

AIR: Complex Instruction Generation via Automatic Iterative Refinement

Anonymous ACL submission

Abstract

With the development of large language models, their ability to follow simple instructions has significantly improved. However, adhering to complex instructions remains a major challenge. Current approaches to generating complex instructions are often irrelevant to the current instruction requirements or suffer from limited scalability and diversity. Moreover, methods such as back-translation, while effective for simple instruction generation, fail to leverage the rich knowledge and formatting in human written documents. In this paper, we propose a novel **Automatic Iterative Refinement (AIR)** framework to generate complex instructions with constraints, which not only better reflects the requirements of real scenarios but also significantly enhances LLMs’ ability to follow complex instructions. The AIR framework consists of two stages: 1) Generate an initial instruction from a document; 2) Iteratively refine instructions with LLM-as-judge guidance by comparing the model’s output with the document to incorporate valuable constraints. Finally, we construct the AIR-10K dataset with 10K complex instructions and demonstrate that instructions generated with our approach significantly improve the model’s ability to follow complex instructions, outperforming existing methods for instruction generation.¹

1 Introduction

Recent advancements in Large Language Models (LLMs) have shown impressive performance across a wide range of tasks (Zhao et al., 2023; Li et al., 2024a). Driven by vast amounts of data and efficient training, most current LLMs are capable of effectively following user instructions and aligning to a certain extent with human preferences (Ouyang et al., 2022; Li et al., 2024b). However, despite these successes, they still face significant chal-

¹Codes and data are anonymously available at anonymous.open.science/r/AIR-0833.

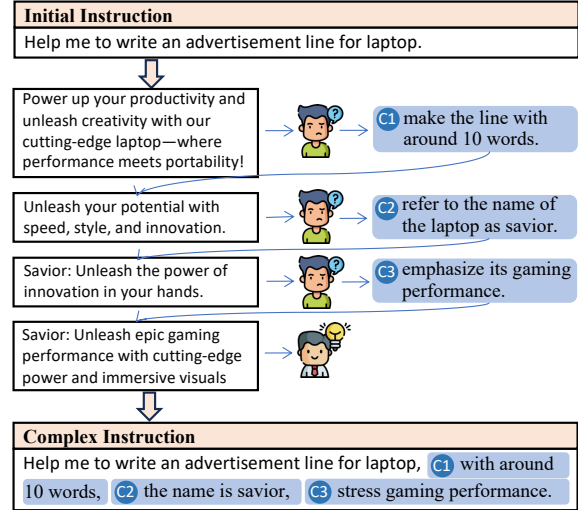


Figure 1: Illustration of how humans iteratively refine instructions to be more complex.

lenges when it comes to following complex instructions (Jiang et al., 2023; Wen et al., 2024).

Existing complex instructions datasets primarily originate from two sources: 1) Curated data from open-source datasets or human annotations (Zhou et al., 2024; Zhang et al., 2024), which are resource-intensive and **lack scalability**, and 2) Transforming simple instructions into complex ones automatically using proprietary LLMs (Xu et al., 2023; Sun et al., 2024). While the automatic transformation improves scalability, the generated constraints are often recombinations of few-shot examples, resulting in **limited diversity**. Moreover, these constraints may have **low relevance** with the target output, failing to reflect real-world scenarios.

Recently, back-translation, which involves translating text from the target side back into the source side, has been proposed to generate scalable and diverse instructions from human-written corpora (Sennrich, 2015; Hoang et al., 2018; Zheng et al., 2024a; Li et al., 2023). However, these methods typically focus on generating **simple instruc-**

tions and have not fully explored the rich knowledge contained in the human corpus.

In this paper, we propose an **Automatic Iterative Refinement (AIR)** framework for generating high-quality complex instructions. Specifically, our approach is based on two key observations. First, human-written documents contain massive human preferences that can be converted to specific constraints, such as formatting conventions in legal documents. Second, human often refine complex instructions iteratively based on feedback from model outputs. As illustrated in Figure 1, simple instructions are progressively adjusted and enriched to better align with human preferences. This iterative process plays a critical role in crafting effective complex instructions.

Therefore, our AIR framework incorporates document-based knowledge and LLM-as-judge to iteratively construct complex instructions. The framework consists of two key steps: 1) **Initial Instruction Generation**, where the model generates initial instructions based on the document content; 2) **Iterative Instruction Refinement**, where instructions are iteratively refined with LLM-as-judge guidance by comparing model outputs with the document, to identify and incorporate valuable constraints. This process enables the framework to generate more challenging instructions that align more closely with real-world scenarios.

In summary, our contributions are as follows:

- To better align with real-world scenarios, we propose the **AIR** framework, which iteratively refines complex instructions with LLM-as-judge guidance by comparing with the document.
- We present a new instruction dataset (AIR-10K) generated using our framework. Experimental results demonstrate that our fine-tuned model significantly outperforms existing methods on instruction-following benchmarks.
- We provide a comprehensive experimental analysis to evaluate the individual components of our framework, validating the contribution of each step to the overall improvement.

2 Related Work

2.1 Instruction Generation

Instruction tuning is essential for aligning Large Language Models (LLMs) with user intentions (Ouyang et al., 2022; Cao et al., 2023). Ini-

tially, this involved collecting and cleaning existing data, such as open-source NLP datasets (Wang et al., 2023; Ding et al., 2023). With the importance of instruction quality recognized, manual annotation methods emerged (Wang et al., 2023; Zhou et al., 2024). As larger datasets became necessary, approaches like Self-Instruct (Wang et al., 2022) used models to generate high-quality instructions (Guo et al., 2024). However, complex instructions are rare, leading to strategies for synthesizing them by extending simpler ones (Xu et al., 2023; Sun et al., 2024; He et al., 2024). However, existing methods struggle with scalability and diversity.

2.2 Back Translation

Back-translation, a process of translating text from the target language back into the source language, is mainly used for data augmentation in tasks like machine translation (Sennrich, 2015; Hoang et al., 2018). Li et al. (2023) first applied this to large-scale instruction generation using unlabeled data, with Suri (Pham et al., 2024) and Kun (Zheng et al., 2024a) extending it to long-form and Chinese instructions, respectively. Nguyen et al. (2024) enhanced this method by adding quality assessment to filter and revise data. Building on this, we further investigated methods to generate high-quality complex instruction dataset using back-translation.

3 Approach

Our approach mainly consists of two steps: 1) Initial Instruction Generation; 2) Iterative Instruction Refinement, as shown in Figure 2. In this section, we will introduce the two steps in detail.

3.1 Initial Instruction Generation (IIG)

Document Collection. Traditional instruction generation methods such as Self-Instruct (Wang et al., 2022) often suffer from limited diversity, as the generated instructions are generally recombinations of the provided few-shot examples. Inspired by the work by Li et al. (2023), we generate initial instructions using back translation based on human-written documents.

To further enhance the diversity of the generated instructions, we implement a density-based sampling mechanism for documents, as shown in Algorithm 1. Specifically, we convert documents into vector representations based on Sentence-Transformers², and perform sampling to maximize

²sentence-transformers/all-MiniLM-L6-v2.

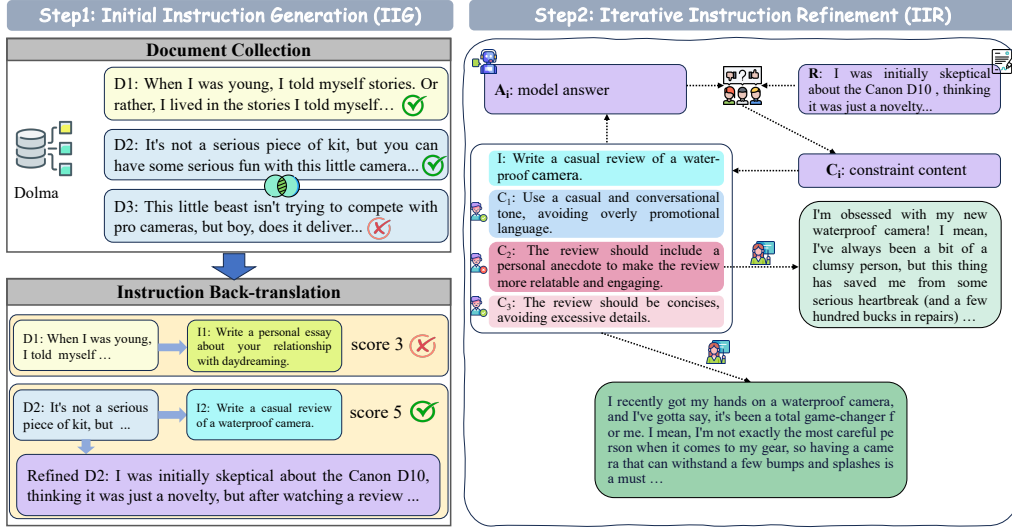


Figure 2: AIR: Automatic Iterative Refinement Framework.

Algorithm 1 Density-based Sampling

Input: Instruction Dataset D with m samples.

Output: Selected Dataset D' with n samples.

- 1: Derive the embeddings for each sample in D .
- 2: Random sample one data point x from D .
- 3: Delete x from D , add x to D' .
- 4: **for** $i = 1, 2, \dots, t$ **do**
- 5: Calculate the cosine similarity score between x and each sample from D .
- 6: Select the least similar sample x' from D' .
- 7: Let $x = x'$.
- 8: Delete x from D , add x to D' .
- 9: **end for**

the density of samples in the representation space.

In this way, we effectively eliminate redundant documents, enhancing the diversity of instructions. Moreover, this approach ensures that the knowledge introduced during instruction fine-tuning is evenly distributed across various domains. This not only prevents the model from overfitting to a specific domain but also mitigates the risk of catastrophic forgetting of fundamental capabilities.

Moreover, to further ensure the quality of the document collection, we filter out documents based on the following criteria: 1) Length: Documents with fewer than 50 words or exceeding 2,048 words are removed. 2) Symbol-to-text ratio: Documents where the proportion of symbols exceeds that of textual content are excluded. 3) Redundancy: Documents containing repetitive paragraphs or excessive symbol repetitions are eliminated.

Instruction Back-translation Based on the sampled documents, we employ the back-translation

method to construct initial instructions. Specifically, we utilize a guidance model to predict an instruction which can be accurately answered by (a portion of) the document³. This enables the generation of new instructions without relying on few-shot examples or pre-designed rules. Moreover, we can further ensure the diversity of the generated instructions by diversifying the documents.

However, despite being constructed from the document, the instruction do not always align well with the document in two key aspects (Nguyen et al., 2024). First, the document is unstructured and does not follow the AI-assistant format. Second, it may contain content irrelevant to the instruction. Therefore, we introduce an additional refinement step to transform the document into response format and remove irrelevant content.

To further ensure the quality of the instructions, we introduce a scoring step to filter out low-quality data. Each instruction is assigned a score on a scale of 1 to 5 by the guidance model, with each point corresponding to a specific aspect. Only instructions with a score greater than (or equal to) 4 are retained for the next step⁴.

3.2 Iterative Instruction Refinement (IIR)

To enhance a model’s ability to follow complex instructions, it is crucial to construct complex instruction data that incorporates multiple constraints. Previous methods typically start with simple instructions and generate complex ones through rewriting or recombination (Xu et al., 2023). However, the constraints generated in this way often do not meet

³Detailed prompt templates are presented in Appendix E.

⁴Results of instruction score are presented in Appendix G.

actual needs or lack diversity.

An effective sample for complex instruction fine-tuning should adhere to two key principles:

1. Whether the model’s response originally misaligns with constraint before it is added;
2. Whether the model’s response still misaligns with the constraint after it is added.

These constraints highlight the model’s weaknesses in handling complex instructions and require further improvement. Conversely, if a constraint does not meet these principles, it indicates that the constraint falls within the model’s current capabilities and does not require additional learning.

Therefore, we introduce constraint generation with LLM-as-judge guidance (Zheng et al., 2023), which mimics the human process of iteratively refining prompts to form complex instructions⁵. As shown in Algorithm 2, during the process of iteration, we obtain the constraints that the model fails to satisfy, which require further fine-tuning.

Algorithm 2 Iterative Instruction Refinement

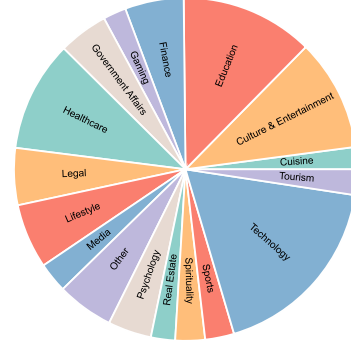
Input: Guidance model M , current model m , refined document R , initial instruction I_0 .

Output: Constraint Sets C_n and C'_n .

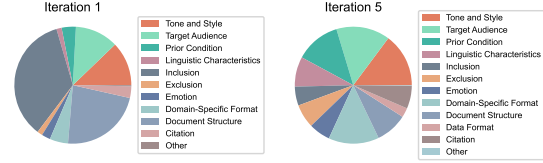
- 1: **for** $i = 1, 2, \dots, n$ **do**
- 2: Use m to generate a response A_i for the previous instruction I_{i-1} .
- 3: Leverage M as the judge, compare A_i with R to identify a new constraint c_i .
- 4: Add c_i to C_n .
- 5: Add c_i to I_{i-1} to form a new instruction I_i .
- 6: Use m to generate a response A'_i for I_i .
- 7: Leverage M as the judge, check whether A'_i satisfies constraint c_i . If not, add c_i to C'_n .
- 8: **end for**

Throughout this process, as the number of constraints increases, the model’s response also improves, making the identification of new constraints more challenging. To uncover constraints that better reflect human preferences, we use the refined document as the reference answer for the judgment process. Human-written documents inherently contain vast amounts of knowledge and formatting conventions that reflect human preferences. Therefore, the derived constraints will also align more closely with human preferences.

⁵Detailed prompt templates are presented in Appendix F.



(a) Distribution of domains



(b) Distribution of constraint types in iteration 1 and 5

Figure 3: Data statistics of AIR-10K.

Finally, the constraint set is merged into a new complex instruction. Notice two constraint sets are derived: the first set C_n satisfies Principle 1, while the second set C'_n , which includes an additional checking step, satisfies both Principle 1 and 2⁶.

While we leverage the refined document as the reference for the judgment process, it should not be used as the target for fine-tuning as in Nguyen et al. (2024), as the document is not refined with the constraints presented explicitly. Therefore, we leverage the guidance model to re-generate the response based on the combined instruction⁷.

3.3 Data Statistics of AIR-10K

We present the real-life scenario-specific domain distribution of our dataset in Figure 3(a). As can be seen, our dataset encompasses nearly 20 different domains in total, demonstrating a high degree of balance across diverse fields. Furthermore, we present the distribution of constraint types during iteration 1 and 5 in Figure 3(b). It is evident that in iteration 1, *Inclusion* and *Document Structure* constraints dominate. However, after four rounds of constraint additions, by iteration 5, the proportions of each constraint type become more uniform⁸.

We also analyze the length distributions of both instructions and responses. As shown in Figure

⁶The effect of the checking step is shown in Section 4.4.

⁷A detailed example illustrating the complete pipeline is provided in Appendix B.

⁸The constraint type definition and complete distributions across all iterations are detailed in Appendix C.

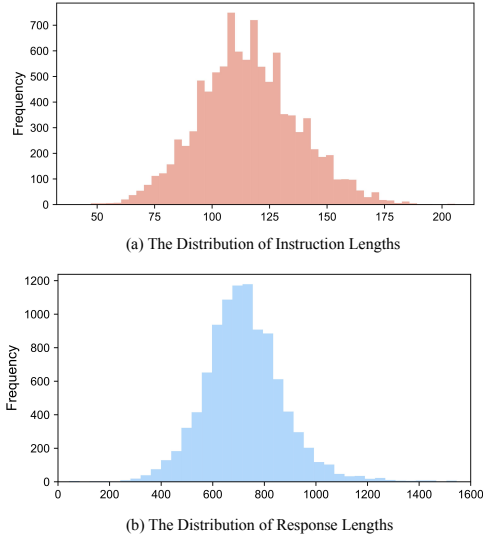


Figure 4: Length distribution of AIR-10K.

4(a) and 4(b), our instructions are of substantial information for capturing complex tasks.

4 Experiments

4.1 Set-up

Data. Following Nguyen et al. (2024), we utilize a subset of Dolma v1.7 (Soldaini et al., 2024) as the document source, which is derived from a collection of web pages and has undergone rigorous quality and content filtering to ensure data quality.

Models. We apply our method on two models, Llama-3-8B and Qwen2.5-7B, and we apply preliminary supervised fine-tuning for both models. The preliminary fine-tuning process is conducted on two general instruction datas, namely ultrachat-200k (Ding et al., 2023) and tulu-330k (Lambert et al., 2024), respectively. For the guidance model to construct the data, we rely on a larger model with the same group to ensure data quality, namely Qwen-2.5-72B-Instruct for Qwen-2.5-7B, and Llama-3-70B-Instruct for Llama-3-8B. We set the maximum number of iterations to 5.

Evaluation. We mainly conduct evaluation on two complex instruction-following benchmarks, **CFBench** (Zhang et al., 2024) and **Follow-Bench** (Jiang et al., 2023), where instructions consist of multiple constraints. We also conduct evaluations on a general instruction benchmark of **AlpacaEval2** (Dubois et al., 2024). Note that all benchmarks require GPT-4 for judgment, and we use GPT-4o-0806⁹ as the evaluator for all of them.

⁹platform.openai.com/docs/models/gp#gpt-4o

We also conduct evaluation on fundamental capability benchmarks, including math, code, and knowledge tasks, and the results are presented in Appendix A due to space limitation.

Baselines. We mainly compare our method with four groups of methods as follows:

1. **Human crafted instruction data:** This includes ShareGPT¹⁰, which is a collections of real human-AI conversations.
2. **Automatic crafted general instruction data:** This includes Self-Instruct (Wang et al., 2022), which leverages few-shot examples to self-generate simple instruction samples.
3. **Automatic rewritten complex instruction data:** This includes Evol-Instruct (Xu et al., 2023), ISHEEP (Liang et al., 2024), Muffin (Lou et al., 2023) and Conifer (Sun et al., 2024), which initiate with simple instructions and progressively construct more complex ones through rewriting or recombination.
4. **Automatic back-translated complex instruction data:** This includes Suri (Pham et al., 2024) and Crab (Qi et al., 2024), which curate the complex instructions and constraints by back-translating the pre-existing response. These methods are the most closest to our work.

Additionally, we also compare with the original back-translation (Cao et al., 2023) and back-and-forth (Nguyen et al., 2024), where IIR is skipped and initial instructions are directly used.

Note that for all constructed datasets, we sample 10k instruction-response pairs for supervised fine-tuning under the same hyper-parameters¹¹.

4.2 Main Results

As shown in Tables 1 and 2, our proposed method achieves the best performance on both complex and general instruction-following benchmarks, demonstrating its effectiveness. In contrast, automatically crafted general instruction data significantly underperform, highlighting the importance of multiple constraints in effective instruction fine-tuning. Automatic rewritten instructions also underperform, as their constructed constraints do not align with real-world practice. Additionally, automatically back-translated instructions underperform as well.

¹⁰huggingface.co/datasets/anon8231489123/ShareGPT_Vicuna_unfiltered

¹¹Detailed hyper-parameters are presented in Appendix D.

Fine-tuned on Llama-3-8B-UltraChat							
Method	CF-Bench			FollowBench		AlpacaEval2	
	CSR	ISR	PSR	HSR	SSR	LC.	Len
Baseline	0.51	0.15	0.22	41.04	57.39	8.86	1,017
back-translation	0.40	0.11	0.15	21.19	33.92	0.96	2,966
back-and-forth	0.58	0.20	0.27	44.65	61.58	10.06	1,440
ShareGPT	0.62	0.22	0.32	40.99	58.59	8.36	1,052
Self-Instruct	0.34	0.08	0.10	12.33	26.92	2.76	384
Evol-Instruct	0.57	0.22	0.28	43.58	59.21	7.15	903
MUFFIN	0.50	0.16	0.22	30.88	48.48	4.51	791
Conifer	0.57	0.22	0.28	47.06	61.32	12.81	1,084
I-SHEEP	0.53	0.17	0.23	34.26	50.28	5.41	838
Suri	0.26	0.05	0.07	3.19	3.83	0.60	29
Crab	0.56	0.18	0.25	39.92	56.83	9.05	1,192
AIR	0.61	0.24	0.31	50.69	63.89	21.00	1,813

Fine-tuned on Qwen-2.5-7B-UltraChat							
Method	CF-Bench			FollowBench		AlpacaEval2	
	CSR	ISR	PSR	HSR	SSR	LC.	Len
Baseline	0.68	0.29	0.40	47.71	64.79	10.87	836
back-translation	0.42	0.14	0.18	21.62	34.86	1.79	3,266
back-and-forth	0.63	0.24	0.34	45.33	60.39	12.59	1,480
ShareGPT	0.69	0.32	0.41	47.67	64.46	10.75	1,028
Self-Instruct	0.39	0.10	0.14	20.10	35.47	2.47	557
Evol-Instruct	0.67	0.30	0.40	46.67	63.98	8.81	964
MUFFIN	0.61	0.26	0.34	45.27	62.45	8.44	880
Conifer	0.70	0.34	0.44	51.65	65.72	19.39	1,024
I-SHEEP	0.63	0.25	0.36	41.96	59.48	6.43	996
Suri	0.31	0.07	0.10	4.55	4.85	0.94	239
Crab	0.62	0.24	0.32	41.48	59.57	9.68	1,102
AIR	0.76	0.41	0.51	59.07	71.35	32.43	1,779

Table 1: Experiment results on Llama-3-8B and Qwen-2.5-7B, with both models fine-tuned with ultrachat-200k (Ding et al., 2023). Llama-3-70B-Instruct and Qwen-2.5-72B-Instruct are used as the guidance models respectively.

Despite the constraints being derived from documents, the documents (even after refinement) suffer from misalignment and should not be direct used as the target for fine-tuning.

4.3 Data Quality Evaluation

To evaluate our dataset’s quality, we employed the Deita scorer (Liu et al., 2024), which utilizes LLM to assess complexity score for instructions and quality score for both instructions and responses. As shown in Figure 5, our approach significantly outperforms human crafted instructions, automatically crafted general instructions, and automatically rewritten complex instructions in terms of both complexity and quality scores. Notably, our method shows marginal improvements over automatic back-translation approaches like Suri and Crab, despite their use of high-quality seed datasets

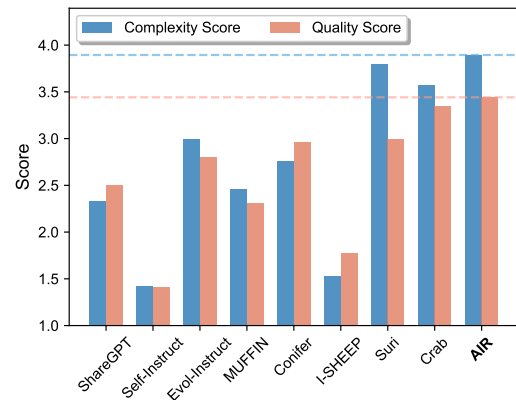


Figure 5: Comparison of averaged complexity and quality scores on different datasets.

(e.g., Alpaca GPT4 for Crab) and advanced models (e.g., GPT-4-turbo for Suri). These results validate the effectiveness of our data generation strategy.

To investigate the effect of iterative refinement,

Fine-tuned on Llama-3-8B-Tulu					
Method	CF-Bench			AlpacaEval2	
	CSR	ISR	PSR	LC.	Len
Baseline	0.50	0.15	0.20	5.20	995
back-trans	0.27	0.07	0.08	1.09	2,263
back&forth	0.47	0.14	0.19	9.04	1,337
ShareGPT	0.61	0.21	0.29	9.00	1,080
Self-Instruct	0.30	0.07	0.09	2.63	378
Evol-Instruct	0.58	0.19	0.27	18.09	991
MUFFIN	0.46	0.15	0.18	5.21	760
Conifer	0.61	0.24	0.32	7.15	903
I-SHEEP	0.49	0.16	0.19	3.11	931
Suri	0.25	0.05	0.06	0.44	151
Crab	0.56	0.19	0.27	8.55	1,221
AIR	0.68	0.28	0.38	22.00	2,097

Table 2: Experiment results on Llama-3-8B, fine-tuned with tul-330k (Lambert et al., 2024), with Llama-3-70B-Instruct as the guidance model.

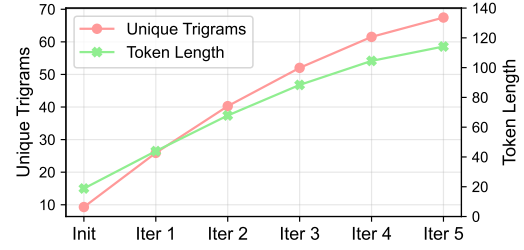
we analyze the variation of average unique trigrams and token lengths across iterations in Figure 6(a). The results demonstrate consistent increases in both instruction length and unique trigrams, indicating that newly added constraints is diverse rather than mere repetition. Furthermore, Figure 6(b) displays the evolution of complexity and quality scores throughout the iterations, showing steady improvement of data quality as the iterations progress.

4.4 Judgment Strategy for Better Constraint

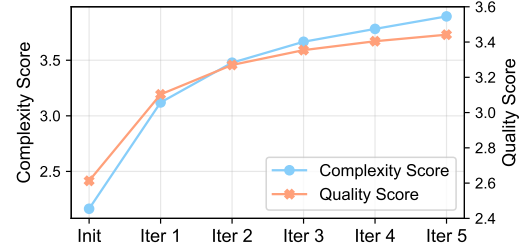
In this section, we investigate the optimal judgment strategy for constraint generation. When humans adjust prompts based on the output, they typically have a pre-expected response as the reference in mind, and constraints are issued to guide the response closer to the reference. Therefore, we compare three judgment settings: 1) No judgment, directly curate constraints; 2) Judge without document as the reference. Instead, use the guidance models’ response as the reference; 3) Judge with the refined document as the reference.

As shown in Table 3, the judgment process is essential for uncovering valuable constraints to improve the complex instruction following ability. LLM-judge can curate constraints that reflects the insufficiency of the model which requires further tuning. Moreover, using document as reference is also essential due to the limited judgment ability of the model, and human-written references aid in more targeted constraint construction.

On the other hand, the additional checking step does not improve complex instruction-following



(a) Diversity: unique trigrams and token length



(b) Complexity and quality score

Figure 6: Variation of quality indicators across iterations. *Init* represents instructions generated through the IIG step and responses from the guidance model.

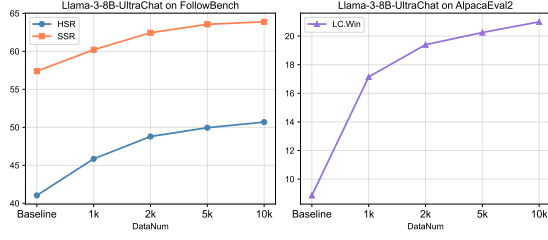
ability, as the checking step would result in fewer constraints. However, we observe improved performance on general-instruction following, indicating there exists a trade-off between general and complex instruction following abilities.

Method	FollowBench		AlpacaEval2	
	HSR	SSR	LC.	Len
<i>Results on Llama-3-8B-UltraChat</i>				
Baseline	41.04	57.39	8.86	1,017
w/o judge	47.15	62.62	19.07	1,706
judge w/o doc	51.24	63.81	20.00	1,717
judge w/ doc	52.34	64.09	19.74	1,408
w/ check	50.69	63.89	21.00	1,813
<i>Results on Llama-3-8B-Tulu</i>				
Baseline	34.91	51.76	5.20	995
w/o judge	47.59	63.60	18.32	2,067
judge w/o doc	50.62	63.69	17.02	2,842
judge w/ doc	54.16	67.52	20.45	1,639
w/ check	51.35	66.09	21.09	2,049

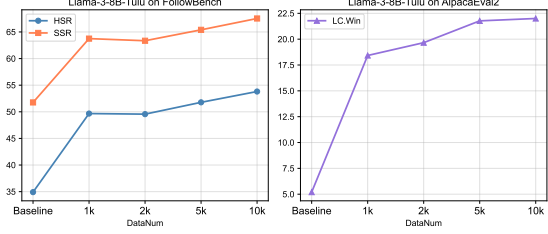
Table 3: Experiment results on Llama-3-8B models with constraints from different judgment strategies.

4.5 Influence of Iterative Judge

In this section, we investigate the effectiveness of iterative judge by examining model performance across different iterations. As shown in Table 4, the iterative judge process demonstrates clear benefits



(a) Llama-3-8B-UltraChat



(b) Llama-3-8B-Tulu

Figure 7: The variation of performance on FollowBench and AlpacaEval2 with the variation of data number.

Iteration	FollowBench		AlpacaEval2	
	HSR	SSR	LC.	Len
Baseline	34.91	51.76	5.20	995
Init	46.37	61.87	17.96	1,602
1	49.75	64.78	21.63	1,994
2	53.82	67.55	21.01	1,829
3	54.46	67.54	20.69	1,722
4	53.97	67.09	22.50	1,672
5	53.30	67.91	20.78	1,599

Table 4: Experiment results on Llama-3-8B-Tulu fine-tuned on different iterations. *Init* represents initial instructions generated through the IIG step.

compared to both the baseline and IIG step.

Specifically, we observe consistent improvements on FollowBench and AlpacaEval2 through the first two iterations. This suggests that the iterative judging process effectively identifies and incorporates increasingly sophisticated constraints that are valuable for complex instruction following. However, improvements tend to plateau after the third iteration. This could be attributed to the fact that the most critical and fundamental constraints have already been discovered in earlier iterations.

4.6 Influence of Data Quantity

In this section, we investigate the impact of data quantity on AIR’s performance. We present the results of models trained with varying amounts of data in Figure 7. As shown, performance on both general and complex instruction tasks improves with increasing data quantity. On the other hand, the model can achieve superior performance with

only 1k training samples, and the performance gains become marginal as more data is added. Therefore, in practical applications, the optimal amount of fine-tuning data can be determined based on available computational resources.

4.7 Influence of Guidance Model Size

Guid. Model	FollowBench		AlpacaEval2	
	HSR	SSR	LC.	Len
Baseline	47.71	64.79	10.87	836
14B	57.72	70.59	29.13	1,501
32B	60.06	71.97	26.39	1,309
72B	59.07	71.35	32.43	1,779

Table 5: Experiment results on Qwen-2.5-7B-UltraChat fine-tuned with different guidance model size.

In Table 5, we investigate the impact of guidance model size on AIR’s performance. We performed experiments with Qwen-2.5-7B-UltraChat as the base model, while varying the guidance model size from 14B to 72B parameters. As shown, all guidance models significantly improve instruction-following ability compared to the baseline, while larger models generally present more improvement. On the other hand, even the 14B guidance model demonstrates remarkable improvement. This scalability across different model sizes highlights the robustness and efficiency of our approach.

5 Conclusion

This paper introduces the Automatic Iterative Refinement (AIR) framework, a novel approach for generating complex instructions that better align with real-world scenarios. The framework employs an iterative refinement process guided by LLM-as-judge to generate high-quality complex constraints. We also construct a complex instruction dataset, AIR-10K, to facilitate the research and application of complex instruction following.

While previous methods for complex instruction following often introduce constraints without clear justification, it is crucial to understand what authentic complex instruction entails. In the future, we will conduct further research on the effectiveness and efficiency of complex instruction data.

Limitations

Our work has several limitations. 1) Although our evaluation includes multiple established benchmarks and metrics, including human evaluation

could further improve its credibility. Due to time and resource limitation, we have to leave this as future work. 2) Despite meticulous preprocessing, the Dolma dataset remains relatively noisy. Incorporating more high-quality documents (for example, judicial documents made public) could provide more knowledge and formality to support constraint construction. 3) The iterative nature of our framework requires multiple rounds of model inference, resulting in higher computational demands. While our ablation studies demonstrate effectiveness even with smaller guidance models and fewer samples, the computational cost remains a challenge for researchers with limited resources.

Ethical Considerations

Our data construction framework primarily leverages proprietary models such as Llama-3-70B-Instruct, which have undergone extensive preference optimization to minimize the likelihood of generating instructions that raise ethical concerns. However, large-scale web corpora—our primary data sources—are uncensored and may contain harmful or toxic content. To address this, we recommend implementing more rigorous and meticulous filtering mechanisms to proactively identify and remove such instances if possible.

While the AIR framework mainly aims to enhance models’ ability to follow complex instructions, it is important to note that some user constraints may conflict with system constraints set by developers. For example, users may request the generation of harmful or toxic content. Although our study does not specifically investigate conflicting constraints, there is a potential risk that the pipeline could prioritize user requests over developer-defined safeguards.

References

Yihan Cao, Yanbin Kang, Chi Wang, and Lichao Sun. 2023. Instruction mining: Instruction data selection for tuning large language models. *arXiv preprint arXiv:2307.06290*.

Mark Chen, Jerry Tworek, Heewoo Jun, Qiming Yuan, Henrique Ponde de Oliveira Pinto, Jared Kaplan, Harri Edwards, Yuri Burda, Nicholas Joseph, Greg Brockman, Alex Ray, Raul Puri, Gretchen Krueger, Michael Petrov, Heidy Khlaaf, Girish Sastry, Pamela Mishkin, Brooke Chan, Scott Gray, Nick Ryder, Mikhail Pavlov, Alethea Power, Lukasz Kaiser, Mohammad Bavarian, Clemens Winter,

Philippe Tillet, Felipe Petroski Such, Dave Cummings, Matthias Plappert, Fotios Chantzis, Elizabeth Barnes, Ariel Herbert-Voss, William Hebgen Guss, Alex Nichol, Alex Paino, Nikolas Tezak, Jie Tang, Igor Babuschkin, Suchir Balaji, Shantanu Jain, William Saunders, Christopher Hesse, Andrew N. Carr, Jan Leike, Josh Achiam, Vedant Misra, Evan Morikawa, Alec Radford, Matthew Knight, Miles Brundage, Mira Murati, Katie Mayer, Peter Welinder, Bob McGrew, Dario Amodei, Sam McCandlish, Ilya Sutskever, and Wojciech Zaremba. 2021. [Evaluating large language models trained on code](#). *Preprint*, arXiv:2107.03374.

Karl Cobbe, Vineet Kosaraju, Mohammad Bavarian, Mark Chen, Heewoo Jun, Lukasz Kaiser, Matthias Plappert, Jerry Tworek, Jacob Hilton, Reiichiro Nakano, Christopher Hesse, and John Schulman. 2021. [Training verifiers to solve math word problems](#). *Preprint*, arXiv:2110.14168.

Ning Ding, Yulin Chen, Bokai Xu, Yujia Qin, Zhi Zheng, Shengding Hu, Zhiyuan Liu, Maosong Sun, and Bowen Zhou. 2023. Enhancing chat language models by scaling high-quality instructional conversations. *arXiv preprint arXiv:2305.14233*.

Yann Dubois, Balázs Galambosi, Percy Liang, and Tatsunori B Hashimoto. 2024. Length-controlled alpacaEval: A simple way to debias automatic evaluators. *arXiv preprint arXiv:2404.04475*.

Hongyi Guo, Yuanshun Yao, Wei Shen, Jiaheng Wei, Xiaoying Zhang, Zhaoran Wang, and Yang Liu. 2024. Human-instruction-free llm self-alignment with limited samples. *arXiv preprint arXiv:2401.06785*.

Qianyu He, Jie Zeng, Wenhao Huang, Lina Chen, Jin Xiao, Qianxi He, Xunzhe Zhou, Jiaqing Liang, and Yanghua Xiao. 2024. Can large language models understand real-world complex instructions? In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 38, pages 18188–18196.

Dan Hendrycks, Collin Burns, Steven Basart, Andy Zou, Mantas Mazeika, Dawn Song, and Jacob Steinhardt. 2021. [Measuring massive multitask language understanding](#). In *International Conference on Learning Representations*.

Cong Duy Vu Hoang, Philipp Koehn, Gholamreza Haffari, and Trevor Cohn. 2018. Iterative back-translation for neural machine translation. In *2nd Workshop on Neural Machine Translation and Generation*, pages 18–24. Association for Computational Linguistics.

Yuxin Jiang, Yufei Wang, Xingshan Zeng, Wanjuan Zhong, Liangyou Li, Fei Mi, Lifeng Shang, Xin Jiang, Qun Liu, and Wei Wang. 2023. Follow-bench: A multi-level fine-grained constraints following benchmark for large language models. *arXiv preprint arXiv:2310.20410*.

Tom Kwiatkowski, Jennimaria Palomaki, Olivia Redfield, Michael Collins, Ankur Parikh, Chris Alberti,

Jiang. 2023. Wizardlm: Empowering large language models to follow complex instructions. *arXiv preprint arXiv:2304.12244*.

Tao Zhang, Yanjun Shen, Wenjing Luo, Yan Zhang, Hao Liang, Fan Yang, Mingan Lin, Yujing Qiao, Weipeng Chen, Bin Cui, et al. 2024. Cfbench: A comprehensive constraints-following benchmark for llms. *arXiv preprint arXiv:2408.01122*.

Wayne Xin Zhao, Kun Zhou, Junyi Li, Tianyi Tang, Xiaolei Wang, Yupeng Hou, Yingqian Min, Beichen Zhang, Junjie Zhang, Zican Dong, et al. 2023. A survey of large language models. *arXiv preprint arXiv:2303.18223*.

Lianmin Zheng, Wei-Lin Chiang, Ying Sheng, Siyuan Zhuang, Zhanghao Wu, Yonghao Zhuang, Zi Lin, Zhuohan Li, Dacheng Li, Eric Xing, Hao Zhang, Joseph E. Gonzalez, and Ion Stoica. 2023. [Judging LLM-as-a-judge with MT-bench and chatbot arena](#). In *Thirty-seventh Conference on Neural Information Processing Systems Datasets and Benchmarks Track*.

Tianyu Zheng, Shuyue Guo, Xingwei Qu, Jiawei Guo, Xinrun Du, Qi Jia, Chenghua Lin, Wenhao Huang, Jie Fu, and Ge Zhang. 2024a. Kun: Answer polishment for chinese self-alignment with instruction back-translation. *arXiv preprint arXiv:2401.06477*.

Yaowei Zheng, Richong Zhang, Junhao Zhang, Yanhan Ye, and Zheyang Luo. 2024b. [LlamaFactory: Unified efficient fine-tuning of 100+ language models](#). In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 3: System Demonstrations)*, pages 400–410, Bangkok, Thailand. Association for Computational Linguistics.

Chunting Zhou, Pengfei Liu, Puxin Xu, Srinivasan Iyer, Jiao Sun, Yuning Mao, Xuezhe Ma, Avia Efrat, Ping Yu, Lili Yu, et al. 2024. Lima: Less is more for alignment. *Advances in Neural Information Processing Systems*, 36.

A Impact on Fundamental Capabilities

Method	MMLU	CQA	NQ	HE	GSM	AVG
<i>Results on Llama-3-8B-UltraChat</i>						
Baseline	64.00	72.97	29.61	30.49	57.47	50.90
AIR	61.64	73.63	30.54	29.88	54.59	50.05
<i>Results on Qwen-2.5-7B-UltraChat</i>						
Baseline	73.64	82.39	25.68	52.20	81.65	63.11
AIR	73.35	82.56	25.76	55.49	84.38	64.30
<i>Results on Llama-3-8B-Tulu</i>						
Baseline	65.43	79.44	32.22	50.61	64.14	58.36
AIR	64.95	79.92	34.62	50.85	63.70	58.80

Table 6: Experiment results on fundamental capabilities.

Previous methods have shown LLMs may suffer from capability degradation during alignment (Ouyang et al., 2022). To evaluate this concern, we tested our AIR method on MMLU (Hendrycks et al., 2021), CommonsenseQA (CQA) (Talmor et al., 2019), Natural Questions (NQ) (Kwiatkowski et al., 2019), HumanEval (HE) (Chen et al., 2021), and GSM8K (GSM) (Cobbe et al., 2021). In Table 6, our method does not have a negative impact on fundamental capabilities. For Qwen-2.5-7B-UltraChat and Llama-3-8B-Tulu, our method even improves the average performance by 1.19 and 0.44 points, respectively. This indicates that instruction constructed from documents with evenly sampled distributions also present even distribution, which would not lead to catastrophic forgetting of fundamental capabilities.

B Case Study for Complete Pipeline

This section presents a detailed end-to-end demonstration of our pipeline in Figure 8. The case study provides a thorough walkthrough of each stage in our instruction generation and refinement process.

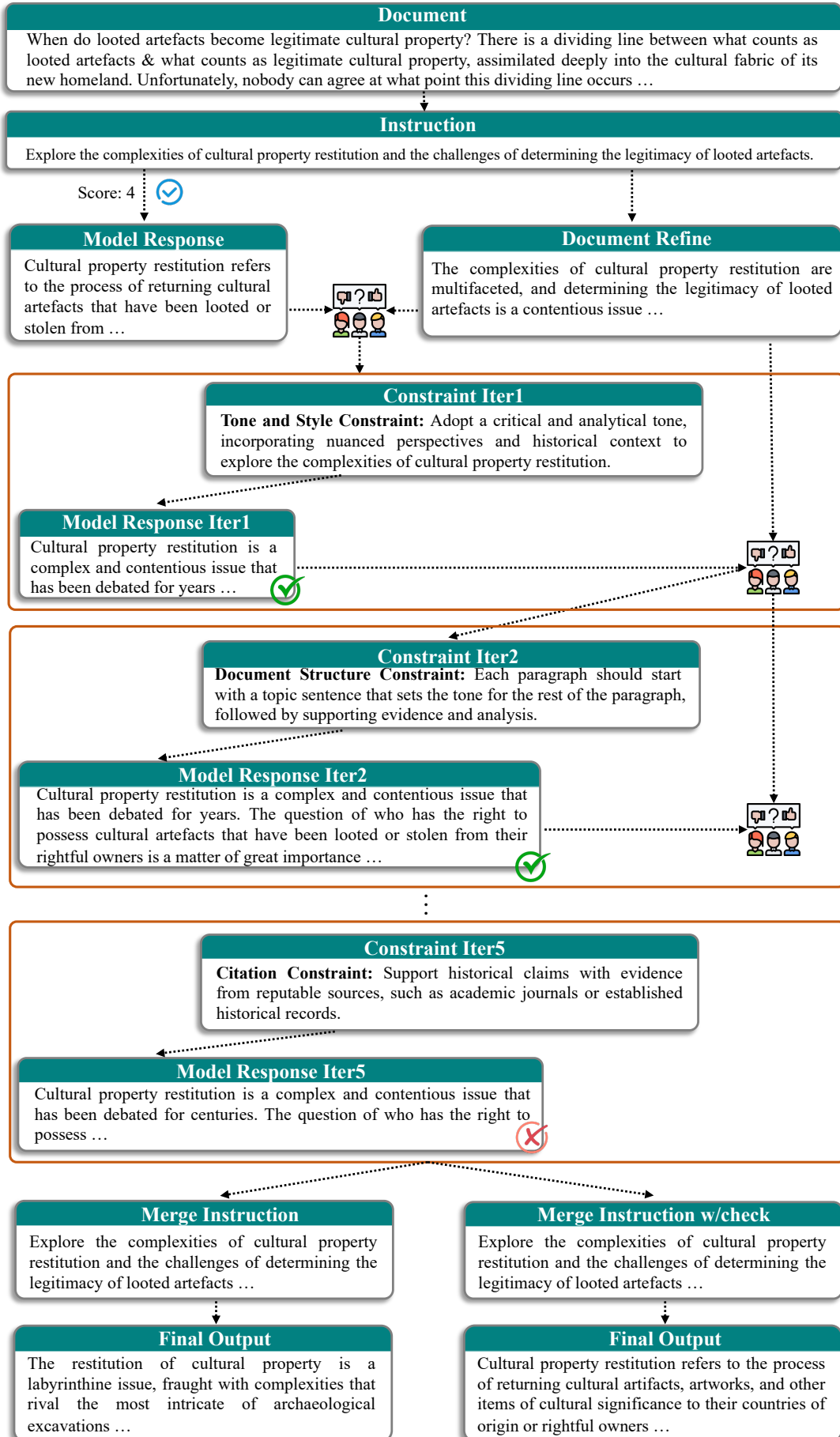


Figure 8: End-to-End Pipeline Implementation Example.

C Constraint Type Taxonomy and Distribution Analysis

This section provides a detailed classification of constraint types, as defined in Table 7. Additionally, we present a comprehensive analysis of constraint type distribution patterns observed across five iterative refinement rounds, as visualized in Figure 9.

Constraint Type	Description
Data Format	The generated content should conform to specific data structure formats, such as JSON, Markdown, Table, CSV, etc.
Document Structure	The generated content should follow specific document organization patterns, including Numbered lists (1, 2, 3 or I, II, III), Bullet points (•, -, *), Custom templates with predefined sections, Tables, Headers, etc.
Domain-Specific Format	Content must follow strict format rules for different industries
Inclusion	Identify and list the specific elements or information that should be included in the generated content
Exclusion	Identify and list the specific elements or information that should not be included in the generated content
Citation	The generated content should include citations to sources, providing reliable sources and literature support; follow specific citation formats or reference styles
Prior Condition	When a specific intention is met, a particular process should be followed to perform an operation or output specific content
Target Audience	The generated content should target a specific audience, which affects the terminology used, the level of detail provided, and the complexity of the content
Tone and Style	The generated content should adopt a specific tone and style, such as formal, polite, academic, concise, literary, romantic, or sci-fi
Emotion	The generated content should express a specific emotion or mood, such as ensuring the content is positive, inspiring, or empathetic
Linguistic Characteristics	Use specific linguistic features, such as metaphors, personification, and other rhetorical devices
Multilingual	The generated content should be written in a specific language, such as English, Mandarin, or Spanish

Table 7: Types of Constraints Used in Dataset Generation.

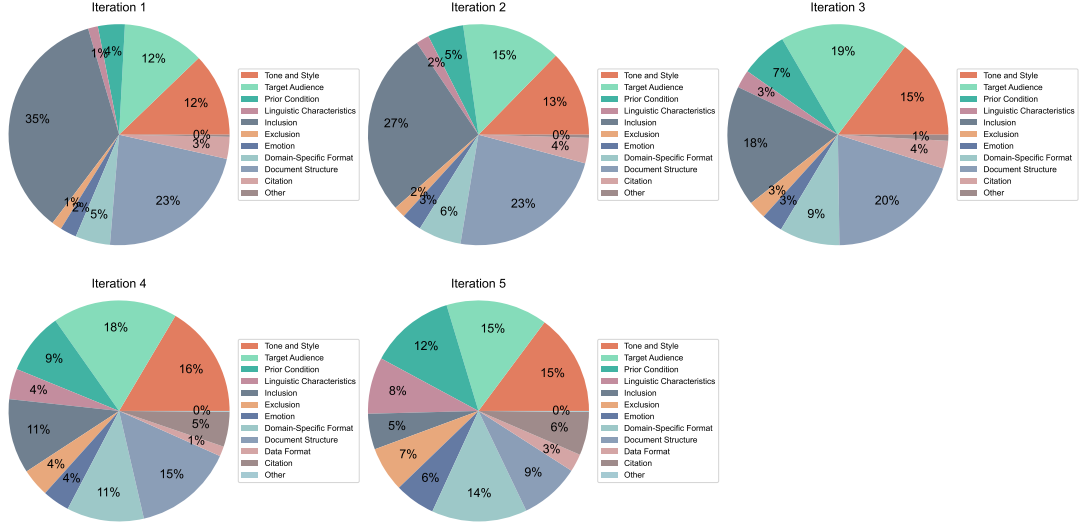


Figure 9: Distribution of constraint types across all iterations.

D Model Training Hyper-parameters

This section details our model training configuration based on the LlamaFactory (Zheng et al., 2024b) framework. We employed Supervised Fine-Tuning (SFT) with hyper-parameters as outlined in Table 8.

Configuration	Llama-3-8B	Qwen-2.5-7B
max length	4096	4096
learning rate	1e-5	1e-5
scheduler	cosine decay	cosine decay
training epochs	3	3
batch size	64	64
flash-attn	fa2	fa2
numerical precision	bf16	bf16
ZeRO optimizer	stage 2	stage 2

Table 8: Hyper-parameters for Supervised Fine-Tuning.

E Prompts for Initial Instruction Generation

This section presents the prompts used in our data generation pipeline in Initial Instruction Generation step. These prompts serve different purposes in our methodology, from initial instruction generation through back-translation (Figure 10) to document refining (Figure 11) and instruction scoring (Figure 12).

F Prompts for Iterative Instruction Refinement

This section presents the prompts used in our data generation pipeline in Iterative Instruction Refinement step. These prompts serve different purposes in our methodology, from constraint generation (Figure 13), constraint verification (Figure 14), and finally combines all elements into refined instructions (Figure 15).

G Instruction Score Examples

This section presents a comprehensive analysis of instruction quality through representative examples. As illustrated in Figure 16, we provide a diverse set of instructions spanning the entire quality spectrum (scores 1-5). Each score category is exemplified by five carefully selected cases, where score 1 represents basic quality and score 5 demonstrates exceptional quality.

Please generate a single instruction that would lead to the given text as a response.

- The instruction should not be a question. Instead, it should be a more general task.
- The instruction should not cover all details of the response. Instead, it should be concise and only focus on the main aspect.

Please generate your instruction based on the text.

Text: {document}

Instruction:

Figure 10: Prompt for generating initial instructions through back-translation.

You are a professional editor. Given an instruction and an original response, your task is to improve the response while ensuring it aligns well with the instruction.

The improvement should focus on:

- Better alignment with the instruction
- Enhanced clarity and coherence
- Aligns with AI assistant response style
- Maintaining the core message while improving expression.

Now, this is your task. Please directly present your modifications, without using ANY headings or prefixes.

Instruction: {instruction}

Original Response: {document}

Enhanced Response:

Figure 11: Prompt for refining document content.

Review the user's instruction using the additive 5-point scoring system described below. Points are accumulated based on the satisfaction of each criterion:

Award 1 point for containing a basic question or task.

Add 1 point if the instruction can be addressed using the language model's existing knowledge base without requiring external resources or current event information.

Add 1 point if the instruction does NOT require analyzing specific texts, documents, or specific person's perspective.

Add 1 point if the instruction effectively communicates both the core question and key preferences, demonstrating clear intent while being self-contained.

Add 1 point if the instruction pertains to general topics or advice that are widely applicable and within the common knowledge base, rather than requiring specialized or niche information about specific individuals or events.

After examining the instruction:

- Briefly justify your total score, up to 100 words.

- Conclude with the score using the format: "Score: <total points>/5"

Example 1:

Instruction: What was the impact of Gary Gilmour's career and his life in the years following his cricketing career?

Answer: The instruction poses a basic question about Gary Gilmour's impact after his cricketing career (1 point). It can be answered using the language model's existing knowledge (1 point). It doesn't require analyzing specific texts, documents, or a specific person's perspective (1 point). The question is clear, self-contained, and demonstrates clear intent (1 point). However, since it involves information about a specific individual, which requires specialized or niche knowledge, the last point is not awarded.

Score: 4/5

Example 2:

Instruction: What's the most helpful advice you have for students who are awaiting their college admission decision?

Answer: The instruction asks for the most helpful advice for students awaiting their college admission decisions, which is a basic question (1 point). It can be answered using the language model's existing knowledge (1 point). It does not require analyzing specific texts, documents, or a specific person's perspective (1 point). The question is clear, self-contained, and demonstrates clear intent (1 point). It pertains to a general topic that is widely applicable and within the common knowledge base (1 point).

Score: 5/5

...

This is your task:

Instruction: {instruction}

Answer:

Figure 12: Prompt for scoring initial instructions.

Based on the provided instruction, I obtained Output1 and Output2 from two different models. Please analyze both outputs carefully to identify the MOST CRITICAL constraint type that Output2 needs to improve to match Output1's quality.

Available Constraint Types:
{constraints_type}

Task Requirements:

1. [Analysis] Compare Output1 and Output2 to identify differences
2. [Selection] Choose the SINGLE most critical constraint type where Output2 shows the biggest gap
3. [Constraint] Create ONE specific constraint that:
 - Addresses ONLY the selected constraint type
 - Exists in Output1 but is missing in Output2
 - Is written in a clear and concise sentence (10-20 words)
 - Avoids references to "Output1" or "Output2"
4. If no significant differences match the available types, specify "None"

Required Response Format:

****Analysis****: [Brief analysis]
****Selected Type****: [Single most critical type]
****Constraint****: [ONE specific constraint]

Context:

#Instruction#
{instruction}

#Output1#
{document_refine}

#Output2#
{model_response}

#Your Response#

Figure 13: Prompt for generating constraints based on judge.

I want you to act as a quality evaluator. You need to evaluate the model answer by combining [User Instructions], [Model Answer], and [Evaluation Criteria] and score with 0-3.

Specifically, [Model Answer] is the response to [User Instructions], and [Evaluation Criteria] defines the points that the model answer should satisfy and needs to be evaluated. You need to strictly score the [Model Answer] according to each evaluation point in [Evaluation Criteria].

Scoring Rules:

- Score 0: Does not meet the evaluation criteria
- Score 1: Meets the evaluation criteria with acceptable response
- Score 2: Meets the evaluation criteria with high quality and comprehensive response
- Score 3: Meets the evaluation criteria with exceptional and flawless response

Output format: 1. Strictly output one line at a time according to the order of evaluation points in [Evaluation Criteria], with lines separated by "\n\n";

2. Each line first outputs the corresponding content in [Evaluation Criteria], then uses "\t" to separate, and outputs the corresponding score (0-3) after it;

3. Please output your evaluation directly without any other content;

4. Note that if a criteria states like "do not include X", the score should be 0 if the answer includes X.

[User Instructions]: {instruction}

[Model Answer]: {model_response}

[Evaluation Criteria]: {constraints}

[Your Evaluation]:

Figure 14: Prompt for verifying model responses against constraints.

You are a skilled writing specialist who excels at blending different elements into cohesive, natural-sounding instructions.

Fusion guidelines:

- Consolidates overlapping constraints and resolves any conflicts
- Craft a cohesive instruction that naturally integrates ALL appropriate constraints
- AVOID expanding constraints

{few_shot}

Now it's your turn. Please merge the following input and constraints, do not output anything else, including response to the merged instruction:

[Original Input]
{instruction}

[Original Constraints]
{constraints}

[Merged Instruction]

Figure 15: Prompt for combining instructions with constraints.

Instruction	
Score: 1 	Conduct an in-depth interview with a standout college basketball player about their career. Write a weekly community newsletter for a small town, covering local news, and opinions. Write a personal account of a company-wide cost reduction. Write a scene where Amato meets with Raith to discuss a new. Write a profile article about a local church and its leadership.
Instruction	
Score: 2 	Conduct an in-depth interview with a professional chef about their career path. Write a review of a recent episode of the TV show Shameless. Review and compare alternative Instagram growth services to Hyper Vote. Provide a progress update on the Pensions Dashboards Programme. Write a personal tribute to a Nigerian politician who has made a positive impression on you.
Instruction	
Score: 3 	Draft a court opinion for the appeal of a grand theft conviction. Write a feature article about the Pac-12's dominance in college athletics. Create an informed consent document for a research study. Write a film review of Top Gun: Maverick. Write a critical analysis of the movie Prometheus, exploring its themes.
Instruction	
Score: 4 	Compile a comprehensive guide to natural remedies for treating yeast infections in women. Write a spiritual reflection on the limitations of human capacity. Write a comprehensive guide about how doctors inform patients about cancer diagnosis. Write a sports article about a football team's creative adjustments due to injuries. Write a comprehensive guide for international students on pursuing MBA program in the UK.
Instruction	
Score: 5 	Write a comprehensive guide to understanding the different types of real estate. Develop a guide for starting a meditation habit. Write a guide on securing valuables and property at home. Develop a guide on leveraging social media stories for business growth. Write an article about the mental health benefits of owning a pet.

Figure 16: Examples of instructions at different score levels (1-5), where each score level is illustrated with five representative cases. Score 1 represents the lowest quality while score 5 represents the highest quality.