Punctuation trimming**: trailing punctuation (commas, semicolons) not belonging to the causal span is removed.

The output of this stage is the final verbatim span $\hat{a}$.

4. Experimental Setup

4.1. Dataset

We evaluate VERSA on the FinCausal 2026 dataset Moreno-Sandoval et al. (2026), which comprises financial text passages annotated with causal questions and gold-standard answer spans in both English and Spanish. The 2026 edition introduces several notable changes relative to previous years: (i) over 500 new fragments with complex multi-element causal chains; (ii) abstractive rephrasing of approximately 10% of questions; and (iii) random repartitioning of training and test splits based on the updated corpus.

Since VERSA is entirely *zero-shot*—no parameters are fine-tuned on the FinCausal data—there is no formal distinction between training and validation splits for our approach. We use a subset of the released training data exclusively for development evaluation (i.e., measuring EM and F1 against gold annotations). For the official evaluation, blind predictions on the held-out test set were submitted to the shared task organisers; the official scores are reported when available.

4.2. Language Model Configuration

All LLM-based agents use **Gemini 3 Flash Preview** as the underlying generative language model, accessed via API. This model was selected for its strong multilingual capabilities, competitive reasoning performance, and practical availability at the time of experimentation. No local or self-hosted

models were explored in this work; the rationale for this decision is discussed in Section 8. Agent-specific temperature settings are as follows: the Causal Analyst and Validator agents use $\tau = 0.1$ to promote deterministic outputs; the Rephraser uses $\tau = 0.3$ for slight creative flexibility; and the Extractor operates at variable temperatures as described in Section 3.3. The Aggregator uses $\tau = 0.2$ to encourage conservative synthesis.

4.3. Evaluation Metrics

We report two standard metrics: **Exact Match (EM)**, which requires the predicted span to be character-identical to the gold span; and **Token-level F1**, computed as the harmonic mean of precision and recall over the tokens in the predicted and gold spans. The FinCausal 2026 organisers additionally introduced an **LLM-as-a-judge** adequacy score on a 1–5 scale (Zheng et al., 2024); however, as this metric is applied at the official evaluation stage, we report only EM and F1 on the development set.

4.4. Ablation Design

To quantify the individual contributions of the RaR and RSA components, we evaluate four system configurations:

Configuration	RaR	RSA
Baseline (zero-shot)	—	—
+ RaR only	✓	—
+ RSA only	—	✓
Full pipeline	✓	✓

Table 1: Ablation configurations. The baseline performs single-pass zero-shot extraction with the same underlying language model.

5. Results and Analysis

5.1. Development Results

Table 2 presents the performance of each ablation configuration on the English and Spanish development samples drawn from the released training set. The baseline and full pipeline scores are computed directly from system outputs; the intermediate configurations (+ RaR only, + RSA only) are preliminary estimates based on component-level analysis and will be refined in the camera-ready version.

The full pipeline achieves 50.0% EM and 87.7% F1 on the English development sample and 33.3% EM and 79.4% F1 on Spanish. Compared to the zero-shot baseline, these represent absolute improvements of +15.0 and +28.3 EM points, respectively. The consistently high F1 scores across all

Configuration	EM (%)	F1 (%)
<i>English (n = 20)</i>		
Baseline (zero-shot)	35.0	82.0
+ RaR only	40.0 [†]	84.5 [†]
+ RSA only	45.0 [†]	86.2 [†]
Full pipeline	50.0	87.7
<i>Spanish (n = 20)</i>		
Baseline (zero-shot)	5.0	75.5
+ RaR only	15.0 [†]	77.8 [†]
+ RSA only	20.0 [†]	78.1 [†]
Full pipeline	33.3	79.4

Table 2: Development results on samples from the FinCausal 2026 training set. EM = Exact Match; F1 = token-level F1. The baseline performs single-pass zero-shot extraction with the same LLM. [†]Preliminary estimates from component-level analysis.

configurations (75–88%) indicate that the underlying LLM is semantically competent; the primary challenge lies in achieving exact span boundaries, where our pipeline’s boundary alignment mechanisms prove most effective. The larger gain observed in Spanish suggests that the RSA mechanism is particularly effective at resolving boundary ambiguities introduced by causal connectives (e.g., “debido a”, “como consecuencia de”) that the baseline frequently includes in the extracted span.

Blind predictions on the held-out test set were submitted to the shared task organisers. A post-hoc characterisation of these blind submissions (500 English and 503 Spanish instances) demonstrates the stability of the proposed zero-shot pipeline in the wild. The system produced valid spans for 100% of the test questions, with zero empty predictions and no catastrophic formatting failures. The extracted English spans had an average length of 25.7 tokens (median 21.0), representing 47.7% of the source context length on average. In Spanish, spans were slightly longer, averaging 31.2 tokens (median 27.0) and covering 43.7% of the context. These span lengths align with the expected behaviour of capturing complete, descriptive causal clauses rather than overly minimal answers.

5.2. Official Evaluation Results

Table 3 reports the official results released by the shared task organisers (Moreno-Sandoval et al., 2026) for the held-out test set, evaluated using the LLM-as-a-judge metric (Zheng et al., 2024) on a 1–5 adequacy scale. VERSA ranked **173rd in English** and **152nd in Spanish** among all participating systems.

Language	LLM Score	Rank
English	4.404	173
Spanish	4.336	152

Table 3: Official test-set results (Moreno-Sandoval et al., 2026). LLM score is the adequacy rating on a 1–5 scale. Rank is the system’s position in the official leaderboard.

5.3. Qualitative Analysis

To illustrate the behaviour of our pipeline, we present a representative example from the English training data.

Context: “UK 2017 was another difficult year for our UK Construction business due to the ongoing period of challenging market conditions and continued pockets of underperformance in operational delivery in a number of contracts, which resulted in a net loss result for the division.”

Question: “What were the reasons for the net loss result in the division?”

In a single-pass zero-shot setting, candidate extractions exhibit two common failure modes: (a) including the causal connective “due to” as part of the answer span, and (b) truncating the initial article “the”, yielding “ongoing period of...” rather than “the ongoing period of...”. Our pipeline addresses both issues. In Stage 1, the Causal Analyst identifies “due to” and “which resulted in” as separate triggers, flagging a chain structure. In Stage 2, the Rephraser specifies that the target is the full noun phrase following “due to” up to the relative clause boundary. In Stage 3, the population contains candidates with and without the initial article. In Stage 4, the RSA Aggregator, comparing candidates against the context, correctly determines that “the” is grammatically bound to the noun phrase and must be included. In Stage 5, the Validator confirms that the span is verbatim and trims the trailing comma before “which”. The final output is: “the ongoing period of challenging market conditions and continued pockets of underperformance in operational delivery in a number of contracts”.

5.4. Effect of RSA Iterations

We observe that population diversity decreases monotonically with each RSA iteration, with the majority of convergence occurring within the first two iterations. Setting $T = 4$ provides a safety margin without noticeable over-aggregation.

6. Conclusion

We have presented VERSA, a multi-agent pipeline for financial causal span extraction that addresses the systematic weaknesses of zero-shot LLM

extraction through two complementary mechanisms. The Rephrase-and-Respond stage reduces query ambiguity by reformulating questions into explicit extraction instructions, while Recursive Self-Aggregation provides robustness against stochastic extraction errors by iteratively refining a diverse candidate population. Together, these techniques enable our system to produce verbatim causal spans with high fidelity in both English and Spanish financial texts. Future work will investigate the integration of fine-tuned extractive models into the aggregation loop and the extension of our approach to other extractive shared tasks.

7. Ethics Statement

Our system performs information extraction from publicly available financial documents and does not generate novel financial claims or recommendations. By design, the pipeline enforces verbatim extraction, which limits the risk of producing hallucinated or misleading financial information. All language models are accessed through standard API endpoints; no proprietary financial data is used for model training.

8. Limitations

The primary limitation of our approach is computational cost. Generating $N = 8$ candidates and performing $T = 4$ RSA iterations, each requiring a full LLM inference call, results in a per-example latency that is approximately $N \times (T + 1) = 40$ times that of a single-pass extraction. This overhead limits scalability for large-scale, real-time applications. Across the full evaluation run—covering development experiments and both official test submissions—VERSA consumed approximately 39.1 million tokens (~56 000 API requests), incurring an estimated cost of \$27.3 USD at standard API rates. While exact CO₂ equivalents are unavailable from the API provider, the energy footprint is comparable to other API-intensive NLP evaluation workflows.

Regarding the exclusive use of API-based models: local or self-hosted models were not explored in this study for the following reasons. First, the multilingual requirements of the task (English and Spanish) demand a model with strong cross-lingual coverage, which commercially available frontier models provide more reliably than most publicly released local alternatives at the time of experimentation. Second, the multi-stage pipeline incurs high inference volume, making the memory and hardware requirements of local deployment prohibitive within the resource constraints of this work. Investigating the substitution of API calls with locally

hosted, quantised models remains an important direction for future work.

Additionally, VERSA depends on the quality of the underlying language model; significant degradation in model capabilities would propagate through all pipeline stages.

9. Acknowledgements

This work has been supported by SOLUCIONES CUATROOCHENTA S.A. through the project “SISTEMA DE GESTIÓN DE ALERTAS DE CIBERSEGURIDAD BASADO EN SISTEMAS DE INTELIGENCIA ARTIFICIAL” (UJI Code: 24I526). The authors thank the FinCausal 2026 organisers for providing the shared task infrastructure and dataset, and for the opportunity to participate in the FNP 2026 workshop.

10. Bibliographical References

- Tom B. Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, et al. 2020. [Language models are few-shot learners](#). *Advances in Neural Information Processing Systems*, 33:1877–1901.
- Alexis Conneau, Karttikeya Khandelwal, Naman Goyal, Vishrav Chaudhary, Guillaume Wenzek, Francisco Guzmán, Edouard Grave, Myles Ott, Luke Zettlemoyer, and Veselin Stoyanov. 2020. [Unsupervised cross-lingual representation learning at scale](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 8440–8451.
- Yihe Deng, Weitong Zhang, Zixiang Chen, and Quanquan Gu. 2023. [Rephrase and respond: Let large language models ask better questions for themselves](#). *arXiv preprint arXiv:2311.04205*.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. [BERT: Pre-training of deep bidirectional transformers for language understanding](#). *arXiv preprint arXiv:1810.04805*.
- John D. Lafferty, Andrew McCallum, and Fernando C. N. Pereira. 2001. Conditional random fields: Probabilistic models for segmenting and labeling sequence data. In *Proceedings of the Eighteenth International Conference on Machine Learning*, pages 282–289.
- Vladimir I. Levenshtein. 1966. Binary codes capable of correcting deletions, insertions, and reversals. *Soviet Physics Doklady*, 10(8):707–710.

- Dominique Mariko, Hanna Abi-Akl, Estelle Labidurie, Stephane Durfort, Hugues De Mazancourt, and Mahmoud El-Haj. 2020. [The financial document causality detection shared task \(FinCausal 2020\)](#). In *Proceedings of the 1st Joint Workshop on Financial Narrative Processing and MultiLing Financial Summarisation*, pages 23–32, Barcelona, Spain (Online). COLING.
- Dominique Mariko, Hanna Abi-Akl, Kim Trottier, and Mahmoud El-Haj. 2022. [The financial causality extraction shared task \(FinCausal 2022\)](#). In *Proceedings of the 4th Financial Narrative Processing Workshop @LREC2022*, pages 105–107, Marseille, France. European Language Resources Association.
- Antonio Moreno Sandoval, Blanca Carbajo Coronado, Jordi Porta Zamorano, Yanco Amor Torterolo Orta, and Doaa Samy. 2025. [The financial document causality detection shared task \(FinCausal 2025\)](#). In *Proceedings of the Joint Workshop of the 9th Financial Technology and Natural Language Processing (FinNLP), the 6th Financial Narrative Processing (FNP), and the 1st Workshop on Large Language Models for Finance and Legal (LLMFinLegal)*, pages 214–221, Abu Dhabi, UAE. Association for Computational Linguistics.
- Antonio Moreno-Sandoval, Jordi Porta, Yanco Torterolo, Alexia Stanescu, Melina Chatzi, and Sofía Roseti. 2026. The Financial Document Causality Detection Shared Task (FinCausal 2026). In *Proceedings of the 7th Financial Narrative Processing Workshop (FNP 2026) at LREC 2026*, Palma de Mallorca, Spain. ELRA. To appear.
- Antonio Moreno-Sandoval, Jordi Porta-Zamorano, Blanca Carbajo-Coronado, Doaa Samy, Dominique Mariko, and Mahmoud El-Haj. 2023. [The financial document causality detection shared task \(FinCausal 2023\)](#). In *2023 IEEE International Conference on Big Data (BigData)*, pages 2855–2860.
- Pranav Rajpurkar, Robin Jia, and Percy Liang. 2018. [Know what you don’t know: Unanswerable questions for SQuAD](#). In *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics*, pages 784–789.
- Siddarth Venkatraman, Vineet Jain, Sarthak Mittal, Vedant Shah, Johan Obando-Ceron, Yoshua Bengio, Brian R. Bartoldson, Bhavya Kaikhura, Guillaume Lajoie, Glen Berseth, Nikolay Malkin, and Moksh Jain. 2025. [Recursive self-aggregation unlocks deep thinking in large language models](#). *arXiv preprint arXiv:2509.26626*.
- Xuezhi Wang, Jason Wei, Dale Schuurmans, Quoc V. Le, Ed H. Chi, Sharan Narang, Aakanksha Chowdhery, and Denny Zhou. 2023. [Self-consistency improves chain of thought reasoning in language models](#). *arXiv preprint arXiv:2203.11171*.
- Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Brian Ichter, Fei Xia, Ed Chi, Quoc Le, and Denny Zhou. 2022. [Chain-of-thought prompting elicits reasoning in large language models](#). *Advances in Neural Information Processing Systems*, 35:24824–24837.
- Lianmin Zheng, Wei-Lin Chiang, Ying Sheng, Siyuan Zhuang, Zhanghao Wu, Yonghao Zhuang, Zi Lin, Zhuohan Li, Dacheng Li, Eric P. Xing, et al. 2024. [Judging LLM-as-a-judge with MT-bench and chatbot arena](#). In *Advances in Neural Information Processing Systems*, volume 36.

11. Language Resource References

- Moreno-Sandoval, Antonio and Torterolo Orta, Yanco Amor and Stanescu, Maria Alexia and Chatzi, Melina. 2026. [The Financial Document Causality Detection Shared Task \(FinCausal 2026\): Dataset](#). e-cienciaDatos.

SpanDiffusion: Flow Matching over Continuous Span Masks for Financial Causal Question Answering

Georg Niess¹, Roman Kern^{1,2}

¹Institute of Machine Learning and Neural Computation, Graz University of Technology

²Know Center Research GmbH

Graz, Austria

{georg.niess, rkern}@tugraz.at

Abstract

We present SpanDiffusion, a continuous diffusion approach to extractive causal question answering for the FinCausal 2026 shared task. SpanDiffusion uses two Gaussian masks, continuous signals with peaks at the answer start and end positions, and learns to denoise them from pure noise through a dedicated transformer conditioned on frozen DeBERTa-v3-large embeddings with LoRA adapters (1.6M parameters). By replacing Denoising Diffusion Probabilistic Models (DDPM) with flow matching (rectified flow), we reduce denoising to only 20 Euler steps at inference. A systematic ablation across six diffusion variants and a span-classification baseline shows that LoRA adaptation is the dominant factor (+34 Exact Match points), followed by flow matching (+5.5 EM). However, the standard span classifier (85.8% EM) outperforms our best diffusion model (83.0% EM), suggesting that the denoiser does not yet justify its added complexity. We discuss tradeoffs between the interpretability of diffusion trajectories and classification accuracy.

Keywords: extractive question answering, diffusion models, flow matching, financial causality, FinCausal

1. Introduction

Detecting causal relationships in financial documents is important for understanding market dynamics and supporting analytical workflows. The FinCausal shared task series (Moreno-Sandoval et al., 2023, 2025) has helped to push progress on this problem since 2020, evolving from span-level BIO tagging to extractive question answering formulations. The 2026 edition further expands its bilingual (English and Spanish) dataset of 4,000 training samples and replaces Exact Match and Semantic Answer Similarity evaluation with an LLM-as-a-judge metric that scores system outputs on a 1-5 adequacy scale (Moreno-Sandoval et al., 2026).

Dominant approaches at FinCausal 2025 relied on fine-tuned LLMs such as Llama 3.1 (Niess et al., 2025) or encoder-based token classification (Devlin et al., 2019). While LLMs achieve strong adequacy scores, they risk hallucinating tokens absent from the source text, a critical failure mode in financial applications that has to be carefully balanced. Extractive models avoid hallucination by construction but lack the capacity to model positional uncertainty over answer boundaries.

We propose **SpanDiffusion** (Figure 1), which frames extractive Q&A as a continuous denoising problem. Instead of classifying each token independently, we diffuse dual Gaussian masks, soft peaks centered at the answer start and end positions, and learn to recover them from noise via a dedicated transformer conditioned on the encoder output. Our contributions are:

1. A novel formulation of extractive Q&A as continuous diffusion over dual Gaussian span masks, with joint start and end prediction through a shared denoising process.
2. Replacing Denoising Diffusion Probabilistic Models (DDPM) with flow matching (rectified flow), achieving simpler training and $2.5\times$ fewer inference steps (20 instead of 50).
3. A systematic ablation across six diffusion variants and a standard span-classification baseline, separating the contributions of the diffusion formulation, encoder adaptation, and training duration.

2. Method

2.1. Task Formulation

Given a context passage $C = (c_1, \dots, c_N)$ and a causal question Q , the task is to extract a contiguous span (s, e) such that the answer $A = (c_s, \dots, c_e)$ addresses the causal relationship in Q . The training data comprises 2,000 English and 2,000 Spanish samples (Moreno-Sandoval et al., 2026). The leaderboard test sets contain 500 and 503 samples, respectively.

2.2. Encoder

We encode the concatenated input $[Q; [\text{SEP}]; C]$ with DeBERTa-v3-large (He et al., 2023), a 434M-parameter Transformer pre-trained with replaced token detection. The encoder weights are frozen,

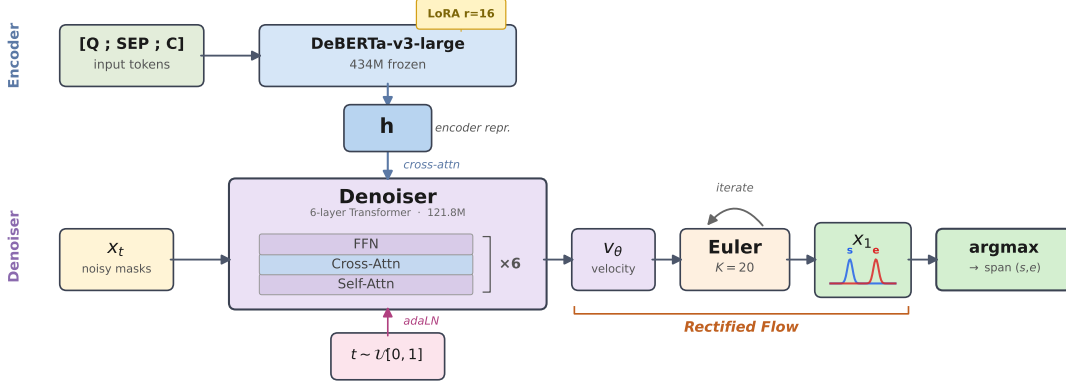


Figure 1: SpanDiffusion architecture. The input is encoded by a frozen DeBERTa-v3-large with LoRA adapters. The denoiser receives noisy dual-peak masks x_t and timestep t , attends to the encoder representations via cross-attention, and predicts the velocity field. At inference, 20 Euler integration steps map noise to clean dual peaks, from which the answer span is extracted via argmax.

and we inject LoRA adapters (Hu et al., 2022) into the `query_proj` and `value_proj` projections of all 24 attention layers. With rank $r=16$, scaling $\alpha=32$, and dropout 0.05, this adds 1.6M trainable parameters (0.3% of the encoder) while enabling domain adaptation to financial text.

2.3. Dual Gaussian Span Masks

Rather than predicting start/end logits independently, we construct a continuous 2-channel target over the L context tokens. For a ground-truth span (s, e) , the target at position i is:

$$x_1^{(c)}(i) = 2 \exp\left(-\frac{(i - p_c)^2}{2\sigma^2}\right) - 1, \quad c \in \{\text{start}, \text{end}\} \quad (1)$$

where $p_{\text{start}} = s$, $p_{\text{end}} = e$, and $\sigma=1.5$ (selected via grid search over $\{0.5, 1.0, 1.5, 2.0\}$). This maps each channel to $[-1, 1]$ with Gaussian peaks centered at the answer boundaries. The soft representation provides a smooth loss landscape for the denoiser, coupling neighboring positions rather than treating each token independently.

2.4. Flow Matching

We adopt rectified flow (Liu et al., 2023; Lipman et al., 2023) instead of DDPM (Ho et al., 2020). Given a source sample $x_0 \sim \mathcal{N}(0, I)$ and target x_1 (the dual-peak mask from Eq. 1), we define a straight interpolation path:

$$x_t = t \cdot x_1 + (1 - t) \cdot x_0, \quad t \in [0, 1] \quad (2)$$

The velocity along this path is constant: $v = x_1 - x_0$. A neural network v_θ is trained to predict this velocity (Eq. 3):

$$\mathcal{L} = \mathbb{E}_{t, x_0} \|v_\theta(x_t, t, h) - (x_1 - x_0)\|_{\mathcal{M}}^2 \quad (3)$$

where h denotes the encoder representations and $\|\cdot\|_{\mathcal{M}}^2$ denotes the masked MSE, $\frac{1}{|\mathcal{M}|} \sum_{i \in \mathcal{M}} (\cdot)_i^2$, where $\mathcal{M}$ is the set of context-token positions (excluding question and special tokens).

At inference, we integrate from $x_0 \sim \mathcal{N}(0, I)$ using Euler steps with $\Delta t = 1/K$ (Eq. 4):

$$x_{t+\Delta t} = x_t + v_\theta(x_t, t, h) \cdot \Delta t \quad (4)$$

We use $K=20$ steps, compared to 50 DDIM steps needed by our DDPM baseline. Final answer positions are extracted as $s = \arg \max x_1^{(\text{start})}$, $e = \arg \max x_1^{(\text{end})}$.

2.5. Denoiser

The velocity predictor (denoiser) is a 6-layer Transformer with hidden dimension $d=1024$ and 16 attention heads. Each layer performs self-attention over the noisy mask representation, followed by cross-attention to the encoder output h and a feed-forward network. Time conditioning follows the adaptive layer normalization (adaLN) scheme from DiT (Peebles and Xie, 2023):

$$\hat{h} = \text{LN}(h) \cdot (1 + s_t) + b_t \quad (5)$$

where (s_t, b_t) are produced by a learnable MLP from a sinusoidal time embedding. The 2-channel noisy mask is projected to d via a 2-layer MLP before entering the Transformer. The denoiser comprises 121.8M trainable parameters.

Variant	Method	LoRA	Start	End	EM
Baseline	Linear	$r=16$	91.5	93.2	85.8
V1 (soft-box)	DDPM	—	—	—	32.5
V2 (dual-peak)	DDPM	—	—	—	42.5
V2-LoRA	DDPM	$r=16$	83.2	89.2	76.5
V3-Flow	Flow	$r=16$	86.2	93.0	82.0
V3-Flow-Long	Flow	$r=16$	87.2	93.5	83.0
V3-LoRA32	Flow	$r=32$	89.8	91.2	82.8

Table 1: Ablation results on the validation set (% accuracy). EM = Exact Match (both start and end correct). Baseline = DeBERTa + LoRA + linear span head (no diffusion). All LoRA variants use $\alpha=32$, dropout 0.05.

3. Experiments

3.1. Experimental Setup

We combine the English and Spanish training splits into a single bilingual set of 4,000 samples and create a stratified 90/10 train/validation split. All models are trained jointly on both languages. We use AdamW with learning rate 3×10^{-4} for the denoiser (or span head) and 3×10^{-5} for LoRA parameters, weight decay 0.01, gradient clipping at 1.0, OneCycleLR with cosine annealing (warmup 0.1), and batch size 8. Training uses fp32 (DeBERTa-v3 overflows under mixed precision). Most variants train for 30 epochs, V3-Flow-Long extends this to 50 with early stopping (patience 10). Each run takes ~ 2 hours on one NVIDIA L40 GPU. Validation performance is measured by Exact Match (EM): the percentage of samples where both predicted start and end positions exactly match the ground truth.

All diffusion variants share the DeBERTa-v3-large encoder and differ in the diffusion formulation, LoRA usage, and training schedule, as detailed in Table 1. We additionally include a standard span-classification baseline (DeBERTa + LoRA + linear head, no diffusion) for reference.

3.2. Ablation Study

Table 1 presents the ablation study. The top row shows a standard span-classification baseline (DeBERTa + LoRA + linear head, no diffusion), which achieves 85.8% EM, outperforming all diffusion variants. We analyze the individual factors below.

Target shape (soft-box vs. dual-peak). As an initial baseline (V1), we tested a ‘soft-box’ target where all tokens within the span are assigned a value of 1 and all others -1. Switching to the dual-peak Gaussian formulation (V2) improved EM from 32.5% to 42.5%, allowing the model to better capture the contiguous nature of the extractive spans.

Submission	EN	ES
SpanDiff. V1 (soft-box)	3.30	—
SpanDiff. V2 (dual-peak)	3.41	—
SpanDiff. V2-LoRA	4.33	4.41
SpanDiff. V3-Flow	4.57	4.63
Span Hybrid V2 [†]	4.08	—
Best competitor	4.81	4.81

Table 2: FinCausal 2026 leaderboard scores (LLM-as-a-judge, 1 to 5 scale). [†]Span Hybrid V2 = RoBERTa-base + flow matching + classifier-free guidance, an earlier architecture abandoned in favor of SpanDiffusion.

LoRA is most influential. Adding LoRA adapters to the frozen encoder yields the largest single improvement in our study. V2 to V2-LoRA increases EM from 42.5% to 76.5%, an increase of +34.0 points. Without adaptation, the frozen DeBERTa representations are poorly aligned with the continuous target space of the denoiser. LoRA with only 1.6M additional parameters (0.3% of the encoder) bridges this gap effectively.

Flow matching outperforms DDPM. Replacing DDPM with rectified flow (V2-LoRA $\rightarrow$ V3-Flow) improves EM by +5.5 points (76.5% $\rightarrow$ 82.0%) while reducing inference from 50 DDIM steps to 20 Euler steps (2.5 $\times$ speedup). The straight interpolation paths of flow matching provide a simpler learning objective, and the constant-velocity targets reduce gradient variance.

Baseline outperforms diffusion. The span classifier (85.8% EM) surpasses our best diffusion variant (V3-Flow-Long, 83.0%) by 2.8 points while requiring only 2.7M total trainable parameters vs. 123.4M for the diffusion model, due to replacing the 121.8M-parameter denoiser with a 1.1M linear span head. Inference is also simplified to a single forward pass. Looking at the training dynamics reveals an interesting contrast: the baseline’s validation cross-entropy loss diverges after epoch 4 (0.38 $\rightarrow$ 1.68 by epoch 21) while EM continues improving (83.5% $\rightarrow$ 85.8%), indicating that the model overfits in probability space but the argmax decision boundary remains correct. In contrast, V3-Flow-Long’s MSE validation loss decreases steadily (0.35 $\rightarrow$ 0.03) but EM saturates at 83.0%, suggesting the continuous regression objective is harder to optimize for discrete position accuracy.

3.3. Competition Results

Table 2 shows the FinCausal 2026 leaderboard scores for our submissions. The progression V1 $\rightarrow$ V2 $\rightarrow$ V2-LoRA $\rightarrow$ V3-Flow mirrors the validation

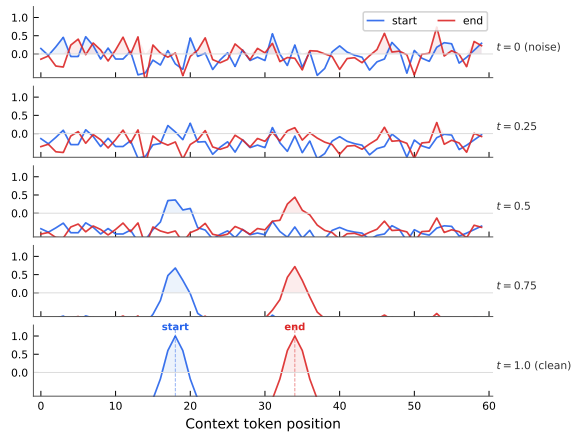


Figure 2: Euler integration from noise ($t=0$) to clean dual peaks ($t=1$) in 20 steps. Blue = start channel, red = end channel. The answer span boundaries sharpen progressively over the trajectory.

ablation. Our best submission (V3-Flow) scores 4.57 (EN) / 4.63 (ES), while the top system achieves 4.81 on both sub-tasks.

The gap to the top system (~ 0.24 on English) may partly reflect evaluation dynamics (Section 5), though the top system may simply produce more accurate answers. The baseline used in Table 1 was not officially submitted, so a direct comparison with SpanDiffusion is not possible under the challenge evaluation measure.

3.4. Diffusion Process Visualization

Figure 2 illustrates flow matching inference on a validation example. At $t=0$ the signal is pure noise; as Euler integration progresses, the start peak (blue) and end peak (red) gradually separate and sharpen, converging to the correct boundaries by $t=1.0$. The intermediate states are interpretable, demonstrating how the model progressively resolves positional uncertainty, a key advantage over single-pass classification.

4. Related Work

Diffusion models for NLP. Diffusion models have been applied to text generation via continuous embeddings (Li et al., 2022; Gong et al., 2023) and discrete masking (Sahoo et al., 2024). Han et al. (2023) propose semi-autoregressive diffusion for controllable generation. While Shen et al. (2023) recently introduced boundary diffusion for Named Entity Recognition, to our knowledge, SpanDiffusion is the first to formulate extractive Q&A as a continuous diffusion process over span boundaries.

Flow matching. Flow matching (Lipman et al., 2023) and rectified flow (Liu et al., 2023) replace the

SDE formulation of DDPM with straight ODE paths, allowing faster sampling. Peebles and Xie (2023) demonstrated their effectiveness with Transformers (DiT). We adapt DiT-style adaLN to 1D positional masks.

Parameter-efficient fine-tuning. LoRA (Hu et al., 2022) injects low-rank updates into frozen weights. Our ablation confirms its critical role: without LoRA, EM drops from 76.5% to 42.5%, the largest single factor.

5. Discussion

LLM-as-a-judge metric. FinCausal 2026 replaced Exact Match with an LLM-as-a-judge metric (Zheng et al., 2023). We hypothesize that exact span extractions may receive lower fluency ratings than paraphrases conveying the same information, though a controlled study is needed to confirm this.

Diffusion vs. classification. The baseline’s advantage (85.8% vs. 83.0%) raises the question of when diffusion-based span prediction is a good choice. SpanDiffusion offers two potential benefits not captured by EM: (1) interpretable intermediate states (Figure 2) showing how the model progressively resolves span boundaries, and (2) the stochastic inference process could in principle yield uncertainty estimates over predictions, though we do not evaluate calibration in this work. However, on the 4,000-sample FinCausal dataset, these do not offset the harder optimization of the diffusion objective.

Limitations. The fixed Gaussian width ($\sigma=1.5$) assumes unimodal boundaries, which may not hold for multi-span answers, since errors concentrate on multi-clause causal chains and very short (1 to 2 token) answers where peaks overlap. Training requires fp32 (DeBERTa-v3 overflows under mixed precision). The 20-step Euler inference is both slower than single-pass classification and stochastic (different random x_0 seeds yield different predictions), introducing inference variance that a deterministic baseline could avoid. We did not quantify this variance and leave it to future work.

6. Conclusion

We presented SpanDiffusion, a continuous diffusion approach to extractive causal question answering that operates over dual Gaussian span masks rather than discrete token labels. Our systematic ablation reveals that LoRA encoder adaptation is the single most important factor (+34 EM points),

followed by the switch from DDPM to flow matching (+5.5 EM with $2.5\times$ fewer inference steps). A standard span classifier with the same encoder outperforms our best diffusion model (85.8% vs. 83.0% EM), indicating that diffusion-based span prediction does not currently justify its added complexity on this task. Nonetheless, our best diffusion model scores competitively on the FinCausal 2026 leaderboard (4.57/4.63 EN/ES). We believe the formulation remains promising: the interpretable denoising trajectory and built-in uncertainty estimation offer qualitative advantages that Exact Match does not capture. Future work includes classifier-free guidance (Ho and Salimans, 2022) for question-conditioned refinement, consistency distillation for single-step inference, and scaling to larger training sets where the capacity of the 121.8M-parameter denoiser may be more effectively utilized.

7. Bibliographical References

- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. BERT: Pre-training of deep bidirectional transformers for language understanding. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (NAACL-HLT)*, pages 4171–4186.
- Shansan Gong, Mukai Li, Jiangtao Feng, Zhiyong Wu, and LingPeng Kong. 2023. DiffuSeq: Sequence to sequence text generation with diffusion models. In *Proceedings of the 11th International Conference on Learning Representations (ICLR)*.
- Xiaochuang Han, Sachin Kumar, and Yulia Tsvetkov. 2023. SSD-LM: Semi-autoregressive simplex-based diffusion language model for text generation and modular control. In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (ACL)*, pages 11575–11596.
- Pengcheng He, Jianfeng Gao, and Weizhu Chen. 2023. DeBERTaV3: Improving DeBERTa using ELECTRA-style pre-training with gradient-disentangled embedding sharing. In *Proceedings of the 11th International Conference on Learning Representations (ICLR)*.
- Jonathan Ho, Ajay Jain, and Pieter Abbeel. 2020. Denoising diffusion probabilistic models. In *Advances in Neural Information Processing Systems (NeurIPS)*, volume 33, pages 6840–6851.
- Jonathan Ho and Tim Salimans. 2022. Classifier-free diffusion guidance. *arXiv preprint arXiv:2207.12598*.
- Edward J. Hu, Yelong Shen, Phillip Wallis, Zeyuan Allen-Zhu, Yanzhi Li, Shanen Wang, Lu Wang, and Weizhu Chen. 2022. LoRA: Low-rank adaptation of large language models. In *Proceedings of the 10th International Conference on Learning Representations (ICLR)*.
- Xiang Lisa Li, John Thickstun, Ishaan Gulrajani, Percy Liang, and Tatsunori B. Hashimoto. 2022. Diffusion-LM improves controllable text generation. In *Advances in Neural Information Processing Systems (NeurIPS)*, volume 35.
- Yaron Lipman, Ricky TQ Chen, Heli Ben-Hamu, Maximilian Nickel, and Matthew Le. 2023. Flow matching for generative modeling. In *Proceedings of the 11th International Conference on Learning Representations (ICLR)*.
- Xingchao Liu, Chengyue Gong, and Qiang Liu. 2023. Flow straight and fast: Learning to generate and transfer data with rectified flow. In *Proceedings of the 11th International Conference on Learning Representations (ICLR)*.
- Antonio Moreno-Sandoval, Jordi Porta, Blanca Carbajo-Coronado, Yanco Torterolo, and Doaa Samy. 2025. The financial document causality detection shared task (FinCausal 2025). In *Proceedings of the Joint Workshop of the 9th Financial Technology and Natural Language Processing (FinNLP), the 6th Financial Narrative Processing (FNP), and the 1st Workshop on Large Language Models for Finance and Legal (LLMFin-Legal)*, pages 214–221, Abu Dhabi, UAE. Association for Computational Linguistics.
- Antonio Moreno-Sandoval, Jordi Porta, Yanco Torterolo, Alexia Stanesco, Melina Chatzi, and Sofía Roseti. 2026. The Financial Document Causality Detection Shared Task (FinCausal 2026). In *Proceedings of the 7th Financial Narrative Processing Workshop (FNP 2026) at LREC 2026*, Palma de Mallorca, Spain. ELRA. To appear.
- Antonio Moreno-Sandoval, Jordi Porta-Zamorano, Blanca Carbajo-Coronado, Doaa Samy, Dominique Mariko, and Mahmoud El-Haj. 2023. The financial document causality detection shared task (FinCausal 2023). In *Proceedings of the 5th Financial Narrative Processing Workshop (FNP 2023) at the 2023 IEEE International Conference on Big Data (IEEE BigData 2023)*, Sorrento, Italy.
- Georg Niess, Houssam Razouk, Stasa Mandic, and Roman Kern. 2025. Addressing hallucination

in causal Q&A: The efficacy of fine-tuning over prompting in LLMs. In *Proceedings of the Joint Workshop of the 9th Financial Technology and Natural Language Processing (FinNLP), the 6th Financial Narrative Processing (FNP), and the 1st Workshop on Large Language Models for Finance and Legal (LLMFinLegal)*.

William Peebles and Saining Xie. 2023. Scalable diffusion models with transformers. In *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*, pages 4195–4205.

Subham Sekhar Sahoo, Marianne Arriola, Yair Schiff, Aaron Gokaslan, Edgar Marroquin, Justin T. Chiu, Alexander Rush, and Aditya Grover. 2024. Simple and effective masked diffusion language models. *Advances in Neural Information Processing Systems (NeurIPS)*, 37.

Yongliang Shen, Kaitao Song, Xu Tan, Dongsheng Li, Weiming Lu, and Yueting Zhuang. 2023. [DiffusionNER: Boundary diffusion for named entity recognition](#). In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 3875–3890, Toronto, Canada. Association for Computational Linguistics.

Lianmin Zheng, Wei-Lin Chiang, Ying Sheng, Siyuan Zhuang, Zhanghao Wu, Yonghao Zhuang, Zi Lin, Zhuohan Li, Dacheng Li, Eric Xing, et al. 2023. Judging LLM-as-a-judge with MT-bench and chatbot arena. In *Advances in Neural Information Processing Systems (NeurIPS)*, volume 36.

8. Language Resource References

Antonio Moreno-Sandoval, Yanco Amor Torterolo Orta, Maria Alexia Stanescu, and Melina Chatzi. 2026. [The Financial Document Causality Detection Shared Task \(FinCausal 2026\): Dataset](#).

Improving Verbatim Financial Causality Extraction with Supervised Fine-Tuning and Prompt Repetition

Sanae Attak, Mohammed Salah Chiadmi, Youssef Lamrani Alaoui

IFE-lab, LERMA, Mohammadia School of Engineers (EMI)

Mohammed V University in Rabat, Morocco

sanae.attak@research.emi.ac.ma, schiadmi@emi.ac.ma, lamrani@emi.ac.ma

Abstract

This paper investigates the application of generative Large Language Models (LLMs) for strict verbatim span extraction. We evaluate our methodology within the FinCausal 2026 shared task. Because generative LLMs optimize next-token probability rather than strict boundaries, they naturally suffer from over-generation and boundary drift in extraction tasks. To address this, we introduce a generalized structural training constraint, extending prompt repetition from a purely inference-time heuristic to a training-time supervision framework. By incorporating duplicated prompts directly into Supervised Fine-Tuning (SFT), we hypothesize that this encourages the model to internalize a form of unidirectional cross-reading behavior, leading to stronger alignment between generated spans and the source context for exact extraction. Evaluating on open-weights (Qwen2.5-14B-Instruct-1M) and proprietary (GPT-4.1-Nano) architectures, we find this soft attention constraint improves Exact Match scores for open models and helps balance cross-lingual performance disparities. Conversely, the proprietary model exhibited sensitivity to prompt duplication, achieving its highest score without repetition. Ultimately, our deterministic SFT approach secured 4th place in the Spanish subtask (4.73) and 6th place in the English subtask (4.70), indicating the viability of structurally simple, natively fine-tuned models compared to complex multi-stage pipelines.

Keywords: Causality Extraction, Generative LLMs, Prompt Repetition, Supervised Fine-Tuning (SFT)

1. Introduction

The increasing complexity of financial documents necessitates advanced methodologies to extract and analyze causality within such texts. The FinCausal 2026 shared task introduces a generative Question-Answering (QA) framework for detecting causal relationships in financial disclosures (Moreno-Sandoval et al., 2026). The task requires models to process abstractive questions regarding causes or effects and answer by extracting verbatim spans directly from the source text.

This generative formulation presents a structural conflict: generative models are inherently designed to synthesize and paraphrase information (Chrysos-tomou et al., 2024), yet the task evaluation strictly penalizes any generative deviation or hallucination from the original text. Because generative LLMs optimize next-token probability rather than strict boundaries, they naturally suffer from over-generation and boundary drift in extraction tasks. To address this, our study investigates the effectiveness of Supervised Fine-Tuning (SFT) coupled with targeted prompting techniques specifically Prompt Repetition to constrain generative models into performing exact span extraction.

In this study, we investigate the following Research Questions (RQs):

- **RQ1:** Can structurally simple SFT constrain generative LLMs to act as verbatim extractors without relying on multi-stage pipelines?
- **RQ2:** Does extending prompt repetition from

an inference trick to a training paradigm improve extraction fidelity?

- **RQ3:** How do open-weight models compare to proprietary models under these extraction constraints across different languages?

Contributions. While prompt repetition has primarily been studied as an inference-time technique, we explore its integration directly into the supervised fine-tuning process for causal extraction tasks. Specifically, this paper makes three core contributions:

1. We propose a generalized structural training constraint for generative extraction. By extending prompt repetition from an inference-time heuristic to a training-time supervision methodology, we hypothesize that generative models internalize context anchoring for strict span extraction.
2. We provide a cross-lingual evaluation (English and Spanish) indicating that this method reduces language-specific extraction biases.
3. Our experiments indicate that mid-scale open-weight models (Qwen2.5-14B-Instruct-1M) can achieve competitive performance relative to proprietary models under strict extraction constraints.

2. Related Work

Causal relationship extraction remains a persistent challenge in financial NLP. Earlier FinCausal shared tasks (Mariko et al., 2022) predominantly framed the problem as a sequence-tagging task, utilizing BIO tagging schemes via token classifiers like BERT or BioBERT to identify causal spans (Saha et al., 2022; Lyu et al., 2022). While token classifiers and pointer networks perform well on small datasets, they often struggle when causal chains span multiple disconnected sentences. Consequently, the field recently shifted toward generative Q&A frameworks (Moreno-Sandoval et al., 2026). To combat the hallucinations inherent in generative architectures, prior approaches utilized complex lexically constrained decoding (Ghosh and Naskar, 2022), explicit pointer-generator copy mechanisms, or passed extractive outputs through LLMs for refinement (Trivedi et al., 2025). Our objective is to approach the strict verbatim fidelity of these classical copy mechanisms natively through generative SFT, without relying on modified decoding algorithms.

2.1. Attention Constraints and Prompt Repetition

Prior work has shown that repeating prompts during inference improves extraction accuracy by reinforcing attention alignment in causal language models (Leviathan et al., 2025). However, this technique has primarily been explored as an inference-time heuristic. In this study, we extend this idea by integrating prompt repetition directly into the supervised fine-tuning stage, enabling the model to internalize cross-reading behavior during training.

Furthermore, evaluations from the FinCausal 2025 shared task demonstrated that standard prompt optimization and few-shot learning are fundamentally insufficient to prevent generative models from hallucinating during causal extraction (Niess et al., 2025). As noted by (Niess et al., 2025), fine-tuning generative architectures is absolutely essential for minimizing boundary drift and enforcing strict extraction constraints. Our work builds directly upon this premise: we not only adopt Supervised Fine-Tuning (SFT) as a baseline necessity, but we further constrain the generative process by introducing prompt repetition as a structural soft-attention mechanism during the SFT phase itself.

3. Methodology

3.1. Dataset and Preprocessing

This study utilizes the official FinCausal 2026 dataset (Moreno-Sandoval et al., 2026). The dataset comprises financial reports sourced from

the UK and Spain, providing 2,000 annotated training samples per language (4,000 samples in total).

To validate our models and conduct internal ablation studies, we employed a two-stage data utilization strategy. First, we merged and randomly split the dataset into an internal Training Set (3,600 samples) and a held-out Development Set (400 samples: 200 English and 200 Spanish). This internal split was exclusively utilized to independently compute Exact Match (EM) and Semantic Answer Similarity (SAS) metrics for our comparative analyses. Second, for the final official blind test submissions, the models were retrained on the entirety of the provided dataset (all 4,000 samples) to maximize domain exposure. Prior to tokenization, all texts underwent a standard normalization pipeline (lowercasing and whitespace removal).

3.2. Internalizing Attention via Prompt Repetition

Unlike prior work which applies prompt repetition only at inference (Leviathan et al., 2025), we incorporate the repeated prompt structure directly during SFT. Generative LLMs, built upon decoder-only transformer architectures (Brown et al., 2020), utilize a lower-triangular causal attention mask to preserve the auto-regressive property; meaning a token t_i can only attend to previous tokens $t_{\leq i}$ (Vaswani et al., 2017). In a standard prompt (*Context + Question*), the question tokens can attend to the context, but the context tokens cannot attend to the question. By duplicating the input (*Context₁ + Question₁ + Context₂ + Question₂*), we fundamentally alter the attention graph: the tokens in *Context₂* are now positioned after *Question₁*, allowing them to compute dense attention over both the source text and the task specification simultaneously. This restructuring of the input sequence may mitigate the limitations imposed by causal masking by allowing later context tokens to attend to both the source text and task instruction. We hypothesize that this structural duplication encourages stronger anchoring of the generated span to the original context.

3.3. Experimental Setup

Exact prompt templates utilized for the experiments are provided in Section A.

Open-Weights Configuration We utilized the Qwen2.5-7B-Instruct-1M and Qwen2.5-14B-Instruct-1M checkpoints. Models were loaded using 4-bit quantization via `unsloth`. We applied Low-Rank Adaptation (QLoRA) targeting all linear modules (`q_proj`, `k_proj`, `v_proj`, `o_proj`, `gate_proj`, `up_proj`, `down_proj`) with a rank of $r = 64$ and $\alpha = 16$.

The models were trained for exactly 1 epoch using the 8-bit AdamW optimizer with a linear learning rate schedule peaking at 2×10^{-4} , a warmup ratio of 0.03, and an effective batch size of 8. Maximum sequence length was set to 2048, and sequence packing was explicitly disabled (`packing = False`) to maintain the structural integrity of the duplicated contexts.

Proprietary Configuration For proprietary comparisons, we applied SFT to the `gpt-4.1-nano` checkpoint via the official API for 3 epochs.

Inference Parameters To isolate the impact of our training methodology, decoding parameters were set to pure greedy decoding. For all architectures, generation was deterministic (`temperature = 0.0` and `top_p = 1.0`). Inference for the open-weights models was executed on a single NVIDIA L40S 48GB GPU. While Prompt Repetition inherently increases the input sequence length, we observed only a moderate computational overhead during the prefill phase.

3.4. Evaluation Metrics and Statistical Testing

To comprehensively evaluate performance, we utilize three distinct metrics. **Exact Match (EM)** measures whether the predicted answer exactly matches the gold reference span. Both the generated predictions and the withheld gold references undergo the exact same normalization pipeline (lowercasing, diacritic stripping, and whitespace trimming) prior to EM calculation. **Semantic Answer Similarity (SAS)** evaluates the semantic equivalence of the answers (Risch et al., 2021). To align precisely with the official evaluation framework established by the FinCausal 2025 organizers (Moreno-Sandoval et al., 2026), we extract 768-dimensional text embeddings using the cross-lingual `paraphrase-multilingual-mpnet-base-v2` Sentence Transformer model and compute their pairwise cosine similarity. Finally, the **LLM-as-a-Judge** assesses answer adequacy on a scale of 1 to 5. We conducted paired bootstrap resampling to verify that improvements in Exact Match are statistically significant ($p < 0.05$).

4. Results

4.1. Validation Set Performance

Because our internal development split differs from the official 2025 test set, directly comparing these scores against historical extractive baselines would be methodologically unsound. Instead, we use this held-out set strictly as an internal ablation to

quantify the absolute impact of Prompt Repetition on textual fidelity.

As shown in Table 1, incorporating Prompt Repetition during SFT yields a consistent improvement in Exact Match for the open-weight models. Internalizing cross-reading behavior increases the `Qwen2.5-14B-Instruct-1M` model’s EM score from 0.8200 to 0.8550 in English. This confirms that generative models, when constrained via SFT and prompt duplication, effectively reduce boundary errors compared to standard zero-shot prompting. Across both model scales and languages, prompt repetition consistently improves Exact Match performance for open-weight architectures, suggesting that simple structural constraints during SFT can improve verbatim span fidelity.

4.2. Official Blind Test Results

For the final evaluation on the official blind test set (where gold references are withheld), our models were retrained on the full 4,000-sample dataset. Because independent calculation of EM and SAS is impossible for these final submissions, Table 2 reports the official standardized LLM-as-a-judge scores provided by the organizers.

5. Discussion and Analysis

1. Language Balance via Prompt Repetition:

An observation from our internal evaluation is the impact of training-time Prompt Repetition on cross-lingual disparities. In the baseline "Simple" setting, the `Qwen2.5-14B-Instruct-1M` model exhibits a slight bias toward Spanish (EM of 0.8350 in ES vs. 0.8200 in EN). Applying the Repeated technique appears to balance this performance (0.8550 EN and 0.8450 ES). The consistent improvements observed for the open-weight models suggest that structural prompt duplication may help stabilize span boundaries during generation, particularly in tasks requiring strict verbatim extraction. This observation reinforces the hypothesis that simple structural constraints applied during supervised fine-tuning can improve the reliability of generative models in high-precision extraction tasks.

2. The Conflict Between RLHF and Structural Constraints:

While open-weights models showed benefits from Prompt Repetition, the proprietary `GPT-4.1-Nano` exhibited divergent behavior. On the internal dev set (Table 1), repetition maintained raw Exact Match extraction boundaries. Yet, under the official blind test evaluated by the LLM-as-a-judge metric (Table 2), applying Prompt Repetition resulted in a noticeable degradation (from ~ 4.73 down to ~ 3.98). We posit this reveals a

Generative Configuration	English (EN)		Spanish (ES)	
	SAS	EM	SAS	EM
Qwen2.5-7B-Instruct-1M (Simple Prompt)	0.9667	0.8000	0.9699	0.7600
Qwen2.5-7B-Instruct-1M (Repeated Prompt)	0.9729	0.8300	0.9723	0.7950
Qwen2.5-14B-Instruct-1M (Simple Prompt)	0.9626	0.8200	0.9749	0.8350
Qwen2.5-14B-Instruct-1M (Repeated Prompt)	0.9730	0.8550*	0.9746	0.8450
GPT-4.1-Nano (Simple Prompt)	0.9578	0.8050	0.9759	0.8450
GPT-4.1-Nano (Repeated Prompt)	0.9683	0.8000	0.9787	0.8600*

Table 1: Semantic Answer Similarity (SAS) and Exact Match (EM) Results evaluated strictly on our 400-sample Internal Development Set. Because this data split differs from historical blind test sets, these scores serve specifically as an internal ablation to isolate the impact of Prompt Repetition. (*) denotes a statistically significant improvement over the Simple baseline for the respective model architecture ($p < 0.05$).

Model	Configuration	EN	ES
Open-Weights SFT (Single Model)			
Qwen2.5-7B-Instruct-1M	Simple Prompt	4.3120	4.0915
Qwen2.5-7B-Instruct-1M	Repeated Prompt	4.3500	4.3877
Qwen2.5-14B-Instruct-1M	Repeated Prompt	4.6720	4.6740
Proprietary SFT (Single Model)			
GPT-4.1-Nano	Zero-Shot	3.9540	3.9841
GPT-4.1-Nano	Repeated Prompt	3.9800	3.9940
GPT-4.1-Nano*	Simple Prompt	4.7040	4.7396
Ablation: Inference-Only Repetition			
GPT-4.1-Nano	Train Simple + Infer Rep.	4.6880	4.6143
Ablation: Complex Pipelines			
Qwen2.5-14B-Instruct-1M	Repeated + Ensemble	4.5800	4.5785
GPT-4.1-Nano	Simple + Ensemble	4.6660	4.6501
GPT-4.1-Nano	Simple + RAG (3-Shot)	4.6600	4.7117
GPT-4.1-Nano	Simple + GPT-4o Judge	4.2560	4.2445

Table 2: Official LLM-as-a-Judge performance on the FinCausal 2026 Blind Test Set (Scored 1 to 5). The ablations indicate that inference-only repetition and multi-stage pipelines (Ensembles, RAG, Correctors) generally resulted in lower scores compared to single-stage, natively fine-tuned models. (*) indicates the official submission.

potential conflict between structural training constraints and Reinforcement Learning from Human Feedback (RLHF). Heavily aligned models optimized for conversational naturalness may penalize or misinterpret highly unnatural, duplicated input structures during generative decoding. This indicates that structural prompting techniques effective on base-aligned open models may conflict with the safety alignment layers of proprietary models.

3. Training-Time vs. Inference-Time Repetition:

To examine whether the performance gains stem from the training paradigm or merely from inference-time context duplication, we conducted a targeted ablation on the blind test set. When the proprietary model was fine-tuned on the *Simple* configuration but evaluated using the *Repeated* prompt during inference, its score decreased (from 4.7396 down to 4.6143 in Spanish). This supports the idea that forcing prompt repetition solely at inference intro-

duces out-of-distribution formatting noise, and that the model benefits from internalizing the attention mechanism during the SFT phase.

4. The Limitations of Multi-Stage Pipelines:

Recent NLP extraction tasks frequently deploy multi-stage pipelines. However, our ablations (Table 2) suggest these architectures can be less effective for strict verbatim extraction. Injecting dynamic context via Retrieval-Augmented Generation (RAG, utilizing 3-shot semantic retrieval for in-context examples) introduced external noise, slightly lowering the score. Similarly, utilizing a powerful meta-model (GPT-4o) as a post-generation corrector via strict formatting prompts caused a decrease in performance (from 4.7040 to 4.2560). Qualitative analysis indicated that the corrector model prioritized grammatical completeness over verbatim copying, disrupting the required span boundaries. Furthermore, Ensemble Voting (majority consensus across

5 high-temperature generations) introduced token-level variations that diluted the Exact Match consistency.

6. Error Analysis and Qualitative Comparison

To understand the mechanism by which Prompt Repetition improves Exact Match (EM) scores, we conducted a qualitative analysis of the residual errors in the baseline models. As demonstrated in Table 4 (Appendix C), the generative formulation introduces specific hallucination patterns.

For instance, consider a boundary error hallucination where the baseline *Simple Prompt* misses the exact starting boundary by appending an introductory connector (e.g., extracting "**As** external threats become more sophisticated" instead of "external threats become more sophisticated"). Furthermore, when faced with causal chains, the baseline model occasionally truncates the extraction, or exhibits generative stutters (Example 1) and mid-generation stops (Example 4).

As shown in Table 4 (Appendix C), internalizing the *Repeated Prompt* during Supervised Fine-Tuning acts as an attention constraint, encouraging the model to respect verbatim spans and reducing these generative tendencies across the tested open-weights architectures. The high verbatim copy-rate achieved by the Repeated configuration supports our hypothesis that the SFT process enforces a unidirectional "cross-reading" mechanism.

7. Conclusion

This study investigated the effectiveness of a structural training constraint combining Supervised Fine-Tuning with prompt repetition for strict span extraction. We demonstrated that extending prompt repetition from an inference-time heuristic to a training-time supervision approach acts as a context anchoring mechanism, allowing generative causal LLMs to internalize cross-reading behaviors. The results suggest that structurally simple, natively fine-tuned generative models can perform reliable verbatim extraction, neutralizing language-specific biases to achieve balanced multilingual performance without relying on complex multi-stage pipelines. Although evaluated in the financial domain, the proposed structural constraint may extend to other high-precision extraction tasks where strict verbatim copying is required.

8. Limitations and Future Work

Our approach demonstrates practical empirical performance but also highlights several avenues for

future research. The models were fine-tuned on a relatively small, domain-specific dataset (3,600 bilingual samples), which was effective in this context; however, extending structural training constraints to larger, multi-domain corpora would help assess broader generalization. While our study focused on the financial sector, adapting prompt repetition for verbatim extraction in other domains, such as biomedical relation extraction or legal evidence analysis, remains an open challenge.

The LLM-as-a-judge metric provides a useful assessment of answer adequacy but may introduce implicit alignment biases. Future work could explore hybrid evaluation frameworks to better relate these subjective judgments to deterministic metrics like Exact Match.

Computationally, prompt repetition imposes only minor overhead. As shown in Table 3 (Appendix B), it increased training time by 0.02 hours, added a negligible +0.0009 kg of CO₂ emissions, and caused only a slight rise in peak VRAM (38.67 GB vs. 39.31 GB), indicating feasibility for mid-scale models.

Finally, while our results show a statistically significant improvement in Exact Match scores (+3.5%), achieving absolute boundary determinism remains a challenge for generative architectures, even under strict SFT constraints. Our interpretation that prompt repetition may promote a unidirectional "cross-reading" mechanism is supported by behavioral evidence, such as reduced generative stutters and boundary offsets. Nonetheless, a formal extraction and visualization of multi-head attention maps could provide further validation and represents a promising direction for future interpretability research.

9. Acknowledgements

We thank the organizers of the FinNLP and FNP workshops for their dedication to the financial NLP community. We acknowledge the creators of the FinCausal 2026 dataset (Moreno-Sandoval et al., 2026), whose bilingual corpus enabled the cross-lingual analyses presented in this research.

10. Bibliographical References

- Tom B Brown, Benjamin Mann, Nick Ryder, et al. 2020. Language models are few-shot learners. In *Advances in Neural Information Processing Systems*, volume 33, pages 1877–1901.
- George Chrysostomou, Zhixue Zhao, Miles Williams, and Nikolaos Aletras. 2024. Investigating hallucinations in pruned large language models for abstractive summarization. *Transactions of the Association for Computational Linguistics*, 12:1163–1181.
- Sohom Ghosh and Sudip Kumar Naskar. 2022. Lipi at fincausal 2022: Mining causes and effects from financial texts. In *Proceedings of the 4th Financial Narrative Processing Workshop*, pages 121–123.
- Yaniv Leviathan, Matan Kalman, and Yossi Matias. 2025. Prompt repetition improves large language models. *arXiv preprint arXiv:2512.14982*.
- Zhiheng Lyu et al. 2022. Dcu-lorcan at fincausal 2022. In *Proceedings of the 4th Financial Narrative Processing Workshop*.
- Dominique Mariko, Hanna Abi-Akl, Kim Trottier, and Mahmoud El-Haj. 2022. The financial causality extraction shared task (fincausal 2022). In *Proceedings of the 4th Financial Narrative Processing Workshop @LREC2022*, pages 105–107.
- Antonio Moreno-Sandoval, Jordi Porta, Yanco Amor Torterolo, Alexia Stanescu, Melina Chatzi, and Sofia Roseti. 2026. The financial document causality detection shared task (fincausal 2026). In *Proceedings of the 7th Financial Narrative Processing Workshop (FNP 2026) at LREC 2026*, Palma de Mallorca, Spain. ELRA. To appear.
- Georg Niess, Houssam Razouk, Stasa Mandic, and Roman Kern. 2025. Addressing hallucination in causal q&a: The efficacy of fine-tuning over prompting in llms. In *Proceedings of the Joint Workshop of the 9th FinNLP, 6th FNP, and 1st LLMFinLegal*, pages 253–258.
- Julian Risch, Timo Möller, Julian Gutsch, and Malte Pietsch. 2021. Semantic answer similarity for evaluating question answering models. In *Proceedings of the 3rd Workshop on Machine Reading for Question Answering*, pages 149–157.
- Anik Saha, Jian Ni, Oktie Hassanzadeh, et al. 2022. Spock at fincausal 2022: Causal information extraction using span-based and sequence tagging models. In *Proceedings of the 4th Financial Narrative Processing Workshop*, pages 108–111.

Avinash Trivedi, Gauri Toshniwal, Sivanesan Sangeetha, and S.R. Balasundaram. 2025. Sarang at fincausal 2025: Contextual qa for financial causality detection combining extractive and generative models. In *Proceedings of the Joint Workshop of the 9th FinNLP, 6th FNP, and 1st LLMFinLegal*, pages 242–247.

Ashish Vaswani, Noam Shazeer, Niki Parmar, et al. 2017. Attention is all you need. In *Advances in Neural Information Processing Systems*, pages 5998–6008.

11. Language Resource References

Language Resources

Antonio Moreno-Sandoval, Yanco Amor Torterolo Orta, Maria Alexia Stanescu, and Melina Chatzi. 2026. [The Financial Document Causality Detection Shared Task \(FinCausal 2026\): Dataset](#).

A. Prompt Templates

We detail the exact prompt structures utilized during both Supervised Fine-Tuning and inference.
Simple Prompt (Zero-Shot SFT Baseline):

```
System:
You are a financial expert participating in FinCausal 2026.
Task: Extract the exact cause or effect from the provided financial text.
Rules:
1. The answer must be a VERBATIM extraction.
2. If the text contains a complex causal chain, extract the full relevant sequence.
3. Do not add introductory words.

User:
Context: {context}
Question: {question}
```

Repeated Prompt (Cross-Reading Configuration):

```
System:
You are a financial expert participating in FinCausal 2026. Task: Extract
the exact cause or effect from the provided financial text.
Rules:
1. The answer must be a VERBATIM extraction.
2. If the text contains a complex causal chain, extract the full relevant sequence.
3. Do not add introductory words.

User:
Context: {context}
Question: {question}
Context: {context}
Question: {question}
```

B. CodeCarbon Environmental Tracking

Model	Setting	Time (h)	CO ₂ (kg)	Peak VRAM (GB)
Qwen2.5-14B-Instruct-1M	Simple	0.20	0.0051	38.67
Qwen2.5-14B-Instruct-1M	Repeated	0.22	0.0060	39.31

Table 3: Empirical computational cost and environmental impact for fine-tuning on a single NVIDIA L40S (48GB) GPU. Total training time, peak VRAM utilization, and CO₂ emissions were directly tracked using the CodeCarbon library. Measurements were taken during a dedicated reproducibility run using identical hyperparameters, data, and hardware.

C. Qualitative Examples

Target Span	Simple Prompt Baseline	Repeated Prompt
Type 1: Boundary Offset & Over-generation		
[Qwen2.5-14B-Instruct-1M - EN] the £9.4m improvement in underlying operating cash flows	the £9.4m improvement in underlying operating cash flows offset by a £2.0m increase in outflows...	the £9.4m improvement in underlying operating cash flows
[GPT-4.1-Nano - EN] external threats become more sophisticated, and the potential impact of service disruption increases	As external threats become more sophisticated, and the potential impact of service disruption increases	external threats become more sophisticated, and the potential impact of service disruption increases
Type 2: Generative Stutter & Typographical Drops		
[Qwen2.5-7B-Instruct-1M EN] Our business serving the grocery sector benefited from several new accounts although the additional business won, combined with a competitive marketplace	Our business serving the grocery sector benefited from several new new accounts	Our business serving the grocery sector benefited from several new accounts although the additional business won, combined with a competitive marketplace
[Qwen2.5-14B-Instruct-1M - ES] no incluyen intereses, dividendos, ganancias o pérdidas procedentes de venta de inversiones o de operaciones de rescate o extinción de deuda	no incluyen intereses, dividendos, ganancias o pérdidas procedentes de venta de inversiones o de operaciones de rescate o extinción de deud	no incluyen intereses, dividendos, ganancias o pérdidas procedentes de venta de inversiones o de operaciones de rescate o extinción de deuda
Type 3: Truncation (Premature Stops)		
[GPT-4.1-Nano - ES] parámetros como el suministro de materias primas, la utilización de los quemadores, los sensores instalados o el balance entre energía fósil y eléctrica pueden ser gestionados de una manera más rápida, moderna y eficiente	parámetros como el s	parámetros como el suministro de materias primas, la utilización de los quemadores, los sensores instalados o el balance entre energía fósil y eléctrica pueden ser gestionados de una manera más rápida, moderna y eficiente
[Qwen2.5-14B-Instruct-1M - EN] The key partnerships established with leading European manufacturers	The key partnerships	The key partnerships established with leading European manufacturers

Table 4: Qualitative comparison of extraction errors. A focused evaluation of six samples illustrates how Prompt Repetition reduces generative stutters, boundary misalignments, and premature truncations across multiple architectures without requiring an external corrector model.

LeedsMEng26: Qwen + Gemini for FinCausal 2026 Causality Detection in Financial Narrative Texts

Ayomide Ivienagbor*, Idrees Asad*, Rijul Shrestha*
Yasemin Bal*, Zahaab Nadeem*, Zaid Shahrouri*

*University of Leeds, Leeds, United Kingdom

sc22ai@leeds.ac.uk, sc22i2a@leeds.ac.uk, sc22r2s@leeds.ac.uk,
sc22y2b@leeds.ac.uk, sc22zn@leeds.ac.uk, sc21z2s@leeds.ac.uk

Abstract

This paper presents the LeedsMEng26 system for the FinCausal 2026 shared task [Moreno-Sandoval et al. \(2026a\)](#) on financial causality detection in narrative texts. The task is formulated as extractive question answering over English and Spanish financial reports, where systems must return a verbatim span from the context that answers an abstractive question about a cause or an effect. We propose a two-stage pipeline consisting of candidate span generation followed by span verification and boundary refinement under a strict extractiveness constraint. We evaluate both an extractive RoBERTa-based baseline and instruction-tuned large language models. Results show that Qwen-2.5-1.5B-Instruct is a stronger candidate generator than the RoBERTa baseline, and that a second-stage verifier further improves answer boundary accuracy and overall adequacy. Our best configuration, Qwen-2.5-1.5B-Instruct with Gemini-2.5-flash refinement, achieved an adequacy score of 4.7000 for English and 4.6143 for Spanish. These findings suggest that a modular generation-and-verification pipeline is effective for extractive financial causality detection.

Keywords: Financial Narrative Processing, Causality Detection, Natural Language Processing, Question Answering, Multilingual NLP, Financial Text Analytics

1. Introduction

Financial narratives such as annual reports, earnings announcements, management commentaries, and regulatory filings play a central role in how firms communicate with investors, regulators, and the public. These documents often describe not only what has happened but also why it occurred. They do this through explicit or implicit causal statements, such as linking changes in performance to macroeconomic events or strategic decisions. Understanding these causal links is essential for interpreting corporate performance, assessing risk, and forming expectations about future developments. Manually analysing narratives at scale is costly and time-consuming, particularly given the volume and growing complexity of financial disclosures. This has motivated increasing interest in applying natural language processing (NLP) to financial text.

Within this emerging area, financial causality detection focuses on identifying cause–effect relations in financial documents. Automatically extracting such relations can support tasks such as risk analysis, fraud detection, forecasting, and explainable decision support. It also enhances transparency by making implicit reasoning patterns in corporate communication more explicit and machine-interpretable. Despite these benefits, financial text presents several challenges: it is often technical, domain-specific, and highly contextual, and causal language may be expressed in subtle, indirect or multi-sentence forms. Moreover, financial narratives frequently combine quantitative data with qualitative explanations, requiring models to integrate

numerical reasoning with linguistic understanding. Recent shared tasks, including the FinCausal track at the Financial Narrative Processing (FNP) workshop, provide benchmark datasets and evaluation protocols for studying these problems in a controlled setting, thereby enabling systematic comparison of modelling approaches.

This paper describes the LeedsMEng26 system for the FinCausal 2026 shared task ([Moreno-Sandoval et al., 2026b](#)), which formulates financial causality detection as question answering over English and Spanish financial texts. Given an abstractive question and a short context paragraph, the system must return an extractive span from the context that answers either the cause or the effect. We study both extractive and instruction-tuned generative approaches under a strict extractiveness constraint, and propose a two-stage pipeline: (i) candidate span generation and (ii) span verification and boundary refinement using a verifier LLM constrained to copy a contiguous substring.

2. Related Work

The FinCausal shared tasks have progressively advanced research on causality detection in financial narratives, moving from span extraction in earlier editions toward question-answering and more generative evaluation settings in recent years. The current edition, FinCausal 2026, continues the multilingual English–Spanish setting and retains the use of annotated contexts paired with abstractive questions and extractive answers. However, it introduces two notable changes: a new random parti-

tioning of the 2026 dataset and an LLM-as-a-judge evaluation metric, which scores responses on a 1–5 adequacy scale and replaces the previous Semantic Answer Similarity (SAS) and Exact Match (EM) metrics used in 2025 (Moreno-Sandoval et al., 2025).

Earlier editions focused more on span- and token-level causality extraction. FinCausal 2023 expanded the task across English and Spanish, using span-level Exact Match (EM) and token-level weighted F1 for evaluation, while encouraging multilingual and prompt-based approaches (Moreno-Sandoval et al., 2023). FinCausal 2022 focused on causality in quantified financial facts, with the winning SPOCK system using an ensemble of RoBERTa-Large sequence-tagging models with the BIO scheme (Mariko et al., 2022). Overall, the FinCausal series shows a progression from structured cause–effect extraction to multilingual reasoning, QA-based formulations, and more flexible generative evaluation.

(Moreno-Sandoval et al., 2025) introduces FinCausal 2025 competition entries and it was used as a guide to see where the most recent advancements for the task were. Participants adopted a range of approaches, spanning discriminative extractive QA models, generative LLMs with prompt engineering (simple, CoT, few-shot), and varying use of fine-tuning and quantization. Notably, strong scores were achieved even without fine-tuning in some cases, highlighting that fine-tuning was not the only route to competitive performance.

The (Trivedi et al., 2025) system paper was useful for our work because it provided a strong, task-aligned example of a hybrid extractive to generative refinement pipeline (RoBERTa-based span prediction followed by Gemma2-9B refinement) that directly targets common QA boundary and coherence errors while remaining competitive on the official SAS/EM metrics.

3. Dataset and Task

The FinCausal 2026 task is formulated as a generative question answering problem over financial annual reports, where systems must identify either the cause or the effect corresponding to an abstractive question. We decided to participate in both sub-tasks, training models on both the English and Spanish texts.

The training data are drawn from the FinCausal 2026 dataset Moreno-Sandoval et al. (2026b), which is provided in CSV format, with each file containing 2000 rows. Each instance was in the following form: (i) **ID** - its unique identifier; (ii) a **context**, corresponding to a paragraph extracted from a financial report; (iii) an **abstractive question**, ask-

ing for either the cause or the effect of a described event; and (iii) an **answer**, which is an extractive span taken precisely from the context. Although the question is abstractive, the expected output is a text span grounded strictly in the provided paragraph.

The 2026 edition introduces a revised and expanded dataset, including more complex causal structures and rephrased questions designed to require deeper reasoning. The training and test sets are randomly partitioned from this updated dataset. The task focuses on explanatory cause–effect relations within the text, especially where specific events result in measurable financial outcomes.

For our submission, we approached the problem from both an extractive and a generative perspective. We first experimented with span-based extraction models to exploit the extractive nature of the gold answers, and subsequently explored fine-tuning and combining large language models to assess whether generative approaches could better capture complex causal structures.

4. Methodology

4.1. Task Formulation

Financial causality extraction is approached as an extractive question answering (QA) task. Each instance consists of a financial context c and a question q specifying a causal relation. The system produces an answer span a that must appear verbatim within c . We apply this extractive constraint strictly to all system variants. This ensures that predictions are always drawn from the provided context.

We report span-level Exact Match (EM) and Semantic Answer Similarity (SAS) metrics. EM evaluates strict boundary correctness. SAS accounts for minor boundary deviations. The official leaderboard score is also tracked where applicable.

4.2. Pipeline Overview

The proposed approach utilises a two-stage framework to enhance boundary accuracy and preserve the extractive constraint:

1. **Candidate generation:** an extractive QA model produces a candidate span $\hat{a}$ from the context.
2. **Span verification and refinement:** an instruction-following LLM receives $(c, q, \hat{a})$ and either confirms $\hat{a}$ or corrects span boundaries, while being constrained to return a verbatim substring of c .

This design distinctly separates relevant evidence retrieval in stage 1 from boundary correction

and consistency verification in stage 2. Preliminary error analysis revealed that boundary truncations, such as missing causal qualifiers, and boundary overruns, such as the inclusion of adjacent clauses, are the most frequent failure modes in pure extractive models. The refinement stage specifically addresses these errors while maintaining strict adherence to the input context.

Figure 1 illustrates the two-stage candidate generation and refinement pipeline used in our system.

All systems employ a unified preprocessing & post-processing pipeline implemented using the Hugging Face `transformers` library. Inputs are constructed by concatenating the question and context with the model-specific tokeniser. For extractive QA models, start and end position labels are derived from the gold answer span.

Predicted spans undergo lightweight normalisation as follows: (i) whitespace trimming and deduplication, and (ii) minor punctuation trimming at span boundaries where it does not alter content. The extractive constraint is enforced through substring verification, requiring the produced answer to appear as a contiguous substring in the given context. If a refinement model outputs text that violates this constraint, the system reverts to the pre-refinement candidate.

4.3. Baseline 1: Extractive QA with RoBERTa

The first baseline employs `question-answering-roberta-base-s-v2`, a RoBERTa-base encoder fine-tuned for extractive QA. Given (c, q) , the model produces probability distributions over context tokens for the start and end indices of the answer span. The highest-scoring (i, j) pair is selected, and the corresponding substring is returned verbatim.

We fine-tuned the model on the official FinCausal English training data. We optimised based on the extractive QA objective, minimising cross-entropy loss over the gold start and end token positions. This baseline serves as a strictly extractive reference point with no generative components.

4.4. Baseline 2: RoBERTa + GPT Refinement

As a result, it was concluded that RoBERTa was sensitive to boundary errors that affected the consistency of meanings. Early truncation and overrun errors were the main issues; consequently, a second-stage refinement step is introduced using GPT-3.5 as a verifier.

The refiner is provided with the context, question, and candidate span predicted by RoBERTa. It is

instructed to:

- Output *only* the final answer span (no explanation),
- Copy the span verbatim from the context (no paraphrasing),
- Keep the candidate unchanged if correct, and
- Correct boundary errors by expanding or contracting to the shortest correct substring when multiple valid spans exist.

The method ensures compliance with the context and achieves improved span boundary precision. If the refined span is not a substring of the context, the original candidate is used.

4.5. Baseline 3: Qwen-2.5-1.5B-Instruct as Candidate Generator

The third baseline replaces the RoBERTa extractor with Qwen-2.5-1.5B-Instruct, a decoder-only instruction-tuned LLM with multilingual capability. Unlike encoder-style QA models that predict token positions, Qwen generates text directly; therefore, constrained prompting is used to enforce extractive behaviour.

The candidate-generation prompt explicitly requires the model to output a verbatim substring from the context and nothing else. Despite these directives, we still observed boundary drift and occasional paraphrasing. These were similar to those seen with RoBERTa when the model attempted to be 'helpful.'

Low-Rank Adaptation (LoRA) is employed, updating only low-rank matrices inserted into selected attention and feed-forward layers while keeping the base weights frozen, enabling Qwen to adapt to the task with limited computational resources and resulting in reduced training memory and cost compared to full fine-tuning.

Reinforcement learning (RL) post-training is explored using a composite reward. The reward encourages extractive correctness and semantic fidelity. It combines: (i) Exact Match (EM), (ii) a semantic similarity component between the prediction and the gold span, and (iii) an auxiliary judge score reflecting span correctness and boundary precision. At inference time, we use conservative decoding (low temperature) to reduce variability and discourage paraphrasing. Qwen-2.5-1.5B-Instruct was also fine-tuned using Low-Rank Adaptation (LoRA) for parameter-efficient fine-tuning. (Schulman and Lab, 2025)

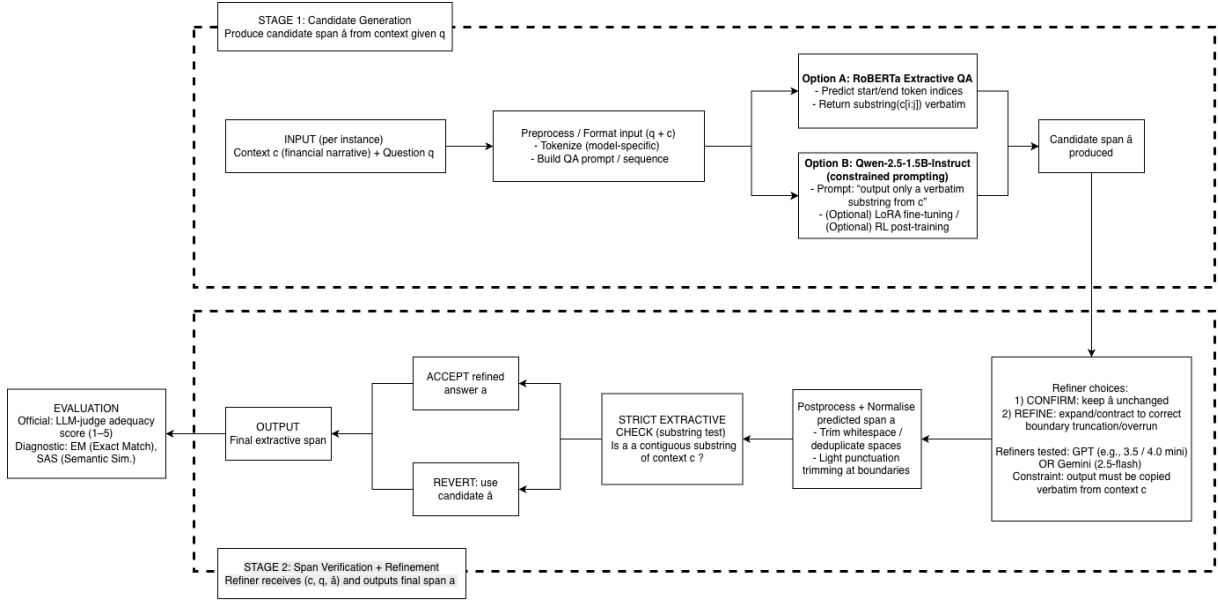


Figure 1: Two-stage pipeline for financial causality extraction.

4.6. Hybrid: Qwen + GPT Refinement

Qwen generates an initial candidate span under strict prompting, and the verifier model performs second-stage verification and boundary correction, replicating the two-stage process used in the RoBERTa + verifier configuration.

We perform substring verification after refinement as before. If GPT-based refinement (GPT-3.5 and GPT-4.0 mini) violates extractiveness, the system reverts to the original Qwen candidate. This hybrid approach uses Qwen’s strong instruction-following together with GPT-3.5 and GPT-4.0 mini for verification, reducing hallucination and improving boundary precision. The following configuration further explores this two-stage strategy by integrating a different refinement model, as outlined next.

4.7. Final System: Qwen + Gemini Refinement

Ultimately, we used Gemini-2.5-Flash as the refinement model. Qwen generates a candidate span using strict extractive prompting, then provides the prediction to Gemini and is instructed to verify correctness and adjust the boundaries if needed.

Similar to previous hybrids, we apply minimal normalisation and substring verification, reverting to the candidate span if refinement breaks extractiveness. This configuration provided the most consistent boundary corrections and the strongest overall performance across English and Spanish settings. The next sections outline the experimental setup supporting these results.

4.8. Refinement Prompt

The refinement stage relies on a constrained prompt designed to correct boundary errors while preserving the extractive constraint.

The verifier model receives the context, question, and candidate span produced by Qwen and is instructed to return only the final corrected answer.

The prompt enforces several rules: (i) the output must be a verbatim span from the context, (ii) no explanations or additional text may be produced, (iii) the candidate span should be returned unchanged if it is already correct, and (iv) if multiple spans are possible, the shortest correct substring should be selected.

A shortened version of the prompt is shown below.

You are correcting a short answer. Output only the final answer. Return a verbatim span from the context. Do not add explanations. If the answer is already correct, return it unchanged. If multiple spans are possible, choose the shortest correct span.

The full prompt, including the in-context examples used during inference, is provided in Appendix A. A Spanish version of the prompt with identical constraints was used for Spanish inputs.

4.9. Experimental Setup

- **Data:** Experiments used the FinCausal 2025 dataset. Training was performed using only the English data. Spanish examples were evaluated without additional training.

- **Data Split:** The English dataset consisted of 2000 instances, which were divided into 75% training and 25% validation data.
- **RoBERTa Training:** The parameters used were learning rate of 3×10^{-5} , weight decay 0.01, and a linear scheduler with warmup ratio 0.1. The batch size was 4 with gradient accumulation of 2 (effective batch size 8). The best model was selected based on validation loss.
- **Qwen Fine-tuning:** Qwen-2.5-1.5B-Instruct was fine-tuned using LoRA with rank $r = 16$, $\alpha = 32$, and dropout 0.0. Bias parameters were not adapted, and gradient checkpointing was enabled using Unsloth.
- **Decoding:** Low-temperature decoding was used during inference. Outputs were constrained to be substrings of the input context to ensure extractive answers.
- **Reproducibility:** Random seeds and preprocessing were kept consistent across all experiments.

5. Experiments and Results

5.1. Evaluation Protocol

The FinCausal 2026 shared task uses an LLM-as-a-judge evaluation protocol. Each system prediction is scored on a 1–5 adequacy scale. The judge’s score rewards answers that fully address the question using evidence from the context. It penalises truncations, overruns, and non-extractive outputs.

In addition to the official adequacy score, we report **Exact Match (EM)** and **Semantic Answer Similarity (SAS)** as *diagnostic* metrics to analyse span boundary quality and semantic correspondence during development. Unless stated otherwise, EM/SAS are computed by comparing predictions against available labelled data and are used for internal evaluation rather than leaderboard ranking.

5.2. English Results

System Configuration	Score
RoBERTa + GPT-3.5	3.8200
Qwen-2.5-1.5B-Instruct + GPT-3.5	4.5080
Qwen-2.5-1.5B-Instruct	4.6060
Qwen-2.5-1.5B-Instruct + GPT-4.0 mini	4.6240
Qwen-2.5-1.5B-Instruct + Gemini-2.5-flash	4.7000

Table 1: Best performing English submission per system configuration (LLM-judge adequacy score).

Table 1 summarises the best-performing English submission for each system configuration. Overall, performance improves monotonically as the candidate generator becomes more instruction-aligned and as verification-based refinement is introduced.

The RoBERTa-based hybrid baseline (**RoBERTa + GPT-3.5**) achieves an adequacy score of **3.8200**. While the extractive QA model reliably returns substrings from the context, qualitative inspection indicates that it repeatedly faces minor span boundary errors (e.g., missing causal qualifiers or including adjacent clauses), which reduces adequacy under judge-based scoring.

Replacing the encoder-based extractor with an instruction-tuned decoder model yielded substantial improvement. **Qwen-2.5-1.5B-Instruct** achieves **4.6060**, indicating that constrained, instruction-following generation is better aligned with the causality-oriented QA prompts and the judge’s emphasis on completeness and adequacy.

Adding a second-stage verifier provides further gains. **Qwen + GPT-4.0 mini** reaches **4.6240**, indicating that refinement helps correct residual boundary variances and improves answer adequacy.

The best-performing English configuration is **Qwen + Gemini-2.5-flash**, achieving **4.7000**. This result supports the effectiveness of a generation verification pipeline in which a strong refiner corrects subtle span boundary errors while preserving the extractive constraint.

For diagnostic analysis, Table 3 reports EM and SAS on labelled English data. Both metrics increase consistently across configurations (EM from **0.3500** to **0.7415**; SAS from **0.8850** to **0.9460**), indicating that adequacy gains are accompanied by improvements in boundary accuracy and semantic correspondence, rather than reflecting superficial formatting differences.

5.3. Spanish Results

System Configuration	Score
RoBERTa + GPT-3.5	3.9264
Qwen-2.5-1.5B-Instruct + GPT-4.0	4.4692
Qwen-2.5-1.5B-Instruct	4.5030
Qwen-2.5-1.5B-Instruct + Gemini-2.5-flash	4.6143

Table 2: Best performing Spanish submission per system configuration (LLM-judge adequacy score).

Spanish results (Table 2) follow similar trends, with the strongest performance again obtained by verification-based refinement using Gemini-2.5-Flash.

The RoBERTa + GPT baseline achieves **3.9264**. The standalone **Qwen-2.5-1.5B-Instruct** model im-

proves adequacy to **4.5030**, demonstrating strong cross-lingual generalisation in Spanish under extractive prompting.

Unlike in English, adding GPT-based refinement slightly reduced performance (**4.4692** vs **4.5030**), suggesting that the refiner may occasionally over-correct boundaries or produce outputs that are less well aligned with the judge’s adequacy preferences for Spanish instances. In contrast, **Qwen + Gemini-2.5-flash** yields the best Spanish score of **4.6143**, indicating that Gemini provides more reliable verification and boundary correction in the bilingual setting.

5.4. Overall Analysis

Across both languages, two observations are deduced. First, instruction-tuned candidate generation (Qwen) better aligns with the task format and produces more adequate spans with limited prompting. Secondly, a second-stage verifier further improves robustness using correcting boundary truncations and overruns, with **Gemini-2.5-flash** providing the most reliable refinement among tested models. RoBERTa achieved lower scores on both the English and Spanish tasks, indicating lower effectiveness than the alternative approaches evaluated.

Overall, the hybrid **Qwen + Gemini-2.5-Flash** system achieved the best results in both English and Spanish, indicating that generation verification pipelines are effective for extractive financial causality tasks evaluated using an adequacy-oriented LLM judge.

Model	EM Accuracy	SAS
RoBERTa + GPT-3.5	0.3500	0.8850
Qwen-2.5-1.5B-Instruct + GPT-3.5	0.6295	0.9155
Qwen-2.5-1.5B-Instruct	0.6855	0.9366
Qwen-2.5-1.5B-Instruct + Gemini-2.5-flash	0.7415	0.9460

Table 3: Diagnostic EM and SAS results for English (computed on labelled data for development).

6. Strengths and Limitations

A key strength of our work is that, despite limited time and resources, we achieved strong competitive performance, placing 7th in the English category with only marginal differences from teams ranked above us. This result demonstrates that our approach was effective and robust even under practical constraints.

Our pipeline also benefits from a clear and structured design that combines the reliability of ex-

tractive question answering with the contextual reasoning benefits of instruction-tuned Large Language Models and verification/refinement steps applied after the initial prediction. By progressing from a strong extractive baseline (RoBERTa) to an instruction-following model (Qwen) and then adding refinement stages (GPT-3.5, GPT-4.0 mini, and Gemini-2.5-Flash), we were able to isolate which components contributed most to performance gains and reduce common QA errors such as span boundary drift. In particular, our two-stage candidate and verifier setup made the system more reliable. The verifier checks the initial span and fixes common issues such as answers being too short (truncated) or too long, while also ensuring the final output stays strictly extractive. This iterative setup also made our experiments more interpretable, since each stage had a clear and measurable role in improving Exact Match and semantic alignment.

Despite the effectiveness of our pipeline, several constraints limited the scope of our experiments and likely capped performance:

- **Time constraints due to schedule clashes.** The shared task timeline overlapped with our university exam period, causing us to begin experimentation later than planned and reducing the time available for broader hyperparameter searches and ablation studies.
- **Limited compute access.** We did not have access to the university’s powerful GPUs, which restricted our ability to train larger models, run longer fine-tuning schedules, or explore more compute-intensive approaches. A single NVIDIA RTX 3050 was used for the reinforcement learning finetuning of Qwen-2.5-1.5B.
- **Restricted access to the newest proprietary LLM APIs.** We were unable to use the latest frontier LLM APIs. Access to stronger models for refinement and verification could plausibly have improved boundary correction and reduced rare failure cases.
- **Training focused only on English.** Our main training and optimisation effort was concentrated on the English dataset rather than fully training separate systems for both English and Spanish. This likely reduced Spanish performance relative to what could be achieved with language-specific fine-tuning and validation.

7. Dataset Feedback

One limitation of the dataset is its relatively small size, which limited the amount of supervision avail-

able for adapting the models to the task. While extractive models remained reasonably stable, larger generative models were more sensitive to the limited supervision and showed less consistent performance. This limitation influenced our decision to adopt a multi-stage approach with a separate verification step, rather than relying solely on direct fine-tuning. A larger training set would likely allow more effective fine-tuning of generative models and lead to more stable and reliable results overall.

8. Conclusion and Future Work

Across the project, we demonstrate that a modular extractive QA pipeline is an effective approach for financial causality extraction, achieving a 7th-place ranking in the English track with only marginal differences from higher-ranked teams. Starting from an extractive baseline (RoBERTa) and progressively incorporating instruction-following modelling (Qwen) and refinement stages (GPT-3.5, GPT-4.0 mini, and Gemini-2.5-Flash), we achieved consistent improvements while keeping outputs strictly extractive. In particular, the candidate-verifier design helped reduce span boundary drift by correcting answers that were too short (truncated) or too long, and the staged structure made it clear which components were responsible for the improvements in Exact Match and semantic alignment.

For future work, this system will serve as a foundation for a Masters-level group project and will be expanded in both scope and capability. We aim to evaluate whether the approach generalises beyond financial reports to other domains such as medical or construction text, and to train and fine-tune stronger models using improved computational resources to further boost performance.

References

Dominique Mariko, Kim Trotter, and Mahmoud El-Haj. 2022. [The financial causality extraction shared task \(fincausal 2022\)](#). In *Proceedings of the 4th Financial Narrative Processing Workshop @ LREC 2022*, pages 105–107, Marseille, France. European Language Resources Association.

Antonio Moreno-Sandoval, Blanca Carbajo Coronado, Jordi Porta Zamorano, Yanco Amor Torterolo Orta, and Doaa Samy. 2025. [The financial document causality detection shared task \(FinCausal 2025\)](#). In *Proceedings of the Joint Workshop of the 9th Financial Technology and Natural Language Processing (FinNLP), the 6th Financial Narrative Processing (FNP), and the*

1st Workshop on Large Language Models for Finance and Legal (LLMFinLegal), pages 214–221, Abu Dhabi, UAE. Association for Computational Linguistics.

Antonio Moreno-Sandoval, Jordi Porta, Yanco Torterolo, Alexia Stanescu, Melina Chatzi, and Sofía Roseti. 2026a. [The Financial Document Causality Detection Shared Task \(FinCausal 2026\)](#). In *Proceedings of the 7th Financial Narrative Processing Workshop (FNP 2026) at LREC 2026*, Palma de Mallorca, Spain. ELRA.

Antonio Moreno-Sandoval, Jordi Porta-Zamorano, Blanca Carbajo-Coronado, Doaa Samy, Dominique Mariko, and Mahmoud El-Haj. 2023. [The financial document causality detection shared task \(fincausal 2023\)](#). In *2023 IEEE International Conference on Big Data (BigData)*, Sorrento, Italy. IEEE.

Moreno-Sandoval, Antonio and Torterolo Orta, Yanco Amor and Stanescu, Maria Alexia and Chatzi, Melina. 2026b. [The Financial Document Causality Detection Shared Task \(FinCausal 2026\): Dataset](#). e-cienciaDatos.

John Schulman and Thinking Machines Lab. 2025. [Lora without regret](#). *Thinking Machines Lab: Connectionism*.

Avinash Trivedi, Gauri Toshniwal, Sivanesan Sangeetha, and S. R. Balasundaram. 2025. [Sarang at FinCausal 2025: Contextual QA for financial causality detection combining extractive and generative models](#). In *Proceedings of the Joint Workshop of the 9th Financial Technology and Natural Language Processing (FinNLP), the 6th Financial Narrative Processing (FNP), and the 1st Workshop on Large Language Models for Finance and Legal (LLMFinLegal)*, pages 242–247, Abu Dhabi, UAE. Association for Computational Linguistics.

A. Appendix

A.1. English Refinement Prompt

The following prompt was used in the second-stage refinement step of the Qwen + Gemini pipeline.

You are correcting a short answer. **Rules:**

- Output **only** the final answer.
- Do **not** add explanations.
- Return a **verbatim span** from the context.
- If the answer is already correct, return it unchanged.
- Prefer the full sentence containing the answer if needed.
- If multiple spans are possible, choose the shortest correct one.
- Goal: maximise Exact Match.

Examples:

Context: Additionally, during the year the Group commenced work, under its EPCC contract with CGE, on four biogas-based power generation plants. As previously stated, due to financial constraints, progress was slower than initially expected and work has been temporarily suspended, awaiting the Company finalising an arrangement with CGE.

Question: What did financial constraints bring about?

Answer: progress was slower than initially expected and work has been temporarily suspended, awaiting the Company finalising an arrangement with CGE

Context: The Directors believe this growth is driven by consumer preferences moving away from chain and branded pubs and towards pubs with an individual identity and an ambience which reflects the local market.

Question: What explains the growth, according to the Directors?

Answer: consumer preferences moving away from chain and branded pubs and towards pubs with an individual identity and an ambience which reflects the local market

Context: The nature of the Group's operations creates an ongoing demand for fuel and therefore the Group is exposed to movements in market fuel prices. The Group enters into commodity derivative instruments to hedge such exposure where it makes commercial and economic sense to do so.

Question: What explains why the Group is exposed to movements in market fuel prices?

Answer: The nature of the Group's operations creates an ongoing demand for fuel

A.2. Spanish Refinement Prompt

For Spanish inputs, a language-adapted version of the refinement prompt was used to ensure consistent instruction following. The structure and constraints remained identical to the English prompt, but the instructions and examples were provided in Spanish.

Estás corrigiendo una respuesta corta.

Reglas:

- Devuelve **solo** la respuesta final.
- No añadas explicaciones.
- La respuesta debe ser un fragmento **verbatim** del contexto.
- Si ya es correcta, devuélvela igual.
- Si es necesario, prefiere la oración completa que contiene la respuesta.
- Si hay múltiples opciones, elige la más corta correcta.
- Objetivo: maximizar Exact Match.

Ejemplos:

Contexto: Debido a restricciones financieras, el progreso fue más lento de lo esperado y el trabajo ha sido suspendido temporalmente hasta que la empresa finalice un acuerdo.

Pregunta: ¿Qué provocaron las restricciones financieras?

Respuesta: el progreso fue más lento de lo esperado y el trabajo ha sido suspendido temporalmente hasta que la empresa finalice un acuerdo

A.3. Model Versions

Table 4: Model versions used in the experiments.

Model	Version / Checkpoint	Date Used
Qwen	Qwen2.5-1.5B-Instruct	March 2026
RoBERTa QA	question-answering-roberta-base-s-v2	March 2026
GPT-3.5	GPT-3.5	March 2026
GPT-4	GPT-4.0	March 2026
Gemini	Gemini 2.5 Flash	March 2026

Financial Causal QA via Instruction and Prompt Tuning of Gemma3-12B

Avinash Trivedi, Chindukuri Mallikarjuna

SRM University-AP,
Amaravati, Andhra Pradesh 522240, India
avinashtrivedi.2008@gmail.com

Abstract

In this paper we present a novel methodology that harnesses the power of prompt tuning applied directly to Gemma3-12B, a state-of-the-art generative large language model to enhance performance on complex natural language processing challenges. Instead of relying solely on extensive retraining, our approach leverages carefully crafted input prompts to steer the pre-trained Gemma-12B towards generating outputs with superior contextual accuracy and interpretability. Our experimental evaluation employed a composite LLM Score metric that quantifies both semantic coherence and relevance; under this framework, our system (Team Name: *Sarang*) achieved a score of 4.54, ranking 9th in the shared task. Furthermore, in the competitive task evaluation, our method demonstrated the potential of prompt tuning as a viable alternative to traditional fine-tuning approaches. This study not only demonstrates the practical benefits of integrating prompt engineering with large language models but also opens avenues for future research aimed at further optimizing model performance in domain-specific applications.

Keywords: FinCausal, Prompt Tuning, LLM Score

1. Introduction

Natural question answering systems which are expected to facilitate decision-making and market research should have a sophisticated sense of cause-and-effect structure embedded in financial narrative. Annual reports, earnings statements, etc., often describe events, antecedents and the consequences of those events such that it requires a keen sense of cause and effect. These linkages are computationally intensive and impractical to extract manually when faced with the large volume of modern financial text corpora. Therefore, automated causal question-answering reduces such limitations, simplifying the processing of financial information and making the produced financial intelligence more readable and understandable. It is on this basis that the creation of systems that can respond to causal questions derived through financial narratives has become one of the critical research areas.

The FinCausal 2026 Shared Task ([Moreno-Sandoval et al., 2026a](#)), which is structured into the Financial Narrative Processing Workshop, is focused on the improvement of the methodologies of causal question-answering in financial texts. The dataset has been carefully designed in terms of triadic settings of context, inquiry, and response whereby the participants are required to make an inference of the missing causal determinant on a financial passage and the query related to that passage. These interrogatives are highly abstracted and hence compel the examinees to identify antecedent or consequent items whereas the intended responses are specific extractive spans

obtained out of the contextual contents. Therefore, the construct can be described as a graceful combination of classical span extraction and question answering with reasoning, and, as a consequence, such an elevated standard of subtle understanding of financial stories. Although modern large language models have made significant progress, it is still an extremely challenging task to identify causal relationships in the financial discourse: causes and effects are incorporated in a systematic way in such texts, specialised terminological registers are used, and long-term causal relationships are cultivated. The 2026 update makes the complexity intrinsic, with increased granularity of causal processes, and re-organization of ways of inquiry, to require a high level of inferential skill. This means that researchers have to come up with mechanisms that can be used to traverse complex causal networks, beyond surface pattern recognition. What adds to these inherent difficulties is the addition of a new judgment measure based on the use of Large Language Models as adjudicators (LLM-as-a-judge). This change of paradigm shifts the focus of emphasis on specific span recall to a general evaluation of the semantic fidelity and reasoning profundity. Correspondingly, analysis will be based on the ability of a system to encode cause-and-effect relationships with probable coherence, and not merely match annotated spans.

2. Literature study

The methodic identification and parsing out of causal relationships in the financial discourse has

become a critical project in the context of financial natural language processing. The FinCausal shared tasks have provided a significant impact on this pattern, providing strict guidelines and increasing complex evaluation models. In 2020, the first FinCausal attempt was launched, introducing the first annotated corpus on financial causality detection, which defines two fundamental subtasks: sentence-level classification and definition of cause-effect spans (Mariko et al., 2020). In 2021 and 2022, further cycles of improvement were made, along with annotation guidelines, increasing the data coverage, and advancing more complex causal patterns, including quantified facts and transformation-based relations (Mariko et al., 2021, 2022). These initial implementations demonstrated the effectiveness of transformer-based encoders in combination with highly structured span-extraction systems, and at the same time, the difficulty in representing implicit causality and cross-sentence reasoning. Based on this background, the 2023 version of FinCausal expanded its criterion to include multilingual tracks and redefined the role of the span-oriented extraction as the part of question answering or sequence generating (Moreno-Sandoval et al., 2023). The 2025 version also better specified its aims by incorporating a multilingual causal question-answer model that may be assessed by Exact Match (EM) and Semantic Answer Similarity (SAS) scores (Moreno-Sandoval et al., 2025). This redefinition marked the end of pure extraction and an incorporative reasoning and generation of answers. FinCausal is closely related to the development of financial question answering research, which places greater emphasis on financial reasoning and multi-step inference, at the cost of more traditional financial reporting (Chen et al., 2021). Taken together, these changes reflect a larger change in causal recognition, that is, surface-level identification to a more abstract financial reasonability that incorporates text with quantitative information. The effectiveness of hybrid architectures which combine extractive precision and generative reasoning is further supported by the recent literature. (Pilault et al., 2020) proposed a model in which an extractive item picks the relevant evidence to modulate a transformer-based generative model and thus showing a better contextual consistency. According to (Luo et al., 2022), extractive and generative QA models have been systematically compared, with the first type proving much more successful in generalisation in limited settings, while the latter type possesses the benefits of abstraction. The NeurIPS EfficientQA competition (Min et al., 2021) unveiled the trade-off between the computation efficiency and results, where well optimised lightweight extractive models can be competitive with state-of-the-art results. Basic transformer models like

RoBERTa (Liu et al., 2019) have also enhanced the abilities of domain adaptation.

3. Dataset for FinCausal2026

This paper uses English version of FinCausal 2026 Question Answering dataset (Moreno-Sandoval et al., 2026b) published as part of FinCausal shared task, which is dedicated to detecting causal links in financial narratives. The task is expected to test systems, which are able to read financial texts and identify the cause or effect of financial events. The dataset is assembled out of the financial reports and corporate disclosures in which causal relations are often found in the descriptions of the company performance, market trends and economic situation. The dataset consists of 2000 training instances and 500 instances for testing. In the training dataset each instance includes *ID*, *Context*, *Question* and *Answer*. Whereas test data contains same attributes except *Answer*.

Considering a financial context and a causal question, the goal is to derive the right cause or effect.

4. Methodology

4.1. Few-shot prompting and Finetuning of Language Model

As a baseline, We tried few-shot prompting of various LLMs. Later started the finetuning of *consciousAI/question-answering-roberta-base-s*, then changed the checkpoint to *deepset/deberta-v3-large-squad2* followed by prompt based enhancement steps as in Fig 1, inspired from (Trivedi et al., 2025) to improve the response received from finetuned model. This technique was giving LLM score of 4.424. We also tried prompt tuning and from there we found our best performing system discussed in section 4.2.

4.2. System Submission

The current section outlines the proposed structure of the FinCausal Question Answering (QA) task. Fig 2 representing the architecture of the model submission. The current research involves a methodology that integrates the few-shot learning, automated prompt optimization and a large language model to identify causal relationships in finance in a systematic manner. The pipeline includes four main elements, dataset utilisation, few-shot prompt construction, prompt tuning, and model inference, which together allow performing sound causal reasoning on financial texts.

```

Prompt

{"role": "system",
"content": "You are a helpful assistant
that provides accurate and improved an-
swers."},
{"role": "user",
"content": """"You are given a Context, a
Question, and an Answer.
1. If the Answer is 100% correct and is ex-
tracted verbatim from the Context, return
the exact same Answer.
2. If the Answer is incorrect or not fully
extracted from the Context, return an im-
proved version of the Answer that is ex-
tracted verbatim from the Context.
Context: {context}
Question: {question}
Answer: {answer} """"}

```

Figure 1: Prompt for enhancement step

4.2.1. Few-Shot Prompt Construction

Few-shot learning has proven to be effective in instructing large language models to perform specialised reasoning (Brown et al., 2020; Wei et al., 2022). In the current methodology, few instances of the dataset are integrated into prompts as demonstrations. Each instance in the dataset is characterized by a financial context, a causal question, and the answer that clearly outlines the cause effect relationship. Such demonstrations offer implicit information to the language model, and it is able to identify how causal relationships are formulated in financial texts and how the appropriate responses can be produced.

4.2.2. Prompt Optimization using MIPROv2

To further enhance timely efficacy, the system will use MIPROv2 prompt optimization through the DSPy teleprompter framework (Khatab et al., 2022, 2024). DSPy provides a programmatic interface that is structured in such a way that it can be optimized and interactions with language models co-ordinated. The MIPROv2 teleprompter automatically optimizes prompts by sequentially sampling a space of instruction and few-shot examples. Through iterative evaluation, the system will pick highly effective prompts, thus improving the ability of the model to identify causal relationship and provide accurate answers for FinCausal QA task.

4.2.3. Model Integration and Inference

Gemma3:12B large language model (Gemma and DeepMind, 2024) is the refined version that performs the prompts and provides the strong language understanding and reasoning skills. The implementation of the model makes use of DSPy, which is used together with Ollama to enable efficient local execution and enable a controlled in-

teraction with the model. In the process of inference, the system obtains a context and a causal query. The optimized prompt along with the selected few-shot examples are sent to the language model using the DSPy framework. The model then performs contextual reasoning on the financial text and then gives out the final answer thus discovering the relevant causal relationship.

Our final model was build on prompt tuning of Gemma3-12B using MIPROv2 of DSPy (Khatab et al., 2022, 2024). The tuned system prompt is in Fig 3 and best parameters are listed in Table 1

Hyperparameter	Value
auto	medium
max_bootstrapped_demos	10
max_labeled_demos	10
model	gemma3:12b

Table 1: MIPROv2 Hyperparameters

5. Experimental Results

The results of our experiments few-shot prompting, prompt tuning and including finetuning of *consciousAI/question-answering-roberta-base-s* and *deepset/deberta-v3-large-squad2* are listed in Table 2.

Technique	Model	LLM Score
Finetuning	roberta based	4.136
Finetuning + Enhancement	roberta based	4.288
Finetuning	deberta based	4.292
Finetuning + Enhancement	deberta based	4.302
Few-shot	gemma2:latest	4.424
Few-shot	qwen3:latest	4.418
Few-shot + light optimization	gemma3:12b	4.448
Few-shot + medium optimization	gemma3:12b	4.54

Table 2: Performance comparison on test set

The experimental findings prove the existence of a performance improvement between the conventional transformer-based fine-tuning strategies and LLM based few-shot strategies augmented with prompt tuning. First, both the RoBERTa-based and DeBERTa-based models perform competitively among the fine-tuning methods. The initial fine-tuned RoBERTa model has a score of 4.136 that is enhanced to 4.288 using prompt based response enhancement step mentioned in Fig 1. In the same manner, the fine-tuned DeBERTa model achieves 4.292, which is slightly higher than RoBERTa and it reaches 4.302 after applying enhancement step. These outcomes demonstrate the idea that the architectural variation (DeBERTa versus RoBERTa) comes with minor benefits, whereas the enhancement strategies are both able to provide incremental improvements to both models.

Conversely, the few-shot LLM-based models have a visible lead over all fine-tuned encoder mod-

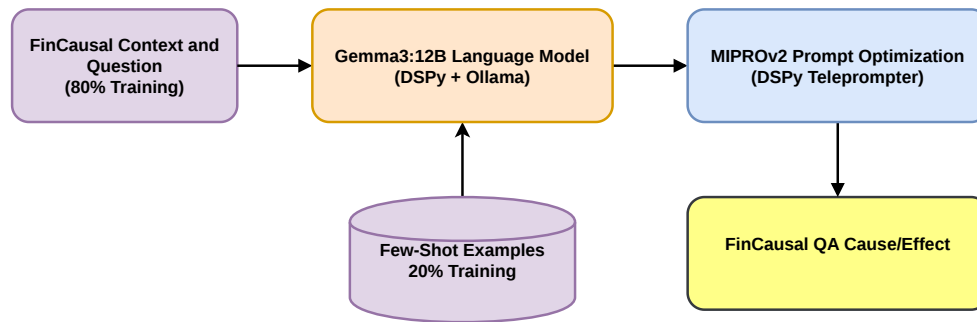


Figure 2: Model architecture

Prompt

Your input fields are:
 1. context (str): Financial report excerpt containing the causal relation
 2. question (str): Question asking for the cause or effect

Your output fields are:
 1. answer (str): Exact causal span copied from the context

All interactions will be structured in the following way, with the appropriate values filled in.

```
[[ ## context ## ]]
{context}
```

```
[[ ## question ## ]]
{question}
```

```
[[ ## answer ## ]]
{answer}
```

```
[[ ## completed ## ]]
```

In adhering to this structure, your objective is:

You are a financial causal reasoning expert.

The answer to the question (cause or effect) is ALWAYS explicitly stated in the provided context.

Extract the exact text span from the context that answers the question.

Rules:

- The answer MUST be copied verbatim from the context.
- Do NOT paraphrase.
- Do NOT add any extra words.
- Return only the precise causal phrase.

Figure 3: Tuned system prompt

els. The few-shot set up using Gemma2 also gives 4.424, and Qwen3 gives 4.418, indicating the high level of generalisation of large instruction-tuned models without task-specific fine-tuning. This indicates that prompt-based learning on modern LLMs performs better at this task environment than in traditional fine-tuning. Lastly, better performance is further improved by optimisation of bigger LLMs. The few-shot plus light optimisation experiment with

Gemma3 (12B) has a result of 4.448 and medium optimisation has the most optimal result of 4.54, which is the best overall result. The development of this process underlines the fact that prompt tuning methods significantly enhance the efficiency of a model. Overall, the outcomes show that there is a definite trend: bigger LLMs with organised prompt tuning outperform classic fine-tuned transformer baselines, achieving the best results in an experimental environment.

6. Conclusions and Future Work

Within this investigation we have delineated several experimental paradigms, including few-shot learning, fine-tuning procedures, and prompt tuning. Our most effective submission i.e. Gemma3-12B prompt tuning achieved an LLM Score of 4.54.

Looking ahead, future work will concentrate on overcoming computational resource constraints, thereby permitting an in-depth exploration of prompt tuning strategies for larger LLMs. Furthermore, investigating data augmentation techniques to fine-tune the deberta based checkpoint represents another promising research direction. Finally, the potential integration of LLM agents remains a viable avenue for subsequent experimental inquiries.

7. Bibliographical References

- Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, et al. 2020. Language models are few-shot learners. *Advances in neural information processing systems*, 33:1877–1901.
- Zhiyu Chen, Wenhui Chen, Charese Smiley, Sameena Shah, Iana Borova, Dylan Langdon, Reema Moussa, Matt Beane, Ting-Hao Huang, Bryan Routledge, and William Yang Wang. 2021. Finqa: A dataset of numerical reasoning over

- financial data. In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 3697–3711. Association for Computational Linguistics.
- Team Gemma and Google DeepMind. 2024. Gemma: Open models based on gemini research and technology. *arXiv preprint arXiv:2403.08295*.
- Omar Khattab, Keshav Santhanam, Xiang Lisa Li, David Hall, Percy Liang, Christopher Potts, and Matei Zaharia. 2022. Demonstrate-search-predict: Composing retrieval and language models for knowledge-intensive NLP. *arXiv preprint arXiv:2212.14024*.
- Omar Khattab, Arnav Singhvi, Paridhi Maheshwari, Zhiyuan Zhang, Keshav Santhanam, Sri Vardhamanan, Saiful Haq, Ashutosh Sharma, Thomas T. Joshi, Hanna Moazam, Heather Miller, Matei Zaharia, and Christopher Potts. 2024. Dspy: Compiling declarative language model calls into self-improving pipelines.
- Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Mandar Joshi, Danqi Chen, Omer Levy, Mike Lewis, Luke Zettlemoyer, and Veselin Stoyanov. 2019. Roberta: A robustly optimized bert pretraining approach. *arXiv preprint arXiv:1907.11692*.
- Man Luo, Kazuma Hashimoto, Semih Yavuz, Zhiwei Liu, Chitta Baral, and Yingbo Zhou. 2022. Choose your qa model wisely: A systematic study of generative and extractive readers for question answering. In *Proceedings of the 1st Workshop on Semiparametric Methods in NLP: Decoupling Logic from Knowledge*, pages 7–22, Dublin, Ireland and Online. Association for Computational Linguistics.
- Dominique Mariko, Hanna Abi-Akl, Estelle Labidurie, Stephane Durfort, Hugues De Mazancourt, and Mahmoud El-Haj. 2020. The financial document causality detection shared task (fincausal 2020). In *Proceedings of the 1st Joint Workshop on Financial Narrative Processing and MultiLing Financial Summarisation*, pages 23–32.
- Dominique Mariko, Hanna Abi-Akl, Kim Trottier, and Mahmoud El-Haj. 2022. The financial causality extraction shared task (fincausal 2022). In *Proceedings of the 4th Financial Narrative Processing Workshop @LREC2022*, pages 105–107, Marseille, France. European Language Resources Association.
- Dominique Mariko, Hanna Abi Akl, Estelle Labidurie, Stephane Durfort, Hugues de Mazancourt, and Mahmoud El-Haj. 2021. [The financial document causality detection shared task \(FinCausal 2021\)](#). In *Proceedings of the 3rd Financial Narrative Processing Workshop*, pages 58–60, Lancaster, United Kingdom. Association for Computational Linguistics.
- Sewon Min et al. 2021. Neurips 2020 efficiency competition: Systems, analyses and lessons learned. In *Proceedings of the NeurIPS 2020 Competition and Demonstration Track*, volume 133 of *PMLR*, pages 86–111.
- Antonio Moreno-Sandoval, Blanca Carbajo Coronado, Jordi Porta Zamorano, Yanco Amor Torterolo Orta, and Doaa Samy. 2025. The financial document causality detection shared task (fincausal 2025). In *Proceedings of the Joint Workshop of the 9th Financial Technology and Natural Language Processing (FinNLP), the 6th Financial Narrative Processing (FNP), and the 1st Workshop on Large Language Models for Finance and Legal (LLMFinLegal)*, pages 214–221, Abu Dhabi, UAE. Association for Computational Linguistics.
- Antonio Moreno-Sandoval, Jordi Porta, Yanco Torterolo, Alexia Stanescu, Melina Chatzi, and Sofía Roseti. 2026a. The Financial Document Causality Detection Shared Task (FinCausal 2026). In *Proceedings of the 7th Financial Narrative Processing Workshop (FNP 2026) at LREC 2026*, Palma de Mallorca, Spain. ELRA.
- Antonio Moreno-Sandoval, Yanco Amor Torterolo Orta, Maria Alexia Stanescu, and Melina Chatzi. 2026b. [The Financial Document Causality Detection Shared Task \(FinCausal 2026\): Dataset](#).
- Antonio Moreno-Sandoval, Jordi Porta Zamorano, Blanca Carbajo-Coronado, Doaa Samy, Dominique Mariko, and Mahmoud El-Haj. 2023. The financial document causality detection shared task (fincausal 2023). In *Proceedings of the 2023 IEEE International Conference on Big Data (Big-Data 2023)*, pages 2855–2860, Sorrento, Italy. IEEE.
- Jonathan Pilault, Raymond Li, Sandeep Subramanian, and Chris Pal. 2020. On extractive and abstractive neural document summarization with transformer language models. In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP 2020)*, pages 9308–9319. Association for Computational Linguistics.
- Avinash Trivedi, Gauri Toshniwal, Sivanesan Sangeetha, and SR Balasundaram. 2025.

Sarang at fincausal 2025: Contextual qa for financial causality detection combining extractive and generative models. In *Proceedings of the Joint Workshop of the 9th Financial Technology and Natural Language Processing (FinNLP), the 6th Financial Narrative Processing (FNP), and the 1st Workshop on Large Language Models for Finance and Legal (LLMFinLegal)*, pages 242–247.

Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Fei Xia, Ed Chi, Quoc V Le, Denny Zhou, et al. 2022. Chain-of-thought prompting elicits reasoning in large language models. *Advances in neural information processing systems*, 35:24824–24837.

QRAFT: QLoRA Retrieval-Augmented Fine-Tuning for Causal Span Extraction in Financial Documents

Bavya Sarda¹, Pulkit Chatwal², Sonal Dabra³

¹Galgotias University, Greater Noida, India

²Rajiv Gandhi Institute of Petroleum Technology, India

³Sony Research, India

{bhavyasarda19, pulkitchatwal, sonaldabral26}@gmail.com

Abstract

Understanding *why* financial outcomes occur is as important as knowing *what* they are. Annual reports and regulatory filings are rich with causal reasoning, yet extracting that reasoning automatically remains a difficult problem — one that sits at the intersection of domain expertise, linguistic nuance, and machine comprehension. In this paper, we describe our participation in the English subtask of the Financial Document Causality Detection shared task, FinCausal 2026, where systems are asked to identify verbatim causal spans from financial paragraphs in response to abstractive causal questions. Our approach is grounded in the intuition that a small, well-adapted model with the right inductive biases can outperform a larger but unfocused one. We fine-tune Qwen3-4B-Instruct-2507 on 2,000 domain-annotated instances using QLoRA, a parameter-efficient technique that enables meaningful adaptation under modest computational resources. Before training, we reformat all instances into the Qwen ChatML instruction template to align the model’s generation behaviour with the verbatim extraction requirement of the task. At inference time, we further guide the model by retrieving the most causally relevant sentence from the context using TF-IDF cosine similarity, providing an explicit local signal before generation. Outputs are produced via greedy decoding to ensure deterministic, source-grounded predictions. Under the official LLM-as-a-judge evaluation framework — which scores responses on a 1–5 adequacy scale based on semantic correctness rather than lexical overlap — our system achieves a score of **4.76 out of 5**, placing **4th out of nine teams** on the English leaderboard. Our results suggest that combining instruction-tuned fine-tuning with lightweight retrieval is a practical and effective strategy for causal reasoning in specialised financial text.

Keywords: causal question answering, financial NLP, QLoRA, parameter-efficient fine-tuning, TF-IDF retrieval, span extraction, LLM-as-a-judge

1. Introduction

Financial documents such as earnings reports, annual filings, and regulatory disclosures do more than report numbers — they explain *why* those numbers changed. Understanding the causal reasoning embedded in these texts is fundamental to building systems that support financial analysis, risk assessment, and automated report generation. However, causality in financial language is rarely straightforward: causal relationships often span multiple sentences, involve compound contributing factors, and are expressed without explicit connectives such as *because* or *therefore* (Cheng et al., 2024; Girju, 2003).

The Financial Document Causality Detection shared task (FinCausal) (Mariko et al., 2020, 2022; Moreno-Sandoval et al., 2023, 2025) directly addresses this challenge by evaluating systems on their ability to identify cause-and-effect relationships in financial text across English and Spanish. The 2026 edition introduces three significant advances over prior years: a revised dataset with richer and more complex causal annotations, the inclusion of multi-hop causal chains involving three or more events, and a new evaluation framework in which an LLM judge scores system responses on a 1–5 adequacy scale (Zheng et al., 2023), re-

placing the earlier exact-match and similarity-based metrics.

We participate in the English subtask and frame it as a **Causal Question Answering (CQA)** problem, where a system must extract the verbatim causal span from a financial paragraph in response to an abstractive question. Our approach consists of three components: (1) reformatting the training data into the Qwen ChatML instruction format, (2) fine-tuning Qwen3-4B-Instruct-2507 using QLoRA for parameter-efficient domain adaptation, and (3) an inference pipeline that combines intra-context TF-IDF retrieval with constrained greedy decoding to enforce verbatim extraction and reduce hallucination.

Our system achieves a score of **4.76 out of 5** on the official test set, demonstrating that a carefully fine-tuned compact language model, paired with a lightweight retrieval signal, can effectively identify causal relationships in complex financial text.

2. Related Work

2.1. Causality Detection in NLP

Early approaches to causality detection relied on lexical cues such as *because*, *therefore*, and *as a result* (Girju, 2003; Blanco et al., 2008), but

struggled with implicit causal expressions. Subsequent work introduced machine learning methods — SVMs, CNNs, and RNNs — that learned patterns directly from data (Do et al., 2011). More recently, transformer-based models such as BERT have become the dominant approach due to their contextual understanding (Cheng et al., 2024). Recent studies have also explored the use of prompt engineering with large language models to enhance causal relationship detection, particularly in domain-specific settings such as finance (Chatwal et al., 2025). Our work extends this line by fine-tuning the Qwen model on financial causal data.

2.2. Causal Question Answering

Extractive QA benchmarks such as SQuAD require answer spans to be identified directly within a context, but do not emphasise causal reasoning. Tasks like COPA (Gordon et al., 2012) target commonsense causality, while FinCausal (Mariko et al., 2020, 2022; Moreno-Sandoval et al., 2023, 2025) focuses specifically on financial text. FinCausal 2026 raises the difficulty further by combining abstractive questions with extractive answers and introducing multi-hop causal chains, requiring models to reason rather than match.

2.3. Financial NLP

Financial text presents unique challenges for general NLP tools. Loughran and McDonald (2011) demonstrated that standard sentiment lexicons perform poorly on financial language, motivating domain-specific models such as FinBERT (Araci, 2019). Causality extraction in financial reports has gained traction through the FinCausal shared task series, which our work directly builds upon.

2.4. Retrieval-Augmented Generation and Parameter-Efficient Fine-Tuning

Retrieval-Augmented Generation (RAG) (Lewis et al., 2020)(Didwania et al., 2024) has shown that grounding generative models with retrieved context improves factual accuracy. We adopt a lightweight variant of this idea, using TF-IDF cosine similarity to surface the most causally relevant sentences before generation. For efficient adaptation, we employ QLoRA (Dettmers et al., 2023), which enables fine-tuning of large language models under 4-bit quantization with minimal performance degradation. Finally, FinCausal 2026 adopts the LLM-as-a-judge evaluation framework (Zheng et al., 2023), which scores responses on semantic adequacy rather than lexical overlap — better reflecting the quality of causal reasoning.

3. Problem Statement

Financial narratives describe measurable changes in economic indicators and corporate performance, yet numerical values alone rarely explain *why* such changes occur. The Financial Document Causality Detection task, FinCausal 2026 (mor; Uniyal et al., 2021), addresses this gap by targeting **text-internal causal relationships** within financial documents. We formalize this as a **Causal Question Answering (CQA)** problem, where each instance is structured as a triplet (C, Q, A) : a financial paragraph $C = \{w_1, w_2, \dots, w_n\}$ serving as context, an abstractive causal question Q , and an extractive answer span $A \subseteq C$ drawn verbatim from the text.

A causal relation is defined as an ordered pair (e_c, e_e) , where e_c is the causal event and e_e is the resulting event, such that $e_c \rightarrow e_e$. The task focuses strictly on how causality is *encoded within the document*, not on the real-world validity of the stated relationships.

3.1. Learning Objective

Given (C, Q) , the goal is to learn a function $f_\theta : (C, Q) \rightarrow \hat{A}$, where $\hat{A} = C[i : j]$ is a contiguous extractive span. Model parameters are optimized by minimizing the negative log-likelihood:

$$\mathcal{L}(\theta) = - \sum_{k=1}^N \log P_\theta(A_k | C_k, Q_k) \quad (1)$$

Despite the extractive constraint, the task is posed in a **generative QA format** to handle complex multi-hop causal chains of the form $e_1 \rightarrow e_2 \rightarrow \dots \rightarrow e_k$, which are a key feature of the 2026 edition.

3.2. Evaluation

FinCausal 2026 replaces the earlier SAS + Exact Match scheme with an **LLM-as-a-judge** framework, where a judge model scores each response on a 1–5 adequacy scale. The overall system score is:

$$\text{Score} = \frac{1}{N} \sum_{k=1}^N J(A_k, \hat{A}_k) \quad (2)$$

This prioritizes **semantic adequacy and reasoning correctness** over strict lexical overlap.

4. Dataset

The FinCausal 2026 English dataset (Moreno-Sandoval et al., 2026) is sourced from UK financial annual reports (2017) compiled by UCREL at Lancaster University, supplemented with excerpts from the 2018 FinT-esp corpus. Relative to prior editions,

the dataset has been substantially revised: ambiguous and trivial instances were removed, and over 500 new fragments featuring complex causal structures — including chains of three or more events — were added. Approximately 10% of questions were rephrased to demand deeper reasoning beyond surface-level lexical matching. Training and test splits were constructed via random partitioning, ensuring uniform distribution of complex examples across both sets. We participated in the **English subtask**; Table 1 reports the split statistics.

Split	Language	Instances
Train	English	2,000
Test	English	500
Total		2,500

Table 1: Dataset statistics for the English subtask of FinCausal 2026.

4.1. Data Format

Each instance is a triplet (C, Q, A) : a financial paragraph C as context, an abstractive causal question Q probing either the cause or the effect, and an extractive answer span A copied verbatim from C . The dataset is provided in CSV format with four fields: *ID*, *context*, *question*, and *answer* — where the answer field is withheld in the test set.

5. Methodology

Our approach to the FinCausal 2026 English subtask consists of two main stages: supervised fine-tuning of a compact instruction-tuned language model, and a retrieval-augmented inference pipeline designed to enforce verbatim causal span extraction. Algorithm 1 provides a high-level overview of the full system.

5.1. Base Model

We select **Qwen3-4B-Instruct-2507** (Team, 2025) as our base model. This model offers a strong balance between parameter efficiency and instruction-following capability, making it well-suited for a constrained extractive QA task on domain-specific financial text. Its compact size also allows fine-tuning and inference within limited computational budgets using quantization.

5.2. Input Formatting

Before fine-tuning, all training instances are reformatted into the **Qwen ChatML** template, which

Algorithm 1 FinCausal 2026 System Pipeline

Require: Context C , Question Q , Fine-tuned model f_θ

Ensure: Predicted causal span $\hat{A}$

- 1: // — **Training Stage** —
- 2: Reformat all (C, Q, A) instances into Qwen ChatML format
- 3: Load base Qwen3-4B-Instruct-2507 in 4-bit quantized mode
- 4: Attach QLoRA adapters to attention and feed-forward projections
- 5: Optimize θ by minimizing $\mathcal{L}(\theta) = -\sum_{k=1}^N \log P_\theta(A_k | C_k, Q_k)$
- 6: // — **Inference Stage** —
- 7: Load fine-tuned f_θ in 4-bit quantized mode
- 8: Tokenize C into sentences $\{s_1, s_2, \dots, s_m\}$
- 9: Compute TF-IDF vectors for each s_i and Q
- 10: $s^* \leftarrow \arg \max_{s_i} \cos(\text{TF-IDF}(s_i), \text{TF-IDF}(Q))$
- 11: Construct prompt using full C , retrieved s^* , and Q
- 12: Generate $\hat{A} \leftarrow f_\theta(C, s^*, Q)$ using greedy decoding
- 13: **return** $\hat{A}$

structures each example as a multi-turn conversation with explicit role markers. Each instance is formatted as follows:

Prompt Template

<|im_start|>system

You are a financial question answering assistant. Given a financial text, extract the answer to the causal question directly and **verbatim** from the context. Do not generate any answer that is not present in the text.

<|im_start|>user

Context: $[C_i]$

Question: $[Q_i]$

<|im_start|>assistant

$[A_i]$

<|im_end|>

This formatting aligns the training distribution with the model’s pre-trained instruction-following behaviour, ensuring that the model learns to respond within the expected conversational structure rather than treating the task as raw text completion.

5.3. QLoRA Fine-Tuning

We fine-tune the model using **QLoRA** (Quantized Low-Rank Adaptation) (Dettmers et al., 2023), which combines 4-bit quantization of the base model weights with low-rank adapter layers inserted into the transformer architecture. This significantly reduces GPU memory requirements while preserv-

ing the model’s ability to adapt to the target domain.

Low-rank adapters are injected into all major projection layers of the transformer, including the attention projections (q, k, v, o) and the feed-forward projections (gate_proj, up_proj, down_proj). Targeting the feed-forward layers in addition to the attention layers has been shown to substantially improve task-specific adaptation (Dettmers et al., 2023). Table 2 summarises the LoRA configuration used.

Hyperparameter	Value
LoRA rank (r)	32
LoRA alpha (α)	64
LoRA dropout	0.05
Bias	none
Task type	Causal LM
Target modules	q, k, v, o, gate, up, down proj
Quantization	4-bit (NF4)

Table 2: QLoRA fine-tuning configuration.

The rank $r = 32$ provides sufficient adapter capacity for capturing financial domain patterns, while $\alpha = 64$ (set to $2r$ following standard practice) controls the scaling of the adapter updates. A dropout of 0.05 is applied to the adapter layers as a lightweight regularisation measure.

5.4. Inference Pipeline

At inference time, we apply a four-step pipeline designed to minimise hallucination and enforce verbatim extraction from the source context.

Step 1 — Model Loading. The fine-tuned model is loaded in **4-bit quantized mode** using BitsAndBytes, reducing memory footprint while maintaining generation quality.

Step 2 — Intra-Context Retrieval. Rather than passing the full context blindly to the model, we apply an **intra-context retrieval** step. Each sentence in the context C is ranked against the question Q using **TF-IDF cosine similarity**. The top-ranked sentence s^* is selected as the most relevant causal evidence:

$$s^* = \arg \max_{s_i \in C} \cos(\text{TF-IDF}(s_i), \text{TF-IDF}(Q)) \quad (3)$$

This retrieved sentence is appended to the prompt alongside the full context, providing the model with an explicit signal about where the causal span is likely to reside.

Step 3 — Constrained Prompt. The model is prompted using a **verbatim extraction format** that explicitly instructs the model to copy the answer word-for-word from the context. Both the full context and the retrieved relevant sentence are included in the prompt, reducing the risk of paraphrased or hallucinated responses.

Step 4 — Greedy Decoding. Generation is performed using **greedy decoding** with temperature set to zero, ensuring fully deterministic outputs. This choice is deliberate: since the task requires precise verbatim spans, stochastic sampling strategies such as top- p or beam search with diversity penalties are counterproductive. Greedy decoding directly maximises the probability of the most likely token at each step, producing stable and reproducible predictions aligned with the source text.

6. Results

6.1. Main Result

We evaluate our system on the FinCausal 2026 English test set comprising 500 instances, using the official LLM-as-a-judge metric. Our system achieves a score of **4.76 out of 5**, ranking **4th** on the official English subtask leaderboard out of nine participating teams. Unlike Exact Match, the LLM judge tolerates minor boundary differences provided the causal meaning is preserved, making it a more faithful measure of reasoning quality in financial text. Table 3 reports the full leaderboard standings.

Rank	Team	LLM Score
1	HSA_CORAL	4.814
1	Sheffield_Causal	4.814
3	Tredence_AICOE	4.812
4	Lab Rats (Ours)	4.760
5	CariMed	4.720
6	EMI	4.704
7	LeedsMeng26	4.700
8	TU Graz Data Team	4.662
9	Sarang	4.540

Table 3: Official English subtask leaderboard for FinCausal 2026.

6.2. Component Analysis

Three design choices contribute to the strong performance of our system.

QLoRA Fine-Tuning. Training on 2,000 expert-annotated instances adapts the Qwen model to the

vocabulary, syntax, and causal patterns of financial reporting. Without this step, the base model lacks the domain grounding needed to distinguish causal spans from surrounding narrative text.

ChatML Instruction Formatting. Structuring each training instance as a ChatML instruction explicitly conditions the model to extract answers verbatim from the context rather than paraphrase or hallucinate plausible-sounding responses.

TF-IDF Intra-Context Retrieval. Ranking sentences by cosine similarity to the question before generation provides the model with an explicit relevance signal. This is particularly effective for instances where the question and causal span share limited lexical overlap, a common characteristic of abstractive causal questions in the 2026 dataset.

7. Limitations

Despite strong overall performance, our system faces two notable limitations. First, predicting long answer spans — particularly those covering multiple clauses — remains challenging, as small boundary errors can reduce semantic fidelity. Second, multi-hop causal chains involving three or more events are difficult to resolve through span extraction alone, as they require discourse-level reasoning that goes beyond identifying a single contiguous passage. Addressing these limitations likely requires models with explicit coreference and discourse structure awareness, rather than relying solely on local retrieval signals.

8. Conclusion

We presented a system for the FinCausal 2026 English subtask that combines parameter-efficient fine-tuning with a lightweight retrieval-augmented inference pipeline. By reformatting training data into the Qwen ChatML instruction format, fine-tuning Qwen3-4B-Instruct-2507 via QLoRA, and augmenting inference with TF-IDF intra-context retrieval and greedy decoding, our system achieves a score of **4.76 out of 5** on the official evaluation.

Our results demonstrate three broader findings. First, compact language models fine-tuned on modest domain-specific datasets can achieve strong performance on financial causal QA when paired with appropriate instruction formatting. Second, lightweight retrieval signals such as TF-IDF remain effective for grounding generation even without dense retrieval infrastructure. Third, the LLM-as-a-judge evaluation framework provides a more meaningful signal than lexical overlap metrics for tasks requiring causal reasoning, rewarding semantic correctness over verbatim matching.

Future work should explore discourse-aware models capable of resolving multi-hop causal chains, as well as dense retrieval methods that better capture semantic similarity between abstractive questions and their corresponding causal spans in financial text.

9. Generative AI Use Disclosure

Generative AI tools were used solely for language editing and paraphrasing during the preparation of this manuscript, including grammar correction and improving the clarity of written expressions. No generative AI tool was used to produce any scientific content, experimental results, analysis, or conclusions presented in this work. All authors are fully responsible and accountable for the content of this paper.

10. Acknowledgements

This research was conducted independently by the authors outside the scope of their professional responsibilities at Sony Research. The views and findings presented in this paper are solely those of the authors and do not reflect the positions or policies of Sony Research.

11. Bibliographical References

- Dogu Araci. 2019. Finbert: Financial sentiment analysis with pre-trained language models. *arXiv preprint arXiv:1908.10063*.
- Eduardo Blanco, Nuria Castell, and Dan I Moldovan. 2008. Causal relation extraction. In *Lrec*, volume 66, page 74.
- Pulkit Chatwal, Amit Agarwal, and Ankush Mittal. 2025. Enhancing causal relationship detection using prompt engineering and large language models. In *Proceedings of the Joint Workshop of the 9th Financial Technology and Natural Language Processing (FinNLP), the 6th Financial Narrative Processing (FNP), and the 1st Workshop on Large Language Models for Finance and Legal (LLMFinLegal)*, pages 248–252.
- Qing Cheng, Zefan Zeng, Xingchen Hu, Yuehang Si, and Zhong Liu. 2024. A survey of event causality identification: Taxonomy, challenges, assessment, and prospects. *arXiv preprint arXiv:2411.10371*.

- Tim Dettmers, Artidoro Pagnoni, Ari Holtzman, and Luke Zettlemoyer. 2023. Qlora: Efficient fine-tuning of quantized llms. In *Advances in Neural Information Processing Systems (NeurIPS)*.
- Krish Didwania, Pratinav Seth, Aditya Kasliwal, and Amit Agarwal. 2024. Agrillm: harnessing transformers for framer queries. In *Proceedings of the Third Workshop on NLP for Positive Impact*, pages 179–187.
- Quang Do, Yee Seng Chan, and Dan Roth. 2011. [Minimally supervised event causality identification](#). In *Proceedings of the 2011 Conference on Empirical Methods in Natural Language Processing*, pages 294–303, Edinburgh, Scotland, UK. Association for Computational Linguistics.
- Roxana Girju. 2003. Automatic detection of causal relations for question answering. In *Proceedings of the ACL 2003 workshop on Multilingual summarization and question answering*, pages 76–83.
- Andrew Gordon, Zornitsa Kozareva, and Melissa Roemmele. 2012. [SemEval-2012 task 7: Choice of plausible alternatives: An evaluation of commonsense causal reasoning](#). In **SEM 2012: The First Joint Conference on Lexical and Computational Semantics – Volume 1: Proceedings of the main conference and the shared task, and Volume 2: Proceedings of the Sixth International Workshop on Semantic Evaluation (SemEval 2012)*, pages 394–398, Montréal, Canada. Association for Computational Linguistics.
- Patrick Lewis, Ethan Perez, Aleksandra Piktus, Fabio Petroni, Vladimir Karpukhin, Naman Goyal, Heinrich Küttler, Mike Lewis, Wen-tau Yih, Tim Rocktäschel, et al. 2020. Retrieval-augmented generation for knowledge-intensive nlp tasks. *Advances in neural information processing systems*, 33:9459–9474.
- Tim Loughran and Bill McDonald. 2011. When is a liability not a liability? textual analysis, dictionaries, and 10-ks. *The Journal of finance*, 66(1):35–65.
- Dominique Mariko, Hanna Abi-Akl, Estelle Labidurie, Stephane Durfort, Hugues De Mazancourt, and Mahmoud El-Haj. 2020. The financial document causality detection shared task (fincausal 2020). In *Proceedings of the 1st Joint Workshop on Financial Narrative Processing and MultiLing Financial Summarisation*, pages 23–32.
- Dominique Mariko, Hanna Abi-Akl, Kim Trottier, and Mahmoud El-Haj. 2022. The financial causality extraction shared task (fincausal 2022). In *Proceedings of the 4th Financial Narrative Processing Workshop@ LREC2022*, pages 105–107.
- Antonio Moreno-Sandoval, Blanca Carbajo-Coronado, Jordi Porta Zamorano, Yanco Amor Torterolo Orta, and Doaa Samy. 2025. The financial document causality detection shared task (fincausal 2025). In *Proceedings of the Joint Workshop of the 9th Financial Technology and Natural Language Processing (FinNLP), the 6th Financial Narrative Processing (FNP), and the 1st Workshop on Large Language Models for Finance and Legal (LLMFinLegal)*, pages 214–221.
- Antonio Moreno-Sandoval, Jordi Porta-Zamorano, Blanca Carbajo-Coronado, Doaa Samy, Dominique Mariko, and Mahmoud El-Haj. 2023. [The financial document causality detection shared task \(fincausal 2023\)](#). In *2023 IEEE International Conference on Big Data (BigData)*, pages 2855–2860.
- Antonio Moreno-Sandoval, Yanco Amor Torterolo Orta, Maria Alexia Stanescu, and Melina Chatzi. 2026. [The Financial Document Causality Detection Shared Task \(FinCausal 2026\): Dataset](#).
- Qwen Team. 2025. [Qwen3 technical report](#).
- Deepak Uniyal, Amit Agarwal, Durga Toshniwal, and Dipanjan Deb. 2021. Dense vector embedding based approach to identify prominent dis-seminators from twitter data amid covid-19 outbreak. *IEEE Transactions on Emerging Topics in Computational Intelligence*, 5(3):308–320.
- Lianmin Zheng, Wei-Lin Chiang, Ying Sheng, Siyuan Zhuang, Zhanghao Wu, Yonghao Zhuang, Zi Lin, Zhuohan Li, Dacheng Li, Eric Xing, et al. 2023. Judging llm-as-a-judge with mt-bench and chatbot arena. *Advances in neural information processing systems*, 36:46595–46623.

Author Index

- Alqarni, Aali Abdullah, 125
Asad, Idrees, 160
Atherly, Rachel, 78
Attak, Sanae, 152
- Bal, Yasemin, 160
Berlanga, Rafael, 139
Blanchon, Hervé, 78
- Chatwal, Pulkit, 175
Chatzi, Melina, 87, 114
Chiadmi, Mohammed Salah, 152
- Dabral, Sonal, 175
Dinarelli, Marco, 78
Dodig, Paula, 49
- Fan, Yimei, 39
- Gautam, Akash Kumar, 1, 132
Gonzalez Saez, Gabriela nicole, 78
- Hamotskyi, Serhii, 1, 132
Hänig, Christian, 1, 132
- Ivienagbor, Ayomide, 160
- Jay, Aldan, 139
- Kabra, Anubha, 39
Kern, Roman, 146
Kim, Katie Jooyoung, 39
Kivimäki, Timo, 98
Koloski, Boshko, 49
Kosai, Yurina, 106
Kou, Zhiwei, 39
Kurfali, Murathan, 28
- Laksito, Arif Dwi, 125
Lamrani Alaoui, Youssef, 152
Lasnier, Théo, 12
Liu, Mo, 28
- Mallikarjuna, Chindukuri, 169
Martinez Vidiri, Gabriel, 39
Mashiku, Melchizedek, 98
Moreno-Sandoval, Antonio, 87, 114
- Moreno, Yoelvis, 139
Mosolova, Anna, 12
Mouilleron, Virginie, 12
- Nadeem, Zahaab, 160
Nakhlé, Mariam, 78
Niess, Georg, 146
- Paredes Amorin, Alvaro, 59
Pollak, Senja, 49
Porta, Jordi, 114
Purver, Matthew, 49
Python, Andre, 59
- Qader, Raheel, 78
- Roseti, Sofia, 114
- Saefken, Benjamin, 98
Sajer, Helene, 39
Santamarta, Vicent, 139
Sarda, Bavya, 175
Schlee, Michael, 98
Seddah, Djamé, 12
Shahrouri, Zaid, 160
Shrestha, Rijul, 160
Sitar Šuštar, Katarina, 49
Stanescu, Alexia, 114
Stevenson, Mark, 125
- Tortero Orta, Yanco Amor, 87, 114
Trivedi, Avinash, 169
Tsuchida, Rikuto, 106
- Utsuro, Takehito, 106
- Weisser, Christoph, 59, 98
- Xie, Yucheng, 106
- Zhu, Tianning, 28