

Table-R1: Inference-Time Scaling for Table Reasoning

Zheyuan Yang^{*1} Lyuhao Chen^{*2} Arman Cohan¹ Yilun Zhao¹

¹Yale NLP Lab ²Carnegie Mellon University

Abstract

In this work, we present the first study to explore inference-time scaling on table reasoning tasks. We develop and evaluate two post-training strategies to enable inference-time scaling: distillation from frontier model reasoning traces and reinforcement learning with verifiable rewards (RLVR). For distillation, we introduce a large-scale dataset of reasoning traces generated by DeepSeek-R1, which we use to fine-tune LLMs into the Table-R1-SFT model. For RLVR, we propose task-specific verifiable reward functions and apply the GRPO algorithm to obtain the Table-R1-Zero model. We evaluate our Table-R1-series models across diverse table reasoning tasks, including short-form QA, fact verification, and free-form QA. Notably, the Table-R1-Zero model matches or exceeds the performance of GPT-4.1 and DeepSeek-R1, while using only a 7B-parameter LLM. It also demonstrates strong generalization to out-of-domain datasets. Extensive ablation and qualitative analyses reveal the benefits of instruction tuning, model architecture choices, and cross-task generalization, as well as emergence of essential table reasoning skills during RL training.

 **Model** huggingface.co/Table-R1
 **Code** github.com/Table-R1

1 Introduction

Reasoning large language models, such as OpenAI’s o-series (Jaech et al., 2024; Pfister and Jud, 2025) and Deepseek’s R1 (Guo et al., 2025), have demonstrated enhanced reasoning capabilities by inference-time scaling, *i.e.*, generating a reasoning chain of tokens that allow the model to “think” before giving the final answer. Building on this success, recent research has extended inference-time scaling to various domains and tasks, including multimodal reasoning (Huang et al., 2025a;

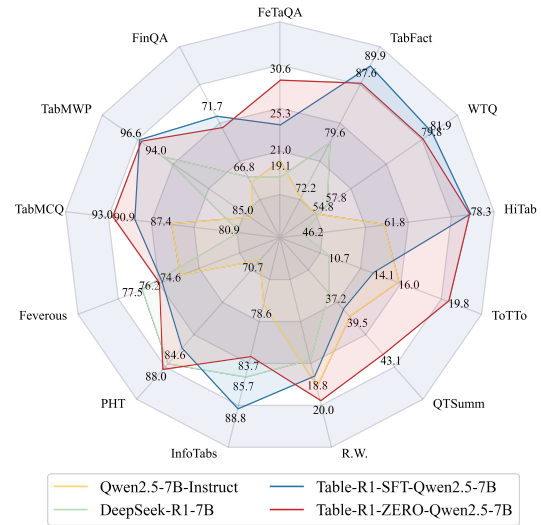


Figure 1: Overall performance comparison between Table-R1 and same-scale baselines on various table reasoning benchmarks. Both Table-R1-SFT and Table-R1-Zero exhibit substantial performance improvements over baselines, showing the effectiveness of our approach across both in- and out-of-domain benchmarks.

Xu et al., 2025), machine translation (Feng et al., 2025b), agent-based tool use (Ouyang et al., 2025; Jin et al., 2025), and information retrieval (Weller et al., 2025; Zhuang et al., 2025).

However, applying inference-time scaling to structure-dependent tasks—particularly table reasoning—remains largely unexplored. Table reasoning presents distinct challenges compared to text-only tasks: it requires interpreting diverse cell contents, aligning data across the table, and performing multi-step reasoning with aggregation and numerical operations (Deng et al., 2024; Wu et al., 2025). These requirements are further complicated by the need to process long and densely structured tabular inputs (Zhao et al., 2023c; Nahid and Rafiei, 2024; Zhang et al., 2025b). Advancing LLMs’ reasoning capabilities over tabular tasks holds significant promise for real-world applications, including

^{*}Equal Contributions.

data analysis (Zhao et al., 2024c), scientific reporting (Liang et al., 2024; Newman et al., 2024), and decision-support systems (Handler et al., 2024).

In this work, we present the first study to explore inference-time scaling on table reasoning tasks. Figure 2 presents the overview of our research. We develop and systematically evaluate two widely used post-training strategies to enable inference-time scaling on table reasoning tasks: (1) distilling from reasoning traces of frontier reasoning models, and (2) reinforcement learning with verifiable rewards (RLVR). For the distillation approach, we curate and open-source a large-scale table reasoning dataset containing reasoning traces generated by DeepSeek-R1 and verified by LLM-based annotators. We fine-tune LLMs on this data to obtain Table-R1-SFT. For the RLVR approach, we design task-specific, verifiable reward functions tailored to table reasoning and apply the Group Relative Policy Optimization (GRPO) algorithm (Shao et al., 2024; Guo et al., 2025) to enable stable and scalable reinforcement learning. This yields the Table-R1-Zero model.

We evaluate Table-R1-series models on a wide range of table reasoning tasks, including short-form table QA, fact verification, and free-form table QA. Our experiments demonstrate the effectiveness of inference-time scaling for table reasoning. The RLVR approach, in particular, exhibits better performance and generalization capabilities, compared to the distillation approach. Notably, our Table-R1-Zero models achieve performance that is competitive with advanced language models such as GPT 4.1 and DeepSeek R1, despite using only a 7B-parameter LLM (*i.e.*, Qwen2.5-7B) as the backbone. We further conduct comprehensive ablation studies on *instruction tuning benefits*, *model family comparisons*, and *cross-task generalization*, providing insights for future applications of inference-time scaling in table reasoning. Our qualitative analysis of model responses reveals that Table-R1-Zero not only acquires multi-step reasoning and reflection abilities like other reasoning models, but also develops essential table-specific reasoning skills such as semantic understanding, information extraction, and arithmetic computation.

2 Related Work

2.1 Inference-Time Scaling

Recently, OpenAI’s o1 has demonstrated that scaling inference-time computation can significantly

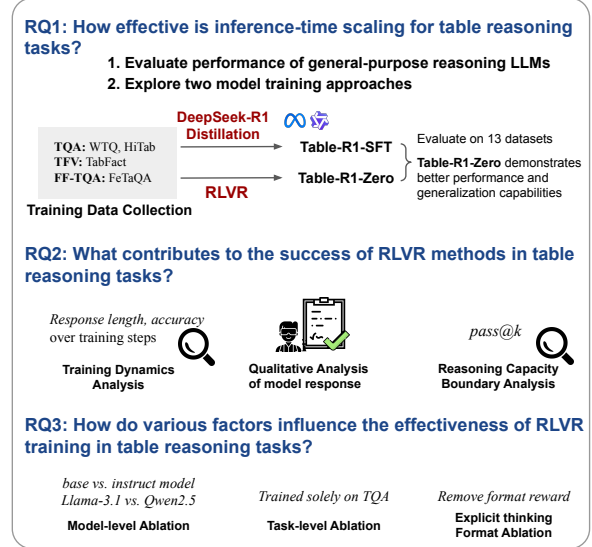


Figure 2: An overview of our research and three research questions investigated in this study.

enhance the reasoning abilities of large language models (LLMs) on complex tasks (Jaech et al., 2024). To leverage this, various inference-time strategies have been explored, including the use of Monte Carlo Tree Search (MCTS) for exploring diverse reasoning trajectories (Feng et al., 2023; Qi et al., 2024; Guan et al., 2025) and process reward models (PRMs) that offer step-level feedback to guide model outputs (Lightman et al., 2023; Yuan et al., 2024). In parallel, supervised fine-tuning (SFT) on reasoning traces has emerged as a practical post-training method, enabling LLMs to better align generation with explicit chain-of-thought reasoning patterns (Wen et al., 2025; Muennighoff et al., 2025; Ye et al., 2025).

Beyond supervised approaches, recent work has introduced reinforcement learning from verifiable rewards (RLVR) as a promising post-training paradigm for LLM reasoning (Guo et al., 2025; Team et al., 2025; Team, 2025). In this setting, models are directly optimized with rule-based rewards, allowing them to autonomously discover effective reasoning strategies without explicit intermediate supervision. Subsequent studies have improved RLVR training by incorporating dynamic sampling, token-level policy gradients, and reward normalization to enhance training stability and sample efficiency (Yu et al., 2025b; Liu et al., 2025; Xia et al., 2025). The RLVR paradigm has demonstrated strong generalization across diverse domains, including mathematical problem solving (Hu et al., 2025; Face, 2025), logical reason-

ing games (Xie et al., 2025), vision-based reasonings (Huang et al., 2025b), and interactive agent scenarios (Wang et al., 2025; Xia and Luo, 2025; Feng et al., 2025a).

2.2 Table Reasoning

Reasoning over tabular data has long attracted attention due to its practical applications in real-world scenarios such as data analysis. It encompasses a variety of tasks, including short-form question answering (Pasupat and Liang, 2015; Cheng et al., 2022; Lu et al., 2023; Zhao et al., 2023d, 2024b; Wu et al., 2025), fact verification (Chen et al., 2020; Gupta et al., 2020), and free-form question answering (Nan et al., 2022; Zhao et al., 2023a). Earlier research primarily focused on fine-tuning smaller language models for specific tasks (Herzig et al., 2020; Zhao et al., 2022; Liu et al., 2022). Recent efforts have advanced the adaptation of LLMs for table reasoning (Zhang et al., 2024a, 2025a; Su et al., 2024; Zha et al., 2023; Su et al., 2024; Deng et al., 2025), enabling more general-purpose capabilities across task types. In parallel, several studies have begun exploring agentic approaches to tackle table reasoning tasks (Ye et al., 2023; Zhao et al., 2024a; Nan et al., 2024; Yu et al., 2025a). Despite these advances, the application and enhancement of inference-time scaling for table reasoning remain largely underexplored. Our study shows that 7B-scale LLMs with inference-time scaling can match the performance of frontier models such as GPT-4.1. This finding highlights a promising direction for advancing table reasoning.

3 Table-R1 Models

To systematically explore inference scaling in table reasoning tasks, we develop two variants of Table-R1 model, each leveraging a widely adopted post-training strategy for inference-time scaling: (1) Table-R1-SFT, trained via supervised fine-tuning on reasoning traces generated by frontier reasoning LLMs, and (2) Table-R1-Zero, trained using our developed RLVR approach tailored for table reasoning tasks. The methodologies for each approach are detailed in the following subsections.

3.1 Training Data Collection

We construct the Table-R1 training dataset by integrating three representative table reasoning tasks, each introducing distinct reasoning challenges: (1) Short-form Table QA, which requires models to provide precise answers to questions grounded in

Task	Dataset	Samples
Short-form QA (TQA)	WTQ (Pasupat and Liang, 2015)	13,706
	HiTab (Cheng et al., 2022)	6,793
Fact Verification (TFV)	TabFact (Chen et al., 2020)	20,740
Free-form QA (FF-TQA)	FeTaQA (Nan et al., 2022)	7,324

Table 1: Overview of datasets collected in Table-R1 training data. For each dataset, we sample examples from its training set.

tabular data; (2) Table Fact Verification, which requires models to determine whether a given claim is entailed by the table content; and (3) Free-form Table QA, which requires models to produce open-ended answers grounded in tabular information. Each task contributes unique reasoning challenges, and we select established datasets (presented in Table 1) to ensure comprehensive coverage. All datasets are preprocessed with the presence of verifiable ground truths for reward computation to align with the RLVR paradigm. More detailed explanation of training data is presented in Appendix A.1.

3.2 Table-R1-SFT via Supervised Finetuning

To enable inference-time scaling in Table-R1-SFT, we curate a new table reasoning dataset featuring long CoT reasoning. Specifically, for each instance in the raw training data described in the previous subsection, we use DeepSeek-R1 (DeepSeek-AI et al., 2025) to generate a long CoT response. The response consists of a step-by-step reasoning process followed by a final answer. We present the prompts for response generation in Appendix C. To ensure the quality and correctness of the training data, we apply automated evaluators (detailed in Section 4.2) to assess the final answers. Examples with incorrect answers are filtered out. After this verification step, we obtain a total of 33,601 high-quality examples for SFT training. This dataset is then used to train the Table-R1-SFT model.

3.3 Table-R1-Zero via RL Training

We describe the RLVR algorithm and the reward design for training Table-R1-Zero as follows.

Reinforcement Learning with Verifiable Rewards (RLVR). We adapt Group Relative Policy Optimization (GRPO) with recent improvements introduced by DAPO (Yu et al., 2025b), including both token-level loss computation and asymmetric (decoupled) clipping. Notably, following recent work (Hu et al., 2025; Liu et al., 2025; Xia et al., 2025), we omit the KL penalty term present in

the original GRPO (Jaech et al., 2024; Pfister and Jud, 2025). We formally define the RL algorithm applied in our study as follows:

For each input (q, a) , the policy π_θ samples a group of G candidate responses $\{o_i\}_{i=1}^G$. Each response receives a reward R_i as described in the next paragraph. The group-normalized advantage for the i -th response at time step t is:

$$\hat{A}_{i,t} = \frac{R_i - \text{mean}(\{R_j\}_{j=1}^G)}{\text{std}(\{R_j\}_{j=1}^G)}. \quad (1)$$

Our objective is optimized at the token level with decoupled, asymmetric clipping:

$$\mathcal{J}_{\text{GRPO}}(\theta) = \mathbb{E}_{(q,a) \sim \mathcal{D}, \{o_i\}_{i=1}^G \sim \pi_{\theta_{\text{old}}}(\cdot|q)} \left[\frac{1}{\sum_{i=1}^G |o_i|} \sum_{i=1}^G \sum_{t=1}^{|o_i|} \min \left(r_{i,t}(\theta) \hat{A}_{i,t}, \right. \right. \\ \left. \left. \text{clip}(r_{i,t}(\theta), 1 - \varepsilon_{\text{low}}, 1 + \varepsilon_{\text{high}}) \hat{A}_{i,t} \right) \right] \quad (2)$$

where the probability ratio $r_{i,t}(\theta)$ is defined as:

$$r_{i,t}(\theta) = \frac{\pi_\theta(o_{i,t} | q, o_{i,<t})}{\pi_{\theta_{\text{old}}}(o_{i,t} | q, o_{i,<t})}. \quad (3)$$

This formulation enables stable and effective RL fine-tuning for table reasoning with LLMs.

Reward Design. To facilitate effective RL training, we design verifiable reward signals tailored to the characteristics of each table reasoning task. Our reward framework consists of two components: *accuracy rewards* and *format rewards*. Accuracy rewards measure the correctness of model outputs. We define task-specific reward functions as follows:

- **TQA:** The ground-truth is a short-answer list, where each element contains several words. We employ exact match to assign a reward of 1 for a correct answer and 0 otherwise.
- **TFV:** The ground-truth is either “entailed” or “refuted”. The reward is 1 if the predicted label matches the ground-truth, and 0 otherwise.
- **FF-TQA:** The ground-truth is a sentence or paragraph. We use the average of normalized BLEU and ROUGE-L scores to reflect semantic overlap between model output and reference answer.

This combination of rule-based and metric-based evaluation ensures reward interpretability and robustness, mitigating instability and reward hacking.

In addition to *accuracy*, we encourage models—especially base models without instruction

tuning—to follow a strict *response format* specified by a system prompt. We introduce a *cumulative format reward* incentivizing outputs that match the required template: `<think>...</think> <answer>...</answer>`, with the `<answer>` block containing a JSON snippet of the form `{"answer": ...}`. The format reward is computed by a deterministic regular expression-based checker, which assigns partial credit as outputs progressively satisfy structural requirements (e.g., tag presence, JSON structure, valid answer types), and awards full credit for strictly conformant outputs. This dense and interpretable reward guides base models to generate well-structured, verifiable responses.

4 Experiment

In this section, we describe the experimental setup and address three central research questions, presenting the findings associated with each.

4.1 Table-R1 Model Training Details

All models are trained using the verl framework. We initialize Table-R1 with Qwen2.5-7B-Instruct and Llama-3.1-8B-Instruct models. For Table-R1-SFT, training is conducted with a batch size of 256, a maximum sequence length of 20,480, and `ulysses_sequence_parallel_size` set to 4. The learning rate is set to 1×10^{-5} , and training is performed for 3 epochs. For Table-R1-Zero, we use a batch size of 256 and 16 rollouts per prompt under the GRPO algorithm. It is trained for 2 epochs. The learning rate is fixed at 1×10^{-6} , sampling temperature is 1.0, with a maximum prompt length of 4096 tokens and a maximum response length of 1024 tokens. The GRPO clipping parameters are set to $\varepsilon_{\text{low}} = 0.2$ and $\varepsilon_{\text{high}} = 0.28$. For validation during training, inference is performed with temperature 0.6 and top-p 1.0. All the experiments are conducted on 4 NVIDIA A100 80GB GPUs.

4.2 Experiment Setup

Evaluation Benchmarks. To address the lack of a unified evaluation framework for table reasoning, we introduce a new benchmark suite encompassing a wide range of datasets. For *in-domain evaluation*, we use test sets from the same distributions as the training data: WTQ (Pasupat and Liang, 2015) and HiTab (Cheng et al., 2022) for TQA; TabFact (Chen et al., 2020) for TFV; and FeTaQA (Nan et al., 2022) for FF-TQA. To assess generalization of Table-R1, we further conduct *out-of-domain evaluation* on datasets that are not seen during training:

Model	In-domain Performance				Out-of-domain Performance								
	FF-TQA		TFV		TQA		FF-TQA			TFV			TQA
	FeTaQA	TabFact	WTQ	HiTab	ToTTo	QTSum	R.W.	InfoTabs	PHT	Feverous	TMCQ	TMWP	FinQA
<i>Proprietary Models</i>													
GPT-4.1	25.1	86.5	68.0	84.7	20.4	45.7	21.0	90.5	88.2	87.7	92.0	77.0	74.0
GPT-4.1 mini	27.2	84.9	69.5	80.7	18.8	46.4	20.0	88.9	86.8	86.1	92.9	86.2	71.4
<i>Open-source Models</i>													
<i>LLMs</i>													
Qwen2.5-7B	21.0	72.2	54.8	61.8	16.0	39.5	19.3	78.6	70.7	74.6	87.4	85.0	66.4
Qwen2.5-32B	21.9	90.3	77.3	79.4	17.8	41.4	<u>20.0</u>	90.5	86.7	79.2	92.1	95.8	<u>77.3</u>
DeepSeek-V3	24.7	92.4	69.9	<u>82.2</u>	19.0	46.2	20.9	91.9	86.2	85.8	87.6	93.4	78.6
<i>Reasoning LLMs</i>													
DeepSeek-R1-Distill-7B	19.1	79.6	57.8	46.2	10.7	37.2	18.0	85.7	87.1	77.5	80.9	94.0	66.8
QwQ-32B	23.8	<u>91.5</u>	85.4	81.6	19.4	41.9	19.6	<u>91.0</u>	87.8	80.1	90.7	99.4	76.2
DeepSeek-R1	26.2	90.8	79.6	82.4	18.5	<u>43.8</u>	19.9	90.4	87.5	76.0	93.3	<u>99.0</u>	75.8
<i>Table-Specific LLMs</i>													
[†] TableBenchLLM (7B)	3.1	27.1	38.8	–	6.2	–	–	–	–	42.3	–	–	–
[†] TableLLM-13B	10.8	69.0	66.3	6.3	5.4	–	–	–	–	21.5	–	–	–
[†] TableLlama (7B)	39.0	82.6	35.0	64.7	20.8	–	–	–	–	73.8	–	–	–
TableGPT2-7B	29.0	77.8	61.4	70.3	14.1	39.0	19.0	85.4	<u>89.1</u>	78.0	77.2	79.7	66.4
[†] *TableGPT2-72B	32.3	85.4	71.5	75.6	22.7	–	–	–	–	76.8	–	–	–
<i>Table-R1 (Ours)</i>													
<i>Llama-3.1 Series</i>													
Llama-3.1-8B-Instruct	21.7	74.1	52.3	58.2	16.5	31.6	18.1	84.1	82.5	78.3	49.5	72.0	57.1
Table-R1-SFT	26.0	91.1	<u>83.8</u>	81.8	13.7	36.6	16.6	89.8	85.8	79.4	90.8	89.0	64.3
Table-R1-Zero	<u>32.7</u>	87.6	81.2	81.4	<u>22.3</u>	30.2	17.7	87.9	91.6	<u>80.2</u>	68.6	84.6	62.3
<i>Qwen2.5 Series</i>													
Qwen2.5-7B-Instruct	21.0	72.2	54.8	61.8	16.0	39.5	19.3	78.6	70.7	74.6	87.4	85.0	66.4
Table-R1-SFT	25.3	89.9	81.9	78.3	14.1	38.8	18.8	88.8	84.6	76.0	90.9	96.6	71.7
Table-R1-Zero	30.6	87.6	79.8	78.1	19.8	43.1	<u>20.0</u>	83.7	88.0	76.2	<u>93.0</u>	96.4	70.8

Table 2: Results on 13 table reasoning benchmarks spanning TQA, TFV, and FF-TQA tasks. For TQA, EM accuracy is reported (with ambiguous cases re-evaluated by GPT-4.1 mini); for TFV, classification accuracy; for FF-TQA, BLEU and ROUGE-L. **Bold** and underlined scores indicate the top-2 performances among *open-source* models. [†]: Due to the context length limitations of most previous table-specific LLMs, it is challenging to conduct a fully fair comparison. Therefore, for these models, we directly use the results as reported in their respective papers, which may be based on sampled or filtered datasets. *: Model weight has not been released.

TabMCQ (Jauhar et al., 2016), TabMWP (Lu et al., 2023), and FinQA (Chen et al., 2021) for TQA; InfoTabs (Gupta et al., 2020), PubHealthTab (Akhtar et al., 2022), and Feverous (Aly et al., 2021) for TFV; ToTTo (Parikh et al., 2020), QTSumm (Zhao et al., 2023b), and RotoWire (Wiseman et al., 2017) for FF-TQA. We provide detailed descriptions of the **evaluated datasets** and **evaluated baseline systems** in Appendix A.2 and Appendix B.

Automated Evaluation System. For each evaluated dataset, we use its test set for evaluation. For the TQA task, we report Exact Match (EM) accuracy. For TFV, we use classification accuracy. For FF-TQA, we measure generation quality with BLEU and ROUGE-L scores. Considering that EM accuracy in short-answer TQA may underestimate

model performance due to formatting variations or semantically equivalent but non-exact matches, we further re-evaluate responses initially marked incorrect by EM using the GPT-4.1 mini model, with the prompt shown in Appendix C.

💡 **RQ1:** How effective is inference-time scaling for table reasoning tasks?

To address RQ1, we conduct a comprehensive evaluation and summarize the key findings from the results presented in Table 2 as follows:

4.3 Main Findings

General-Purpose LLMs vs. Reasoning LLMs.

We evaluate three pairs of general-purpose

LLMs—with and without inference-time scaling capabilities—on table reasoning tasks: Qwen2.5-7B vs. DeepSeek-R1-Distill-7B, Qwen2.5-32B vs. QwQ-32B, and DeepSeek-V3 vs. DeepSeek-R1. While reasoning LLMs tend to outperform their counterparts on the TQA benchmark, their results on TFV and FF-TQA are mixed and not consistently superior. This indicates that inference-time scaling alone, without table-specific training, does not provide a clear advantage for table reasoning tasks. These findings underscore the importance of specialized adaptation strategies, such as Table-R1, for effective performance in this domain.

Table-R1 In-Domain Performance. Both Table-R1-SFT and Table-R1-Zero achieve substantial improvements across all tasks. Specifically, Table-R1-Zero-8B obtains a BLEU score of 32.7 on Fe-TaQA for FF-TQA, significantly surpassing the best among other models of 26.2 from Deepseek-R1; Table-R1-SFT-8B reaches 91.1 accuracy on TabFact for TFV, closely matching the leading 91.9; and for TQA, our models achieve 83.8 and 81.8 on WTQ and HiTab, respectively, which are comparable to the best scores among other models of 85.4 and 84.7. These results consistently demonstrate that Table-R1 models deliver robust gains and competitive performance across diverse table reasoning scenarios, validating the effectiveness and versatility of both SFT and RLVR training strategies.

Table-R1 Out-of-Domain Performance. Table-R1 models exhibit strong generalization capabilities on out-of-domain datasets. Across most out-of-domain benchmarks, our models consistently outperform their respective initial baselines, demonstrating the effectiveness of both SFT and R1-Zero training methods for table reasoning. Notably, Table-R1-Zero-7B achieves the best overall generalization among all variants. In contrast, we observe that supervised fine-tuning (SFT) leads to weaker generalization compared to the R1-Zero method, and models initialized from Llama tend to generalize less effectively than those based on Qwen. These results highlight the advantage of our approach in improving table reasoning robustness beyond the training distribution.

💡 RQ2: What contributes to the success of RLVR methods in table reasoning tasks?

To address this research question, we present a



Figure 3: Response length during Table-R1 training across different models.

comprehensive analysis of Table-R1-Zero in the following three subsections: the training dynamics, a qualitative assessment of model responses, and an exploration of the reasoning capacity boundaries.

4.4 Analysis of Training Dynamics

We conduct a detailed analysis of the training dynamics exhibited by our Table-R1 models across various model backbones, including Qwen2.5 7B and Llama-3.1 8B, under both *base* and *instruct* configurations.

Figure 3 presents the evolution of *response length* throughout reinforcement learning. Notably, the base models consistently start with longer responses compared to their instruct counterparts. During the initial stage of RL training, we observe a sharp drop in response length, corresponding to a phase of format acquisition—where the model learns to produce outputs adhering to the expected answer format. Subsequently, response length gradually increases, with base models exhibiting a more pronounced growth trajectory than instruct models. Among all four model variants, Table-R1-Zero-Qwen2.5-7B-Instruct demonstrates the greatest stability, showing a smooth and moderate increase in response length. In contrast, Table-R1-Zero-Llama-3.1-8B displays considerable instability, ultimately failing to acquire the desired response format. We attribute this to the weaker instruction-following capabilities of the Llama-3.1-8B base model.

Figure 4 illustrates the progression of *model accuracy* over the course of RL training. We find that instruct models consistently achieve higher accuracy than their base versions throughout training. While base models start from a lower performance baseline, they undergo a phase of format adaptation, after which their accuracy trends converge with those of instruct models. Interestingly,

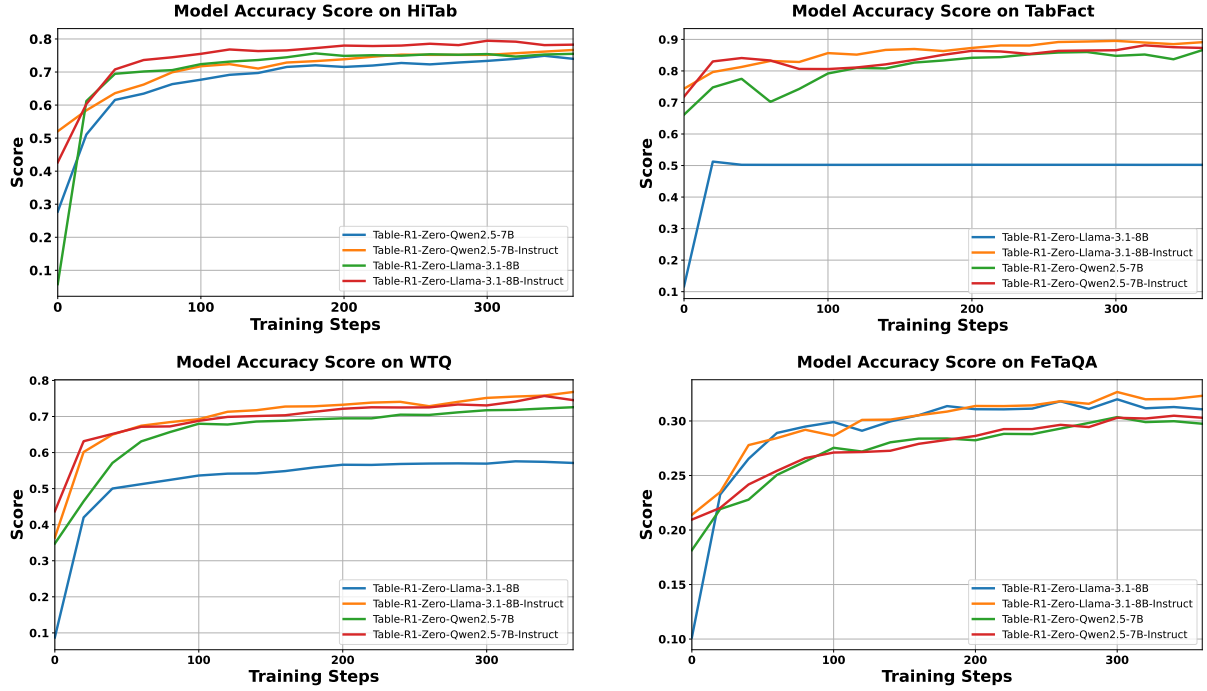


Figure 4: Accuracy and BLEU score dynamics across four table reasoning datasets during RLVR training. Results are shown for all four Table-R1-Zero models, which are trained from Qwen2.5 7B or Llama-3.1 8B as initialization.

although Table-R1-Zero-Llama-3.1-8B exhibits unstable changes in response length, it still demonstrates performance gains on most tasks, with the exception of TabFact—a binary classification task.

Overall, these observations reveal that instruction tuning provides significant benefits in both stability and performance during RLVR training for table reasoning tasks. The results further suggest that initial format alignment and instruction-following capability play critical roles in the successful adaptation of large language models to structured reasoning scenarios.

4.5 Reasoning Capacity Boundary

Understanding the upper limits of a model’s reasoning ability is crucial for evaluating the true impact of RLVR on tabular reasoning tasks. Inspired by prior studies on the boundaries of RLVR method (Yue et al., 2025), we employ the $pass@k$ metric to quantify the model’s capacity. Specifically, $pass@k$ measures the probability that at least one out of k generated responses is correct, given a fixed input prompt. This metric is particularly well-suited for our setting, as it captures not only the model’s accuracy, but also its ability to produce diverse and plausible reasoning trajectories within a limited number of attempts.

We systematically evaluate $pass@k$ for k up to

32, both before and after RLVR training, on the HiTab and WTQ datasets. As illustrated in Figure 8, RLVR training leads to a notable increase in $pass@k$ values throughout the evaluated range. This improvement is consistent across different datasets and model architectures, demonstrating that RLVR enhances not only the likelihood of obtaining a correct answer on first attempt ($pass@1$), but also the breadth of valid reasoning paths the model can explore in a small sampling budget.

4.6 Qualitative Analysis

To gain deeper insights into how RLVR shapes model behavior, we conduct a qualitative analysis of model responses throughout the training process. By examining the same set of representative prompts across different RL training steps, we observe that the model not only internalizes general R1-style reasoning characteristic, but also acquires table-specific reasoning abilities critical for tabular tasks, as illustrated in Figure 5 and Appendix D.1.

On the **reasoning** axis, we observe clear progress toward sophisticated, multi-step reasoning: after RL training, the Table-R1-Zero model decomposes complex queries into sequential sub-tasks, explicitly outlining intermediate steps and sometimes “rethinking” earlier conclusions to check or correct errors. Such reflective patterns,

Model-level Ablation. We first analyze the effect of model initialization by comparing base and instruct variants. Across all settings, instruct models consistently outperform their base counterparts. We attribute this to the enhanced instruction-following ability inherent to instruct models, which enables faster adaptation to the explicit reasoning formats required by table tasks and leads to more stable training dynamics.

We next compare representative model architectures—Qwen2.5 and Llama-3.1. Under identical training regimes, Llama-3.1-8B-Instruct achieves superior in-domain performance on table reasoning benchmarks compared to Qwen2.5-7B-Instruct, indicating a stronger capacity for learning table-specific reasoning. However, Qwen2.5 demonstrates better out-of-domain generalization, suggesting that model architecture and pretraining data may influence the balance between in-domain accuracy and cross-domain robustness.

Furthermore, we evaluate the effect of distillation data by comparing our SFT models—fine-tuned on table-specific DeepSeek-R1 distillation data—against the official DeepSeek-R1 distilled models. Our models not only surpass the official versions, but also outperform larger-scale distilled models. This underscores the effectiveness of domain-specific fine-tuning and the importance of high-quality, task-aligned training data.

Task-level Ablation. To investigate cross-task generalization, we train Table-R1 exclusively on the TQA dataset and evaluate its performance on TFV and FF-TQA tasks. Interestingly, models trained solely on TQA exhibit notable performance gains on TFV, indicating that the reasoning capabilities required for TFV are closely aligned with those developed for short-answer TQA. In contrast, no significant improvement is observed on FF-TQA, likely due to the distinct reasoning and answer generation skills required for free-form responses, which are not adequately covered by TQA training. These results highlight the varying degrees of transferability among table reasoning tasks and emphasize the need for targeted training to achieve robust generalization.

Format Ablation. We assess the role of explicit reasoning format supervision by removing the format reward during training. This ablation reduces training stability and slightly lowers in-domain performance. More notably, generalization suffers: while TFV scores may improve, performance on

short-answer TQA and FF-TQA declines, sometimes even below baseline, indicating that format supervision is crucial for transferable reasoning.

5 Conclusion

This work presents the first comprehensive study on applying inference-time scaling to table reasoning tasks. Through extensive evaluation across 13 diverse benchmarks, we demonstrate that inference-time scaling enables substantial improvements in reasoning quality, with RLVR methods yielding stronger generalization to out-of-domain tasks. Ablation studies confirm the benefits of instruction tuning, model architecture choice, and task composition in enhancing training effectiveness. Qualitative analysis reveals that RLVR fosters the emergence of structured, multi-step reasoning and table-specific capabilities. We hope this work paves the way for future research in structured reasoning.

Limitations

Several limitations remain that warrant future investigation: The SFT data was generated exclusively using DeepSeek-R1. Additionally, the data verification and filtering processes may have inadvertently removed difficult or high-quality training examples. Future research could explore incorporating outputs from other reasoning LLMs, such as QwQ-32B, to enhance distillation performance and data diversity. Furthermore, in Section 4.4, we observe that models initialized from the LLaMA-3.1-8B backbone exhibit unstable training dynamics during RLVR fine-tuning, including inconsistent acquisition of the desired output format and significant fluctuations in response length. While instruct-tuned variants mitigate some of these issues, the underlying causes of instability—such as sensitivity to initialization, reward sparsity, or optimization hyperparameters—remain underexplored. Future work could investigate strategies to improve the robustness and generalizability of RLVR for structured reasoning tasks.

Acknowledgments

We would like to thank Zijie Zhou for making the TQA-Distill-R1 dataset¹ publicly available, which we used in our early experiments.

¹<https://huggingface.co/datasets/jared-zhou/TQA-Distill-R1>

References

- Mubashara Akhtar, Oana Cocarascu, and Elena Simperl. 2022. [PubHealthTab: A public health table-based dataset for evidence-based fact checking](#). In *Findings of the Association for Computational Linguistics: NAACL 2022*, pages 1–16, Seattle, United States. Association for Computational Linguistics.
- Rami Aly, Zhijiang Guo, Michael Sejr Schlichtkrull, James Thorne, Andreas Vlachos, Christos Christodoulopoulos, Oana Cocarascu, and Arpit Mittal. 2021. [FEVEROUS: Fact extraction and VERification over unstructured and structured information](#). In *Thirty-fifth Conference on Neural Information Processing Systems Datasets and Benchmarks Track (Round 1)*.
- Wenhu Chen, Hongmin Wang, Jianshu Chen, Yunkai Zhang, Hong Wang, Shiyang Li, Xiyu Zhou, and William Yang Wang. 2020. [Tabfact: A large-scale dataset for table-based fact verification](#). In *International Conference on Learning Representations*.
- Zhiyu Chen, Wenhu 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*, pages 3697–3711, Online and Punta Cana, Dominican Republic. Association for Computational Linguistics.
- Zhoujun Cheng, Haoyu Dong, Zhiruo Wang, Ran Jia, Jiaqi Guo, Yan Gao, Shi Han, Jian-Guang Lou, and Dongmei Zhang. 2022. [HiTab: A hierarchical table dataset for question answering and natural language generation](#). In *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 1094–1110, Dublin, Ireland. Association for Computational Linguistics.
- DeepSeek-AI, Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma, Peiyi Wang, Xiao Bi, Xiaokang Zhang, Xingkai Yu, Yu Wu, Z. F. Wu, Zhibin Gou, Zhihong Shao, Zhuoshu Li, Ziyi Gao, and 181 others. 2025. [Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning](#). *Preprint*, arXiv:2501.12948.
- Naihao Deng, Zhenjie Sun, Ruiqi He, Aman Sikka, Yulong Chen, Lin Ma, Yue Zhang, and Rada Mihalcea. 2024. [Tables as texts or images: Evaluating the table reasoning ability of llms and mllms](#). *arXiv preprint arXiv:2402.12424*.
- Naihao Deng, Sheng Zhang, Henghui Zhu, Shuaichen Chang, Jiani Zhang, Alexander Hanbo Li, Chung-Wei Hang, Hideo Kobayashi, Yiqun Hu, and Patrick Ng. 2025. [Towards better understanding table instruction tuning: Decoupling the effects from data versus models](#). *Preprint*, arXiv:2501.14717.
- Hugging Face. 2025. [Open r1: A fully open reproduction of deepseek-r1](#).
- Jiazhan Feng, Shijue Huang, Xingwei Qu, Ge Zhang, Yujia Qin, Baoquan Zhong, Chengquan Jiang, Jinxin Chi, and Wanjun Zhong. 2025a. [Retool: Reinforcement learning for strategic tool use in llms](#). *arXiv preprint arXiv:2504.11536*.
- Xidong Feng, Ziyu Wan, Muning Wen, Stephen Marcus McAleer, Ying Wen, Weinan Zhang, and Jun Wang. 2023. [Alphazero-like tree-search can guide large language model decoding and training](#). *arXiv preprint arXiv:2309.17179*.
- Zhaopeng Feng, Jiahan Ren, Jiayuan Su, Jiamei Zheng, Zhihang Tang, Hongwei Wang, and Zuozhu Liu. 2025b. [Mt-rewardtree: A comprehensive framework for advancing llm-based machine translation via reward modeling](#). *arXiv preprint arXiv:2503.12123*.
- Xinyu Guan, Li Lyna Zhang, Yifei Liu, Ning Shang, Youran Sun, Yi Zhu, Fan Yang, and Mao Yang. 2025. [rstar-math: Small llms can master math reasoning with self-evolved deep thinking](#). *arXiv preprint arXiv:2501.04519*.
- Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma, Peiyi Wang, Xiao Bi, and 1 others. 2025. [Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning](#). *arXiv preprint arXiv:2501.12948*.
- Vivek Gupta, Maitrey Mehta, Pegah Nokhiz, and Vivek Srikumar. 2020. [INFOTABS: Inference on tables as semi-structured data](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 2309–2324, Online. Association for Computational Linguistics.
- Abram Handler, Kai R Larsen, and Richard Hackathorn. 2024. [Large language models present new questions for decision support](#). *International Journal of Information Management*, 79:102811.
- Jonathan Herzig, Pawel Krzysztof Nowak, Thomas Müller, Francesco Piccinno, and Julian Eisenschlos. 2020. [TaPas: Weakly supervised table parsing via pre-training](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 4320–4333, Online. Association for Computational Linguistics.
- Jingcheng Hu, Yinmin Zhang, Qi Han, Daxin Jiang, Xiangyu Zhang, and Heung-Yeung Shum. 2025. [Open-reasoner-zero: An open source approach to scaling up reinforcement learning on the base model](#). *arXiv preprint arXiv:2503.24290*.
- Wenxuan Huang, Bohan Jia, Zijie Zhai, Shaosheng Cao, Zheyu Ye, Fei Zhao, Zhe Xu, Yao Hu, and Shaohui Lin. 2025a. [Vision-r1: Incentivizing reasoning capability in multimodal large language models](#). *Preprint*, arXiv:2503.06749.

- Wenxuan Huang, Bohan Jia, Zijie Zhai, Shaosheng Cao, Zheyu Ye, Fei Zhao, Zhe Xu, Yao Hu, and Shaohui Lin. 2025b. Vision-r1: Incentivizing reasoning capability in multimodal large language models. *arXiv preprint arXiv:2503.06749*.
- Aaron Jaech, Adam Kalai, Adam Lerer, Adam Richardson, Ahmed El-Kishky, Aiden Low, Alec Helyar, Aleksander Madry, Alex Beutel, Alex Carney, and 1 others. 2024. Openai o1 system card. *arXiv preprint arXiv:2412.16720*.
- Sujay Kumar Jauhar, Peter Turney, and Eduard Hovy. 2016. [Tabmcq: A dataset of general knowledge tables and multiple-choice questions](#). *Preprint*, arXiv:1602.03960.
- Bowen Jin, Hansi Zeng, Zhenrui Yue, Jinsung Yoon, Sercan Arik, Dong Wang, Hamed Zamani, and Jiawei Han. 2025. [Search-r1: Training llms to reason and leverage search engines with reinforcement learning](#). *Preprint*, arXiv:2503.09516.
- Weixin Liang, Yaohui Zhang, Zhengxuan Wu, Haley Lepp, Wenlong Ji, Xuandong Zhao, Hancheng Cao, Sheng Liu, Siyu He, Zhi Huang, and 1 others. 2024. Mapping the increasing use of llms in scientific papers. *arXiv preprint arXiv:2404.01268*.
- Hunter Lightman, Vineet Kosaraju, Yuri Burda, Harrison Edwards, Bowen Baker, Teddy Lee, Jan Leike, John Schulman, Ilya Sutskever, and Karl Cobbe. 2023. Let’s verify step by step. In *The Twelfth International Conference on Learning Representations*.
- Qian Liu, Bei Chen, Jiaqi Guo, Morteza Ziyadi, Zeqi Lin, Weizhu Chen, and Jian-Guang Lou. 2022. [TAPEX: Table pre-training via learning a neural SQL executor](#). In *International Conference on Learning Representations*.
- Zichen Liu, Changyu Chen, Wenjun Li, Penghui Qi, Tianyu Pang, Chao Du, Wee Sun Lee, and Min Lin. 2025. Understanding r1-zero-like training: A critical perspective. *arXiv preprint arXiv:2503.20783*.
- Pan Lu, Liang Qiu, Kai-Wei Chang, Ying Nian Wu, Song-Chun Zhu, Tanmay Rajpurohit, Peter Clark, and Ashwin Kalyan. 2023. [Dynamic prompt learning via policy gradient for semi-structured mathematical reasoning](#). In *The Eleventh International Conference on Learning Representations*.
- Niklas Muennighoff, Zitong Yang, Weijia Shi, Xiang Lisa Li, Li Fei-Fei, Hannaneh Hajishirzi, Luke Zettlemoyer, Percy Liang, Emmanuel Candès, and Tatsunori Hashimoto. 2025. s1: Simple test-time scaling. *arXiv preprint arXiv:2501.19393*.
- Md Mahadi Hasan Nahid and Davood Rafiei. 2024. Tabsqlify: Enhancing reasoning capabilities of llms through table decomposition. *arXiv preprint arXiv:2404.10150*.
- Linyong Nan, Chiachun Hsieh, Ziming Mao, Xi Victoria Lin, Neha Verma, Rui Zhang, Wojciech Kryściński, Hailey Schoelkopf, Riley Kong, Xiangru Tang, Mutethia Mutuma, Ben Rosand, Isabel Trindade, Renusree Bandaru, Jacob Cunningham, Caiming Xiong, Dragomir Radev, and Dragomir Radev. 2022. [FeTaQA: Free-form table question answering](#). *Transactions of the Association for Computational Linguistics*, 10:35–49.
- Linyong Nan, Ellen Zhang, Weijin Zou, Yilun Zhao, Wenfei Zhou, and Arman Cohan. 2024. [On evaluating the integration of reasoning and action in LLM agents with database question answering](#). In *Findings of the Association for Computational Linguistics: NAACL 2024*, pages 4556–4579, Mexico City, Mexico. Association for Computational Linguistics.
- Benjamin Newman, Yoonjoo Lee, Aakanksha Naik, Pao Siangliulue, Raymond Fok, Juho Kim, Daniel S Weld, Joseph Chee Chang, and Kyle Lo. 2024. Arxivdigestables: Synthesizing scientific literature into tables using language models. *arXiv preprint arXiv:2410.22360*.
- OpenAI. 2024. [Gpt-4o system card](#). *arXiv preprint arXiv:2410.21276*.
- Jie Ouyang, Ruiran Yan, Yucong Luo, Mingyue Cheng, Qi Liu, Zirui Liu, Shuo Yu, and Daoyu Wang. 2025. [Training powerful llm agents with end-to-end reinforcement learning](#).
- Ankur Parikh, Xuezhi Wang, Sebastian Gehrmann, Manaal Faruqui, Bhuwan Dhingra, Diyi Yang, and Dipanjan Das. 2020. [ToTTo: A controlled table-to-text generation dataset](#). In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 1173–1186, Online. Association for Computational Linguistics.
- Panupong Pasupat and Percy Liang. 2015. [Compositional semantic parsing on semi-structured tables](#). In *Proceedings of the 53rd Annual Meeting of the Association for Computational Linguistics and the 7th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, pages 1470–1480, Beijing, China. Association for Computational Linguistics.
- Rolf Pfister and Hansueli Jud. 2025. Understanding and benchmarking artificial intelligence: Openai’s o3 is not agi. *arXiv preprint arXiv:2501.07458*.
- Zhenting Qi, Mingyuan Ma, Jiahang Xu, Li Lyna Zhang, Fan Yang, and Mao Yang. 2024. Mutual reasoning makes smaller llms stronger problem-solvers. *arXiv preprint arXiv:2408.06195*.
- Zhihong Shao, Peiyi Wang, Qihao Zhu, Runxin Xu, Junxiao Song, Xiao Bi, Haowei Zhang, Mingchuan Zhang, Y. K. Li, Y. Wu, and Daya Guo. 2024. [Deepseekmath: Pushing the limits of mathematical reasoning in open language models](#). *Preprint*, arXiv:2402.03300.

- Aofeng Su, Aowen Wang, Chao Ye, Chen Zhou, Ga Zhang, Gang Chen, Guangcheng Zhu, Haobo Wang, Haokai Xu, Hao Chen, Haoze Li, Haoxuan Lan, Jiaming Tian, Jing Yuan, Junbo Zhao, Junlin Zhou, Kaizhe Shou, Liangyu Zha, Lin Long, and 14 others. 2024. [Tablegpt2: A large multi-modal model with tabular data integration](#). *Preprint*, arXiv:2411.02059.
- Kimi Team, Angang Du, Bofei Gao, Bowei Xing, Changjiu Jiang, Cheng Chen, Cheng Li, Chenjun Xiao, Chenzhuang Du, Chonghua Liao, and 1 others. 2025. Kimi k1. 5: Scaling reinforcement learning with llms. *arXiv preprint arXiv:2501.12599*.
- Qwen Team. 2024. [Qwen2.5: A party of foundation models](#).
- Qwen Team. 2025. [Qwq-32b: Embracing the power of reinforcement learning](#).
- Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier Martinet, Marie-Anne Lachaux, Timothée Lacroix, Baptiste Rozière, Naman Goyal, Eric Hambro, Faisal Azhar, and 1 others. 2023. Llama: Open and efficient foundation language models. *arXiv preprint arXiv:2302.13971*.
- Zihan Wang, Kangrui Wang, Qineng Wang, Pingyue Zhang, Linjie Li, Zhengyuan Yang, Kefan Yu, Minh Nhat Nguyen, Licheng Liu, Eli Gottlieb, and 1 others. 2025. Ragen: Understanding self-evolution in llm agents via multi-turn reinforcement learning. *arXiv preprint arXiv:2504.20073*.
- Orion Weller, Kathryn Ricci, Eugene Yang, Andrew Yates, Dawn Lawrie, and Benjamin Van Durme. 2025. [Rank1: Test-time compute for reranking in information retrieval](#). *Preprint*, arXiv:2502.18418.
- Liang Wen, Yunke Cai, Fenrui Xiao, Xin He, Qi An, Zhenyu Duan, Yimin Du, Junchen Liu, Lifu Tang, Xiaowei Lv, and 1 others. 2025. Light-r1: Curriculum sft, dpo and rl for long cot from scratch and beyond. *arXiv preprint arXiv:2503.10460*.
- Sam Wiseman, Stuart Shieber, and Alexander Rush. 2017. [Challenges in data-to-document generation](#). In *Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing*, pages 2253–2263, Copenhagen, Denmark. Association for Computational Linguistics.
- Xianjie Wu, Jian Yang, Linzheng Chai, Ge Zhang, Jiaheng Liu, Xeron Du, Di Liang, Daixin Shu, Xianfu Cheng, Tianzhen Sun, and 1 others. 2025. Tablebench: A comprehensive and complex benchmark for table question answering. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 39, pages 25497–25506.
- Bingquan Xia, Bowen Shen, Dawei Zhu, Di Zhang, Gang Wang, Hailin Zhang, Huaqiu Liu, Jiebao Xiao, Jinhao Dong, Liang Zhao, and 1 others. 2025. Mimo: Unlocking the reasoning potential of language model—from pretraining to posttraining. *arXiv preprint arXiv:2505.07608*.
- Xiaobo Xia and Run Luo. 2025. Gui-r1: A generalist r1-style vision-language action model for gui agents. *arXiv preprint arXiv:2504.10458*.
- Tian Xie, Zitian Gao, Qingnan Ren, Haoming Luo, Yuqian Hong, Bryan Dai, Joey Zhou, Kai Qiu, Zhirong Wu, and Chong Luo. 2025. Logic-r1: Unleashing llm reasoning with rule-based reinforcement learning. *arXiv preprint arXiv:2502.14768*.
- Guowei Xu, Peng Jin, Hao Li, Yibing Song, Lichao Sun, and Li Yuan. 2025. [Llava-cot: Let vision language models reason step-by-step](#). *Preprint*, arXiv:2411.10440.
- Yixin Ye, Zhen Huang, Yang Xiao, Ethan Chern, Shijie Xia, and Pengfei Liu. 2025. Limo: Less is more for reasoning. *arXiv preprint arXiv:2502.03387*.
- Yunhu Ye, Binyuan Hui, Min Yang, Binhua Li, Fei Huang, and Yongbin Li. 2023. [Large language models are versatile decomposers: Decomposing evidence and questions for table-based reasoning](#). In *Proceedings of the 46th International ACM SIGIR Conference on Research and Development in Information Retrieval, SIGIR '23*, page 174–184, New York, NY, USA. Association for Computing Machinery.
- Peiying Yu, Guoxin Chen, and Jingjing Wang. 2025a. [Table-critic: A multi-agent framework for collaborative criticism and refinement in table reasoning](#). *Preprint*, arXiv:2502.11799.
- Qiyang Yu, Zheng Zhang, Ruofei Zhu, Yufeng Yuan, Xiaochen Zuo, Yu Yue, Tiantian Fan, Gaohong Liu, Lingjun Liu, Xin Liu, and 1 others. 2025b. Dapo: An open-source llm reinforcement learning system at scale. *arXiv preprint arXiv:2503.14476*.
- Lifan Yuan, Wendi Li, Huayu Chen, Ganqu Cui, Ning Ding, Kaiyan Zhang, Bowen Zhou, Zhiyuan Liu, and Hao Peng. 2024. Free process rewards without process labels. *arXiv preprint arXiv:2412.01981*.
- Yang Yue, Zhiqi Chen, Rui Lu, Andrew Zhao, Zhaokai Wang, Shiji Song, and Gao Huang. 2025. Does reinforcement learning really incentivize reasoning capacity in llms beyond the base model? *arXiv preprint arXiv:2504.13837*.
- Liangyu Zha, Junlin Zhou, Liyao Li, Rui Wang, Qingyi Huang, Saisai Yang, Jing Yuan, Changbao Su, Xiang Li, Aofeng Su, Tao Zhang, Chen Zhou, Kaizhe Shou, Miao Wang, Wufang Zhu, Guoshan Lu, Chao Ye, Yali Ye, Wentao Ye, and 6 others. 2023. [Tablegpt: Towards unifying tables, nature language and commands into one gpt](#). *Preprint*, arXiv:2307.08674.
- Tianshu Zhang, Xiang Yue, Yifei Li, and Huan Sun. 2023. Tablellama: Towards open large generalist models for tables. *arXiv preprint arXiv:2311.09206*.
- Tianshu Zhang, Xiang Yue, Yifei Li, and Huan Sun. 2024a. [TableLlama: Towards open large generalist models for tables](#). In *Proceedings of the 2024*

- Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)*, pages 6024–6044, Mexico City, Mexico. Association for Computational Linguistics.
- Xiaokang Zhang, Sijia Luo, Bohan Zhang, Zeyao Ma, Jing Zhang, Yang Li, Guanlin Li, Zijun Yao, Kangli Xu, Jinchang Zhou, Daniel Zhang-Li, Jifan Yu, Shu Zhao, Juanzi Li, and Jie Tang. 2025a. [Tablellm: Enabling tabular data manipulation by llms in real office usage scenarios](#). *Preprint*, arXiv:2403.19318.
- Xiaokang Zhang, Sijia Luo, Bohan Zhang, Zeyao Ma, Jing Zhang, Yang Li, Guanlin Li, Zijun Yao, Kangli Xu, Jinchang Zhou, and 1 others. 2024b. Tablellm: Enabling tabular data manipulation by llms in real office usage scenarios. *arXiv preprint arXiv:2403.19318*.
- Xuanliang Zhang, Dingzirui Wang, Longxu Dou, Qingfu Zhu, and Wanxiang Che. 2025b. A survey of table reasoning with large language models. *Frontiers of Computer Science*, 19(9):199348.
- Yilun Zhao, Lyuhao Chen, Arman Cohan, and Chen Zhao. 2024a. [TaPERA: Enhancing faithfulness and interpretability in long-form table QA by content planning and execution-based reasoning](#). In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 12824–12840, Bangkok, Thailand. Association for Computational Linguistics.
- Yilun Zhao, Hongjun Liu, Yitao Long, Rui Zhang, Chen Zhao, and Arman Cohan. 2024b. [Financemath: Knowledge-intensive math reasoning in finance domains](#). In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 12841–12858, Bangkok, Thailand. Association for Computational Linguistics.
- Yilun Zhao, Yitao Long, Hongjun Liu, Ryo Kamoi, Linyong Nan, Lyuhao Chen, Yixin Liu, Xiangru Tang, Rui Zhang, and Arman Cohan. 2024c. [DocMath-eval: Evaluating math reasoning capabilities of LLMs in understanding long and specialized documents](#). In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 16103–16120, Bangkok, Thailand. Association for Computational Linguistics.
- Yilun Zhao, Linyong Nan, Zhenting Qi, Rui Zhang, and Dragomir Radev. 2022. [ReasTAP: Injecting table reasoning skills during pre-training via synthetic reasoning examples](#). In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, pages 9006–9018, Abu Dhabi, United Arab Emirates. Association for Computational Linguistics.
- Yilun Zhao, Zhenting Qi, Linyong Nan, Lorenzo Jaime Flores, and Dragomir Radev. 2023a. [LoFT: Enhancing faithfulness and diversity for table-to-text generation via logic form control](#). In *Proceedings of the 17th Conference of the European Chapter of the Association for Computational Linguistics*, pages 554–561, Dubrovnik, Croatia. Association for Computational Linguistics.
- Yilun Zhao, Zhenting Qi, Linyong Nan, Boyu Mi, Yixin Liu, Weijin Zou, Simeng Han, Ruizhe Chen, Xiangru Tang, Yumo Xu, Dragomir Radev, and Arman Cohan. 2023b. [QTSumm: Query-focused summarization over tabular data](#). In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 1157–1172, Singapore. Association for Computational Linguistics.
- Yilun Zhao, Haowei Zhang, Shengyun Si, Linyong Nan, Xiangru Tang, and Arman Cohan. 2023c. [Investigating table-to-text generation capabilities of large language models in real-world information seeking scenarios](#). In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing: Industry Track*, pages 160–175, Singapore. Association for Computational Linguistics.
- Yilun Zhao, Chen Zhao, Linyong Nan, Zhenting Qi, Wenlin Zhang, Xiangru Tang, Boyu Mi, and Dragomir Radev. 2023d. [RobuT: A systematic study of table QA robustness against human-annotated adversarial perturbations](#). In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 6064–6081, Toronto, Canada. Association for Computational Linguistics.
- Shengyao Zhuang, Xueguang Ma, Bevan Koopman, Jimmy Lin, and Guido Zuccon. 2025. [Rank-r1: Enhancing reasoning in llm-based document rerankers via reinforcement learning](#). *Preprint*, arXiv:2503.06034.

A Dataset Details

A.1 Detailed Explanation of Training Data

Our training data originates from four benchmark datasets: WTQ, HiTab, TabFact, and FeTaQA.

Training Data for RLVR. For WTQ, HiTab, and FeTaQA, we used their complete training sets. For TabFact, which has a very large training set ($> 90,000$ samples), we implemented a difficulty-based filtering process to create a more focused and challenging subset. Specifically, we used Qwen2.5-7B-Instruct to generate 16 candidate responses for each sample. We retained all samples with a pass rate of $\leq 7/16$ and randomly sampled 10,000 instances from those with a pass rate between $8/16$ and $15/16$. This resulted in a final set of 20,740 samples for TabFact. After constructing the prompts, we filtered out any instances where the prompt length exceeded 4096 tokens. This final, combined dataset (48,563 samples) was used for RLVR training.

Training Data for SFT. We used the RLVR dataset as a starting point. We then generated reasoning traces for each sample using DeepSeek-R1 and applied a stringent quality-control filter to retain only high-quality examples. For WTQ, HiTab, and TabFact, we kept only the samples where the generated answer was an exact match to the ground truth. For FeTaQA, we kept samples where the average of BLEU and ROUGE-L scores was ≥ 0.35 . This resulted in a high-quality SFT dataset of 33,601 samples.

The data volumes at each stage are summarized below, as shown in Table 3.

Dataset	Raw Samples	Samples for RLVR	Samples for SFT (Filtered)
WTQ	14,152	13,706	9,524
HiTab	7,301	6,793	4,170
TabFact	92,283	20,740	16,006
FeTaQA	7,326	7,324	3,901
Total	121,062	48,563	33,601

Table 3: The data volumes at each stage.

Table 4 provides detailed statistics on the SFT data. Notably, the quality filtering significantly reduces the average response-to-prompt length ratio (from 1.73 to 1.49 overall). This aligns with the common observation that incorrect or lower-quality reasoning traces are often longer and more convoluted, reinforcing the effectiveness of our filtering strategy.

Dataset / Length (Chars)	Min Prompt	Max Prompt	Avg. Prompt	Min Resp.	Max Resp.	Avg. Resp.	Avg. Resp/Prompt Ratio (Before/After Filtering)
WTQ	625	14,192	2,114.67	465	28,023	2,356.26	1.78 / 1.45
HiTab	694	10,264	3,251.17	526	32,630	2,827.42	1.57 / 1.12
TabFact	624	8,906	1,767.64	529	28,762	2,474.04	1.76 / 1.61
FeTaQA	473	13,793	1,656.14	546	28,661	1,982.57	1.71 / 1.53
Overall	473	14,192	2,037.17	465	32,630	2,427.45	1.73 / 1.49

Table 4: Detailed statistics on the SFT data.

A.2 Information of Evaluated Table Reasoning Datasets

Task Category	Task Name	Dataset	Task Description	Metric
Training	Table QA	WTQ, HiTab	QA over flat or hierarchical tables	Acc.
	Free-form QA	FeTaQA	Generate long-form answers from tables	BLEU
	Table Fact Verification	TabFact	Verify factual correctness based on a table	Acc.
In-Domain	Table QA	WTQ, HiTab	In-domain table QA evaluation	Acc.
	Table Fact Verification	TabFact	In-domain fact verification evaluation	Acc.
	Free-form QA	FeTaQA	Evaluate long-form generation quality	BLEU
Out-of-Domain	Table QA	TabMCQ, TMWP, FinQA	Multiple-choice and word problem QA	Acc.
	Numerical Reasoning QA	FinQA, TABMWP	Financial and numerical table QA	Acc.
	Table Fact Verification	InfoTabs, PHT, Feverous	Health and factual judgment from tables	Acc.
	Free-form QA	ToTTo, Qtsumm, R.W.	Table-to-text generation from highlights	BLEU/R-L

Table 5: Overview of datasets used in training and evaluation.

B Experiment Setup

B.1 Baseline Systems

We benchmark our approach against a comprehensive set of strong baselines, encompassing both proprietary and open-source models. Among proprietary models, we include GPT-4.1 and GPT-4.1 mini. For open-source baselines, we evaluate general purpose LLMs from the Qwen2.5 and LLaMA3 series, reasoning models such as Deepseek-R1, its official distilled variant, our SFT model distilled from Deepseek-R1 on table reasoning data (as described in Section 3.2), and QwQ-32B, as well as table-oriented models including TableLlama, TableLLM, and TableBenchLLM.

B.2 Inference Setup

For all open-source models, inference is performed using the vLLM framework, while for closed-source models, the official OpenAI API is utilized. The maximum output length is set to 2048 tokens for most models. However, for reasoning models, this limit is increased to 18,000 tokens to accommodate their long chain-of-thought generation. The temperature is set to 0.6 and the top-p value to 0.95. All inference processes are conducted on four NVIDIA A100-80G GPUs.

B.3 Evaluated Model Configuration

Model	Citation	Version
GPT-4.1	OpenAI (2024)	gpt-4.1-2025-04-14
GPT-4.1 mini	OpenAI (2024)	gpt-4.1-mini-2025-04-14
Qwen2.5-7B	Team (2024)	Qwen/Qwen2.5-7B-Instruct
Qwen2.5-14B	Team (2024)	Qwen/Qwen2.5-14B-Instruct
Qwen2.5-32B	Team (2024)	Qwen/Qwen2.5-32B-Instruct
QwQ-32B	Team (2025)	Qwen/QwQ-32B
Llama-3.1-8B	Touvron et al. (2023)	meta-llama/Llama-3.1-8B-Instruct
DeepSeek-R1-7B	Guo et al. (2025)	deepseek-ai/DeepSeek-R1-Distill-Qwen-7B
DeepSeek-R1-14B	Guo et al. (2025)	deepseek-ai/DeepSeek-R1-Distill-Qwen-14B
DeepSeek-R1-32B	Guo et al. (2025)	deepseek-ai/DeepSeek-R1-Distill-Qwen-32B
DeepSeek-R1	Guo et al. (2025)	deepseek-ai/DeepSeek-R1
TableGPT2-7B	Su et al. (2024)	TableGPT/TableGPT2-7B
TableLLM-13B	Zhang et al. (2024b)	RUCKBReasoning/TableLLM-13B
TableLlama-7B	Zhang et al. (2023)	osunlp/TableLlama
TableBenchLLM	Zhang et al. (2024b)	Multilingual-Multimodal-NLP/TableLLM-Llama3.1-8B

Table 6: Model List.

C Prompts

System prompt used to guide structured response generation

A conversation between User and Assistant. The user asks a question, and the assistant solves it. The assistant first thinks about the reasoning process in the mind and then provides the user with the answer. The reasoning process and answer are enclosed within `<think>` `</think>` and `<answer>` `</answer>` tags, respectively, i.e., `<think>` reasoning process here `</think>` `<answer>` answer here `</answer>`.

Prompt Template for TQA

Instruction: This is a short-answer table QA task. Answer the question based on the provided table.

Table

Table Title: {table_title}

Table Content: {table_repr (markdown / html)}

Question: {question}

Answer Format:

The final answer should be concise and use the following format:

```
```json
{
 "answer": ["answer1", "answer2", ...]
}
```
```

Prompt Template for TFV

Instruction: This is a table fact verification task. The goal is to determine whether the given statement is entailed or refuted by the table.

Table

Table Title: {table_title}

Table Content: {table_repr (markdown / html)}

Statement: {statement}

Answer Format:

The final answer should be either "entailed" or "refuted" and use the following format:

```
```json
{
 "answer": "entailed" or "refuted"
}
```
```

Prompt Template for Free-Form TQA

Instruction: This is a free-form table QA task. Answer the question based on the provided table.

Table

Table Title: {table_title}

Table Content: {table_repr (markdown / html)}

Question: {question}

Answer Format:

The final answer should be a sentence and use the following format:

```
```json
{
 "answer": "your_generated_sentence_here"
}
```
```


Prompt Template for LLM-as-a-Judge

You are given two answers for a short-answer Table QA task: **response** and **ground_truth**.

- **response**: This is the LLM's answer to the task. It may include reasoning steps and a final answer.
- **ground_truth**: A list of short answers, typically 2-3 word noun phrases or numbers.

Your task is to determine whether the response is fully correct, using these rules:

- **Noun phrases**: Considered correct if meaning matches ground_truth regardless of wording.
- **Numbers**: Considered correct if numerically close (tolerance < 0.01).
- Every ground_truth item must be matched in the response. Order doesn't matter.

Your output must be in the following format:

```
```json
{
 "judgement": "correct" or "incorrect"
}
```
```

Do not provide any explanation or additional output.

Input:

Response: {response}

Ground_truth: {ground_truth}

Evaluate and output the judgement.

D Analysis

D.1 Qualitative Analysis Cases

D.2 Reasoning Capacity Boundary

D.3 Ablation Studies

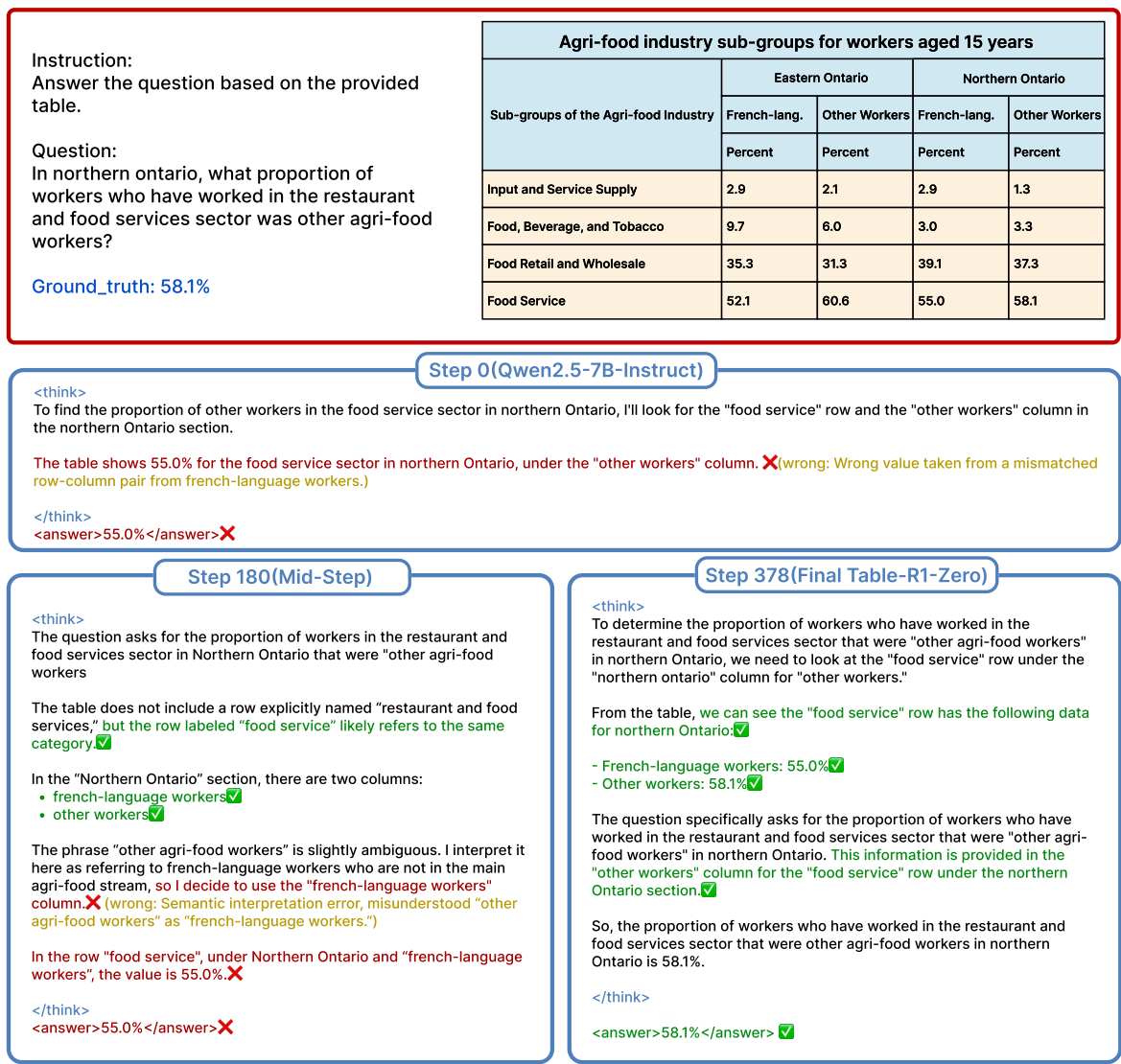


Figure 6: Qualitative Analysis Case Example for TQA tasks

Instruction:

This is a table fact verification task. The goal of this task is to distinguish whether the given statement is entailed or refuted by the given table.

Statement:

riann van zyl have the shortest time span on the united state national rugby union team and tie with 3 others for the smallest number of tries

Ground_truth: entailed

| United States National Rugby Union Team | | | | | | |
|---|-------------|-------|-------|------|------|------|
| Player | Span | Start | Tries | Conv | Pens | Drop |
| Vaea Anitoni | 1992 - 2000 | 44 | 26 | 0 | 0 | 0 |
| Paul Emerick | 2003 - 2012 | 49 | 17 | 0 | 0 | 0 |
| Todd Clever | 2003 - | 51 | 11 | 0 | 0 | 0 |
| Philip Eloff | 2000 - 2007 | 34 | 10 | 0 | 0 | 0 |
| Takudzwa Ngwenya | 2007 - | 27 | 10 | 0 | 0 | 0 |
| Chris Wyles | 2007 - | 35 | 10 | 14 | 22 | 1 |
| David Fee | 2002 - 2005 | 28 | 9 | 0 | 0 | 0 |
| Mike Hercus | 2002 - 2009 | 45 | 9 | 90 | 76 | 4 |
| Riaan van Zyl | 2003 - 2004 | 12 | 9 | 0 | 0 | 0 |

Step 0(Qwen2.5-7B-Instruct)

<think>

First, let's check the duration span for Riaan van Zyl. According to the table, his span is 2003-2004, which is 1 year. Next, let's compare this with other players. Todd Clever (2003-) and Chris Wyles (2007-) have ongoing spans, so they are not the shortest.

Now, let's check the number of tries for Riaan van Zyl. He has 9 tries. ✓

For other players, Philip Eloff (10 tries), Chris Wyles (10 tries), and Takudzwa Ngwenya (10 tries) also have 10 tries each. So, Riaan van Zyl does have the shortest span, but he does not tie with 3 others for the smallest number of tries, as 3 others have more tries than him. ✗(wrong: Incorrect count of players with the fewest tries.)

</think>

<answer>refuted</answer>✗(wrong: Conclusion error.)

Step 378(Final Table-R1-Zero)

<think>

To verify the statement "riann van zyl have the shortest time span on the united state national rugby union team and tie with 3 others for the smallest number of tries," we need to analyze the "span" and "tries" columns for each player in the table.

First, let's look at the "span" column to determine the shortest time span:

1. vaea anitoni: 1992 - 2000 (8 years)
2. paul emerick: 2003 - 2012 (9 years)
3. todd clever: 2003 - (ongoing, so we consider the start year 2003)
4. philip eloff: 2000 - 2007 (7 years)
5. takudzwa ngwenya: 2007 - (ongoing, so we consider the start year 2007)
6. chris wyles: 2007 - (ongoing, so we consider the start year 2007)
7. david fee: 2002 - 2005 (3 years)
8. mike hercus: 2002 - 2009 (7 years)
9. riaan van zyl: 2003 - 2004 (1 year)

From the above, the shortest time span is 1 year, which is for riaan van zyl. ✓The next shortest spans are:

- david fee: 3 years
- vaea anitoni: 8 years
- philip eloff: 7 years
- mike hercus: 7 years
- others: 9 years or more

So, riaan van zyl has the shortest time span of 1 year, and ties with david fee for the second shortest span of 3 years.

Next, let's look at the "tries" column to determine the smallest number of tries:

1. vaea anitoni: 26 tries
2. paul emerick: 17 tries
3. todd clever: 11 tries
4. philip eloff: 10 tries
5. takudzwa ngwenya: 10 tries
6. chris wyles: 10 tries
7. david fee: 9 tries ✓
8. mike hercus: 9 tries ✓
9. riaan van zyl: 9 tries ✓

The smallest number of tries is 9, which is shared by david fee, mike hercus, and riaan van zyl. Therefore, riaan van zyl ties with 3 others (david fee, mike hercus) for the smallest number of tries. ✓

Based on the analysis, the statement "riann van zyl have the shortest time span on the united state national rugby union team and tie with 3 others for the smallest number of tries" is correct.

</think>

<answer>entailed</answer>✓

Figure 7: Qualitative Analysis Case Example for TFV tasks

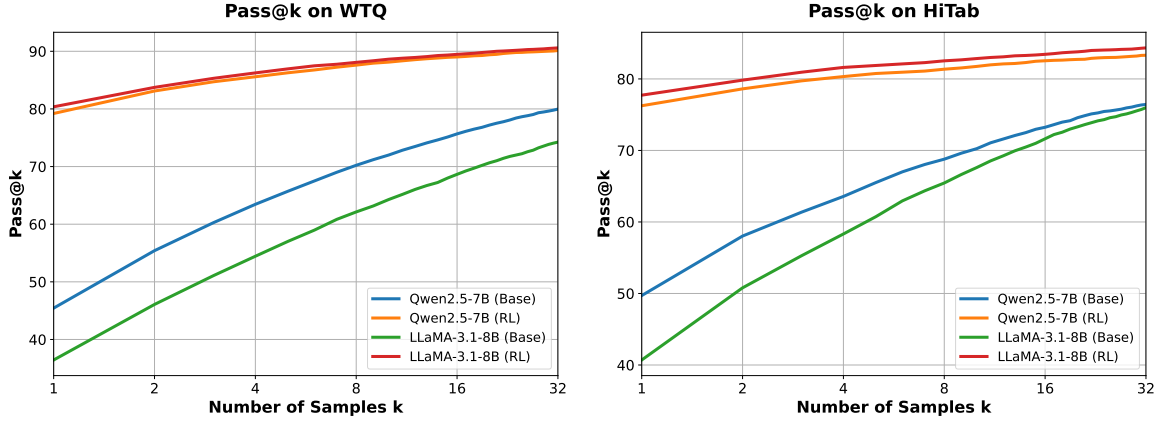


Figure 8: Pass@k performance on WTQ and HiTab.

| Model | In-domain Performance | | | | Out-of-domain Performance | | | | | | | | | |
|--------------------------------|-----------------------|---------|------|-------|---------------------------|--------|------|----------|------|----------|------|------|-------|--|
| | FF-TQA | | TFV | TQA | | FF-TQA | | | TFV | | | TQA | | |
| | FeTaQA | TabFact | WTQ | HiTab | ToTTo | QTSum | R.W. | InfoTabs | PHT | Feverous | TMCQ | TMWP | FinQA | |
| | | | | | | | | | | | | | | |
| Model-level Ablation | | | | | | | | | | | | | | |
| Base vs. Instruct | | | | | | | | | | | | | | |
| Qwen2.5-7B | 18.2 | 67.7 | 51.0 | 55.2 | 9.4 | 37.8 | 16.4 | 74.5 | 75.0 | 74.1 | 78.7 | 80.3 | 58.4 | |
| Table-R1-Qwen2.5-7B | 29.8 | 87.3 | 78.5 | 75.2 | 15.9 | 38.1 | 18.2 | 87.9 | 91.6 | 78.4 | 84.8 | 94.4 | 61.2 | |
| Qwen2.5-7B-Instruct | 21.0 | 72.2 | 54.8 | 61.8 | 16.0 | 39.5 | 19.3 | 78.6 | 70.7 | 74.6 | 87.4 | 85.0 | 66.4 | |
| Table-R1-Qwen2.5-7B-Instruct | 30.6 | 87.6 | 79.8 | 78.1 | 19.8 | 43.1 | 20.0 | 83.7 | 88.0 | 76.2 | 93.0 | 96.4 | 70.8 | |
| Model Architecture Comparison | | | | | | | | | | | | | | |
| Llama-3.1-8B | 10.2 | 21.2 | 35.0 | 44.7 | 6.8 | 18.6 | 10.9 | 21.6 | 26.8 | 19.4 | 68.9 | 43.1 | 21.9 | |
| Table-R1-Llama-3.1-8B | 30.4 | 50.3 | 61.5 | 76.3 | 19.8 | 33.5 | 18.0 | 50.0 | 69.7 | 50.6 | 93.1 | 56.8 | 19.3 | |
| Llama-3.1-8B-Instruct | 21.7 | 74.1 | 52.3 | 58.2 | 16.5 | 31.6 | 18.1 | 84.1 | 82.5 | 78.3 | 49.5 | 72.0 | 57.1 | |
| Table-R1-Llama-3.1-8B-Instruct | 32.7 | 87.6 | 81.2 | 81.4 | 22.3 | 30.2 | 17.7 | 87.9 | 91.6 | 80.2 | 68.6 | 84.6 | 62.3 | |
| SFT on Domain-Specific Data | | | | | | | | | | | | | | |
| DeepSeek-R1-7B | 19.1 | 79.6 | 57.8 | 46.2 | 10.7 | 37.2 | 18.0 | 85.7 | 87.1 | 77.5 | 80.9 | 94.0 | 66.8 | |
| Table-SFT-Qwen2.5-7B | 25.3 | 89.9 | 81.9 | 78.3 | 14.1 | 38.8 | 18.8 | 88.8 | 84.6 | 76.0 | 90.9 | 96.6 | 71.7 | |
| Table-SFT-Llama-3.1-8B | 26.0 | 91.1 | 83.8 | 81.8 | 13.7 | 36.6 | 16.6 | 89.8 | 85.8 | 79.4 | 90.8 | 89.0 | 64.3 | |
| Task-level Ablation | | | | | | | | | | | | | | |
| Cross-task Generalization | | | | | | | | | | | | | | |
| Table-R1-TQA-Qwen2.5-7B | 14.1 | 86.2 | 79.2 | 77.4 | 16.5 | 36.4 | 19.8 | 89.7 | 89.9 | 82.7 | 93.5 | 96.1 | 65.9 | |
| Table-R1-TQA-Llama-3.1-8B | 12.3 | 83.9 | 81.7 | 81.0 | 17.2 | 26.9 | 19.2 | 87.3 | 87.2 | 77.1 | 80.3 | 82.8 | 58.6 | |
| Table-SFT-TQA-Qwen2.5-7B | 13.3 | 87.7 | 79.8 | 77.0 | 13.4 | 27.6 | 18.1 | 88.5 | 87.1 | 78.0 | 90.1 | 95.7 | 69.0 | |
| Table-SFT-TQA-Llama-3.1-8B | 9.6 | 88.8 | 82.4 | 80.5 | 10.3 | 20.6 | 12.6 | 89.1 | 90.4 | 82.1 | 89.9 | 86.9 | 62.9 | |
| Format Ablation | | | | | | | | | | | | | | |
| Effect of Format Reward | | | | | | | | | | | | | | |
| Table-R1-Explicit-Qwen2.5-7B | 29.0 | 85.8 | 76.0 | 73.2 | 13.6 | 37.2 | 14.8 | 87.5 | 88.8 | 81.0 | 52.5 | 91.3 | 49.6 | |
| Table-R1-Qwen2.5-7B | 30.6 | 87.6 | 79.8 | 78.1 | 19.8 | 43.1 | 20.0 | 83.7 | 88.0 | 76.2 | 93.0 | 96.4 | 70.8 | |

Table 7: Ablation study results on model-level, task-level, and formats.