

RKEFino1: A Regulation Knowledge-Enhanced Large Language Model

1st Yan Wang*

Yale University

New Haven, CT, USA

yan.wang.yw937@yale.edu

2nd Yueru He

Columbia University

New York, NY, USA

yh3507@columbia.edu

3rd Ruoyu Xiang

New York University

New York, NY, USA

rx2306@nyu.edu

4th Jeff Zhao

The University of Texas at Austin

Austin, Texas, USA

jeffzhao@utexas.edu

Abstract—Recent advances in large language models (LLMs) hold great promise for financial applications but introduce critical accuracy and compliance challenges in Digital Regulatory Reporting (DRR). To address these issues, we propose RKEFino1, a regulation knowledge-enhanced financial reasoning model built upon Fino1, fine-tuned with domain knowledge from XBRL, CDM, and MOF. We formulate two QA tasks—knowledge-based and mathematical reasoning—and introduce a novel Numerical NER task covering financial entities in both sentences and tables. Experimental results demonstrate the effectiveness and generalization capacity of RKEFino1 in compliance-critical financial tasks.

Index Terms—Financial knowledge, Digital Regulatory Reporting, Large Language Model, XBRL, Fine-tune

I. INTRODUCTION

The financial industry increasingly leverages reinforcement learning (RL) techniques, giving rise to the interdisciplinary field known as Financial Reinforcement Learning (FinRL), which includes applications in portfolio management, algorithmic trading, and option pricing. Recent advancements in large language models (LLMs) have further accelerated the growth of open finance by providing scalable, affordable, and personalized solutions such as enhanced financial search and robo-advisory services.

Despite these promising developments, the application of LLMs to Digital Regulatory Reporting (DRR) introduces significant challenges due to the complexity and precision required by financial regulations. Financial reporting standards such as the eXtensible Business Reporting Language (XBRL) commonly experience high error rates, highlighting critical vulnerabilities in existing reporting mechanisms. Additionally, managing financial transactions through the Common Domain Model (CDM) demands accurate lifecycle tracking, while the Model Openness Framework (MOF) requires comprehensive transparency and completeness from machine learning models. Issues with misinformation and inaccuracies—hallucinations—that plague current LLMs pose severe risks, potentially leading to regulatory non-compliance, financial losses, and diminished trust among stakeholders.

To address these critical challenges, we introduce an RKEFino1, a regulation knowledge-enhanced Fino1 model [1], fine-tuned specifically with domain knowledge from XBRL, CDM, and MOF. This enhanced model aims to significantly improve interpretability, compliance accuracy, and reliability

in digital regulatory reporting tasks, thereby contributing effectively to market integrity and regulatory compliance.

II. RELATED WORKS

Large Language Models (LLMs) have demonstrated significant success across various tasks in the general domain, such as information extraction [2], retrieval [3], ranking [4], and summarization [5]. Inspired by these promising outcomes, recent research has progressively extended the application of LLMs into finance-specific scenarios. Initial efforts have focused on adapting existing models for financial applications, demonstrating their capabilities in financial text analysis [6], market sentiment prediction [7], and automated trading strategies [8]. Building upon these findings, recent advancements have led to the development of specialized financial LLMs, including BloombergGPT [9], FinGPT [7], PIXIU [10], OpenFinLLMs [11], and Plutus [12]. These finance-specific models leverage pretraining on financial corpora, significantly enhancing their ability to understand and accurately interpret domain-specific terminology and contexts. Empirical evaluations confirm that these adapted LLMs effectively handle complex financial tasks, such as financial market sentiment analysis, earnings call analysis, regulatory compliance monitoring, fraud detection, portfolio optimization, and risk management.

III. METHODOLOGY

In this section, we primarily focus on introducing our RKEFino1 model, defining our fine-tuning tasks, collecting training data, and setting model parameters.

A. RKEFino1: regulation knowledge-enhanced Fino1

Fino1 [1] is a lightweight financial reasoning model built upon LLaMA-3.1-8B-Instruct, designed for accurate and efficient deployment in real-world financial scenarios. It adopts a two-stage training strategy combining supervised fine-tuning and reinforcement learning, guided by curated reasoning paths. In the first stage, Fino1 learns domain-specific reasoning patterns from financial question-answer pairs, with an emphasis on step-by-step interpretability. In the second stage, the model leverages verifier-guided reinforcement learning to iteratively refine its outputs, enhancing correctness and coherence through techniques such as backtracking, verification, and correction.

This design enables Fino1 to achieve strong performance on financial mathematical reasoning tasks. However, it cannot analyze Digital Regulatory Reporting (DRR) scenarios, which require a nuanced understanding of regulatory structures and compliance logic. To address this limitation, we further fine-tune Fino1 on curated DRR-related datasets, incorporating domain-specific knowledge from financial regulations. The resulting model, RKEFino1 (Regulation Knowledge-Enhanced Fino1), is equipped with enhanced reasoning abilities tailored to regulatory reporting and compliance tasks.

B. Task formulation

To enhance regulatory knowledge in the Fino1 model, we define a series of Digital Regulatory Reporting (DRR) question-answering tasks based on CDM, MOF, and XBRL data, specifically categorized into knowledge-based QA and mathematical reasoning QA. For the knowledge-based QA task, given a regulation-related question q , covering licenses, abbreviations, terminologies, and tags. The model generates answers a by determining license applicability, expanding abbreviations, explaining terminologies, and identifying tags based on its inherent knowledge.

$$a = f(q) \quad (1)$$

For the mathematical reasoning QA task, given a mathematical question q about a financial report, a financial formulation t , and descriptions d of terms in the formulation, the model calculates a numerical answer a by leveraging its comprehensive regulatory knowledge, financial understanding, and inference capabilities. So,

$$a = f(q, t, d) \quad (2)$$

C. Data collection

We collect our training data mainly from 3 data source: the CDM documentations¹ for CDM knowledge, the Open Source Initiative (OSI) website² for MOF knowledge, and the U.S. Securities and Exchange Commission (SEC) website³ for XBRL knowledge. Moreover, we also collect extra training data from existing datasets such as XBRL terminology⁴. Therefore, as shown in Table I, we collected 9,898 training samples for our model.

D. Model training settings

We fine-tuned the Fino1 model [1] for our DRR task using supervised instruction tuning. The model was trained with a block size of 4096 tokens and a maximum context length of 8192 tokens, enabling it to handle long-form document reasoning. To enable efficient training on limited GPU memory, we applied parameter-efficient fine-tuning (PEFT) via LoRA, with rank $r = 64$, scaling factor $\alpha = 128$, and a dropout rate

¹<https://cdm.finos.org/>

²<https://opensource.org/licenses>

³<https://www.sec.gov/>

⁴https://huggingface.co/datasets/KirkHan/XBRL_Terminology

TABLE I
THE STATISTIC OF TRAINING DATA.

Task	Domain	#Samples	Source
K-QA ^a	CDM	478	CDM documentation
	MOF	258	OSI website
	XBRL	8,052	SEC and XBRL Terminology
MR-QA ^b	XBRL	1,110	SEC

^a K-QA indicates the knowledge-based QA task.

^b MR-QA denotes the mathematical reasoning QA task.

of 0.05. We enabled int4 quantization and applied LoRA to all linear modules.

The training was conducted for 10 epochs with a batch size of 1, using gradient accumulation over 4 steps to simulate a larger effective batch size. We adopted the AdamW optimizer with a learning rate of 3e-5, a cosine learning rate scheduler, and a warmup ratio of 1%. Training was performed with bf16 mixed-precision on 4 NVIDIA H100 GPUs.

IV. EXPERIMENT AND RESULT

A. Evaluation dataset

To better evaluate our RKEFino1 model, we used evaluation data from the regulation challenge of FinNLP-FNP-LLMFinLegal-2025 shared task⁵ [13], with the statistical details summarized in Table II. For the Knowledge-based QA (K QA) task, we constructed 150 test samples, and for the Mathematical Reasoning QA (MR QA) task, we collected 50 test samples. To assess the model’s generalizability, we introduced a new task called Numerical NER, inspired by FiNER [14] and FNXL [15]. Unlike previous work, our Numerical NER task focuses on identifying five entity types (Integer Item Type, Monetary Item Type, Per Share Item Type, Percent Item Type, and Shares Item Type) instead of us gaap tags, and it processes both sentences and tables rather than sentences only.

TABLE II
THE STATISTIC OF EVALUATION DATA.

Task	Domain	#Samples	Source
K-QA	CDM	126	CDM documentation
	MOF	161	OSI website
	XBRL	700	SEC and XBRL Terminology
MR-QA	XBRL	1,000	SEC
Numerical NER	XBRL	3,638	SEC

B. Evaluation metrics

To keep consistent with the competition, we also leveraged Accuracy and FactScore to evaluate our model’s performance on K-QA and MR-QA tasks. Furthermore, we used F1-scores to evaluate performance on the numerical NER task.

- Accuracy is primarily used for questions that require precise answers, such as full expansions of abbreviations,

⁵https://github.com/Open-Finance-Lab/Regulations_Challenge_COLING_2025

TABLE III
COMPARISON OF FINO1 AND RKEFINO1 ON FINE-GRAINED TASKS.

Task	Fine-grained Task	Metrics	Fino1	RKEFino1
K-QA	CDM-QA	FActScore (%)	36.76	42.58
	MOF Abbreviation	Accuracy (%)	0.00	12.23
	MOF Approval	Accuracy (%)	0.00	62.58
	MOF Details	FActScore (%)	27.13	40.56
	XBRL Domain	FActScore (%)	20.08	45.87
	XBRL Tag	Accuracy (%)	0.00	16.02
	XBRL Term	FActScore (%)	26.22	50.28
MR-QA	XBRL Math	Accuracy (%)	56.87	70.69
NER	Numerical NER	F1-score (%)	14.99	26.62

yes-or-no questions, and financial mathematical reasoning.

- FactScore is mainly used for question-answering scenarios, including CDM documentation, MOF detailed QA, and explanations of XBRL terms.
- F1-scores, commonly used in named entity recognition, is the harmonic mean of precision and recall, reflecting overall performance and generalization.

C. Initial results

To quantify the effectiveness of domain knowledge integration, we conduct a detailed comparison between the baseline Fino1 model and its enhanced variant RKEFino1, which is trained with additional knowledge. The results are presented in Table III.

RKEFino1 significantly outperforms Fino1 across all three evaluation tasks. In knowledge-based QA, it achieves much higher accuracy on MOF Approval (62.58% vs. 0.00%) and XBRL Tag (16.02% vs. 0.00%), and shows notable FactScore gains on CDM-QA, MOF Details, and XBRL Term, highlighting the benefits of knowledge injection for factual correctness and relevance.

In the MR-QA task, which involves mathematical reasoning over financial data, RKEFino1 achieves 70.69% accuracy, clearly surpassing Fino1’s 56.87%, suggesting better grounding through domain knowledge.

For the newly proposed Numerical NER task, RKEFino1 reaches an F1-score of 26.62%, compared to Fino1’s 14.99%, demonstrating improved generalization and structured information extraction across both textual and tabular inputs.

V. CONCLUSION AND FUTURE WORK

We present RKEFino1, an enhanced version of Fino1 that integrates regulatory knowledge to better address challenges in DRR. Fine-tuned on curated data from XBRL, CDM, and MOF, it performs domain-specific reasoning via supervised learning guided by a verifier. To evaluate its capabilities, we design two QA tasks—one on regulatory knowledge, the other on mathematical reasoning—and a numerical NER task to test generalization across regulatory data. Experiments show RKEFino1 significantly outperforms Fino1 in all DRR tasks and excels at recognizing numerical entities. In future work, we plan to expand the dataset for MOF abbreviation and

XBRL tag to further improve the model’s performance on this task.

REFERENCES

- [1] L. Qian, W. Zhou, Y. Wang, X. Peng, J. Huang, and Q. Xie, “Fino1: On the transferability of reasoning enhanced llms to finance,” 2025.
- [2] V. K. Keloth, Y. Hu, Q. Xie, X. Peng, Y. Wang, A. Zheng, M. Selek, K. Raja, C. H. Wei, Q. Jin *et al.*, “Advancing entity recognition in biomedicine via instruction tuning of large language models,” *Bioinformatics*, vol. 40, no. 4, p. btae163, 2024.
- [3] Y. Wang, J. Huang, H. He, V. Zhang, Y. Zhou, X. Hao, P. Ram, L. Qian, Q. Xie, R.-L. Weng *et al.*, “Cdemapper: Enhancing nih common data element normalization using large language models,” *arXiv preprint arXiv:2412.00491*, 2024.
- [4] Y. Wang, L. Qian, X. Peng, J. Huang, and D. Feng, “Ordrankben: A novel ranking benchmark for ordinal relevance in nlp,” *arXiv preprint arXiv:2503.00674*, 2025.
- [5] Y. Zhang, H. Jin, D. Meng, J. Wang, and J. Tan, “A comprehensive survey on process-oriented automatic text summarization with exploration of llm-based methods,” *arXiv preprint arXiv:2403.02901*, 2024.
- [6] Y. Yang *et al.*, “Finbert: A pre-trained financial language representation model for financial text mining,” *International Joint Conference on Artificial Intelligence*, 2020.
- [7] H. Yang, X.-Y. Liu, and C. D. Wang, “Fingpt: Open-source financial large language models,” *arXiv preprint arXiv:2306.06031*, 2023.
- [8] Z. Kou, H. Yu, J. Peng, and L. Chen, “Automate strategy finding with llm in quant investment,” *arXiv preprint arXiv:2409.06289*, 2024.
- [9] S. Wu *et al.*, “Bloomberggpt: A large language model for finance,” *arXiv preprint arXiv:2303.17564*, 2023.
- [10] Q. Xie, W. Han, X. Zhang, Y. Lai, M. Peng, A. Lopez-Lira, and J. Huang, “Pixiu: a large language model, instruction data and evaluation benchmark for finance,” in *Proceedings of the 37th International Conference on Neural Information Processing Systems*, 2023, pp. 33 469–33 484.
- [11] Q. Xie, D. Li, M. Xiao, Z. Jiang, R. Xiang, X. Zhang, Z. Chen, Y. He, W. Han, Y. Yang *et al.*, “Open-finllms: Open multimodal large language models for financial applications,” *arXiv preprint arXiv:2408.11878*, 2024.
- [12] X. Peng, T. Papadopoulos, E. Soufleri, P. Giannouris, R. Xiang, Y. Wang, L. Qian, J. Huang, Q. Xie, and S. Ananiadou, “Plutus: Benchmarking large language models in low-resource greek finance,” 2025.
- [13] K. Wang, J. Patel, C. Shen, D. Kim, A. Zhu, A. Lin, L. Borella, C. Osborne, M. White, S. Yang *et al.*, “Finnlp-fnp-llmfinlegal-2025 shared task: Regulations challenge,” in *Proceedings of the Joint Workshop of the 9th FinNLP, the 6th Financial Narrative Processing (FNP), and the 1st Workshop on Large Language Models for Finance and Legal (LLMFinLegal)*, 2025, pp. 363–370.
- [14] L. Loukas, M. Fergadiotis, I. Chalkidis, E. Spyropoulou, P. Malakasiotis, I. Androutopoulos, and P. George, “Finer: Financial numeric entity recognition for xbrl tagging,” in *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (ACL 2022)*, 2022.
- [15] S. Sharma, S. Khatuya, M. Hegde, A. Shaikh, K. Dasgupta, P. Goyal, and N. Ganguly, “Financial numeric extreme labelling: A dataset and benchmarking,” in *Findings of the Association for Computational Linguistics: ACL 2023*, 2023, pp. 3550–3561.