

Pragyaan: Designing and Curating High-Quality Cultural Post-Training Datasets for Indian Languages

Neel Prabhanjan Rachamalla, Aravind Konakalla, Gautam Rajeev,
Ashish Kulkarni, Chandra Khatri, and Shubham Agarwal

Krutrim AI, Bangalore, India

Contact: {neel.rachamalla1, ashish.kulkarni, shubham.agarwal1}@olakrutrim.com

Abstract

The effectiveness of Large Language Models (LLMs) depends heavily on the availability of high-quality post-training data, particularly instruction-tuning and preference-based examples. Existing open-source datasets, however, often lack multilingual coverage, cultural grounding, and suffer from task diversity gaps that are especially pronounced for Indian languages. We introduce a human-in-the-loop pipeline that combines translations with synthetic expansion to produce reliable and diverse Indic post-training data. Using this pipeline, we curate two datasets: *Pragyaan-IT* (22.5K) and *Pragyaan-Align* (100K) across 10 Indian languages covering 13 broad and 56 sub-categories, leveraging 57 diverse datasets. Our dataset protocol incorporates several often-overlooked dimensions and emphasize task diversity, multi-turn dialogue, instruction fidelity, safety alignment, and preservation of cultural nuance, providing a foundation for more inclusive and effective multilingual LLMs.

1 Introduction

Recent developments around Large Language Models (LLMs) (Touvron et al., 2023; Grattafiori et al., 2024; Abidin et al., 2025; Guo et al., 2025) have demonstrated that post-training data, comprising both instruction-tuning and preference data, plays a critical role in enhancing model alignment, task generalization, and usability (Ouyang et al., 2022; Bai et al., 2022b; Chung et al., 2022). Particularly in a multilingual and multicultural landscape, like India, the availability of high-quality, culturally grounded post-training data is crucial to address performance gaps in low-resource languages that often arise from the scarcity of relevant and representative training data (Joshi et al., 2020).

While several open-source datasets for post-training (Longpre et al., 2023; Wang et al., 2022b; Bercovich et al., 2025a) exist, they are predominantly English-centric and often suffer from limita-

tions such as inconsistent quality, restricted coverage, insufficient task complexity, and limited multilingual coverage. These challenges extend to Indic post-training data as well, focus of our work.

Direct translations of existing English post-training datasets are prone to translation biases, errors (Hartung et al., 2023; Savoldi et al., 2021; Muennighoff et al., 2022) and loss of cultural grounding (Wang et al., 2022a; Pudjiati et al., 2022). For instance, a prompt like “Tell me about a small herb to plant in backyard” might yield Western herbs such as *thyme* or *rosemary*, whereas Indian users would expect culturally familiar options like *tulsi* (holy basil), *pudina* (mint), or *curry leaves*. Similarly, when asked “What is a good comfort meal for a rainy day?”, English-centric answers such as *tomato soup* or *grilled cheese* overlook Indian preferences like *masala chai with pakoras* or *khichdi with ghee*. Even in wellness contexts, “Recommend a workout routine for beginners” may default to *squats and push-ups*, neglecting practices like *Surya Namaskar* or *yoga exercises*. Such mismatches highlight the need for post-training data that reflects not just language, but also local traditions and cultural context.

The recent popularity of LLM-based synthetic data generation (Wang et al., 2022c) for creating post-training datasets, while promising, still suffers in quality due to linguistic inaccuracies, grammatical inconsistencies, and reduced fluency, especially in multilingual settings, that could degrade the performance of models trained on them. Moreover, the lack of fine-grained control over output complexity and the potential for hallucination can lead to the generation of low-quality, unreliable data.

With the aim of addressing these gaps, we present an approach to curate high-quality post-training datasets, especially in multilingual settings. Our curation approach combines the above techniques with post-hoc manual editing, leading to a scalable human-in-the-loop pipeline, with spe-

cific focus on several aspects of quality, like task coverage, multilingual representation, task complexity, culture, multi-turns, reasoning, and others. We leverage our approach to curate high-quality Indic post-training datasets: *Pragyaan-IT* comprising 22.5K instruction tuning examples and *Pragyaan-Align*, a dataset of 100K preference examples in 10 Indian languages covering 56 task categories. Our contributions could thus be summarized as follows:

- We present a scalable pipeline for curating high-quality post-training data. Our approach emphasizes a human-in-the-loop (HITL) that is more efficient, reliable and ensures higher quality than direct synthetic generation or translation of existing English datasets.
- We introduce high-quality, manually curated, and culturally-inclusive post-training *Pragyaan* dataset series, consisting of 1) 22.5K *Pragyaan-IT* and 2) 100K *Pragyaan-Align*, designed for aligning LLMs to the diverse Indian cultural context.
- Our dataset includes a broad spectrum of instruction-following tasks with varying levels of complexity, ensuring the resulting models can handle a wide range of real-world scenarios. We provide a detailed analysis of the dataset’s characteristics, including its language distribution and domain representation, also showcasing its suitability for robust instruction-following capabilities through a small-scale pilot experiment.

2 Related Work

Post-training is a key step in aligning large language models (LLMs) with human intent, commonly achieved through instruction tuning (Wei et al., 2022a) and preference tuning (Bai et al., 2022a) datasets, constructed in several ways. Task template based resources such as Flan 2021 (Wei et al., 2022a), Flan 2022 (Longpre et al., 2023), and P3 (Sanh et al., 2022) adapt NLP datasets into instruction–response format. Human-authored datasets like Open Assistant (Köpf et al., 2023), Dolly (Conover et al., 2023), and LIMA (Zhou et al., 2023) demonstrate the value of curated instructions but face scalability challenges. To overcome this, synthetic generation approaches leverage LLMs to expand from small human-annotated *seeds*, with efforts such as Self-Instruct (Wang et al., 2022c), Alpaca (Taori et al., 2023), and Guanaco (Joseph Cheung, 2023), often distilling

knowledge (Hinton et al., 2015) from stronger teacher LLM models. Advanced pipelines like Evol-Instruct (Xu et al., 2023) iteratively increase instruction complexity, while later works extend these methods to reasoning and code generation (Luo et al., 2023; Gunasekar et al., 2023). Complementing these, user-contributed datasets such as InstructionWild (Ni et al., 2023) and ShareGPT¹ provide naturally occurring conversational data, and Unnatural Instructions (Honovich et al., 2022) show how seed tasks can be scaled into diverse synthetic corpora. Subsequent work expanded into specialized domains, including dialogue systems (Köpf et al., 2023), structured knowledge grounding (Xie et al., 2022), and chain-of-thought reasoning (Wei et al., 2022b; Kim et al., 2023). More recently, the Magpie dataset (Xu et al., 2024) introduced a fine-grained taxonomy spanning creative writing, math, role-playing, planning, and data analysis, emphasizing the importance of broad coverage in post-training resources. Preference datasets such as UltraFeedback (Cui et al., 2023) and Tulu3 (Lambert et al., 2025) comprises human and synthetic preference pairs for LLM alignment. Building on these advances, we construct our approach and dataset tailored to Indian languages and cultural contexts leveraging manual annotations in complement with synthetic generation.

While large-scale post-training datasets have become increasingly available, they remain predominantly English-centric, with limited coverage for other languages. A few exceptions incorporate some proportions of multilingual data (Köpf et al., 2023; Longpre et al., 2023; Muennighoff et al., 2023; Zhuo et al., 2024; Nguyen et al., 2023), but they remain limited in cultural and linguistic diversity compared to English resources. Prior efforts to extend post-training resources beyond English have typically followed three strategies: (1) translating English datasets into additional languages (Li et al., 2023a; Khan et al., 2024), (2) generating template-based datasets (Yu et al., 2023b; Gupta et al., 2023), and (3) manually curating instruction datasets in non-English languages (Li et al., 2023b; Wang et al., 2022d). Amongst these, template-based efforts such as xP3 (Muennighoff et al., 2022) extend the P3 taxonomy with 28 multilingual datasets. However, xP3 relies on uniform templates across languages, leading to limited task diversity and frequent repetition. Translation-based

¹<https://sharegpt.com/>

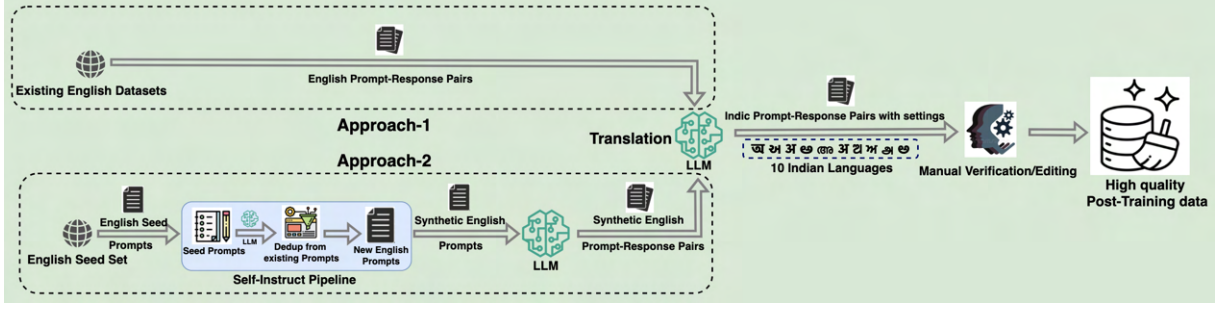


Figure 1: Workflow for building Indian language post-training data: English prompts are either translated or expanded via modified self-instruct pipeline to generate synthetic prompts. In both cases, responses are then produced with an LLM, translated into one of the 10 Indian languages, and manually refined (Section 3).

approaches face similar limitations, such as Bactrian (Li et al., 2023a), which translated Alpaca (Taori et al., 2023) and Dolly (Conover et al., 2023) into 52 languages. In contrast, our work introduces human-edited datasets across 10 Indian languages, addressing issues of redundancy and cultural grounding while providing a more diverse and representative resource for multilingual alignment.

3 Methodology

We present our multi-stage post-training dataset creation process that encompasses a variety of task categories at both broad and fine-grained levels, critical settings (complexity, interaction depth, constraints, safety, indian context, thinking trails) and leverages human annotators alongside curated data sources to ensure quality and coverage. While our approach itself is generic, we discuss how we used it to create high-quality Indic post-training datasets that we collectively refer to as *Pragyaan-IT* and *Pragyaan-Align*.

3.1 Data Construction Approaches

We employ two complementary approaches (Figure 1) that both combine translation and synthetic generation with post-hoc manual editing to ensure linguistic accuracy, fluency, and cultural appropriateness in our datasets.

3.1.1 Approach 1: Translation with Human Refinement

Here, we directly source English prompt-response pairs from existing English datasets (more details later in Section 3.5).

Prompts: We begin with English prompts which are first translated into Indic languages using an LLM, then refined by human annotators. During verification, annotators correct linguistic errors,

improve readability, and adapt expressions where needed to reflect Indian cultural norms. This results in two categories, i.e. 1) *Indic Generic Prompts*: direct translations of the English originals. 2) *Indic Context Prompts*: culturally adapted and edited versions incorporating Indian references and contexts by human annotators.

Responses: Corresponding English responses undergo a similar pipeline independently, with LLM-based translation into Indic languages, followed by human editing for grammar, relevance, length, and cultural appropriateness. Thus, we have 1) *Indic Generic Responses*: literal translations of the English outputs. 2) *Indic Context Responses*: refined versions adapted to Indian discourse norms.

3.1.2 Approach 2: Synthetic Expansion with Human Refinement

This approach introduces an additional intermediate stage of synthetic prompt expansion in English.

Prompts: Starting with a seed set of English prompts (sourced or created), we use the Self-Instruct pipeline (Wang et al., 2022b) to iteratively expand this set into a larger synthetic pool. While in the original pipeline they generate new prompts for classification and non-classification types via different strategies, in our adaptation, we use the same prompt template for both the cases. The resulting synthetic English prompts are then translated into Indic languages with LLMs and refined by human annotators to ensure correctness, clarity, and cultural grounding. This yields 1) *Synthetic Indic Generic Prompts*: literal translations of synthetic English prompts. 2) *Synthetic Indic Context Prompts*: culturally enriched translations.

Responses: For each synthetic English prompt, we generate English responses using an LLM independently. These are then translated into Indic

languages and refined through human editing. Annotators correct factual or linguistic errors, polish style, and, when appropriate, enrich with cultural nuances. This process produces 1) *Synthetic Indic Generic Responses*: faithful translations of the English responses. 2) *Synthetic Indic Context Responses*: culturally adapted versions aligned with the Indian local context.

While our first approach ensures fidelity through the translation of existing English datasets, it remains constrained in both scope and diversity. The second approach complements this by introducing synthetic expansion prior to translation, enabling broader task coverage, richer cultural representation, and greater scalability with reduced dependency on large English resources. Together, the two approaches strike a balance between reliability and diversity, yielding multilingual datasets that are both high-quality and contextually rich.

3.2 Task Categories

Building on the design principles of several existing datasets, we curate a broad set of task categories that combine core language tasks such as reasoning (limited to CoT and self-thinking), inference, natural language understanding (NLU) and generation, question answering (QA), dialogue and interaction, information extraction, mathematics, coding, function calling, and instruction following, while also extending into culturally grounded domains like Indian states, religions, geo-political questions, etc. To promote robustness and responsible deployment, we additionally include safety and non-compliance, Indian contentious content, and self-identity tasks. Collectively, these categories establish a structured yet comprehensive framework that spans diverse sub-categories, enabling richer and more inclusive post-training curation. A detailed breakdown of sub-categories is provided in Table 1.

3.3 Task Settings

For each task category, we additionally define several task settings that encourage diversity of prompt complexity, interaction depth, instruction-following, safety considerations, cultural grounding, and explicit reasoning trails. We provide systematic descriptions of these settings next.

3.3.1 Complexity

We categorize tasks by complexity to ensure models are trained for both simple and challenging scenarios, with two primary levels: 1) *Easy* prompts

are direct and clearly defined, usually requiring minimal reasoning (e.g. a single factual or descriptive query). 2) *Hard* prompts feature greater structural complexity, often embedding multiple sub-questions within a single query, requiring nuanced reasoning and fine-grained understanding. Importantly, complexity is defined within each task category, enabling fair assessment across heterogeneous task types (see Figures 6–11 in Appendix).

3.3.2 Multi-Turn Interactions

Multi-turn settings capture tasks where contextual continuity is critical, such as dialogue, planning, or role-play. These scenarios require models to maintain memory of prior turns while generating coherent and adaptive responses. We consider three levels of interaction depths: 1) *Single-turn (1 turn)*: A response to an isolated prompt; 2) *Short multi-turn (3 turns)*: Three back-and-forth exchanges, ensuring local continuity; 3) *Extended multi-turn (5 turns)*: Five exchanges, for long-range memory and coherence in extended conversations (e.g. planning a festival with evolving constraints).

3.3.3 Instruction Following

We categorize instruction-following into three levels, defined by the number and type of constraints imposed on the response such as “*answer in 100 words*”, “*respond in json format*”, etc. (Figure 12 in Appendix). This ensures coverage of tasks that range from loosely guided prompts to highly structured outputs. Different combinations of these constraints are applied depending on the nature of the prompt: 1) *Simple instruction following*: prompts include minimal or no explicit constraints on the format or content of the response; 2) *Medium instruction following*: prompts introduce two to three explicit constraints, requiring the model to accommodate multiple conditions at once; 3) *Complex instruction following*: prompts impose several simultaneous constraints, demanding precise control and structured outputs.

3.3.4 Safety

We define safety settings to ensure that models behave responsibly when faced with sensitive, controversial, or potentially harmful content in a real-world setting. This dimension helps guide appropriate responses while maintaining ethical standards. We include 1) *Safe*: prompts are neutral and non-controversial, allowing the model to provide direct answers without ethical or policy concerns; 2) *Non-*

Broad Category	Sub Categories	Broad Category	Sub Categories
Reasoning & Inference	Non-Math Reasoning Math Reasoning Code Reasoning Indian Relationships Inference Data Analysis	Natural Language Understanding & Generation	Named Entity Recognition Text Classification Grammar Correction Translation Creative Writing Paraphrase Identification Paraphrase Generation Text Summarization Headline Generation Question Generation Sentiment Analysis
			Multi Turn Conversation Role Playing Advice Seeking Planning Brainstorming
Question Answering	General Question Answering Fact Check	Interaction & Dialogue	Sanskrit Festival Greetings Sanskrit Auspicious Day / Other Occasions Sanskrit Subhashitas (Quotes) Sanskrit Captions and Mottos Sanskrit Person's Name Sanskrit Building / Institution / Company Name Sanskrit Product Name
Information Extraction	Information Seeking Indian Cultural Context Comprehension	Sanskrit Cultural & Creative Usage	Code Generation Code Debugging Code Editing Code Explanation Code Translation Unit Test Generation Code Theory Code Review Repository level Code Generation
Mathematics	Math QA Math Instruction Tuning Math Proofs	Coding	Instruction Following
Function Calling	Function Calling	Instruction Following	Indian Geo Political Indian Politicians Indian States Indian Languages Indian Religions
Safety & Non-Compliance	Safety & Non-Compliance	Indian Contentious Questions	
Self-Identity	Model-name Person-based		

Table 1: Taxonomy of NLP task categories and sub-categories for creating post-training datasets in Section 3.2.

safe: prompts involve sensitive or harmful material, where the model is expected to either refuse politely (e.g., “Sorry, I cannot assist with that ...”) or generate a safe response.

3.3.5 Thinking Trails

We define ‘thinking trail’ settings to capture the role of explicit reasoning in model responses, ensuring that outputs range from direct answers to more reflective reasoning styles. 1) *Normal*: direct response generation without intermediate reasoning traces; 2) *Chain-of-Thought (CoT)*: step-by-step reasoning (Wei et al., 2022b) articulated explicitly before the final answer; 3) *Self-Thinking*: inspired by recent “deep thinking” paradigms (Guo et al., 2025; Bercovich et al., 2025b; Abdin et al., 2025), where models produce more elaborate, self-reflective reasoning trails prior to the final response.

3.3.6 Indian Cultural Context

Given the centrality of Indic languages and cultural alignment in our framework, we explicitly model contextual grounding through three progressively richer levels. 1) *IC-1* represents generic prompts leading to generic responses, with no explicit India related anchoring (e.g., Prompt: “Suggest some breakfast items.”, Response: “Pancakes, cereal, toast, scrambled eggs.”). Such responses

are accurate but remain culturally neutral, with no particular alignment to cultural settings. 2) *IC-2* represents generic prompts that nonetheless yield Indic-grounded responses. For instance, the same prompt above, in this setting, would elicit responses such as “Idli, dosa, paratha, poha”, which are the most popular breakfast items in India. 3) *IC-3* involves prompts that are themselves explicitly Indic, thereby eliciting fully Indic-based responses. For example, the prompt itself mentions “Suggest some Indian breakfast items” with a similar response as in the IC-2 setting. This setting encourages grounding and diversity of responses with respect to Indian cultural context that is required for training Indic focused LLMs.

3.4 Human-In-The-Loop (HITL) Refinement

While synthetic generation and automated translation provides the backbone of our dataset creation pipeline, human annotators play an equally central role in shaping its final form. Each prompt–response pair, once generated and assigned a configuration of settings (e.g., *easy*, *1-Turn*, *Simple-IF*, *Safe*, *IC-3*, *Normal* (*No Thinking Trails*)), enters a stage of manual intervention where annotators act not merely as reviewers, but as curators of the data. If a pair does not fully align

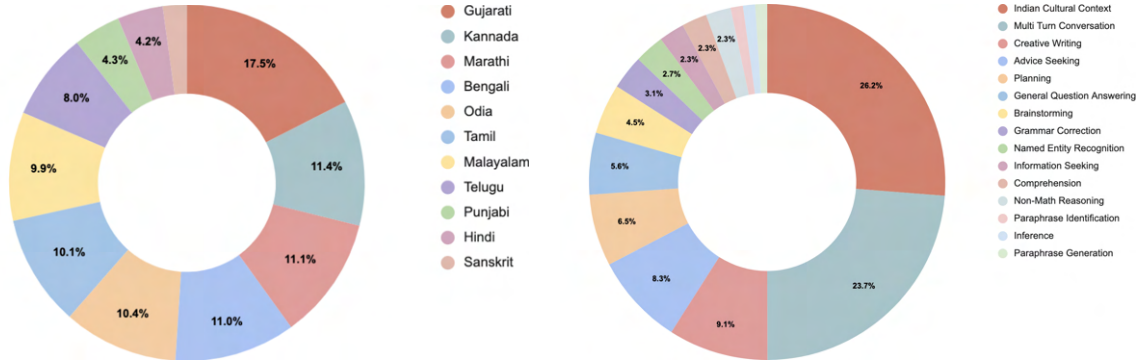


Figure 2: Distribution of *Pragyaan-IT* (Instruction-Tuning) data across languages (left) and categories (right).

with its designated configuration, annotators may either (i) adapt the configuration to better reflect the pair, or (ii) create a new prompt–response pair that correctly conforms to the specified configuration. This decision balances efficiency with fidelity to the framework. In cases where manually generating a new response is especially time-intensive, the annotators flag the prompt for regeneration and they undergo another iteration through the pipeline.

Crucially, manual intervention goes beyond mechanical verification. Annotators conduct linguistic quality checks ensuring fluency, grammatical accuracy, syntactic correctness, and appropriate response length, but are also encouraged to exercise creative judgment. This includes refining awkward phrasings, restructuring unclear outputs, or enriching responses with culturally relevant details. Even when a pair formally satisfies all defined constraints, annotators may modify it to adjust tone, improve readability, contextual appropriateness or pedagogical value as well as elevate the overall communicative value of the response. These interventions ensure that the dataset is not only consistent with the defined settings, but also meaningful and robust for deployment in real-world post-training scenarios.

3.5 Dataset Curation

Our curation process draws from a broad pool of existing resources while systematically adapting them for Indic languages and tasks. In total, we considered 57 diverse language datasets, together with their relevant splits, spanning different categories and task families. These resources serve either as direct candidates for translation into Indic languages or as seed data for synthetic expansion through a modified Self-Instruct pipeline (Wang et al., 2022b) described earlier. By combining translation and generation in a complementary manner,

we ensure that the curated data covers not only core Natural Language Processing (NLP) tasks but also culturally grounded and contextually relevant dimensions. Table 3 in the Appendix provides an overview of the corresponding candidate datasets associated with each task sub-category.

Setting	Configuration	%
Complexity	Easy	62.32
	Hard	37.68
Multi Turn	1-Turn	91.66
	3-Turn	6.76
	5-Turn	1.58
Instruction Following	Simple IF	96.95
	Medium IF	2.49
	Complex IF	0.56
Safety	Safe	92.51
	Non-Safe	7.49
Indian Context	IC-1	32.04
	IC-2	10.16
	IC-3	57.80
Thinking Trails	Normal	99.98
	CoT	0.01
	Self Thinking	0.01

Table 2: Distribution of instances in *Pragyaan-IT* across different task settings and configurations in Section 3.3.

4 Pragyaan: Indic Post-training Datasets

As part of this work, we construct high-quality post-training datasets that explicitly target 10 Indian languages (Bengali, Gujarati, Hindi, Kannada, Malayalam, Marathi, Oriya, Punjabi, Tamil, and Telugu), yielding two complementary resources:

i) *Pragyaan-IT* (22.5K): an instruction-tuning dataset designed to enhance a model’s ability to follow diverse prompts across multiple domains, ensuring that models can generalize well to everyday user interactions.

ii) *Pragyaan-Align* (100K): a preference dataset curated for Reinforcement Learning (RL)-based

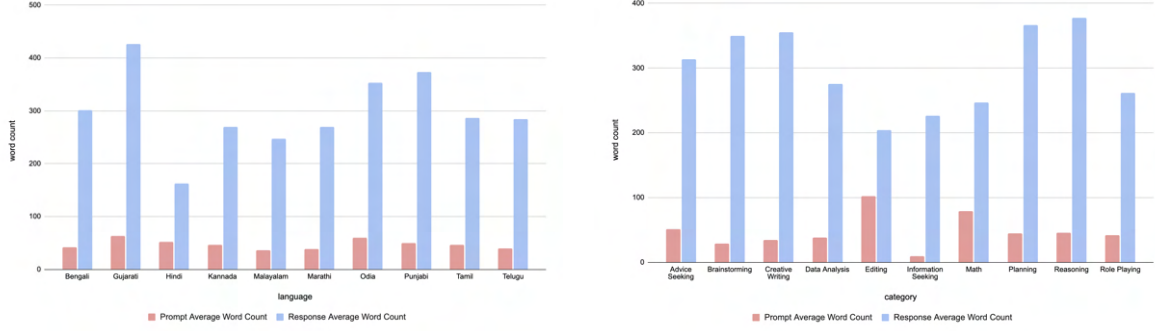


Figure 3: Average word counts of *Pragyaan-Align* alignment data across languages (left) and categories (right).

alignment methods, emphasizing preference learning, safety, and cultural grounding, allowing models to align more closely with the user intent.

5 Analysis

We begin with an assessment of raw synthetic generations refined through human annotation, followed by a broader dataset-level analysis.

5.1 Human Annotation Refinement

We evaluate the underlying LLM’s performance for both generation and translation tasks across 5 dimensions, each on a scale of 1-5, for a small subset (see Section A.5 in the Appendix for details on evaluation). Particularly, for English generations, grammatical accuracy remains slightly lower in comprehension (3.20) and creative writing (3.93) tasks, while receiving high scores for other tasks. For translation, the model shows the strongest performance for Hindi (Figure 14). Telugu and Gujarati exhibit moderate Lexical Diversity (3.43 and 3.30), while grammatical accuracy remains modest in Telugu (3.62), Hindi (3.60), and Punjabi (3.55). Thus, human refinements were critical for converting raw synthetic output into culturally grounded, linguistically accurate, and task-aligned data, directly underpinning the reliability of our data curation framework.

5.2 Dataset analysis

Figure 2 shows the distribution of *Pragyaan-IT* across 10 languages and 15 categories, while Table 3 lists the candidate datasets used in its construction. As seen in Figure 2, Indian Cultural Context (26.2%) and Multi-Turn Conversation (23.7%) dominate, while reasoning and paraphrasing remain limited (1–3%), forming targets for future expansion. Language coverage is led by Gujarati (17.7%), Kannada (11.4%), Marathi (11.1%),

and Odia (10.8%). In the current setting (Table 2), we have 62.3% ‘Easy’ tasks, single-turn interactions (91.7%), simple instruction-following (96.9%), safe content (92.5%) with Indian context well covered (IC-3: 57.8%). While this analysis reflects the status of the *Pragyaan-IT* dataset at the time of writing, the dataset is under active curation and will be more comprehensive across task categories and settings as described in earlier sections.

Word count analysis highlights linguistic and task-level variation (Figure 4). Gujarati and Odia are most verbose in both prompts (~110 words) and responses (~320 words), whereas Hindi mostly remain concise. At the task level, Multi-turn conversation, Advice Seeking, and Creative Writing yield the longest responses (550–620 words). Conversely, QA, Indian cultural context, and Comprehension tasks are consistently brief. Verbosity thus correlates with conversational and creative tasks, while simpler or context-specific settings produce shorter outputs.

Pragyaan-Align, the preference dataset, has an equal representation across all languages and 10 categories that helps promote fairness and mitigate bias. All instances in *Pragyaan-Align* follow a standardized configuration: single-turn interactions with instruction following and safe responses. Our analysis shows variation in text lengths across categories and languages. Average word counts for prompts range from 36–63 words, preferred responses 137–539, while rejected responses range from 154–466 words reflecting differences in task complexity and elaboration. Detailed language-wise as well as category-wise trends are also presented in Figure 3.

Overall, our *Pragyaan* datasets for post-training provide a broad task and Indian language coverage. The various task settings and our curation

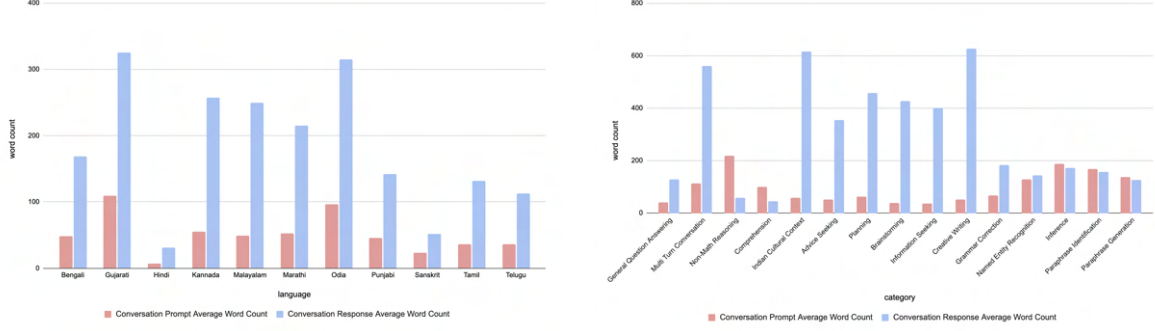


Figure 4: Average word counts of *Pragyaan-IT* data across languages (left) and categories (right).

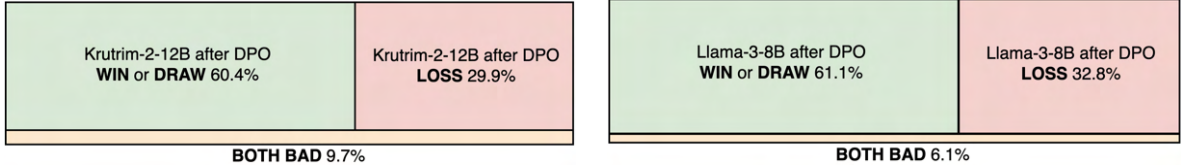


Figure 5: Win rates of the *Krutrim-2-12B* (left) and *Llama-3-8B* (right) models after DPO, compared against their respective pre-DPO versions.

process leveraging LLMs along with human-in-the-loop help ensure high quality. Our future iterations will expand complex reasoning capabilities and enhance representation for under-covered low-resource languages.

5.3 Downstream performance

To evaluate the quality of our curated dataset, we conduct a pilot study with Direct Preference Optimization (DPO) (Rafailov et al., 2023) based alignment on the *Pragyaan-Align* dataset using two open-weight models which supports the languages under consideration: *Krutrim-2-12B Instruct* (Kallappa et al., 2025) and *Llama-3-8B Instruct* (Grattafiori et al., 2024).

For evaluation, we use the recently released *Updesh* dataset², which covers similar categories and languages. We sample around 100 examples from nine relevant categories, balanced across ten languages. Responses from the models are scored against the ground truth on a scale of 1–5 using an LLM-as-a-judge. We provide more information about the hyperparameter configuration and other experimental details in Section A.6 and the corresponding prompts used for evaluation in Section B of the Appendix.

As shown in Figure 5, *Krutrim-2-12B* after DPO either wins or draws in 60.4% of cases, loses in 29.9%, and both pre- and post-DPO responses

score poorly (<2) in 9.7%. A similar trend holds for *Llama-3-8B* (61.1% wins or draws), confirming the promising potential of our curated dataset for alignment across different categories and multiple Indian languages.

6 Conclusion

This work addresses the scarcity of high-quality post-training data for multilingual LLMs by developing a human-in-the-loop pipeline to ensure diversity, quality and cultural grounding. Through this approach, we construct two datasets: *Pragyaan-IT* and *Pragyaan-Align* covering 10 Indian languages and multiple task categories. The datasets highlight inclusion of local cultural context, task diversity, multi-turn dialogue, and safety alignment, overcoming the limitations of naive translations and low-quality synthetic resources. Although designed for Indian languages, the pipeline is readily adaptable to other multilingual contexts. We present a comprehensive analysis of the dataset’s characteristics, covering language distribution and domain coverage, and further demonstrate its effectiveness for alignment through a small-scale pilot study. Future efforts will expand language coverage and further annotation quality refinement. We aim for this work to support broader efforts in building culturally inclusive resources that strengthen LLM applicability in multilingual contexts.

²https://huggingface.co/datasets/microsoft/Updesh_beta

Limitations

Our dataset is part of an ongoing effort, with plans to continually expand post-training instances. While our human-in-the-loop framework mitigates many issues, challenges such as minor linguistic inaccuracies, fluency variation across languages, and potential annotation subjectivity may still persist. Moreover, our research prioritized Indian languages, and the generalization of findings to other multilingual settings remain currently unexplored. Extending the pipeline to new cultural and linguistic contexts will require additional validation. We view these limitations as avenues for future work towards broader applicability and refinement of our proposed framework.

Ethics Statement

This study focuses on curating large-scale post-training datasets for Indian languages, encompassing diverse tasks and cultural contexts. The pipeline combines synthetic generation with human-in-the-loop refinement to ensure quality, safety, and cultural fidelity. We provide proper attribution to all source datasets and tools through citations. Human involvement was limited to annotation and quality control; no personally identifiable or sensitive information was collected. We engage a team of 50 in-house annotators for dataset creation. All contributors were clearly informed that their work supports LLM training and were compensated fairly at locally prevailing market rates. This study did not require formal IRB approval. Throughout the process, we prioritized preserving cultural nuances while avoiding harmful, biased, or unsafe content. The resulting dataset is designed to advance the development of multilingual and culturally inclusive LLMs.

Acknowledgements

We thank the leadership at Krutrim for their support in carrying out this research. We also thank the Data Annotation Team for their meticulous efforts especially Sanmathi P N, Dr. Salman Alam, Rupakatha Mukherjee, Vinod Kumar, Shivaram Prasad Putta, Divyashree K, Arati V Thakur, Vaibhav M Sutariya, Mitali Pradhan, Payal Yadav, Sneha Aniyani, Hemant P Rajopadhye, Joel Johnson and Srinidhi S. Additionally, we also thank Guduru Manoj and Souvik Rana for the helpful discussions.

References

- Marah Abdin, Sahaj Agarwal, Ahmed Awadallah, Vidhisha Balachandran, Harkirat Behl, Lingjiao Chen, Gustavo de Rosa, Suriya Gunasekar, Mojan Javaheripi, Neel Joshi, et al. 2025. Phi-4-reasoning technical report. *arXiv preprint arXiv:2504.21318*.
- Yuntao Bai, Andy Jones, Kamal Ndousse, Amanda Askell, Anna Chen, Nova DasSarma, Dawn Drain, Stanislav Fort, Deep Ganguli, Tom Henighan, Nicholas Joseph, Saurav Kadavath, Jackson Kernion, Tom Conerly, Sheer El-Showk, Nelson Elhage, Zac Hatfield-Dodds, Danny Hernandez, Tristan Hume, Scott Johnston, Shauna Kravec, Liane Lovitt, Neel Nanda, Catherine Olsson, Dario Amodei, Tom Brown, Jack Clark, Sam McCandlish, Chris Olah, Ben Mann, and Jared Kaplan. 2022a. [Training a helpful and harmless assistant with reinforcement learning from human feedback](#).
- Yuntao Bai, Saurav Kadavath, Sandhini Kundu, Amanda Askell, et al. 2022b. Constitutional ai: Harmlessness from ai feedback. *arXiv preprint arXiv:2212.08073*.
- Samuel Barham, , Weller, Michelle Yuan, Kenton Murray, Mahsa Yarmohammadi, Zhengping Jiang, Siddharth Vashishtha, Alexander Martin, Anqi Liu, Aaron Steven White, Jordan Boyd-Graber, and Benjamin Van Durme. 2023. [Megawika: Millions of reports and their sources across 50 diverse languages](#).
- Akhiad Bercovich, Itay Levy, Izik Golan, Mohammad Dabbah, Ran El-Yaniv, Omri Puny, Ido Galil, Zach Moshe, Tomer Ronen, Najeeb Nabwani, et al. 2025a. Llama-nemotron: Efficient reasoning models. *arXiv preprint arXiv:2505.00949*.
- Akhiad Bercovich, Itay Levy, Izik Golan, Mohammad Dabbah, Ran El-Yaniv, Omri Puny, Ido Galil, Zach Moshe, Tomer Ronen, Najeeb Nabwani, et al. 2025b. Llama-nemotron: Efficient reasoning models. *arXiv preprint arXiv:2505.00949*.
- Ning Bian, Hongyu Lin, Yaojie Lu, Xianpei Han, Le Sun, and Ben He. 2023. Chatacpa: A multi-turn dialogue corpus based on alpaca instructions. <https://github.com/cascip/ChatAlpaca>.
- Daniel Borkan, Lucas Dixon, Jeffrey Sorensen, Nithum Thain, and Lucy Vasserman. 2019. Nuanced metrics for measuring unintended bias with real data for text classification. In *Companion Proceedings of The 2019 World Wide Web Conference*.
- Hyung Won Chung, Le Hou, Shayne Longpre, Barret Zoph, et al. 2022. Scaling instruction-finetuned language models. *arXiv preprint arXiv:2210.11416*.
- 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. *arXiv preprint arXiv:2110.14168*.

- Arman Cohan, Franck Dernoncourt, Doo Soon Kim, Trung Bui, Seokhwan Kim, Walter Chang, and Nazli Goharian. 2018. A discourse-aware attention model for abstractive summarization of long documents. In *Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 2 (Short Papers)*.
- Alexis Conneau, Ruty Rinott, Guillaume Lample, Adina Williams, Samuel R. Bowman, Holger Schwenk, and Veselin Stoyanov. 2018. Xnli: Evaluating cross-lingual sentence representations. In *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing*. Association for Computational Linguistics.
- Mike Conover, Matt Hayes, Ankit Mathur, Jianwei Xie, Jun Wan, Sam Shah, Ali Ghodsi, Patrick Wendell, Matei Zaharia, and Reynold Xin. 2023. [Free Dolly: Introducing the World’s First Truly Open Instruction-Tuned LLM](#).
- Ganqu Cui, Lifan Yuan, Ning Ding, Guanming Yao, Bingxiang He, Wei Zhu, Yuan Ni, Guotong Xie, Ruobing Xie, Yankai Lin, et al. 2023. Ultrafeedback: Boosting language models with scaled ai feedback. *arXiv preprint arXiv:2310.01377*.
- Sumanth Doddapaneni, Rahul Aralikkatte, Gowtham Ramesh, Shreya Goyal, Mitesh M. Khapra, Anoop Kunchukuttan, and Pratyush Kumar. 2023. [Towards leaving no Indic language behind: Building monolingual corpora, benchmark and models for Indic languages](#). In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 12402–12426, Toronto, Canada. Association for Computational Linguistics.
- Aaron Grattafiori, Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Alex Vaughan, et al. 2024. The llama 3 herd of models. *arXiv preprint arXiv:2407.21783*.
- Suriya Gunasekar, Yi Zhang, Jyoti Aneja, Caio César Teodoro Mendes, Allie Del Giorno, Sivakanth Gopi, Mojan Javaheripi, Piero Kauffmann, Gustavo de Rosa, Olli Saarikivi, Adil Salim, Shital Shah, Harkirat Singh Behl, Xin Wang, Sébastien Bubeck, Ronen Eldan, Adam Tauman Kalai, Yin Tat Lee, and Yuanzhi Li. 2023. [Textbooks are all you need](#).
- Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma, Peiyi Wang, Xiao Bi, et al. 2025. Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning. *arXiv preprint arXiv:2501.12948*.
- Himanshu Gupta, Kevin Scaria, Ujjwala Ananthaswaran, Shreyas Verma, Mihir Parmar, Saurabh Arjun Sawant, Chitta Baral, and Swaroop Mishra. 2023. [Targen: Targeted data generation with large language models](#).
- Kai Hartung, Aaricia Herygers, Shubham Kurlekar, Khabbab Zakaria, Taylan Volkan, Sören Gröttrup, and Munir Georges. 2023. [Measuring sentiment bias in machine translation](#).
- Dan Hendrycks, Collin Burns, Steven Basart, Andy Zou, Mantas Mazeika, Dawn Song, and Jacob Steinhardt. 2020. Measuring massive multitask language understanding. *arXiv preprint arXiv:2009.03300*.
- Dan Hendrycks, Collin Burns, Saurav Kadavath, Akul Arora, Steven Basart, Eric Tang, Dawn Song, and Jacob Steinhardt. 2021. Measuring mathematical problem solving with the math dataset. *NeurIPS*.
- Geoffrey Hinton, Oriol Vinyals, and Jeff Dean. 2015. Distilling the knowledge in a neural network. *arXiv preprint arXiv:1503.02531*.
- Or Honovich, Thomas Scialom, Omer Levy, and Timo Schick. 2022. [Unnatural instructions: Tuning language models with \(almost\) no human labor](#). *arXiv preprint arXiv:2212.09689*.
- Joseph Cheung. 2023. [GuanacoDataset \(Revision 8cf0d29\)](#).
- Mandar Joshi, Eunsol Choi, Daniel S Weld, and Luke Zettlemoyer. 2017. Triviaqa: A large scale distantly supervised challenge dataset for reading comprehension. *arXiv preprint arXiv:1705.03551*.
- Pratik Joshi, Sebastin Santy, Amar Budhiraja, Kalika Bali, and Monojit Choudhury. 2020. The state and fate of linguistic diversity and inclusion in the nlp world. *arXiv preprint arXiv:2004.09095*.
- Aditya Kallappa, Palash Kamble, Abhinav Ravi, Akshat Patidar, Vinayak Dhruv, Deepak Kumar, Raghav Awasthi, Arveti Manjunath, Shubham Agarwal, Kumar Ashish, Gautam Bhargava, and Chandra Khatri. 2025. [Krutrim llm: Multilingual foundational model for over a billion people](#).
- Mohammed Safi Ur Rahman Khan, Priyam Mehta, Ananth Sankar, Umashankar Kumaravelan, Sumanth Doddapaneni, Sparsh Jain, Anoop Kunchukuttan, Pratyush Kumar, Raj Dabre, Mitesh M Khapra, et al. 2024. Indicllmsuite: A blueprint for creating pre-training and fine-tuning datasets for indian languages. *arXiv preprint arXiv:2403.06350*.
- Seungone Kim, Se June Joo, Doyoung Kim, Joel Jang, Seonghyeon Ye, Jamin Shin, and Minjoon Seo. 2023. The cot collection: Improving zero-shot and few-shot learning of language models via chain-of-thought fine-tuning. *arXiv preprint arXiv:2305.14045*.
- Andreas Köpf, Yannic Kilcher, Dimitri von Rütte, Sotiris Anagnostidis, Zhi-Rui Tam, Keith Stevens, Abdullah Barhoum, Nguyen Minh Duc, Oliver Stanley, Richárd Nagyfi, et al. 2023. Openassistant conversations—democratizing large language model alignment. *arXiv preprint arXiv:2304.07327*.

- Aman Kumar, Himani Shrotriya, Prachi Sahu, Amogh Mishra, Raj Dabre, Ratish Puduppully, Anoop Kunchukuttan, Mitesh M. Khapra, and Pratyush Kumar. 2022. [IndicNLG benchmark: Multilingual datasets for diverse NLG tasks in Indic languages](#). In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, pages 5363–5394, Abu Dhabi, United Arab Emirates. Association for Computational Linguistics.
- Tom Kwiatkowski, Jennimaria Palomaki, Olivia Redfield, Michael Collins, Ankur Parikh, Chris Alberti, Danielle Epstein, Illia Polosukhin, Matthew Kelcey, Jacob Devlin, Kenton Lee Toutanova, Llion Jones, Ming-Wei Chang, Andrew Dai, Jakob Uszkoreit, Quoc Le, and Slav Petrov. 2019. Natural questions: a benchmark for question answering research. *TACL*.
- Nathan Lambert, Jacob Morrison, Valentina Pyatkin, Shengyi Huang, Hamish Ivison, Faeze Brahman, Lester James V. Miranda, Alisa Liu, Nouha Dziri, Shane Lyu, Yuling Gu, Saumya Malik, Victoria Graf, Jena D. Hwang, Jiangjiang Yang, Ronan Le Bras, Oyvind Tafford, Chris Wilhelm, Luca Soldaini, Noah A. Smith, Yizhong Wang, Pradeep Dasigi, and Hannaneh Hajishirzi. 2025. [Tulu 3: Pushing frontiers in open language model post-training](#).
- Haonan Li, Fajri Koto, Minghao Wu, Alham Fikri Aji, and Timothy Baldwin. 2023a. Bactrian-x: A multilingual replicable instruction-following model with low-rank adaptation. *arXiv preprint arXiv:2305.15011*.
- Lei Li, Yuwei Yin, Shicheng Li, Liang Chen, Peiyi Wang, Shuhuai Ren, Mukai Li, Yazheng Yang, Jingjing Xu, Xu Sun, Lingpeng Kong, and Qi Liu. 2023b. [M³it: A large-scale dataset towards multi-modal multilingual instruction tuning](#).
- Yang Liu, Xiaotian Zhang, Haoze Sun, Yuzhen Huang, Wei Wu, Yiliang Li, Hao Yu, Jun Liu, Yueqi Li, Zhicheng Guo, Xiaoguang Li, Yongfeng Huang, Guilin Li, Yujun Zhou, Yufeng Chen, Chenyan Jia, and Xing-Jian Jiang. 2024. Mmlu-pro: a more advanced and challenging multi-task evaluation for llms. *arXiv preprint arXiv:2406.01574*.
- Shayne Longpre, Le Hou, Tu Vu, Albert Webson, et al. 2023. The flan collection: Designing data and methods for effective instruction tuning. *arXiv preprint arXiv:2301.13688*.
- Ziyang Luo, Can Xu, Pu Zhao, Qingfeng Sun, Xubo Geng, Wenxiang Hu, Chongyang Tao, Jing Ma, Qingwei Lin, and Daxin Jiang. 2023. [Wizardcoder: Empowering code large language models with evol-instruct](#).
- Arnav Mhaske, Harshit Kedia, Sumanth Doddapaneni, Mitesh M. Khapra, Pratyush Kumar, Rudra Murthy, and Anoop Kunchukuttan. 2022. [Naamapadam: A large-scale named entity annotated data for indic languages](#).
- Todor Mihaylov, Peter Clark, Tushar Khot, and Ashish Sabharwal. 2018. A new dataset for open book question answering. In *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing*.
- Niklas Muennighoff, Qian Liu, Armel Zebaze, Qinkai Zheng, Binyuan Hui, Terry Yue Zhuo, Swayam Singh, Xiangru Tang, Leandro von Werra, and Shayne Longpre. 2023. Octopack: Instruction tuning code large language models. *arXiv preprint arXiv:2308.07124*.
- Niklas Muennighoff, Thomas Wang, Lintang Sutawika, Adam Roberts, Stella Biderman, Teven Le Scao, M Saiful Bari, Sheng Shen, Zheng-Xin Yong, Hailey Schoelkopf, et al. 2022. Crosslingual generalization through multitask finetuning. *arXiv preprint arXiv:2211.01786*.
- Niklas Muennighoff, Zitong Yang, Weijia Shi, Xiang Lisa Li, Li Fei-Fei, Hannaneh Hajishirzi, Luke Zettlemoyer, Percy Liang, and Emmanuel. 2025. [s1: Simple test-time scaling](#).
- Huu Nguyen, Sameer Suri, Ken Tsui, and Christoph Schuhmann. 2023. The open instruction generalist (oig) dataset. <https://laion.ai/blog/oig-dataset/>.
- Jinjie Ni, Fuzhao Xue, Kabir Jain, Mahir Hitesh Shah, Zangwei Zheng, and Yang You. 2023. Instruction in the wild: A user-based instruction dataset. <https://github.com/XueFuzhao/InstructionWild>.
- NVIDIA. 2025. Llama-nemotron-post-training dataset: A comprehensive collection of instruction tuning and alignment data. *arXiv preprint arXiv:2505.00949*.
- Long Ouyang, Jeff Wu, Xu Jiang, Diogo Almeida, et al. 2022. Training language models to follow instructions with human feedback. *arXiv preprint arXiv:2203.02155*.
- Guilherme Penedo, Anton Lozhkov, Hynek Kydlíček, Loubna Ben Allal, Edward Beeching, Agustín Piñeros Lajarán, Quentin Gallouédec, Nathan Habib, Lewis Tunstall, and Leandro von Werra. 2025. Codeforces cots. <https://huggingface.co/datasets/open-r1/codeforces-cots>.
- Danti Pudjiati, Ninuk Lustyantie, Ifan Iskandar, and Tira Nur Fitria. 2022. Post-editing of machine translation: Creating a better translation of cultural specific terms. *Language Circle: Journal of Language and Literature*, 17(1):61–73.
- Ratish Puduppully, Himani Shrotriya, Anoop Kunchukuttan, and Mitesh M. Khapra. 2024. [Mathinstruct: A high-quality math instruction dataset](#).
- Rafael Rafailov, Archit Sharma, Eric Mitchell, Christopher D Manning, Stefano Ermon, and Chelsea Finn. 2023. Direct preference optimization: Your language model is secretly a reward model. *Advances in neural information processing systems*, 36:53728–53741.

- Nazneen Rajani, Lewis Tunstall, Edward Beeching, Nathan Lambert, Alexander M. Rush, and Thomas Wolf. 2023. No robots. https://huggingface.co/datasets/HuggingFaceH4/no_robots.
- Samyam Rajbhandari, Jeff Rasley, Olatunji Ruwase, and Yuxiong He. 2020. Zero: Memory optimizations toward training trillion parameter models. In *Proceedings of the International Conference for High Performance Computing, Networking, Storage and Analysis (SC)*.
- Pranav Rajpurkar, Robin Jia, and Percy Liang. 2018. Know what you don’t know: Unanswerable questions for squad. *arXiv preprint arXiv:1806.03822*.
- David Rein, Betty Li Hou, Asa Cooper Stickland, Jackson Petty, Richard Yuanzhe Pang, Julien Dirani, Julian Michael, and Samuel R. Bowman. 2024. **GPQA: A graduate-level google-proof q&a benchmark**. In *First Conference on Language Modeling*.
- ServiceNow AI Research. 2024. R1-distill-sft: A dataset for instruction tuning with a focus on distillation from large language models. *arXiv preprint arXiv:2403.06350*.
- Pritika Rohera, Chaitrali Ginimav, Akanksha Salunke, Gayatri Sawant, and Raviraj Joshi. 2024. L3cube-indicquest: A benchmark question answering dataset for evaluating knowledge of llms in indic context. *arXiv preprint arXiv:2405.18789*.
- Christian Rothermund, David Vögele, Lucas Roth, Philipp Schmutz, Tobias Wiegand, and Sebastian Aschenbrenner. 2025. Finemedlm-o1: Enhancing medical knowledge reasoning ability of llm from supervised fine-tuning to test-time training. *arXiv preprint arXiv:2501.09213*.
- Victor Sanh, Albert Webson, Colin Raffel, Stephen Bach, Lintang Sutawika, Zaid Alyafeai, Antoine Chaffin, Arnaud Stiegler, Arun Raja, Manan Dey, M Saiful Bari, Canwen Xu, Urmish Thakker, Shanya Sharma Sharma, Eliza Szczechla, Taewoon Kim, Gunjan Chhablani, Nihal Nayak, Debajyoti Datta, Jonathan Chang, Mike Tian-Jian Jiang, Han Wang, Matteo Manica, Sheng Shen, Zheng Xin Yong, Harshit Pandey, Rachel Bawden, Thomas Wang, Trishala Neeraj, Jos Rozen, Abheesht Sharma, Andrea Santilli, Thibault Fevry, Jason Alan Fries, Ryan Teehan, Teven Le Scao, Stella Biderman, Leo Gao, Thomas Wolf, and Alexander M Rush. 2022. **Multi-task prompted training enables zero-shot task generalization**. In *International Conference on Learning Representations*.
- Beatrice Savoldi, Marco Gaido, Luisa Bentivogli, Matteo Negri, and Marco Turchi. 2021. **Gender bias in machine translation**.
- Harman Singh, Nitish Gupta, Shikhar Bharadwaj, Dinesh Tewari, and Partha Talukdar. 2024a. **Indicgen-bench: A multilingual benchmark to evaluate generation capabilities of llms on indic languages**.
- Shivalika Singh, Freddie Vargus, Daniel Dsouza, Börje F. Karlsson, Abinaya Mahendiran, Wei-Yin Ko, Herumb Shandilya, Jay Patel, Deividas Mataciunas, Laura OMahony, Mike Zhang, Ramith Hettiarachchi, Joseph Wilson, Marina Machado, Luisa Souza Moura, Dominik Krzemiński, Hakimeh Fadaei, Irem Ergün, Ifeoma Okoh, Aisha Alaagib, Oshan Mudannayake, Zaid Alyafeai, Vu Minh Chien, Sebastian Ruder, Surya Guthikonda, Emad A. Alghamdi, Sebastian Gehrmann, Niklas Muennighoff, Max Bartolo, Julia Kreutzer, Ahmet Üstün, Marzieh Fadaee, and Sara Hooker. 2024b. **Aya dataset: An open-access collection for multilingual instruction tuning**.
- Zafir Stojanovski, Oliver Stanley, Joe Sharratt, Richard Jones, Abdulhakeem Adefoye, Jean Kaddour, and Andreas Köpf. 2025. **Reasoning gym: Reasoning environments for reinforcement learning with verifiable rewards**.
- Alon Talmor, Jonathan Herzig, Nicholas Lourie, and Jonathan Berant. 2019. CommonsenseQA: A question answering challenge targeting commonsense knowledge. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*.
- Rohan Taori, Ishaan Gulrajani, Tianyi Zhang, Yann Dubois, Xuechen Li, Carlos Guestrin, Percy Liang, and Tatsunori B Hashimoto. 2023. Stanford alpaca: An instruction-following llama model.
- Shubham Toshniwal, Ivan Moshkov, Sean Narenthiran, Daria Gitman, Fei Jia, and Igor Gitman. 2024. Openmathinstruct-1: A 1.8 million math instruction tuning dataset. *arXiv preprint arXiv: 2402.10176*.
- Foutse Touileb, Xavier Longpre, Adrien L’Heureux, Tom Houlisby, Duy-Kien Le, Vincent Aribaud, Hugo Le Scao, Tommaso Pasquini, Marco Miaschi, and Patrick Muennighoff. 2024. Magpie: A broad-coverage, fine-grained multitask instruction dataset. *arXiv preprint arXiv:2403.00010*.
- Hugo Touvron, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmine Babaei, Nikolay Bashlykov, Soumya Batra, Prajjwal Bhargava, Shruti Bhosale, Dan Bikel, Lukas Blecher, Cristian Canton Ferrer, Moya Chen, Guillem Cucurull, David Esiobu, Jude Fernandes, Jeremy Fu, Wenyin Fu, Brian Fuller, Cynthia Gao, Vedanuj Goswami, Naman Goyal, Anthony Hartshorn, Saghar Hosseini, Rui Hou, Hakan Inan, Marcin Kardas, Viktor Kerkez, Madian Khabsa, Isabel Kloumann, Artem Korenev, Punit Singh Koura, Marie-Anne Lachaux, Thibaut Lavril, Jenya Lee, Diana Liskovich, Yinghai Lu, Yuning Mao, Xavier Martinet, Todor Mihaylov, Pushkar Mishra, Igor Molybog, Yixin Nie, Andrew Poulton, Jeremy Reizenstein, Rashi Rungta, Kalyan Saladi, Alan Schelten, Ruan Silva, Eric Michael Smith, Ranjan Subramanian, Xiaoqing Ellen Tan, Binh Tang, Ross Taylor, Adina Williams, Jian Xiang Kuan, Puxin Xu,

- Zheng Yan, Iliyan Zarov, Yuchen Zhang, Angela Fan, Melanie Kambadur, Sharan Narang, Aurelien Rodriguez, Robert Stojnic, Sergey Edunov, and Thomas Scialom. 2023. [Llama 2: Open foundation and fine-tuned chat models](#).
- Jun Wang, Benjamin Rubinstein, and Trevor Cohn. 2022a. Measuring and mitigating name biases in neural machine translation. In *ACL*.
- Yizhong Wang, Yeganeh Kordi, Swaroop Mishra, Alisa Liu, Noah A. Smith, Daniel Khashabi, and Hannaneh Hajishirzi. 2022b. Self-instruct: Aligning language model with self generated instructions.
- Yizhong Wang, Yeganeh Kordi, Swaroop Mishra, Alisa Liu, Noah A. Smith, Daniel Khashabi, and Hannaneh Hajishirzi. 2022c. [Self-instruct: Aligning language model with self generated instructions](#). *ArXiv preprint*, abs/2212.10560.
- Yizhong Wang, Yeganeh Li, Niklas Michael, Hila Gonen, Jesse He, Yi Zhao, Ziyi Lin, Yejin Shafi, Karan Singh, Foutse Touileb, and et al. 2022d. Super-naturalinstructions: A generalization beyond english language. *arXiv preprint arXiv:2204.05367*.
- Jason Wei, Maarten Bosma, Vincent Zhao, Kelvin Guu, Adams Wei Yu, Brian Lester, Nan Du, Andrew M. Dai, and Quoc V Le. 2022a. [Finetuned language models are zero-shot learners](#). In *International Conference on Learning Representations*.
- Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Fei Xia, Ed Chi, Quoc V Le, Denny Zhou, et al. 2022b. Chain-of-thought prompting elicits reasoning in large language models. *Advances in neural information processing systems*, 35:24824–24837.
- Sean Welleck, Ronan Le Bras, Hannaneh Hajishirzi, and Yejin Choi. 2021. Naturalproofs: Mathematical theorem proving in natural language. *arXiv preprint arXiv:2104.01112*.
- Tianbao Xie, Chen Henry Wu, Peng Shi, Ruiqi Zhong, Torsten Scholak, Michihiro Yasunaga, Chien-Sheng Wu, Ming Zhong, Pengcheng Yin, Sida I Wang, et al. 2022. Unifiedskg: Unifying and multi-tasking structured knowledge grounding with text-to-text language models. *arXiv preprint arXiv:2201.05966*.
- Can Xu, Qingfeng Sun, Kai Zheng, Xiubo Geng, Pu Zhao, Jiazhan Feng, Chongyang Tao, and Daxin Jiang. 2023. [Wizardlm: Empowering large language models to follow complex instructions](#).
- Zhangchen Xu, Fengqing Jiang, Luyao Niu, Yuntian Deng, Radha Poovendran, Yejin Choi, and Bill Yuchen Lin. 2024. [Magpie: Alignment data synthesis from scratch by prompting aligned llms with nothing](#).
- Zhangchen Xu, Yang Liu, Yueqin Yin, Mingyuan Zhou, and Radha Poovendran. 2025. KodCode: A diverse, challenging, and verifiable synthetic dataset for coding. In *Findings of the Association for Computational Linguistics: ACL 2025*, pages 6980–7008.
- Yi Yang, Wen-tau Yih, and Christopher Meek. 2015. Wikiqa: A challenge dataset for open-domain question answering. In *Proceedings of the 2015 Conference on Empirical Methods in Natural Language Processing*.
- Zhilin Yang, Peng Qi, Saizheng Zhang, Yoshua Bengio, William Cohen, Ruslan Salakhutdinov, and Christopher D. Manning. 2018. [HotpotQA: A dataset for diverse, explainable multi-hop question answering](#). In *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing*, pages 2369–2380, Brussels, Belgium. Association for Computational Linguistics.
- Yixin Ye, Zhen Huang, Yang Xiao, Ethan Chern, Shijie Xia, and Pengfei Liu. 2025. [Limo: Less is more for reasoning](#).
- Xiang Yu, Haidong Liu, Ning Yang, Qiaoli Wang, Jianye Wang, Jia Li, Li Zhang, and Jian Zhang. 2023a. Metamath: Bootstrap your own mathematical reasoning dataset. *arXiv preprint arXiv:2309.12284*.
- Yue Yu, Yuchen Zhuang, Jieyu Zhang, Yu Meng, Alexander Ratner, Ranjay Krishna, Jiaming Shen, and Chao Zhang. 2023b. [Large language model as attributed training data generator: A tale of diversity and bias](#).
- Tianyu Zhang, Yang Liu, Yuxuan Xu, Yujia Wang, Haoyu Chen, Hongwei Li, Chen Li, Jiawei Xie, and Chen Zhang. 2024. Open-thoughts: A large-scale dataset for learning and evaluating llm thinking processes. *arXiv preprint arXiv:2406.04178*.
- Chunting Zhou, Pengfei Liu, Puxin Xu, Srinu Iyer, Jiao Sun, Yuning Mao, Xuezhe Ma, Avia Efrat, Ping Yu, Lili Yu, et al. 2023. Lima: Less is more for alignment. *arXiv preprint arXiv:2305.11206*.
- Terry Yue Zhuo, Armel Zebaze, Nitchakarn Supattarachai, Leandro von Werra, Harm de Vries, Qian Liu, and Niklas Muennighoff. 2024. Astraios: Parameter-efficient instruction tuning code large language models. *arXiv preprint arXiv:2401.00788*.

A Appendix

A.1 Dataset Curation

For each sub-category, we curated multiple publicly available candidate datasets and evaluated their applicability to our data collection objectives, assessing how they align with the chosen curation approaches. Table 3 summarizes representative categories along with corresponding datasets incorporated into our pipeline. It is worth noting that data collection for certain sub-categories remains an ongoing effort.

A.2 Data Construction Approaches

We further provide more details about the prompt and responses for both the approaches in Table 4.

Sub Category	Candidate Datasets	Sub Category	Candidate Datasets
Non-Math Reasoning	MMLU (Hendrycks et al., 2020) MMLU-Pro (Liu et al., 2024) GPQA (Rein et al., 2024) medical-o1-reasoning-SFT (Rothermund et al., 2025) OpenThoughts (Zhang et al., 2024)	Sentiment Analysis	IndicSentiment (Doddapaneni et al., 2023) pietrollesci/imdb
Math Reasoning	OpenThoughts (Zhang et al., 2024) Llama-Nemotron-Posttraining-Dataset (NVIDIA, 2025) ServiceNow-AI/R1-Distil-SFT (Research, 2024) LIMO (Ye et al., 2025) s1K (Muennighoff et al., 2025)	General Question Answering	OpenBookQA (Mihaylov et al., 2018) TriviaQA (Joshi et al., 2017) CommonSenseQA (Talmor et al., 2019) WikiQA (Yang et al., 2015) HotPotQA (Yang et al., 2018) MegaWika (Barham et al., 2023)
Code Reasoning	OpenThoughts (Zhang et al., 2024) Llama-Nemotron-Posttraining-Dataset (NVIDIA, 2025) KodCode (Xu et al., 2025) open-r1/codeforces-cots (Penedo et al., 2025)	Fact Check	NaturalQuestions (Kwiatkowski et al., 2019) TriviaQA (Joshi et al., 2017) Indic Quest (Rohera et al., 2024)
Indian Relationships	Reasoning Gym (Stojanovski et al., 2025)	Multi Turn Conversation	ChatAlpaca (Bian et al., 2023) NoRobots (Rajani et al., 2023) Opus Samantha
Inference	XNLI (Conneau et al., 2018)	Role Playing	Public Domain Alpaca LimaRP Magpie (Touileb et al., 2024)
Data Analysis	Magpie (Touileb et al., 2024)	Advice Seeking	Magpie (Touileb et al., 2024)
Named Entity Recognition	Naamapadam (Mhaske et al., 2022) Bharath Bench	Planning	Magpie (Touileb et al., 2024)
Text Classification	pietrollesci/civilcomments-wilds (Borkan et al., 2019) pietrollesci/wikitoxic pietrollesci/hyperpartisan_news_detection Bharath Bench Aya Collection (Singh et al., 2024b)	Brainstorming	Magpie (Touileb et al., 2024)
Grammar Correction	Bharath Bench	Information Seeking	Magpie (Touileb et al., 2024)
Translation	Flores-IN (Singh et al., 2024a) IN-22	Indian Cultural Context	Bharath Bench
Creative Writing	Magpie (Touileb et al., 2024) Poetry	Comprehension	TriviaQA (Joshi et al., 2017) SQUAD (Rajpurkar et al., 2018) SQUAD 2.0 (Rajpurkar et al., 2018) IndicQA
Paraphrase Identification	IndicXParaphrase (Doddapaneni et al., 2023)	Math QA	OpenMathInstruct 2 (Toshniwal et al., 2024) GSM8K (Cobbe et al., 2021) MATH (Hendrycks et al., 2021) Tulu3 Persona Math (Lambert et al., 2025) Tulu3 Persona GSM8K (Lambert et al., 2025)
Text Summarization	CrossSumIN (Singh et al., 2024a) Indic Sentence Summarization (Kumar et al., 2022) arxiv-summarization (Cohan et al., 2018) news-summarization	Math Instruction Tuning	MetaMathQA (Yu et al., 2023a) Math Instruct (Puduppully et al., 2024)
Headline Generation	Indic Headline Generation (Kumar et al., 2022) NewSHead	Math Proofs	Natural Proofs (Welleck et al., 2021)
Question Generation	Indic Question Generation (Kumar et al., 2022)		

Table 3: Representative datasets curated for different sub-categories that are used as candidate datasets in our data curation approaches.

Approaches	Prompts	Responses
Approach 1: Translation + Human Refinement	English prompts → LLM translation → human verification/adaptation. <i>Outputs:</i> Indic Generic Prompts, Indic Context Prompts.	English responses → LLM translation → human verification/adaptation. <i>Outputs:</i> Indic Generic Responses, Indic Context Responses.
Approach 2: Synthetic Expansion + Human Refinement	Seed English prompts → LLM expansion (Self-Instruct) → LLM translation → human verification/adaptation. <i>Outputs:</i> Synthetic Indic Generic Prompts, Synthetic Indic Context Prompts.	Synthetic English responses → LLM translation → human verification/adaptation. <i>Outputs:</i> Synthetic Indic Generic Responses, Synthetic Indic Context Responses.

Table 4: Comparison of data construction approaches, showing pipelines for prompts and responses with resulting categories.

A.3 Complexity Definitions

Notably, complexity is evaluated relative to each task category, ensuring that difficulty levels are interpreted within the specific context of that category. Figures 6–11 illustrate how complexity is defined across different task categories.

A.4 Instruction Following

Various configurations of Instruction Following incorporate different combinations of constraints within the prompt. Figure 12 outlines the types of constraints considered, accompanied by illustrative examples for clarity while human annotation.

A.5 Human Annotation Refinement

Data annotators initially assessed the performance of the underlying LLM on both generation and translation tasks across multiple dimensions.

For generation, raw LLM responses were assessed along dimensions such as Response Relevance, Grammatical Accuracy, Cohesion and Coherence, Rationality, and Completeness (see Table 5 for the guidelines). As shown in Figure 13, English generations perform consistently well: tasks such as Indian Cultural Context, Advice Seeking, Information Seeking, Named Entity Recognition, Inference, Paraphrase Identification, Paraphrase Generation, and Headline Generation achieve near-perfect scores in Response Relevance (5.00) and Completeness (4.93–5.00). Multi-Turn Conversation also scores high in Cohesion and Coherence (5.00) and Response Relevance (4.70). Grammatical Accuracy remains strong overall (4.00–4.40), though slightly lower in Comprehension (3.20) and Creative Writing (3.93).

For translation, annotators evaluated Lexical Diversity, Coherence and Cohesion, Completeness, Grammatical Accuracy, and Named Entity Handling (see Table 6 for the guidelines). Hindi shows the strongest performance, achieving the highest Completeness (4.92) and Named Entity Handling (4.85). Telugu and Gujarati exhibit strong Lexical Diversity (3.43 and 3.30), while Grammatical Accuracy is highest in Telugu (3.62), Hindi (3.60), and Punjabi (3.55). Further detailed evaluation scores are illustrated in Figure 14.

A.6 Implementation

For training, we adopt a distributed DeepSpeed-based setup with ZeRO-3 optimization (Rajbhandari et al., 2020) across 2 H100 nodes (8 GPUs

per node) to efficiently handle large-scale model fine-tuning. The Llama-3-8B and Krutrim-2-12B instruct models are optimized using the Direct Preference Optimization (DPO) objective with a Beta parameter of 0.3, controlling the strength of preference alignment. Training is performed on sequences up to 4096 tokens, with a maximum prompt length of 2048 tokens to accommodate complex multi-turn instructions. We employ the AdamW optimizer with a learning rate of 5×10^{-7} , weight decay of 0.01, and a cosine learning rate scheduler with a warming ratio of 10% warmup ratio for stable convergence and run for 1 epoch. The batch configuration consists of 4 samples per device with gradient accumulation over 2 steps, yielding an effective batch size suitable for large-scale distributed training setup.

Post-trained Krutrim-2-12B and Llama-3-8B models are evaluated on the Updesh dataset across 10 languages and nine categories using LLM-as-a-Judge scoring on a scale of 1-5. Post-DPO, Krutrim achieves 31.7% wins, 29.9% losses, 28.6% draws, and 9.8% both-bad cases, while Llama records 35.0% wins, 26.1% losses, 27.9% draws, and 11.0% both-bad cases.

B Prompts Used

Different prompts were crafted for each category and complexity level during prompt generation using the self-instruct pipeline. For instance, Figures 15 and 16 illustrate example prompts for the Indian Cultural Context category under easy and hard settings, respectively. Similar design considerations were applied across other categories to address their specific requirements.

Figure 17 presents the prompt template employed for translating English prompt-response pairs via LLMs. The Krutrim-2-12B and Llama-3-8B Instruct models, after being fine-tuned using DPO on 100K examples from the *Pragyaan-Align* dataset, were evaluated on the Updesh dataset. The evaluation utilized an LLM-as-a-Judge framework for scoring, with the corresponding prompt design shown in Figure 18.

C Guidelines For Manual Annotation

The data annotation team follows a set of standardized guidelines designed to maintain consistency and uniformity throughout the annotation process. These guidelines include precise definitions for each category and setting, which are elaborated

Sub Category	Complexity	Sub Category	Complexity
Non-Math Reasoning	1. Number of reasoning steps Easy: Typically require a single fact or step. Hard: Requires multiple connected inferences. 2. Obviousness of needed knowledge Easy: Questions rely on common knowledge. Hard: May require less intuitive or more obscure background knowledge. 3. Clarity of distractors Easy: Questions have clearly incorrect options. Hard: Questions have plausible distractors that require careful elimination.	Data-Analysis	1. Data Complexity Easy: Small, clean datasets (few rows/columns). Hard: Larger, noisier, or irregular datasets (missing values, inconsistent formatting). 2. Reasoning Depth Easy: Single-step reasoning (e.g., identify max/min, simple difference). Hard: Multi-step reasoning (e.g., grouping, computing moving averages, interpreting trends, conditional filtering). 3. Type of Question Easy: Direct lookup or comparison. Hard: Involves synthesis across rows/columns, or combining values to reach a conclusion.
Math Reasoning	1. Number of reasoning steps Easy: Typically require a single fact or step. Hard: Requires multiple connected inferences. 2. Obviousness of needed knowledge Easy: Questions rely on common knowledge. Hard: May require less intuitive or more obscure background knowledge. 3. Clarity of distractors Easy: Questions have clearly incorrect options. Hard: Questions have plausible distractors that require careful elimination.	Named Entity Recognition	1. Entity Density and Type Variety Easy: Few entities, mostly common types (Person, Location). Hard: Multiple entities of varied or domain-specific types (e.g., Organization, Event, Product), possibly nested. 2. Language and Script Complexity Easy: Monolingual, clean, formal sentences in a standard script. Hard: Code-mixed, noisy, informal language, or regional scripts. 3. Ambiguity and Context Requirement Easy: Clearly named, well-known entities. Hard: Requires context to disambiguate entities or identify non-obvious named terms (e.g., "Amazon" as a company vs. river).
Code Reasoning	1. Number of reasoning steps Easy: Typically require a single fact or step. Hard: Requires multiple connected inferences. 2. Obviousness of needed knowledge Easy: Questions rely on common knowledge. Hard: May require less intuitive or more obscure background knowledge. 3. Clarity of distractors Easy: Questions have clearly incorrect options. Hard: Questions have plausible distractors that require careful elimination.	Text Classification	1. Clarity of Class Indicators Easy: Clear, surface-level cues that strongly correlate with the label. Hard: Subtle, context-dependent cues; overlapping vocabulary across classes. 2. Length and Structure Easy: Short, focused texts with one intent. Hard: Longer, noisy, or multi-intent texts requiring disambiguation. 3. Label Set Complexity Easy: Binary classification (e.g., spam vs. ham). Hard: Multi-class (e.g., news category: politics, sports, tech) or multi-label (e.g., tags for movie reviews: drama, romance).
Indian Relationships	1. Number of relations involved Easy: 2-3 relations, mostly linear. Hard: 4+ relations, often nested or cross-generational. 2. Clarity of phrasing Easy: Straightforward statements with direct relationships. Hard: Implicit or linguistically ambiguous statements, possibly requiring gender inference or cultural interpretation. 3. Relation type Easy: Direct blood relations (e.g., mother, brother). Hard: In-law relations, culturally specific terms, or multigenerational links.	Grammar Correction	1. Error Type Easy: Simple subject-verb agreement, plural/singular forms, or article usage. Hard: Tense consistency, word order, idiomatic usage, or multiple error types in one sentence. 2. Error Density Easy: Single, localized error. Hard: Multiple interrelated errors, possibly requiring rephrasing. 3. Contextual Dependence Easy: Error is identifiable without additional context. Hard: Requires deeper understanding of meaning or nuance to correct properly.
Inference	1. Lexical Overlap Easy: High overlap; direct wording or synonyms. Hard: Low overlap; requires paraphrasing or world knowledge. 2. Inference Type Easy: Simple entailment or direct contradiction. Hard: Implicit entailment, pragmatic reasoning, or nuanced neutrality. 3. Context Dependence Easy: No external knowledge needed. Hard: Requires commonsense, temporal, or cultural knowledge.	Translation	1. Syntactic Complexity Easy: Simple sentence structures (subject-verb-object). Hard: Complex or nested clauses, passive voice, or long-distance dependencies. 2. Lexical Ambiguity & Idioms Easy: Literal language with direct word equivalents. Hard: Idioms, metaphors, or polysemous words requiring interpretation. 3. Cultural Context Easy: Universally understood concepts. Hard: Culturally specific references, wordplay, or region-specific expressions.

Figure 6: Complexity Definitions for each sub-category.

Dimension	Explanation
Response Relevance	Measures how well the model output addresses the user query or task instruction. Why it matters for Indic/Multilingual Data: Responses must stay on-topic and satisfy user intent; irrelevant outputs reduce usability and may confuse readers. Errors in this area: Off-topic responses, inclusion of unrelated information, misinterpretation of the prompt.
Grammatical Accuracy	Correct use of grammar rules including syntax, tense, agreement, punctuation, and morphology. Why it matters for Indic/Multilingual Data: Proper grammar ensures readability and clarity; morphologically rich languages are prone to agreement and inflection errors. Errors in this area: Wrong tense, subject-verb disagreement, missing auxiliaries, incorrect case markings, improperly inflected words.
Cohesion and Coherence	Cohesion = linguistic devices (connectors, pronouns, conjunctions) linking sentences; Coherence = logical flow of ideas. Why it matters for Indic/Multilingual Data: Maintains well-structured and easily understandable responses; lack of cohesion or coherence leads to fragmented outputs. Errors in this area: Abrupt topic shifts, disconnected sentences, missing references, inappropriate pronoun usage.
Rationality	Logical correctness and factual consistency of the response, including reasoning and alignment with real-world knowledge. Why it matters for Indic/Multilingual Data: Ensures trustworthiness and usefulness of outputs, particularly for reasoning or factual tasks. Errors in this area: Contradictory statements, illogical conclusions, factually incorrect assertions, hallucinations.
Completeness	The extent to which the response fully addresses the user prompt or includes all necessary information. Why it matters for Indic/Multilingual Data: Partial answers reduce usefulness, especially for multi-step reasoning or detailed explanations. Errors in this area: Missing steps in reasoning, skipped entities, truncated explanations, insufficient coverage of subtopics.

Table 5: Key dimensions for assessing LLM outputs in Indic languages, highlighting relevance and frequent pitfalls.

Sub Category	Complexity	Sub Category	Complexity
Creative Writing	1. Prompt Specificity Easy: Constrained or highly specific prompts that guide the plot. Hard: Open-ended or abstract prompts requiring full narrative invention.	Question Generation	1. Contextual Simplicity Easy: Clear, straightforward context with a direct relationship between entities. Hard: Ambiguous or nuanced context, requiring more reasoning or interpretation to generate a meaningful question.
	2. Narrative Complexity Easy: Linear plot with one event or idea. Hard: Multi-layered plot, character development, or world-building.		2. Question Specificity Easy: Basic, factual questions asking for a specific detail (e.g., "Who?", "What?"). Hard: Complex or abstract questions that require synthesis, comparisons, or deeper knowledge about the context.
	3. Language & Style Easy: Simple, straightforward language. Hard: Use of metaphors, poetic devices, shifts in tone or voice.		3. Knowledge Requirement Easy: No prior knowledge beyond the given text is needed. Hard: Requires background knowledge, inference, or reasoning to generate a relevant question.
Paraphrase Identification	1. Lexical & Syntactic Similarity Easy: Minor reordering or synonym substitution with near-identical structure. Hard: Significant rephrasing, different vocabulary or sentence structure, yet same meaning.	Sentiment Analysis	1. Sentiment Clarity Easy: Clear, unambiguous sentiment (positive, negative, or neutral). Hard: Sentiment expressed indirectly, or a mixture of sentiments in the same text.
	2. Semantic Inference Easy: Direct, surface-level paraphrases. Hard: Requires understanding of implied meaning, paraphrased idioms, or abstract concepts.		2. Contextual Nuance Easy: Simple, straightforward expressions of sentiment with no hidden or subtle implications. Hard: Sentiment expressed in a nuanced, sarcastic, or contextual manner that requires understanding beyond surface-level words.
	3. Ambiguity or Pragmatic Context Easy: No ambiguity or external context needed. Hard: Paraphrase identification depends on pragmatic/contextual cues or disambiguation.		3. Length and Depth of the Text Easy: Short and simple sentences expressing clear sentiment. Hard: Longer, more complex sentences with mixed feelings or abstract expressions.
Paraphrase Generation	1. Lexical and Syntactic Simplicity Easy: Minor changes in wording, using synonyms or simple reordering of sentence components. Hard: Significant changes in sentence structure, more complex rewording, or use of idiomatic expressions.	General Question Answering	1. Directness of the Question Easy: Questions that are straightforward and directly answerable from the context without the need for additional reasoning. Hard: Questions that require interpretation, inference, or reasoning based on the provided context.
	2. Semantic Consistency Easy: Paraphrases that are almost identical in meaning, with small vocabulary swaps. Hard: Paraphrases that need to capture more subtle or nuanced meanings, requiring deeper understanding of the original context.		2. Contextual Clarity Easy: Clear, unambiguous context where the answer is explicitly stated. Hard: Ambiguous or vague context, requiring deeper understanding or synthesis of multiple pieces of information.
	3. Creativity and Stylistic Variety Easy: Simple, direct rewording with minimal stylistic variation. Hard: Creative paraphrases that involve stylistic or tone shifts, or those that include complex syntactic changes while preserving the original meaning.		3. Factual vs. Inferential Answers Easy: The answer is directly stated or easily extracted from the text. Hard: The answer involves synthesizing information, making inferences, or requires external knowledge.
Text Summarization	1. Text Length Easy: Short input text that already highlights key points. Hard: Long, complex text with multiple subtopics or dense information.	Fact Check	1. Factual Specificity Easy: Questions asking for easily verifiable, specific facts, such as dates, names, or locations. Hard: Questions that require detailed verification or fact-checking from multiple sources or with some level of ambiguity.
	2. Content Density Easy: Clear, well-structured information where the main idea is obvious. Hard: Dense or technical content requiring understanding of the details to create a meaningful, concise summary.		2. Contextual Complexity Easy: The answer can be found in a single piece of well-known, widely agreed-upon information. Hard: The answer requires analyzing complex data, multiple interpretations, or understanding a situation in context to determine accuracy.
	3. Context and Subtlety Easy: Little to no background knowledge required for the summary. Hard: Summarization requires inference, understanding context, or synthesizing information from different parts of the text.		3. Source Dependence Easy: Clear, straightforward facts that can be verified from commonly known or easily accessible sources (e.g., official records, universally recognized events). Hard: The fact being checked is less universally agreed upon, has evolving data, or involves conflicting reports.

Figure 7: Complexity Definitions for each sub-category.

Sub Category	Complexity	Sub Category	Complexity
Headline Generation	1. Text Length Easy: Short article with a single main idea, easy to extract the core message. Hard: Long article with multiple ideas or subtopics, requiring discernment of the most critical information.	Multi Turn Conversation	1. Contextual Continuity Easy: The conversation is straightforward, with clear transitions and a simple back-and-forth interaction. Hard: The conversation involves complex or abstract ideas, requiring the model to recall, understand, and reference multiple pieces of context across several turns.
	2. Content Type Easy: Straightforward, factual content with a clear subject and action. Hard: Complex content with nuanced language, needing a creative yet accurate headline.		2. Response Generation Easy: Responses require little reasoning and are based on direct information from previous turns. Hard: Responses require inferences, deeper reasoning, or nuanced understanding of prior context and the user's implied needs.
	3. Information Distillation Easy: Clear event or announcement that can be summarized in a few words. Hard: Abstract, multifaceted articles requiring summarization of broader themes or mixed topics.		3. Domain Knowledge Easy: The conversation stays within a simple domain, requiring basic information retrieval and straightforward interaction. Hard: The conversation involves specialized knowledge or abstract topics, requiring the system to adapt dynamically based on the evolving context.
Role-Playing	1. Character Consistency Easy: The character or persona is straightforward, with clear and simple attributes. There is little deviation in how the character should behave. Hard: The character requires nuanced understanding, deep role adaptation, or responding to complex scenarios while staying in-character.	Indian Cultural Context	1. Depth of Cultural Understanding Easy: The query requires a basic, widely understood cultural concept or tradition that does not involve nuanced details or historical context. Hard: The query delves into more specialized or nuanced cultural practices, traditions, or beliefs, which might require historical or region-specific knowledge.
	2. Creativity and Imagination Easy: The role-play is based on established norms, and responses follow predictable patterns. Hard: The role requires creative input, improvisation, or handling abstract or unpredictable situations while maintaining authenticity in the persona.		2. Specificity of the Query Easy: The question is about a widely known and commonly practiced cultural topic or festival that is universal across India. Hard: The question pertains to a specific regional tradition, a lesser-known festival, or cultural practice that might differ significantly across regions or communities.
	3. Contextual Depth Easy: The role-playing context is simple, and only surface-level understanding of the role is necessary. Hard: The role requires understanding deep philosophical, historical, or emotional aspects of the character, possibly over multiple turns or complex scenarios.		3. Integration with Broader Context Easy: The query asks for a general explanation without requiring integration of cultural beliefs, practices, and their social implications. Hard: The query requires an explanation that integrates various aspects of Indian culture, including its philosophical, religious, or social implications, and might ask for comparative or cross-regional analysis.
Advice-Seeking	1. Depth of the Problem Easy: The user's issue is straightforward, commonly understood, and doesn't require deep insight or complex problem-solving. Hard: The problem is more complex, requiring a nuanced approach, multiple perspectives, or deeper understanding to provide meaningful advice.	Comprehension	1. Depth of Information Extraction Easy: The task requires directly extracting information that is clearly stated in the passage, with no need for inference. Hard: The task requires inference, where the answer is implied but not directly stated in the passage, or requires synthesizing information from multiple parts of the text.
	2. Personalization and Context Easy: The user's issue is general, and the advice can be applied broadly without needing to account for specific personal factors. Hard: The user's query involves personal, emotional, or contextual factors, and the advice must be tailored to their unique circumstances.		2. Clarity and Directness of the Passage Easy: The passage is simple, with clear and direct statements that make it easy to extract answers. Hard: The passage is complex, contains nuanced language, or presents a more abstract concept, requiring deeper understanding or analysis.
	3. Emotional Sensitivity Easy: The user's issue does not involve strong emotional components or stress. Hard: The user's situation requires empathetic, emotionally aware responses to ensure the advice is sensitive and supportive.		3. Complexity of the Question Easy: The question asks for a straightforward fact or detail that is easily accessible from the passage. Hard: The question requires critical thinking, such as cause-effect relationships or understanding of multiple layers of context.

Figure 8: Complexity Definitions for each sub-category.

Sub Category	Complexity	Sub Category	Complexity
Planning	1. Scope of the Plan Easy: The plan is narrow in scope, with a limited number of steps or elements to consider. Hard: The plan involves multiple variables, such as many tasks, time constraints, or various dependencies that must be addressed. 2. Level of Detail Easy: The plan requires only basic details like locations or general activities with little need for deep customization. Hard: The plan requires detailed steps, including specific times, types of activities, resources, and possibly even contingency plans for unexpected situations. 3. Customization and Constraints Easy: The user's preferences or constraints are minimal or common, and the plan can be generated with general knowledge. Hard: The plan must account for specific user preferences, such as dietary restrictions, mobility concerns, budget, or local availability of services.	Math QA	1. Mathematical Operation Complexity Easy: The problem involves simple arithmetic operations like addition, subtraction, multiplication, or division that can be solved quickly. Hard: The problem requires more complex operations like algebraic manipulation, geometry, or multi-step calculations involving fractions, percentages, or systems of equations. 2. Problem Complexity (Word Problem Difficulty) Easy: The problem is straightforward, with clear and simple wording. No ambiguity or need for interpretation exists. Hard: The problem may involve convoluted or indirect phrasing that requires parsing or understanding multiple pieces of information to solve. 3. Number of Steps Required Easy: The problem can be solved with one simple operation, often involving direct calculation. Hard: The problem involves multiple steps or a combination of operations to arrive at the final answer.
Brainstorming	1. Creativity and Originality Easy: The task requires simple, commonly known ideas or solutions that do not need much innovation or deep thought. Hard: The task demands unique, novel, or unconventional ideas that require extensive creativity or thinking outside the box. 2. Scope and Specificity of the Problem Easy: The task is relatively broad and general, allowing for a wide range of ideas without needing to consider detailed constraints or limitations. Hard: The task involves more specific or narrow areas, requiring ideas to fit within defined parameters, which could include specific industries, goals, or limitations. 3. Depth of Domain Knowledge Required Easy: The task is general enough that no specific domain knowledge is required to generate ideas, or the task is on a topic familiar to a wide audience. Hard: The task requires specialized knowledge of a particular domain (e.g., technology, healthcare, or economics) to generate ideas that are realistic and feasible.	Math Instruction Tuning	1. Mathematical Operation Complexity Easy: The math problem involves basic arithmetic operations such as addition, subtraction, multiplication, or division. The instruction is simple and clear. Hard: The problem involves advanced concepts like algebra, geometry, calculus, or multi-step problem solving, requiring deeper understanding or manipulation of equations. 2. Instruction Complexity (Reasoning Required) Easy: The instructions are straightforward and involve a single step to solve. Hard: The instructions require several intermediate steps or involve solving multi-step equations with additional reasoning, like factoring, completing the square, or applying specific mathematical theorems. 3. Number of Steps in Solution Easy: The solution involves few steps—often just a direct application of one operation. Hard: The solution involves multiple steps, such as simplifying expressions, factoring, or breaking down complex terms.
Information Seeking	1. Depth of Information Easy: The question seeks basic, straightforward information that is commonly available and well-known. Hard: The question requires more nuanced, detailed, or in-depth information, possibly involving multiple sources or complex concepts. 2. Specificity of the Query Easy: The question is general and can be answered with a clear, direct response. Hard: The question is specific or highly detailed, requiring an elaborate response or synthesis of various aspects. 3. Domain or Contextual Knowledge Easy: The information sought is familiar or common knowledge, requiring little specialized knowledge to answer. Hard: The information involves a specialized domain, or understanding the context or background knowledge is necessary to provide a comprehensive answer.	Math Proofs	1. Theorem Complexity Easy: The theorem is simple and straightforward, typically involving basic set theory, arithmetic, or algebra. Hard: The theorem involves advanced concepts such as topology, number theory, or higher-level algebra. The logic required to prove the theorem might be more complex and abstract. 2. Proof Structure Easy: The proof is short, involving direct logical steps with no intermediate lemmas or deep reasoning. Hard: The proof involves multiple intermediate steps, possibly requiring auxiliary lemmas, the use of multiple mathematical properties, or breaking down the proof into several cases. 3. Mathematical Concepts Involved Easy: The proof relies on simple mathematical concepts, such as basic set operations or arithmetic properties. Hard: The proof uses complex mathematical ideas, such as induction, contrapositive reasoning, or advanced theorems, and may require careful justification of each step.

Figure 9: Complexity Definitions for each sub-category.

Sub Category	Complexity	Sub Category	Complexity
Code Generation	1. Problem Simplicity Easy: Involves basic control structures or common patterns. Hard: Requires algorithmic thinking, multi-step planning, or integration with libraries/APIs. 2. Code Length Easy: One-liner or small function. Hard: Multiple functions, modular structure, or boilerplate setup. 3. Conceptual Complexity Easy: Uses elementary constructs like loops, conditionals. Hard: Requires concepts like recursion, concurrency, or complex data structures. 4. Edge Cases and Testing Easy: Few or no edge cases. Hard: Needs robust handling of invalid inputs, errors, or corner cases.	Unit Test Generation	1. Input Complexity Easy: Function accepts primitive types and has predictable output. Hard: Function has multiple inputs, nested structures, or dynamic behavior. 2. Output Behavior Easy: Deterministic and consistent output. Hard: Output depends on side effects, randomness, or state. 3. Edge Case Coverage Easy: Few or obvious edge cases. Hard: Many possible failure modes or boundary conditions. 4. Format Required Easy: Just input/output pairs or informal descriptions. Hard: Tests in specific frameworks with setup/teardown, mocking, etc.
Code Debugging	1. Error Obviousness Easy: Bug is a syntax error or obvious logic mistake (e.g., == vs =). Hard: Bug is subtle, involves side effects, off-by-one logic, or misunderstood APIs. 2. Code Scope Easy: Small functions with limited state or logic. Hard: Multiple functions, state tracking, or external dependencies involved. 3. Required Knowledge Easy: General debugging, print/log inspection suffices. Hard: Requires domain-specific understanding, complex reasoning, or tool-based debugging. 4. Impact of Fix Easy: Fix is localized and doesn't affect other logic. Hard: Fix has ripple effects across the codebase and must maintain correctness globally.	Code Theory	1. Concept Familiarity Easy: Widely taught or intuitive topics (e.g., loops, arrays, O(n)). Hard: Niche or abstract topics (e.g., amortized complexity, monads). 2. Explanation Depth Easy: Surface-level description of how something works. Hard: Deep dive into trade-offs, limitations, and formal reasoning. 3. Required Background Easy: Assumes basic CS knowledge. Hard: Requires formal math or algorithmic training. 4. Application Scope Easy: Explains theory in isolation. Hard: Tied to real-world design or performance decisions.
Code Editing	1. Edit Type Easy: Cosmetic or small structural change (e.g., rename, add a check). Hard: Deep refactor or major logic change across files. 2. Context Required Easy: Local context is enough to apply the edit. Hard: Requires understanding of broader dependencies or interfaces. 3. Behavior Change Easy: No change in behavior (e.g., renaming, reformatting). Hard: Behavior is intentionally changed, requiring correctness guarantees. 4. Risk Level Easy: Low-risk edit, unlikely to break anything. Hard: High-risk edit, may affect performance, correctness, or security.	Code Review	1. Code Quality Easy: Well-structured code with clear intent. Hard: Poorly written, unstructured, or ambiguous code. 2. Scope of Review Easy: Small diff or single function. Hard: Large PR with multiple files and architectural impact. 3. Standards Compliance Easy: Trivial formatting or naming violations. Hard: Requires domain knowledge, security awareness, or architectural judgment. 4. Reviewer Responsibility Easy: Only comments on style or clarity. Hard: Needs to verify correctness, performance, and maintainability.
Code Explanation	1. Code Length Easy: Few lines with straightforward logic. Hard: Long code with deeply nested structures. 2. Abstraction Level Easy: Explains "what it does." Hard: Needs to explain "why," covering design decisions or trade-offs. 3. Concept Depth Easy: Basic logic with familiar constructs. Hard: Involves complex algorithms, libraries, or domain-specific behavior. 4. Audience Adaptation Easy: Target audience already knows the language and topic. Hard: Requires tailoring to beginners or non-technical audiences.	Repository level Code Generation	1. Architecture Complexity Easy: Scaffold for a simple project (e.g., a CRUD app). Hard: Full-stack system with frontend, backend, DB, and CI/CD config. 2. Component Integration Easy: Independent modules or templates. Hard: Components must interoperate with consistent API contracts. 3. Language/Framework Mastery Easy: One language, minimal config. Hard: Polyglot codebase (e.g., Python backend + React frontend + Docker). 4. Customization Required Easy: Follows standard patterns (e.g., Flask starter). Hard: Needs tailoring to specific domain, logic, or 3rd-party APIs.

Figure 10: Complexity Definitions for each sub-category.

Sub Category	Complexity	Sub Category	Complexity
Code Translation	1. Language Similarity Easy: Translation between similar paradigms (e.g., Python ↔ JavaScript). Hard: Translation between very different paradigms (e.g., Python ↔ Haskell). 2. Feature Mapping Easy: All features map cleanly between languages. Hard: Requires emulating behavior (e.g., list comprehensions, async models). 3. Code Size Easy: Short, self-contained function. Hard: Large module with classes, decorators, or closures. 4. Idiomatic Translation Easy: Direct one-to-one line conversion. Hard: Needs idiomatic rewriting to follow target language conventions.	Function Calling	1. Complexity of Function Signatures Easy: The function signatures are straightforward, with simple parameter types (e.g., integers, strings) and no additional constraints. Hard: The function signatures involve complex parameters, such as nested structures, multiple required fields, or specialized types like objects, arrays, or custom types (e.g., qubits, Fourier components, etc.). 2. Interpretation of User Requirements Easy: The user request is simple and directly maps to the required parameters without ambiguity. There are minimal or no special considerations for things like optional arguments or conflicting values. Hard: The user's request is complex, involving multiple steps of reasoning or context-dependent choices, requiring careful extraction of parameters and appropriate default values for missing information. 3. Parameter Mapping and Constraints Easy: The parameters have clear, unambiguous mappings (e.g., integers for quantities, strings for identifiers), and there are no special conditions or dependencies between parameters. Hard: The parameters have complex dependencies (e.g., the number of Fourier components must align with the number of qubits) or constraints that must be satisfied (e.g., time complexity requirements, control operations). 4. Function Logic and Output Generation Easy: The function call results in simple, easily interpretable outputs that don't require additional processing or understanding of underlying logic. Hard: The function call involves sophisticated logic that results in complex outputs (e.g., quantum circuits), or the output requires further interpretation or refinement to be useful in the context.
Instruction Following	1. Task Complexity Easy: The task involves simple, straightforward instructions with minimal structure or specific formatting requirements. Hard: The task requires multi-step reasoning, adherence to multiple specific formatting guidelines, and generation of complex or detailed content. 2. Output Structure and Formatting Easy: Instructions involve generating basic content with minimal formatting or structural constraints (e.g., short answers or paragraphs without complex guidelines). Hard: The instructions demand precise formatting (headings, bullet points, specific word count limits, etc.), and the content must be logically organized and adhere to multiple constraints. 3. Level of Detail and Precision Easy: The instructions require a general response with little emphasis on fine-grained details, and minor flexibility is allowed. Hard: The instructions demand highly detailed responses with precise limits (e.g., word count per section), strict content divisions (subsections), and comprehensive coverage of a topic. 4. Interpretation of Instructions Easy: The instructions are unambiguous, with clear expectations and no need for interpretation beyond direct execution. Hard: The instructions involve multiple layers of interpretation (e.g., understanding how to balance word count with content depth, determining how to break down sections logically, or ensuring that all key points are covered).	Safety & Non - Compliance	1. Clear vs. Ambiguous Harmful Requests Easy: The harmful or unethical request is clear and direct, with no ambiguity or need for interpretation. Hard: The request is more subtle, requiring the model to understand context, infer potential harm, or identify violations that may not be immediately obvious. 2. Politeness and Firmness Easy: The model simply needs to reject the request without much elaboration or any need to soften the response. Hard: The model needs to balance politeness, firmness, and clarity, ensuring it is both respectful and firm in refusal while explaining why the request is not permissible. 3. Contextual Decision-Making Easy: The request is easily recognized as violating ethical norms or policies, so the response is straightforward. Hard: The model must assess context, which might involve determining whether the request is borderline or masked behind polite language, requiring careful judgment and nuanced response. 4. Complexity of Policy Interpretation Easy: The model only needs to reject straightforward violations, such as promoting illegal activities or encouraging harm. Hard: The model must navigate complex or gray areas where there may be multiple interpretations of policies or ethics, requiring it to consider indirect implications.

Figure 11: Complexity Definitions for each sub-category.

Constraint Type	Description	Examples
Length Constraints	Defines how long or concise a response should be.	"Answer in one sentence", "Summarize in 50 words", "Explain in two paragraphs"
Structural Formatting	Governs the organization of the output.	Paragraphs, bullet points, numbered lists, tables ("Compare A and B in a table"), JSON, XML, HTML
Tone and Style	Specifies tone, voice, or target audience.	Formal, conversational, academic, friendly, motivational, child-friendly, expert-level
Persona / Role Play	Instructs the model to adopt a specific identity.	"Act as a doctor/teacher/software engineer", "Answer like Albert Einstein or a pirate"
Reasoning Style	Defines how the model should arrive at the answer.	Chain-of-thought, step-by-step explanation, final answer only, show intermediate steps
Time / Historical Context	Specifies a temporal or historical frame.	"Explain from a 19th-century perspective", "What would a Roman say?", "Imagine you are in 2100"
Output Type / Genre	Determines the nature of the response.	Poem, story, dialogue, screenplay, essay, report, memo, tweet, press release
Comparative / Multi-perspective	Requests multiple viewpoints or contrasts.	"Compare pros and cons", "Show both sides of the argument", "Write from X's and Y's perspective"

Figure 12: Instruction following: constraint types and examples

Dimension	Explanation
Lexical Diversity	The variety and richness of words used in the text. Higher lexical diversity means using a wide range of vocabulary instead of repetitive or generic terms. Why It Matters for Indic Data: Indic languages have rich vocabulary and multiple synonyms; poor diversity makes translations monotonous and unnatural. Errors in this area: Overuse of common words, failure to use synonyms, repetitive phrasing.
Coherence and Cohesion	Coherence: Logical flow and overall sense of the text. Cohesion: Use of linguistic devices (connectors, pronouns, conjunctions) to link sentences smoothly. Why It Matters for Indic Data: Many Indic languages use connectives and honorific markers that impact cohesion; direct translation from English often breaks these links. Errors in this area: Disconnected sentences, abrupt topic shifts, missing conjunctions or pronouns.
Completeness	Whether the output contains all necessary information from the source without omissions or additions. Why It Matters for Indic Data: When translating long or complex Indic sentences, models often skip certain parts (e.g., verb phrases or subordinate clauses). Errors in this area: Missing phrases, dropped entities, truncated sentences, or extra hallucinated details.
Grammatical Accuracy	Correct use of grammar rules (syntax, tense, agreement, case, morphology). It affects fluency and correctness of the output. Why It Matters for Indic Data: Indic languages have complex inflectional morphology and word order; errors often occur in case endings, gender/number agreement, and verb conjugations. Errors in this area: Wrong tense, missing auxiliary verbs, subject-verb disagreement, wrong case marking.
Named Entity Handling	Correct recognition and rendering of named entities (persons, places, organizations, dates, currencies, etc.) across languages.

Table 6: Key dimensions for assessing LLM translations in Indic languages, highlighting their significance and common errors.

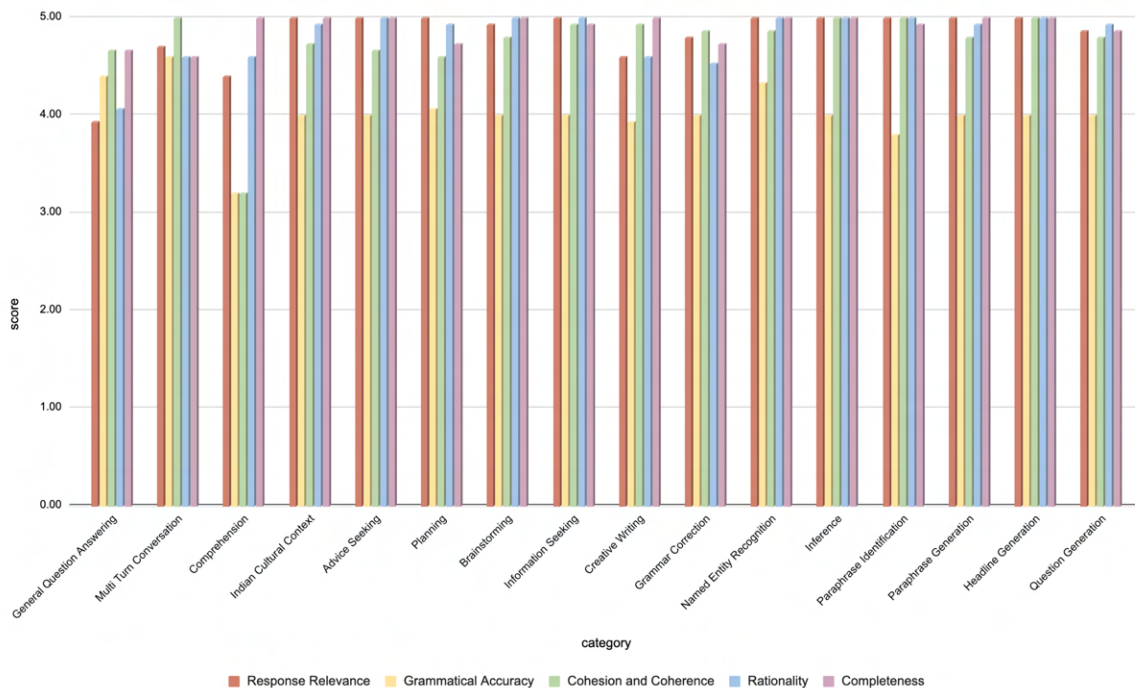


Figure 13: LLM generation quality evaluation scores across various categories.

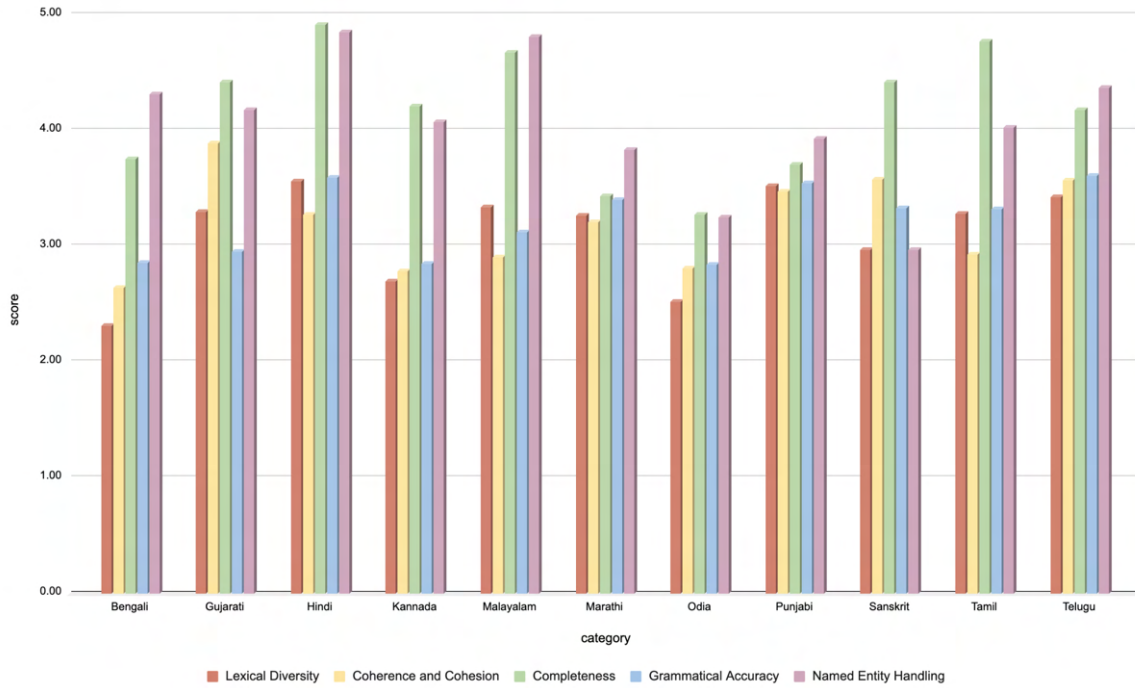


Figure 14: LLM translation quality evaluation scores across 11 Indian languages.

in Section 3.3. In addition to these specifications, dedicated frameworks for quality assurance are provided, encompassing both language quality verification and content quality verification, as illustrated in Figures 19 and 20, respectively.

D Examples

We provide examples of both instruction and preference tuning datasets in Figures 21-30.

Prompt used for Generation of Easy Indian Cultural Context English Prompts in Self Instruct Pipeline
System Prompt You are tasked with generating easy-level, culturally diverse, and descriptive questions that explore various aspects of Indian culture. Your questions should avoid being too generic or frequently asked. Instead, focus on less commonly discussed but still culturally meaningful topics that are easy to understand and answer. Guidelines for Generating Questions: 1. Difficulty Level – Easy: The question should be simple enough for someone with basic cultural awareness to answer in a few sentences. Avoid complex or academic phrasing. Each question should address only one clear aspect or topic — do not combine multiple elements in a single question. 2. Question Type – Descriptive: The answer should require a short paragraph, not a one-word or one-line response. Avoid yes/no or multiple-choice style formats. 3. Scope – Specific but Not Niche: Focus on fine-grained yet non-obscure topics. Do not ask for broad essays like “Describe a festival” or over-complicated questions like “What is the significance of X across different regions with its variations and philosophical meanings?” 4. Diversity – Pan-Indian Coverage: Ensure questions represent diverse states, communities, religious groups, art forms, customs, and practices. Cover a variety of domains like traditional food, clothing, local art, customs, symbols, minor rituals, or lesser-known celebrations. Use the given questions as examples of how to frame the questions. Generate the instructions in the json format: <pre>{ ... Question i: ith Question, ... }</pre> Note that the instructions must be in english language. Return only the json format response, do not return anything other than json. User Prompt Here are some example questions: <8 Easy Indian Cultural Context English Prompts as Few Shots>

Figure 15: Prompt used for generation of easy Indian cultural context English prompt in Self-Instruct pipeline.

Prompt used for Generation of Hard Indian Cultural Context English Prompts in Self Instruct Pipeline
System Prompt You are tasked with generating hard-level, culturally diverse, descriptive questions that explore complex, nuanced, or region-specific aspects of Indian culture. These questions should go beyond basic or commonly asked topics, aiming to challenge the responder's depth of cultural understanding. Guidelines for Generating Questions: 1. Difficulty Level – Hard: The questions should require a detailed and informed response, often involving historical context, philosophical meaning, or regional variation. Expect the responder to draw from less commonly known knowledge, comparative insights, or interconnected traditions. Avoid overly academic phrasing, but do not oversimplify the question. 2. Question Type – Descriptive: The answer should require at least a short essay-style paragraph, explaining the underlying cultural, religious, philosophical, or social dimensions. Do not create multi-part or overly long questions — keep one core focus per question, but explore it deeply. 3. Scope – Specialized or Regionally Rooted: Focus on less mainstream topics, regional traditions, or philosophical/religious practices that are not typically covered in general knowledge. Explore cultural beliefs, historical movements, ritualistic practices, sectarian philosophies, regional art forms, and symbolic meanings. 4. Diversity – Pan-Indian and Cross-Traditional: Ensure a wide range of questions from different Indian states, languages, communities, castes, and religious traditions. Include topics from Hindu, Buddhist, Jain, Sikh, Islamic, Christian, tribal, and other regional traditions. Use the given questions as examples of how to frame the questions. Generate the instructions in the json format: <pre>{ ... Question i: ith Question, ... }</pre> Note that the instructions must be in english language. Return only the json format response, do not return anything other than json. User Prompt Here are some example questions: <8 Hard Indian Cultural Context English Prompts as Few Shots>

Figure 16: Prompt used for generation of hard Indian cultural context English prompt in Self-Instruct pipeline.

Prompt used for Translation via LLM
System Prompt Translate the following text into {lang}, targeting educated native speakers in an academic or professional setting. Translation Guidelines: 1. Exclusive Translation: Translate the provided text only. No answers, solutions, interpretations, or commentary. 2. Preserve Original Structure: Maintain the input's original structure, including equations, formulas, syntax, and logical structures. Retain original text if direct translation alters meaning. 3. Prioritize Accuracy: Ensure accurate translation of technical terms and concepts using widely accepted academic terminology. 4. Maintain Meaning and Intent: Preserve the original meaning, intent, and challenge. Maintain syntax and variable names for code and math. 5. Translate Instructions/Metadata Verbatim: Translate instructions (e.g., "Enter your answer here:", "Return your final response within \boxed{.}") and metadata (e.g., file names, timestamps) exactly as they appear, without attempting to execute or interpret them. 6. Translation Requirements (Reinforced): Perform translations only; do not add new text. Make no changes beyond accurate word-for-word translation. Leave code, formulas, placeholders, and instructions untouched. User Prompt Text to Translate: {text}

Figure 17: Prompt used for translation via LLM.

Prompt used for Evaluating models DPOed on 100K Pragmaan-Align Data via LLM as Judge
System Prompt You are an impartial evaluator. Your task is to judge a model response against a given ground truth answer for a prompt. Evaluate the response with a real number score between 1-5, using **one decimal place only**. Never output integers; always use a decimal number with one digit (e.g., 3.7, 4.2). Consider the category of the prompt when scoring. Categories and their focus are: Reasoning – Logical correctness and clarity of explanation. Role Playing – Staying in character, creativity, and consistency. Math – Accuracy of calculations and clarity of steps. Editing – Improving or revising text while preserving meaning. Information Seeking – Factual accuracy and relevance. Advice Seeking – Practicality, usefulness, and clarity. Data Analysis – Correct interpretation of data and logical conclusions. Creative Writing – Originality, coherence, style, and engagement. Planning – Clear, feasible, and well-structured steps. Brainstorming – Creativity, diversity, and relevance of ideas. Provide scores for each response separately. Output must strictly follow this format: "Response score: [[x.y]] <Explanation of your scoring, highlighting strengths and weaknesses>." Notes: Focus on correctness, relevance, and quality over verbosity. Be critical but fair. Base scores only on alignment with the category requirements and the ground truth. Always provide an explanation after the score. User Prompt < Prompt > {prompt} < Ground Truth > {ground_truth} < Response > {response}

Figure 18: Prompt used for evaluating post-trained model on *Pragmaan-Align* data via LLM-as-a-Judge.

1. Language Quality Verification Guidelines
The evaluation of Prompt–Response Pairs (PRPs) focuses on key linguistic and stylistic parameters to ensure high-quality outputs across Indic languages. These include fluency, grammatical accuracy, coherence and cohesion, sentence structure, lexical diversity, and accurate handling of named entities. The following principles guide this process: 1.1 Preserve Naturalness and Meaning Translations should read naturally in the target language while preserving the original intent. Avoid literal, word-for-word translations. Instead, adapt phrasing and sentence structure in accordance with the linguistic norms of the respective Indic language. Ensure cultural appropriateness in expressions and idiomatic usage. 1.2 Localize Editing Instructions Adapt instructions so they align with how editing tasks are conventionally expressed in Indic languages. Example: English: "Rewrite the sentence in Active Voice." Hindi: "वाक्य को सक्रिय वाच्य में पुनः लिखें।" 1.3 Maintain Consistency in Terminology Use widely accepted equivalents for domain-specific terms (e.g., AI, computing, mathematics, medicine, law). Where no direct equivalent exists, apply transliteration while ensuring clarity of meaning. Examples: English: "Refactor the code for better readability." Hindi: "कोड को बेहतर पठनीयता के लिए पुनः व्यवस्थित करें।" 1.4 Follow Proper Formatting and Punctuation Apply punctuation rules specific to the target Indic language rather than English conventions. Ensure proper spacing and script usage. Avoid using Latin punctuation marks such as . where language-specific marks like । (Hindi) or । (Bangla) should be used. Example: Incorrect: "वाक्य को सक्रिय वाच्य में पुनः लिखें." Correct: "वाक्य को सक्रिय वाच्य में पुनः लिखें।" 1.5 Handle Untranslatable Terms Appropriately For terms such as brand names, technical concepts, or programming keywords, use transliteration in Indic script where possible. Maintain semantic clarity in context. Example: English: "Replace the variable with a meaningful name." Hindi: "चर/वैरिएबल को एक अर्थपूर्ण नाम से बदलें।" 1.6 Code Snippets and Symbols Retain code snippets in their original form without translation. Example: English: "Change int x = 10; to int y = 20;" Hindi: "int x = 10; को int y = 20; में बदलें।" 1.7 Localization of Editable Text When localizing content that requires editing, the translated text should reflect the specific requirement of the instruction. For instance, if the instruction involves correcting grammatical errors, the translated text should intentionally include such errors before editing. Example: English: "Correct the grammatical mistakes in the following text: 'I have took the medicines.'" Hindi: "दिए गए पाठ में व्याकरण संबंधी गलतियों को सुधारें। 'मैं दवा खाता।'"

Figure 19: Comprehensive language quality guidelines outlining key linguistic dimensions such as grammar, fluency, clarity, and naturalness for ensuring high-quality annotated data.

2. Content Quality Verification Guidelines
Each Prompt–Response Pair (PRP) undergoes a structured evaluation process to ensure alignment with six core dimensions: complexity, multi-turn interaction, instruction-following capability, safety considerations, cultural relevance to the Indian context, and inclusion of reasoning or thinking trails. PRPs that fully meet these criteria are accepted without modification, while those exhibiting minor inconsistencies undergo targeted revisions to ensure compliance. PRPs with significant deviations are comprehensively rewritten to maintain both linguistic quality and adherence to the specified configuration requirements. Definitions for each category and setting are explicitly documented and must be strictly followed. Beyond compliance with language and setting-specific requirements, the following general principles guide annotation and validation: 2.1 Relevance to the Prompt The response must accurately and completely address the intent of the prompt. Extraneous details, digressions, or unrelated content should be avoided. Responses that partially satisfy the prompt but omit essential information should be deprioritized. 2.2 Fluency and Naturalness Responses should exhibit grammatical correctness, structural coherence, and logical flow. Language must read naturally in the target Indic language, adhering to its specific linguistic conventions and avoiding literal or awkward phrasing. Outputs containing major syntactic or grammatical errors are considered suboptimal. 2.3 Factual Accuracy For prompts requiring factual information, responses must be accurate and verifiable. Any response containing incorrect, misleading, or unverifiable claims should be rejected. When multiple responses are factually correct, preference should be given to those that are concise, precise, and clearly articulated. 2.4 Cultural and Linguistic Appropriateness Responses should conform to cultural norms and linguistic standards of the target Indic language. Avoid regionally inappropriate idioms, metaphors, or expressions unless explicitly required by the prompt. Ensure contextual relevance for the intended audience of the target language. 2.5 Style and Formality Tone and level of formality must align with the prompt’s context (e.g., academic, conversational, or instructional). Responses that fail to maintain the appropriate tone or deviate from the intended style should be deprioritized. 2.6 Avoiding Bias and Stereotypes Responses must not reinforce stereotypes or exhibit bias related to gender, caste, religion, ethnicity, or socio-economic status. Use neutral and inclusive language wherever applicable. If all available responses contain bias, flag them for further review. 2.7 Filtering Harmful or Offensive Content Reject responses containing hate speech, explicit or offensive material, political bias, or harmful language. Any content promoting violence, discrimination, or misinformation should not be selected. All outputs must comply with ethical AI principles and community guidelines. 2.8 Handling Subjectivity in Sensitive Topics For prompts concerning controversial or sensitive domains (e.g., politics, religion, social issues), responses should maintain a neutral and balanced tone. Avoid presenting subjective opinions as objective facts. Favor informative, fact-based responses over emotionally charged or opinionated ones. 3. Final Considerations Consistency in annotation is critical; all guidelines must be applied uniformly across evaluations. Assess responses solely on their intrinsic quality—model attribution (e.g., ChatGPT, Krutrim) should not influence judgment. Prefer responses that deliver essential details succinctly, avoiding unnecessary verbosity.

Figure 20: Content quality guidelines detailing principles for maintaining factual accuracy, contextual relevance, safety, and adherence to task-specific requirements during data annotation.

Category	Prompt	Preferred Response	Rejected Response
Advice Seeking	माझ्या मनेजरला वाटत की कामाच्या ठिकाणी माझी संवाद शैली असंवेदनशील आणि संघटित आहे. त्यांना वाटत की मी माझ्या सहकाऱ्यांशी विचारपूर्वक संवाद साधायचा. मला कधीच वाटत नाही की मी असंवेदनशील असेन, म्हणून मला हे ऐकून खरेतर आश्चर्य वाटतं. मी फक्त थेट आणि मुद्द्याचं बोलीतो. मी प्रामाणिकपणे, परंतु थोडक्यात बोलताना इतरांच्या भावना लक्षात ठेवून त्यांच्याशी संवाद कसा साधू शकतो?	आपल्या मनेजरच्या टिप्पण्या ऐकून आपल्याला आश्चर्य वाटणे स्वाभाविक आहे. संवादाची शैली वैयक्तिक असू शकते, आणि आपले थेट आणि मुद्देसुद्द दिसणे काही वेळा इतरांमधील संवेदनशीलता कमी करू शकते. इतरांच्या भावना लक्षात घेतल्यास, संवादाची शैली सुधारण्यासाठी काही टिप्स येथे दिल्या आहेत: 1. “सकारात्मक आरंभ”: आपले बोलणे सुरू करताना प्रत्येक गोष्टीत काही सकारात्मकता समाविष्ट करा. उदाहरणार्थ, “मला तुमचं विचार आवडले, पण...” अशी शैली वापरल्यास, आपली अभिप्राय थोडी सौम्य आणि सकारात्मक निर्माण होईल. 2. “सवाल विचारणे”: संवादत समालिप्त झालेल्या व्यक्तीच्या विचारांची आणि भावना जाणून घेण्यासाठी प्रश्न विचारा. हे संवादाला संवादात्मक बनवतो आणि टीममधील सहकार्य वाढवतो. 3. “आवाजाचा टोन”: आपल्या आवाजाचा टोन आणि शरीराची भाषा महत्वाची आहे. थेट बोलताना सॉफ्ट टोन वापरणे आणि थोडा हसत बोलणे, इतरांनाही आरामदायक बनवू शकतात. 4. “त्यांच्या भावना समजून घेणे”: आपण जेव्हा कोणाशी बोलता, तेव्हा त्यांची स्थिती आणि भावना लक्षात घेणे महत्वाचे आहे. “तुमचं मत मोजलं जातं आणि मला तुमचं विचार समजून घ्यायला आवडेल,” असा टुप्चिकोन ठेवता येईल. 5. “परतफेड देणे”: कोणताही अभिप्राय दिल्यावर, त्यात सकारात्मकता ठेवा. जर कुणी काही विचार केले असेल, तर त्याचा आदर करा आणि त्यावर चर्चा करा. 6. “संपूर्णता”: बोलताना थोडी थांबा, जेणेकरून इतरांनाही आपल्या विचारांची पूर्णता समजून घेता येईल. हे विशेषतः गुंतागुंतीच्या मुद्द्यांवर लागू शकते. 7. “प्रतिक्रिया विचारणे”: त्यावर चर्चा केल्यानंतर, आपल्या सहकाऱ्यांकडून आपली संवाद शैली कशी वाटली याबद्दल काही फीडबॅक घ्या, हे आपल्याला पुढील संवाद सुधारण्यास मदत करेल. आपण आपल्या संवाद शैलीत सुधारणा करण्याचे प्रयत्न करत असताना, नेहमी लक्षात ठेवा की प्रत्येक जण वेगळा असतो आणि संवाद करण्याची शैली देखील वेगळी असते. आपल्या संवादाच्या जागरूकतेने आणि संवेदनशीलतेने आपले सहकारी आणि टीमसह चांगले संबंध निर्माण होण्यास मदत करेल.	तुम्हाला तुमच्या व्यवस्थापकाकडून अभिप्राय मिळाल्यानंतर तुम्ही धक्का बसले आहेत याचा अर्थ समजून घेण्यासारखा आहे. कोणत्याही कार्यालयीन वातावरणात प्रभावी संवाद महत्त्वपूर्ण आहे आणि तो अनेकदा वैयक्तिक पसंतीवर अवलंबून बदलतो. अधिक चिंताग्रस्त किंवा सूक्ष्म टुप्चिकोन स्वीकारण्यासाठी स्वतःची पुनर्निर्माणित करण्याचा प्रयत्न करा; प्रत्येक संभाषणासाठी योग्य शब्द निवडा आणि आवश्यकतेनुसार तपशीलवार स्पष्टता देऊन तुमचे संदेश सुधारण्याचा सराव करा. जर एखाद्या विचारावर मत व्यक्त केले तर ते अशाद्धिक भाषेच्या माध्यमातून देखील करता येते-उदाहरणार्थ, सौहार्दपूर्ण स्मितहास्याने किंवा भौतिक पूर्ण हाताभंगाने. जरी तुम्ही नेहमीचे “थेट” मार्ग अनुसरलात तरीही या छोट्या बदलांनी तुमचा सहकारी संबंध दृढ होईल आणि एकंदर संघटनात्मक संस्कृती सुधारेल.
Brainstorming	मैं नया पदोन्नति के लिए 100 विचारों की सूची बनाना चाहता हूँ। इन विचारों में से 10 सबसे बेहतरीन विचारों को चुनना चाहता हूँ।	आपका आइडिया पदोन्नति के लिए 100 विचारों की सूची बनाने के लिए एक बेहतरीन विचार है। आपका विचार 100 विचारों की सूची बनाने के लिए एक बेहतरीन विचार है। 1. नया विचार: नया विचार 100 विचारों की सूची बनाने के लिए एक बेहतरीन विचार है। नया विचार 100 विचारों की सूची बनाने के लिए एक बेहतरीन विचार है। 2. नया विचार: नया विचार 100 विचारों की सूची बनाने के लिए एक बेहतरीन विचार है। नया विचार 100 विचारों की सूची बनाने के लिए एक बेहतरीन विचार है। 3. नया विचार: नया विचार 100 विचारों की सूची बनाने के लिए एक बेहतरीन विचार है। नया विचार 100 विचारों की सूची बनाने के लिए एक बेहतरीन विचार है। 4. नया विचार: नया विचार 100 विचारों की सूची बनाने के लिए एक बेहतरीन विचार है। नया विचार 100 विचारों की सूची बनाने के लिए एक बेहतरीन विचार है। 5. नया विचार: नया विचार 100 विचारों की सूची बनाने के लिए एक बेहतरीन विचार है। नया विचार 100 विचारों की सूची बनाने के लिए एक बेहतरीन विचार है। यदि यह बेहतरीन विचार है, तो आपका विचार 100 विचारों की सूची बनाने के लिए एक बेहतरीन विचार है।	नया विचार 100 विचारों की सूची बनाने के लिए एक बेहतरीन विचार है। नया विचार 100 विचारों की सूची बनाने के लिए एक बेहतरीन विचार है। 1. नया विचार: नया विचार 100 विचारों की सूची बनाने के लिए एक बेहतरीन विचार है। नया विचार 100 विचारों की सूची बनाने के लिए एक बेहतरीन विचार है। 2. नया विचार: नया विचार 100 विचारों की सूची बनाने के लिए एक बेहतरीन विचार है। नया विचार 100 विचारों की सूची बनाने के लिए एक बेहतरीन विचार है। 3. नया विचार: नया विचार 100 विचारों की सूची बनाने के लिए एक बेहतरीन विचार है। नया विचार 100 विचारों की सूची बनाने के लिए एक बेहतरीन विचार है। 4. नया विचार: नया विचार 100 विचारों की सूची बनाने के लिए एक बेहतरीन विचार है। नया विचार 100 विचारों की सूची बनाने के लिए एक बेहतरीन विचार है। 5. नया विचार: नया विचार 100 विचारों की सूची बनाने के लिए एक बेहतरीन विचार है। नया विचार 100 विचारों की सूची बनाने के लिए एक बेहतरीन विचार है। यदि यह बेहतरीन विचार है, तो आपका विचार 100 विचारों की सूची बनाने के लिए एक बेहतरीन विचार है।

Figure 21: *Pragyaan-Align* alignment data examples in Advice Seeking and Brainstorming categories with *Easy*, *1-Turn*, *Simple-IF*, *Safe*, *IC-3*, *Normal (No Thinking Trails)* settings configuration.

Category	Prompt	Preferred Response	Rejected Response
Creative Writing	మ మొదటి స్టోరీ ప్లాంట్ ఈ క్రింది విధంగా ఉంది: కొంతమంది స్నేహితుల బృందం చీకటి చరిత్ర కలిగిన సరస్సుకి క్యాంపింగ్ ట్రిప్ కు వెళ్ళారు. సరస్సు యొక్క అందమైన దృశ్యం కనపించే ప్రదేశంలో వారు తమ టెంట్ ను ఏర్పాటు చేసుకున్నారు, కానీ వారి ప్రశాంతత స్వల్పకాలికం. వారు అడవిలో ఒంటరిగా లేదని త్వరలోనే గ్రహించారు. స్నేహితులు ఎదుర్కొనే రహస్యమైన ఉనికి వరంగా తగిన కథాంశ అభివృద్ధి ఏమిటని మీరు సూచిస్తారు?	ఒక ఆసక్తికరమైన మరియు ఉత్తరభరితమైన కథ కోసం, ఇక్కడ ఒక సంభావ్య కథాంశ అభివృద్ధిని నేను సూచించుచున్నాను: టైటిల్: "ఎకేన్ ఆఫ్ ద పాస్ట్" కథాంశం: స్నేహితుల సమూహం మట్టుపట్టులను అనేషియన్, సరస్సు యొక్క అందాన్ని ఆస్వాదిస్తూ ఉండగా, వారిలో కొందరు నీటిలో ఆనాధారణమైన ప్రతిధ్వనులను గమనించడం మొదలుపెడతారు. వారు మొదట ఇది వాతావరణం కారణంగా అని భావిస్తారు, కానీ క్షణాలు ఎక్కువ బలంగా మారుతాయి, కేవలం క్రమంలో ఆత్యయముకొనక పాస్టరీకలు లాగా ధ్వనిస్తుంది. వారు సమీపంలోని గ్రామస్థుడైన జ్ఞాన సంపన్నుడు అయిన ఓల్డ్ మాన్ రాజును కలుస్తారు, అతను స్థానిక పురాణాలను వారికి చెప్తాడు: అనేక సంవత్సరాలగా సరస్సులో మొనిగి మరణించిన వ్యక్తి యొక్క ఆత్మ గురించి. ఆత్యయముకొనక తెలుసుకోవడానికి నిర్ణయించబడిన స్నేహితులు లోతుగా తీవ్రమైన ప్రారంభిస్తారు. వారు సరస్సుకు సంబంధించిన చీకటి చరిత్ర గురించి తెలుసుకుంటారు, ఇది వరపు రేజులను తనతో పాటు తీసుకువెళ్ళిందని పేరుగాంచింది. వారు సరస్సు నుండి బయటకు వచ్చినట్లు కనపించే వంతెన సంతకాలను కూడా గమనించడం ప్రారంభిస్తారు, మరియు వాటిలో కొన్ని తాము మాత్రమే వినగలుగుతున్నాయని నమ్ముతారు. అని మరింత లోతుగా తీవ్రమౌతుండగా, వారు సరస్సు యొక్క చరిత్ర మరియు దాని మట్టు ఉన్న సంఘటనలకు మధ్య ఆత్యయకరమైన సంబంధాలను కనుగొంటారు. వారు సరస్సు యొక్క ప్రవాహాలలో దాచబడి ఉన్న ఒక రహస్య దృష్టిని కనుగొంటారు, ఇది మట్టుపట్టుల ప్రాంతానికి ఉన్నది వీలుకోని దృష్టిని అందిస్తుంది. దీనివల్ల, వారు ఒక పురాతన ఆలయం యొక్క అవశేషాలను కనుగొంటారు, ఇది సరస్సు యొక్క చరిత్ర మరియు దాని మట్టు జరిగిన మరణాల గురించి అంతర్దృష్టులను బహిర్గతం చేస్తుంది. దీనిని యొక్క అనేకరూప వేస్తున్నప్పుడు, వారు ఆలయంలో దాక్కుని ఉన్న ఒక రహస్య గది పైకి వస్తారు. గదిలో, వారు ఒక డైరీని కనుగొంటారు, ఇది సరస్సు యొక్క చరిత్ర గురించి వివరిస్తుంది మరియు దాని ఆలయ కేంద్ర దాచబడి ఉన్న రహస్యాలను వివరిస్తుంది. డైరీలో సరస్సు యొక్క ఆత్మ యొక్క ఉనికిని మరియు దాని కోపానికి కారణాన్ని వివరిస్తుంది. స్నేహితులు డైరీని అధ్యయనం చేస్తున్నప్పుడు, వారు సరస్సు యొక్క ఆత్మ తమతో నేరుగా సంభాషించడానికి ప్రయత్నిస్తారని అని గ్రహిస్తారు. ఆత్మ తన కోపానికి కారణం - సరస్సు యొక్క అందం మరియు మహిమను నాశనం చేయడానికి మానవ కార్యకలాపాలు. స్నేహితులు సరస్సు యొక్క ఆత్మ యొక్క అభ్యర్థనను అంగీకరించాలని నిర్ణయించుకుంటారు, మరియు దాని రక్షణ కోసం వారు తమ వంతెన కృషి చేస్తారు. వారు తమ ఇంటికి తిరిగి వచ్చినప్పుడు, వారు సరస్సు యొక్క ఆత్మ యొక్క ప్రతిధ్వనులను ఇంకా వింటారు, కానీ ఇప్పుడు అది నింతోషకరమైనది మరియు కృతజ్ఞతతో ఉంటాయి. సరస్సు యొక్క ఆత్మ వారిని కాపాడినందుకు వారికి కృతజ్ఞతలు తెలుపుతుంది, మరియు వారు దాని అందం మరియు మహిమను సంగ్రహించడానికి వారి ప్రతిజ్ఞను పూర్తి చేశారు. స్నేహితులు తమ క్యాంపింగ్ ట్రిప్ నుండి దూరంగా వెళ్ళేటప్పుడు, వారు సరస్సు యొక్క ఆత్మ యొక్క పాతాలను మరియు సహజ ప్రపంచంతో సామరస్యంగా జీవించడం యొక్క ప్రాముఖ్యతను ఎప్పటికీ గుర్తుంచుకోవాలని ప్రతిజ్ఞ చేస్తారు.	చీకటి చరిత్ర కలిగిన సరస్సులో క్యాంపింగ్ చేస్తున్న స్నేహితుల బృందం గురించిన కథకు కొన్ని సూచనలు ఇక్కడ ఉన్నాయి: "1. దున్ను ఆత్మలు లేదా దెయ్యాలు:" సరస్సు చీకటి చరిత్ర కలిగి ఉందని చెప్పబడింది కాబట్టి, ఆ సరస్సులో మునిగిపోయినవారి ఆత్మలు లేదా దెయ్యాలు తిరుగుతున్నాయని చూపించవచ్చు. స్నేహితులు వారి ఉనికికి సంకేతాలను గమనించవచ్చు, ఉదాహరణకు, రాత్రి అరుపులు, అసహజమైన నీడలు లేదా ఎక్కువల స్థానచలనం. ఈ ఆత్మలు స్నేహితులను వేధించడం ప్రారంభించవచ్చు, వారిని భయపెట్టడానికి లేదా వారిని దాడి చేయడానికి ప్రయత్నించవచ్చు. "2. దాగి ఉన్న ప్రాణాలు:" సరస్సులో లేదా దాని మట్టుపట్టుల దాగి ఉన్న ప్రమాదకరమైన బంతువులు లేదా పొరాలతో ప్రాణాలు ఉండవచ్చు. ఉదాహరణకు, ఒక పెద్ద జలవర ప్రాణి, ఒక పాము లేదా ఒక ప్రమాదకరమైన వన్యప్రాణి. ఈ ప్రాణాలు స్నేహితులను వేటాడవచ్చు, వారిని దాడి చేయవచ్చు లేదా వారిని ఒక ప్రమాదంలో ఇరికించవచ్చు. "3. పురాతన శాపం:" సరస్సుకు పురాతన శాపం ఉండవచ్చు, ఇది ఆ ప్రదేశానికి వచ్చే ప్రతి ఒక్కరిని ప్రభావితం చేస్తుంది. ఈ శాపం కోరికలకు తర్వాత స్నేహితులను దాడి చేయవచ్చు లేదా వారి మనస్సులను ప్రభావం చూపవచ్చు, దాన్ని కనుగొనడానికి లేదా అధిగమించడానికి వారిని బలవంతం చేస్తుంది. "4. మానవ దురాచారాలు:" సరస్సు చరిత్రలో జరిగిన మానవ దురాచారాల ఫలితంగా ఈ దృశ్యం ఏర్పడి ఉండవచ్చు. ఉదాహరణకు, పాత, డిబ్బి లేదా క్షణ. ఈ దురాచారాలకు కారణమైనవారు లేదా వారి ఆత్మలు ఇప్పటికీ ఆ ప్రదేశంలో తిరుగుతున్నాయి. స్నేహితులు సరస్సు చరిత్ర గురించి తెలుసుకోని వారి అనుభవాలతో వ్యవహరించవలసి వస్తుంది. "కథాంశ అభివృద్ధికి ఉదాహరణ:" స్నేహితులు ఒక రాత్రి సరస్సు సమీపంలో ఉన్న అడవిలో ఒక ఆనాధారణ కట్టించినవచ్చు, ప్రారంభంలో వారు ఏమి పట్టించుకోరు, కానీ క్షణాలు మరింత తీవ్రతరమవుతాయి. వారు అడవిలోకి వెళ్తే, వారు ఒక పెద్ద, చీకటి వృక్షం చూస్తారు, అది వారిని దాడి చేస్తుంది. స్నేహితులు పారిపోతారు, కానీ వారు ప్రమాదకర ప్రాణి నుండి తప్పించుకోవడానికి ప్రయత్నిస్తున్నప్పుడు, వారు సరస్సు చరిత్రను కనుగొంటారు, మరియు వారు దానిలో చిక్కుకున్నారు. వారు సరస్సు చరిత్ర నుండి తప్పించుకోవడానికి మరియు ప్రాణాలను కాపాడుకోవడానికి ప్రయత్నించే కథ ఇది. ఈ సూచనలు కథాంశాన్ని అభివృద్ధి చేయడంలో మిమ్మల్ని సహాయపడతాయి అని నేను ఆశిస్తున్నాను.

Figure 22: Pragmaan-Align alignment data example in Creative Writing category with Easy, 1-Turn, Simple-IF, Safe, IC-3, Normal (No Thinking Trails) settings configuration.

Category	Prompt	Preferred Response	Rejected Response												
Data Analysis	મૂંઠકન ગણેશક-એકાદિ નવન સાર ગાણેર દુઃખિત કવતા પ્રજાન જાલ, જા નિયે ગાવથના કરાજન। તારા ૧૦દિ ગાણેર સાર પ્રયોગ કરેન એવં ૬ સપ્તાહ પર દુઃખિત રાર પરિમાર્પ કરેન। તાનેર કાણે ૧૦દિ ગાણેર-એકાદિ કવ્વોલ ક્રમ૩ આજ, યેગુલિતે સાર લેગ્યા હયનિ। મૂતિ ક્રમ્પર દુઃખિત રાર નિયક્રમ: સાર પ્રયોગક્રૂત ક્રમ: ૧૦, ૧૨, ૧૬, ૧૮, ૨૦, ૨૧, ૨૬, ૨૦, ૩૦, ૩૦, ૩૦, ૩૦ કવ્વોલ ક્રમ: ૬, ૭, ૯, ૮, ૭, ૧૦, ૧૨, ૧૬, ૧૮, ૨૦ પ્રતિદિ ક્રમ્પર મિન, મિડિયાન એવં યોગેર દુઃખિત રાર કવ?	નવન આપરા પ્રથમ પ્રતિદિ ક્રમ્પર જના મિન, મિડિયાન એવં યોગેર વેર કાર। **સાર પ્રયોગક્રૂત ક્રમ:** **મિન:** સર્વનિમ્ન માન હલ ૧૦। **મિડિયાન:** માનગુલાક ક્રમવર્ધમાન ક્રમ સાક્ષિય (૧૦, ૧૨, ૧૬, ૧૮, ૨૦, ૨૧, ૨૬, ૨૦, ૩૦) આપરા મધમ મૂતિ માન (૨૦ એવં ૨૧) પારિ। મિડિયાન હલ એઈ મૂતિ માનેર ગર: (૨૦ + ૨૧) / ૨ = ૨૦.૬ **યોગ:** કોન માનરે એકાદિકરાર આપનિ, જાઈ એઈ લેટા લેટે કોન યોગ લેઈ। **કવ્વોલ ક્રમ:** **મિન:** સર્વનિમ્ન માન હલ ૬। **મિડિયાન:** માનગુલાક ક્રમવર્ધમાન ક્રમ સાક્ષિય (૬, ૭, ૯, ૮, ૭, ૧૦, ૧૨, ૧૬, ૧૮, ૨૦) આપરા મધમ મૂતિ માન (૭ એવં ૧૦) પારિ। મિડિયાન હલ એઈ મૂતિ માનેર ગર: (૭ + ૧૦) / ૨ = ૭.૬ **યોગ:** કોન માનરે એકાદિકરાર આપનિ, જાઈ એઈ લેટા લેટે કોન યોગ લેઈ। સૂતરાં, સરવ્વક્રમ: <table> <tr> <td>ક્રમ</td><td>મિન</td><td>મિડિયાન</td><td>યોગ</td></tr> <tr> <td>સાર પ્રયોગક્રૂત ક્રમ</td><td>૧૦</td><td>૨૦.૬</td><td>૨૦.૬</td></tr> <tr> <td>કવ્વોલ ક્રમ</td><td>૬</td><td>૭.૬</td><td>૭.૬</td></tr> </table>	ક્રમ	મિન	મિડિયાન	યોગ	સાર પ્રયોગક્રૂત ક્રમ	૧૦	૨૦.૬	૨૦.૬	કવ્વોલ ક્રમ	૬	૭.૬	૭.૬	મિન (મિન): કવ્વોલ ક્રમ્પર જના, સરવેર (જોઈ માન હલ ૬)। સારમૂત ક્રમ્પર જના, સરવેર (જોઈ માન હલ ૧૦)। મેડિયાન (મિન): ઉચ્ચ નિર્દેશ માધ્યમિક અવશ્યાન થાકા સંસ્થા ચૂંકાત હલ। કવ્વોલ ક્રમ્પર (જોઈ, એઈ હલ (8 + 9)/2 = 8.5, કિન્તુ આપાનેર કાણે પૂર્વસંસ્થાર લેટા રોગેજ જાઈ આપરા એકિક કાંઠાકાંઈ પૂર્વસંસ્થાર રાઉંડ કરાત પારિ ના કારણ આપાનેર કાણે એમન કોન૩ મધા માન લેઈ યા મૂંઠકાર પ્રદર્શિત હલ। જર, એઈ સમસ્યાઈ સમાધાનેર આરેકાઈ ઉપાય હલ યે યથન સંસ્થાગુલિ (જોઈ થાકે તથન માધ્યમિક માનગુલિ સરાસરિ ગર્પના કરા યથ। યદિ૩ આપાનેર કાણે 10 ટિ સંસ્થા રોગેજ, યાર મધા 2 ટિ મધમ માન રોગેજ - 7 ૩ 8 ૧૬ ૩ 17 એર મળા નય (યેથાને વ્રકૂત મધમ માન હલ સેઈ મૂતિર ગર)। આપાનેર (જોઈ, આપાનેર કાણે 10 ટિ સંસ્થા રોગેજ, જાઈ આપરા જાનિ (યે મધમ માન હલ 5મ એવં ૬ઠ્ઠ સંસ્થાર ગર)। એકઈજાવે, સારમૂત ક્રમ્પર જના, આપરા 5મ એવં ૬ઠ્ઠ સંસ્થાર ગર મૂંઠક (વેર કરવા)। કવ્વોલ ક્રમ્પર જના, 5મ એવં ૬ઠ્ઠ સંસ્થા યથાક્રમ 8 એવં 9; સૂતરાં, મધાવર્તી માન હલ (8+9)/2 = 8.5। સારમૂત ક્રમ્પર જના, 5મ એવં ૬ઠ્ઠ સંસ્થા યથાક્રમ 18 એવં 20; સૂતરાં, મધાવર્તી માન હલ (18+20)/2 = 19। યોગ: યોગ હલ એમન એકાદિ માન યા લેટાલેટે યથન યથાઈ। કવ્વોલ ક્રમ્પર જના, પ્રતિદિ સંસ્થા સુધુપાઠ એકવાર લેથા લેથ, જાઈ (યોગ) જનના નય। એકઈજાવે, સારમૂત ક્રમ્પર જના૩, પ્રતિદિ સંસ્થા સુધુપાઠ એકવાર પ્રદર્શિત હૃગાર કારણ યોગ) જનના નય। ઉત્તર: - કવ્વોલ ક્રમ્પર મિન, મિડિયાન એવં યોગ યથાક્રમ 10.3, 8.5, એવં જનના। - સારમૂત ક્રમ્પર મિન, મિડિયાન એવં યોગ યથાક્રમ 20.1, 19, એવં જનના।
ક્રમ	મિન	મિડિયાન	યોગ												
સાર પ્રયોગક્રૂત ક્રમ	૧૦	૨૦.૬	૨૦.૬												
કવ્વોલ ક્રમ	૬	૭.૬	૭.૬												
Editing	લેખિત સામગ્રીને પરિશ્કૃત કરતી વખતે તમારા મુખ્ય સંપાદન ઉદ્દેશ્યો શું છે?	લેખિત સામગ્રીને પરિશ્કૃત કરતી વખતે, મારા મુખ્ય સંપાદન ઉદ્દેશ્યો આ પ્રમાણે છે: - સ્પષ્ટતા અને સુસંગતતા સુધારવી: હું લેખકના સંદેશને અસરકારક રીતે પહોંચાડવા માટે સામગ્રીને સ્પષ્ટ, સંક્ષિપ્ત અને અનુસરવામાં સરળ બનાવવાનું લક્ષ્ય રાખું છું. - ભાષા અને શૈલીને વધારવી: હું સામગ્રીની એકદર ગુણવત્તા વધારવા માટે ચોક્કસ, સંયોજક અને આકર્ષક ભાષાનો ઉપયોગ કરવાની ખાતરી કરું છું. - વ્યાકરણ અને મિર્મિકસને સુધારવું: હું વ્યાકરણની ભૂલો, વિસંગતતાઓ અને ફોર્મેટિંગની ભૂલોને સુધારે છે જેથી સામગ્રી લુલ્લ-ખુલ્લ અને વ્યાવસાયિક હોય. - માળખું અને સંગઠનને સુધારવું: હું સામગ્રીને સ્પષ્ટ માળમા સાથે સુવ્યવસ્થિત અને તાર્કિક બનાવવા માટે તેને ફરીથી ગોઠવું છું, જેનાથી વાચકોને તેને સમજવું અને અનુસરવું સરળ બને છે. - સ્વર અને અવાજને વધારવું: હું લેખકના ઘચિત પ્રેક્ષકો અને હેતુને અનુરૂપ સામગ્રીનો સ્વર અને અવાજ સુસંગત અને અસરકારક છે તેની ખાતરી કરું છું. - ચોક્કસ અને તથ્યસાદર ચોક્કસાઈને યકાસવી: હું સામગ્રીની ચોક્કસાઈ અને વિશ્વસનીયતાની ખાતરી કરવા માટે હકીકતો, આંકડાઓ અને માહિતીની ચોક્કસાઈને યકાસું છું. - સંવેદનશીલતા અને સાંસ્કૃતિક યોગ્યતાની સમીક્ષા કરવી: હું એ સુનિશ્ચિત કરું છું કે સામગ્રી સાંસ્કૃતિક રીતે સંવેદનશીલ, આદરપૂર્ણ અને વિવિધ પ્રેક્ષકો માટે યોગ્ય હોય. - સંપાદનની ધોરણો અને માર્ગદર્શિકાઓનું પાલન કરવું: હું સંપાદનની ધોરણો અને માર્ગદર્શિકાઓનું પાલન કરું છું જે લાગુ હોય, જેમ કે શિક્ષણ મેન્યુઅલ ઓફ સ્ટાઇલ અથવા અપી સ્ટાઇલગુક, જેથી સામગ્રી સુસંગત અને ઉદ્યોગના ધોરણોને અનુરૂપ હોય. - લેખકની અનન્ય અવાજ અને દ્રષ્ટિને જાળવી રાખવી: જ્યારે હું સામગ્રીને સુધારું છું અને તેમાં સુધારો કરું છું, ત્યારે હું લેખકની અનન્ય અવાજ અને દ્રષ્ટિને જાળવી રાખવાનું અને તેને જાળવી રાખવાનું લક્ષ્ય રાખું છું, જેથી સામગ્રી પ્રમાણિક અને અધિકૃત લાગે. - પ્રતિસાદ અને સમર્થન આપવું: હું લેખકને તેમના કાર્યને સુધારવામાં અને સુધારવામાં મદદ કરવા માટે મારા સંપાદનની સ્પષ્ટતા અને સમર્થન કરું છું.	**સ્પષ્ટતા**: લેખકની ભૂલોને ઠીક કરવી અને વાંચનને સરળ બનાવવું, જેથી વિચાર સ્પષ્ટ રહે. **સંરચના**: લેખની રચનાને મજબૂત બનાવવી, જેવા કે પેરાગ્રાફ અને ઉપવિષયોનું યોગ્ય વિભાજન. **ભાષાશુદ્ધતા**: કોઈ પણ વ્યવાકરણ, વ્યાકરણ, અને ઉપભાષા સંબંધિત ભૂલોને સુધારવી. **અર્થવૈભવ**: અનાવશ્યક શબ્દો અને પંક્તીઓને દૂર કરી, પીછીને લાગણીસાં અને અસરકારક બનાવવું. **લક્ષ્ય સમજૂતી**: લેખના હેતુ અને લક્ષ્ય સમજૂતીઓને કાર્યક્રમમાં સમાવિષ્ટ કરવું, જેથી તે પ્રેક્ષક સાથે જોડાય. **અર્થ બંધનશીલતા અને ઉદ્દેશ્યસર**: પુસ્તકમાં કોઈપણ અસંવાદિતા કે બાકી રહેવાના મુદ્દાઓને દૂર કરવા. આ બાબતોને ધ્યાનમાં રાખીને, સંપાદક સમૂહ અને સંક્ષિપ્ત સામગ્રી તૈયાર કરે છે, જે વાંચક માટે સ્પષ્ટ અને ઉપયોગી હોય.												

Figure 23: *Pragyaan-Align* alignment data examples in Data Analysis and Editing categories with *Easy*, *1-Turn*, *Simple-IF*, *Safe*, *IC-3*, *Normal (No Thinking Trails)* settings configuration.

Category	Prompt	Preferred Response	Rejected Response
Information Seeking	ತಾರಾ ಅಥವಾ ಗ್ರಹ ಯಾವುದೇ ಬೆಳಕು ಹೊಮ್ಮುತ್ತದೆ?	ತಾರಾ ಅಥವಾ ಗ್ರಹಗಳು ಹೊಮ್ಮುವ ಬೆಳಕು ಹೊಮ್ಮುತ್ತದೆ. ಇದು ಅವುಗಳ ಸ್ವಾಭಾವಿಕ ಗುಣವಾಗಿದೆ. ತಾರಾ ಅಥವಾ ಗ್ರಹಗಳು ಹೊಮ್ಮುವ ಬೆಳಕು ಹೊಮ್ಮುತ್ತದೆ. ಇದು ಅವುಗಳ ಸ್ವಾಭಾವಿಕ ಗುಣವಾಗಿದೆ.	ತಾರಾ ಅಥವಾ ಗ್ರಹಗಳು ಹೊಮ್ಮುವ ಬೆಳಕು ಹೊಮ್ಮುತ್ತದೆ. ಇದು ಅವುಗಳ ಸ್ವಾಭಾವಿಕ ಗುಣವಾಗಿದೆ. ತಾರಾ ಅಥವಾ ಗ್ರಹಗಳು ಹೊಮ್ಮುವ ಬೆಳಕು ಹೊಮ್ಮುತ್ತದೆ. ಇದು ಅವುಗಳ ಸ್ವಾಭಾವಿಕ ಗುಣವಾಗಿದೆ.
Math	ಒಂದು ಸರಣಿಯಲ್ಲಿ 100 ಸಂಖ್ಯೆಗಳಿವೆ. ಅವುಗಳ ಮೊತ್ತವು 1000 ಆಗಿದೆ. ಈ ಸರಣಿಯಲ್ಲಿ 10 ಸಂಖ್ಯೆಗಳನ್ನು ತೆಗೆದುಕೊಂಡು, ಅವುಗಳ ಮೊತ್ತವು 100 ಆಗಿದೆ. ಉಳಿಕೆ 90 ಸಂಖ್ಯೆಗಳ ಮೊತ್ತವು ಎಷ್ಟು?	ಒಂದು ಸರಣಿಯಲ್ಲಿ 100 ಸಂಖ್ಯೆಗಳಿವೆ. ಅವುಗಳ ಮೊತ್ತವು 1000 ಆಗಿದೆ. ಈ ಸರಣಿಯಲ್ಲಿ 10 ಸಂಖ್ಯೆಗಳನ್ನು ತೆಗೆದುಕೊಂಡು, ಅವುಗಳ ಮೊತ್ತವು 100 ಆಗಿದೆ. ಉಳಿಕೆ 90 ಸಂಖ್ಯೆಗಳ ಮೊತ್ತವು 900 ಆಗಿದೆ.	ಒಂದು ಸರಣಿಯಲ್ಲಿ 100 ಸಂಖ್ಯೆಗಳಿವೆ. ಅವುಗಳ ಮೊತ್ತವು 1000 ಆಗಿದೆ. ಈ ಸರಣಿಯಲ್ಲಿ 10 ಸಂಖ್ಯೆಗಳನ್ನು ತೆಗೆದುಕೊಂಡು, ಅವುಗಳ ಮೊತ್ತವು 100 ಆಗಿದೆ. ಉಳಿಕೆ 90 ಸಂಖ್ಯೆಗಳ ಮೊತ್ತವು 900 ಆಗಿದೆ.
Reasoning	ಒಂದು ಸರಣಿಯಲ್ಲಿ 100 ಸಂಖ್ಯೆಗಳಿವೆ. ಅವುಗಳ ಮೊತ್ತವು 1000 ಆಗಿದೆ. ಈ ಸರಣಿಯಲ್ಲಿ 10 ಸಂಖ್ಯೆಗಳನ್ನು ತೆಗೆದುಕೊಂಡು, ಅವುಗಳ ಮೊತ್ತವು 100 ಆಗಿದೆ. ಉಳಿಕೆ 90 ಸಂಖ್ಯೆಗಳ ಮೊತ್ತವು ಎಷ್ಟು?	ಒಂದು ಸರಣಿಯಲ್ಲಿ 100 ಸಂಖ್ಯೆಗಳಿವೆ. ಅವುಗಳ ಮೊತ್ತವು 1000 ಆಗಿದೆ. ಈ ಸರಣಿಯಲ್ಲಿ 10 ಸಂಖ್ಯೆಗಳನ್ನು ತೆಗೆದುಕೊಂಡು, ಅವುಗಳ ಮೊತ್ತವು 100 ಆಗಿದೆ. ಉಳಿಕೆ 90 ಸಂಖ್ಯೆಗಳ ಮೊತ್ತವು 900 ಆಗಿದೆ.	ಒಂದು ಸರಣಿಯಲ್ಲಿ 100 ಸಂಖ್ಯೆಗಳಿವೆ. ಅವುಗಳ ಮೊತ್ತವು 1000 ಆಗಿದೆ. ಈ ಸರಣಿಯಲ್ಲಿ 10 ಸಂಖ್ಯೆಗಳನ್ನು ತೆಗೆದುಕೊಂಡು, ಅವುಗಳ ಮೊತ್ತವು 100 ಆಗಿದೆ. ಉಳಿಕೆ 90 ಸಂಖ್ಯೆಗಳ ಮೊತ್ತವು 900 ಆಗಿದೆ.

[illegible]

[illegible]

Figure 27: *Pragyaan-IT* data examples in categories Planning - (*Hard, 1-Turn, Simple-IF, Safe, IC-3, Normal (No Thinking Trails)*) and General Question Answering - (*Easy, 1-Turn, Simple-IF, Safe, IC-1, Normal (No Thinking Trails)*) settings configuration.

[illegible]

Category	Prompt	Response
Information Seeking	ಸ್ವಯಂ ಇಚ್ಛಾಶಕ್ತಿಯ ಕುರಿತು ಭಾರತದ ಐತಿಹಾಸಿಕ ಮತ್ತು ತಾಂತ್ರಿಕ ವಾದಗಳನ್ನು ವಿಶ್ಲೇಷಿಸಿ. ಸೈತಿಕ ಹೋಣೆಗಾರಿಕೆ, ಭಾರತೀಯ ಕಾನೂನು ವ್ಯವಸ್ಥೆಗಳು, ಮತ್ತು ವೈಯಕ್ತಿಕ ಗುರುತುಗಳ ಮೇಲೆ ಇದರ ಪರಿಣಾಮಗಳೇನು? ಸಾಮರಸ್ಯವಾದಿ ಮತ್ತು ಅಸಾಮರಸ್ಯವಾದಿ ದೃಷ್ಟಿಕೋನಗಳನ್ನು ಚರ್ಚಿಸಿ. ನಿರ್ಣಾಯಕತೆ ಮತ್ತು ಏಜೆನ್ಸಿಗೆ ಸಂಬಂಧಿಸಿದ ನರವಿಜ್ಞಾನ ಹಾಗೂ ತಾಂತ್ರಿಕ ವಾದಗಳ ಬಲ ದುರ್ಬಲತೆಗಳನ್ನು ಮೌಲ್ಯಮಾಪನ ಮಾಡಿ.	ಸ್ವಯಂ ಇಚ್ಛಾಶಕ್ತಿ ಅಂದರೆ ನಮ್ಮ ಆಯ್ಕೆಗಳು ಮತ್ತು ನಿರ್ಧಾರಗಳು ನಿಜಕ್ಕೂ ನಮ್ಮ ಸ್ವಂತದ್ದೇ ಅಥವಾ ಅವು ಪೂರ್ವನಿರ್ಧಾರಿತ ಅಂಶಗಳಿಂದ ನಿಯಂತ್ರಿಸಲ್ಪಟ್ಟವೆಯೇ ಎಂಬ ಪ್ರಶ್ನೆ ಪ್ರಶ್ನೆ ಮನುಷ್ಯರ ನಡುವೆ ನಡೆದ ಚರ್ಚೆಯಲ್ಲಿದೆ. ಇದು ಸೈತಿಕ ಹೋಣೆಗಾರಿಕೆ, ಕಾನೂನು ವ್ಯವಸ್ಥೆಗಳು ಮತ್ತು ವೈಯಕ್ತಿಕ ಗುರುತುಗಳಂತಹ ಮೂಲಭೂತ ಪರಿಕಲ್ಪನೆಗಳ ಮೇಲೆ ಆಳವಾದ ಪರಿಣಾಮಗಳನ್ನು ಬೀರುತ್ತದೆ. ಭಾರತದ ಐತಿಹಾಸಿಕ ಮತ್ತು ತಾಂತ್ರಿಕ ವಾದಗಳು: ಭಾರತೀಯ ತತ್ವಶಾಸ್ತ್ರದಲ್ಲಿ ಸ್ವಯಂ ಇಚ್ಛಾಶಕ್ತಿಯ ಕುರಿತು ಚರ್ಚೆಗಳು ಬಹಳ ಪ್ರಾಚೀನ ಮತ್ತು ವೈವಿಧ್ಯಮಯವಾಗಿವೆ. ಕರ್ಮ ಸಿದ್ಧಾಂತ: ಬಹುಶಃ ಭಾರತೀಯ ಆಧ್ಯಾತ್ಮಿಕ ಮತ್ತು ಧಾರ್ಮಿಕ ಸಂಪ್ರದಾಯಗಳಲ್ಲಿ, ವಿಶೇಷವಾಗಿ ಹಿಂದೂ ಧರ್ಮ, ಬೌದ್ಧಧರ್ಮ ಮತ್ತು ಜೈನಧರ್ಮದಲ್ಲಿ, ಕರ್ಮ ಸಿದ್ಧಾಂತವು ಪ್ರಮುಖ ಸ್ಥಾನ ಪಡೆದಿದೆ. ಇದು ನಾವು ಮಾಡುವ ಪ್ರತಿಯೊಂದು ಕ್ರಿಯೆಯೂ (ಕರ್ಮ) ಭವಿಷ್ಯದ ಪರಿಣಾಮಗಳನ್ನು ಹೊಂದಿರುತ್ತದೆ ಎಂದು ಹೇಳುತ್ತದೆ. ಇಲ್ಲಿ ಸ್ವಯಂ ಇಚ್ಛಾಶಕ್ತಿಗೆ ಮಹತ್ವದ ಪಾತ್ರವಿದೆ, ಏಕೆಂದರೆ ವ್ಯಕ್ತಿಯು ತನ್ನ ಕರ್ಮಗಳನ್ನು ಆಯ್ಕೆ ಮಾಡುವ ಸಾಮರ್ಥ್ಯವನ್ನು ಹೊಂದಿದ್ದಾನೆ ಮತ್ತು ಆ ಆಯ್ಕೆಗಳಿಗೆ ಹೊಣೆಗಾರನಾಗಿರುತ್ತಾನೆ. ಭಗವದ್ಗೀತೆಯು ಈ ಕರ್ಮ ಸಿದ್ಧಾಂತವನ್ನು ಸ್ಪಷ್ಟವಾಗಿ ವಿವರಿಸುತ್ತದೆ. ಅಲ್ಲಿ ಕ್ರಿಶ್ಚಿಯಾನಿಟಿಯ ಅರ್ಥಾನುಗುಣ ಕರ್ಮ ಮಾಡಿ, ಫಲದ ನಿರೀಕ್ಷೆ ಇಲ್ಲದೆ ಎಂದು ಬೋಧಿಸುತ್ತದೆ. ಇದು ಆಯ್ಕೆಯ ಸ್ವಾತಂತ್ರ್ಯ ಮತ್ತು ಅದರ ಫಲವಾಗಿಯನ್ನು ಒತ್ತಿಹೇಳುತ್ತದೆ. ವ್ಯಕ್ತಿಯ ಆಯ್ಕೆಗಳು ಅವನ ಭವಿಷ್ಯವನ್ನು ರೂಪಿಸುತ್ತವೆ ಎಂಬುದನ್ನು ಇದು ಸೂಚಿಸುತ್ತದೆ. ಆದರೆ ಅದೇ ಸಮಯದಲ್ಲಿ ಪೂರ್ವಜನ್ಯ ಕರ್ಮಗಳ ಪ್ರಭಾವವನ್ನು ಒಪ್ಪಿಕೊಳ್ಳುತ್ತದೆ. ಯೋಗಸೂತ್ರಗಳು: ಪತಂಜಲಿಯ ಯೋಗಸೂತ್ರಗಳು ಮನಸ್ಸಿನ ನಿಯಂತ್ರಣ ಮತ್ತು ಅಳವಡ ಸ್ವಯಂ ಅರಿವಿನ ಮೂಲಕ ವ್ಯಕ್ತಿಯು ತನ್ನ ಪ್ರಜ್ಞೆಯನ್ನು ಉನ್ನತೀಕರಿಸಬಹುದು ಎಂದು ಹೇಳುತ್ತವೆ. ಇಲ್ಲಿಯೂ, ವ್ಯಕ್ತಿಯು ತನ್ನ ಆಲೋಚನೆಗಳು ಮತ್ತು ಕಾರ್ಯಗಳ ಮೇಲೆ ಹಿಡಿತ ಸಾಧಿಸುವ ಸಾಮರ್ಥ್ಯವನ್ನು ಹೊಂದಿದ್ದಾನೆ ಎಂಬ ಅಂಶವು ಸ್ವಯಂ ಇಚ್ಛಾಶಕ್ತಿಗೆ ಅವಕಾಶ ನೀಡುತ್ತದೆ. ಸೈತಿಕ ಹೋಣೆಗಾರಿಕೆ, ಕಾನೂನು ವ್ಯವಸ್ಥೆಗಳು ಮತ್ತು ವೈಯಕ್ತಿಕ ಗುರುತುಗಳ ಮೇಲೆ ಪರಿಣಾಮಗಳು: ಸ್ವಯಂ ಇಚ್ಛಾಶಕ್ತಿಯ ಅಸ್ತಿತ್ವದಲ್ಲಿದೆ ಎಂಬ ನಂಬಿಕೆಯು ನಮ್ಮ ಸೈತಿಕ ಮತ್ತು ಕಾನೂನು ವ್ಯವಸ್ಥೆಗಳ ಆಡಿಪಾಯವಾಗಿದೆ. ಸೈತಿಕ ಹೋಣೆಗಾರಿಕೆ: ನೀವು ಒಂದು ಕೆಲಸ ಮಾಡಲು ನಿಮ್ಮದೇ ಇಚ್ಛೆಯಿಂದ ನಿರ್ಧರಿಸಿದರೆ, ಅದರ ಒಳ್ಳೆಯ ಅಥವಾ ಕೆಟ್ಟ ಪರಿಣಾಮಗಳಿಗೆ ನೀವೇ ಜವಾಬ್ದಾರಿಯು. ಉದಾಹರಣೆಗೆ: ನೀವು ಸ್ವತಃ ನಿರ್ಧರಿಸಿ ಒಬ್ಬರಿಗೆ ಸಹಾಯ ಮಾಡಿದರೆ, ಆ ಒಳ್ಳೆಯ ಕಾರ್ಯಕ್ಕೆ ನೀವು ಕಡುಪ್ರಾಣಿರಿ ಪಡೆಯುತ್ತೀರಿ. ಆದರೆ, ಯಾರಾದರೂ ನಿಮ್ಮನ್ನು ಒತ್ತಾಯಿಸಿ ಆ ಸಹಾಯ ಮಾಡಿದರೆ, ಅದರ ಜವಾಬ್ದಾರಿ ನಿಮಗೆ ಸೇರುವುದಿಲ್ಲ. ಅದೇ ರೀತಿ, ನೀವು ನಿಮ್ಮ ಸ್ವಂತ ಇಚ್ಛೆಯಿಂದ ತಪ್ಪು ಮಾಡಿದರೆ, ಅದರ ದೂಷಣೆ ಅಥವಾ ಶಿಕ್ಷೆಯನ್ನು ಎದುರಿಸಬೇಕಾಗುತ್ತದೆ. ಕಾನೂನು ವ್ಯವಸ್ಥೆಗಳು: ನಮ್ಮ ಭಾರತೀಯ ನ್ಯಾಯಾಂಗ ವ್ಯವಸ್ಥೆಯು, ಅವರಾದ ಪ್ರಕರಣಗಳಲ್ಲಿ, ಒಬ್ಬ ವ್ಯಕ್ತಿ ತನ್ನ ಕೆಲಸವನ್ನು ತಾನಾಗಿಯೇ, ಉದ್ದೇಶಪೂರ್ವಕವಾಗಿ ಮಾಡಿದ್ದಾನೆ ಎಂದು ಸಾಮಾನ್ಯವಾಗಿ ಪರಿಗಣಿಸುತ್ತದೆ. ಯಾರಾದರೂ ಮಾಡಿದ ಕೆಲಸಕ್ಕೆ ಅವರೇ ಕಾರಣ ಎಂದು ನಂಬಿದಾಗ ಮಾತ್ರ ಶಿಕ್ಷೆ ನೀಡುವುದು ಸರಿಯಾಗುತ್ತದೆ. ಮನ್ವರಿಯಾ (ಅಪರಾಧ ಮಾಡುವ ಉದ್ದೇಶ) ಎಂಬ ನಿಯಮ, ಒಬ್ಬರು ಮಾಡುವ ಅಪರಾಧದ ಹಿಂದೆ ಅವರ ಅಂತರ ಅಥವಾ ಇಚ್ಛೆ ಇದೆಯೇ ಎಂದು ಪರಿಗಣಿಸುತ್ತದೆ. ವೈಯಕ್ತಿಕ ಗುರುತು: ನಮ್ಮನ್ನು ನಾವು ಯಾರು ಎಂದು ಗುರುತಿಸಿಕೊಳ್ಳುವಲ್ಲಿ ಸ್ವಯಂ ಇಚ್ಛಾಶಕ್ತಿ ಮುಖ್ಯ. ನಮ್ಮ ಆಯ್ಕೆಗಳು, ನಿರ್ಧಾರಗಳು ಮತ್ತು ಅವುಗಳ ಪರಿಣಾಮಗಳ ಮೂಲಕವೇ ನಾವು ನಮ್ಮ ವ್ಯಕ್ತಿತ್ವವನ್ನು ರೂಪಿಸಿಕೊಳ್ಳುತ್ತೇವೆ. ಸಾಮರಸ್ಯವಾದಿ (Compatibilist) ಮತ್ತು ಅಸಾಮರಸ್ಯವಾದಿ (Incompatibilist) ದೃಷ್ಟಿಕೋನಗಳು: ಸ್ವಯಂ ಇಚ್ಛಾಶಕ್ತಿ ಮತ್ತು ನಿರ್ಣಾಯಕತೆಯ ನಡುವಿನ ಸಂಬಂಧದ ಬಗ್ಗೆ ಎರಡು ಪ್ರಮುಖ ದೃಷ್ಟಿಕೋನಗಳಿವೆ: ಸಾಮರಸ್ಯವಾದಿ (Compatibilism): ಈ ದೃಷ್ಟಿಕೋನವು ಸ್ವಯಂ ಇಚ್ಛಾಶಕ್ತಿ ಮತ್ತು ನಿರ್ಣಾಯಕತೆ (Determinism - ಎಲ್ಲ ಘಟನೆಗಳು ಪೂರ್ವನಿರ್ಧಾರಿತವಾಗಿವೆ) ಪರಸ್ಪರ ಹೊಂದಾಣಿಕೆಯಾಗುತ್ತವೆ ಎಂದು ವಾದಿಸುತ್ತದೆ. ನಾವು ನಮ್ಮ ಆಯ್ಕೆಗಳನ್ನು ಮಾಡುವಾಗ ಬಾಹ್ಯ ಒತ್ತಡದಿಂದ ಮುಕ್ತರಾಗಿದ್ದರೆ, ನಮ್ಮ ನಿರ್ಧಾರಗಳು ಹಿಂದಿನ ಕಾರಣಗಳಿಂದ ಪ್ರಭಾವಿತವಾಗಿದ್ದರೂ ಸಹ, ನಾವು ಮುಕ್ತರು ಎಂದು ಸಾಮರಸ್ಯವಾದಿಗಳು ವಾದಿಸುತ್ತಾರೆ. ಇಲ್ಲಿ, ಸ್ವಾತಂತ್ರ್ಯ ಎಂದರೆ ಒತ್ತಾಯವಿಲ್ಲದಿರುವುದು, ಅನಿವಾರ್ಯತೆಯಲ್ಲ. ಅಸಾಮರಸ್ಯವಾದಿ (Incompatibilism): ಈ ದೃಷ್ಟಿಕೋನವು ನಿರ್ಣಾಯಕತೆ ಮತ್ತು ಸ್ವಯಂ ಇಚ್ಛಾಶಕ್ತಿ ಪರಸ್ಪರ ಹೊಂದಾಣಿಕೆಯಾಗುವುದಿಲ್ಲ ಎಂದು ವಾದಿಸುತ್ತದೆ. ನಿರ್ಣಾಯಕತೆ ನಿಜವಾಗಿದ್ದರೆ, ನಮ್ಮ ಆಯ್ಕೆಗಳು ನಿಜವಾಗಿಯೂ ಸ್ವಯಂ ಸ್ಪಂದವಾಗಿಲ್ಲವು. ಏಕೆಂದರೆ ಅವು ಹಿಂದಿನ ಘಟನೆಗಳ ಸರಣಿಯ ಫಲಿತಾಂಶಗಳಾಗಿರುತ್ತವೆ. ಇಲ್ಲಿ ಎರಡು ಉಪ ವಿಭಾಗಗಳಿವೆ: ಮುಕ್ತ ಆಯ್ಕೆ ಸಿದ್ಧಾಂತ : ಇದು ಸ್ವಯಂ ಇಚ್ಛಾಶಕ್ತಿ ನಿಜವಾಗಿದೆ ಮತ್ತು ನಿರ್ಣಾಯಕತೆ ಸುಳ್ಳು ಎಂದು ವಾದಿಸುತ್ತದೆ ನಿರಾಕಾರವಾದ : ಇದು ನಿರ್ಣಾಯಕತೆ ನಿಜವಾಗಿದೆ ಮತ್ತು ಸ್ವಯಂ ಇಚ್ಛಾಶಕ್ತಿ ಒಂದು ಭ್ರಮೆ ಎಂದು ವಾದಿಸುತ್ತದೆ. ನಿರ್ಣಾಯಕತೆ ಮತ್ತು ಏಜೆನ್ಸಿಗೆ ಸಂಬಂಧಿಸಿದ ನರವಿಜ್ಞಾನ ಮತ್ತು ತಾಂತ್ರಿಕ ವಾದಗಳು: ಅಧುನಿಕ ನರವಿಜ್ಞಾನವು ಈ ಚರ್ಚೆಗೆ ಹೊಸ ಆಯಾಮವನ್ನು ಸೇರಿಸಿದೆ. ನರವಿಜ್ಞಾನದ ವಾದಗಳು: ಕೆಲವು ನರವಿಜ್ಞಾನದ ಆಧ್ಯಯನಗಳು ಜನರು ತಮ್ಮ ಕ್ರಿಯೆಗಳನ್ನು ನಿರ್ಧರಿಸುವ ಮೊದಲು ಮೆದುಳು ಈಗಾಗಲೇ ಪ್ರತಿಕ್ರಿಯೆಯನ್ನು ಪ್ರಾರಂಭಿಸಿರುತ್ತದೆ ಎಂದು ಸೂಚಿಸುತ್ತವೆ. ಇದು ಸ್ವಯಂ ಇಚ್ಛಾಶಕ್ತಿಯು ಒಂದು ಭ್ರಮೆ ಇರಬಹುದೇ ಎಂಬ ಪ್ರಶ್ನೆಯನ್ನು ಹುಟ್ಟುಹಾಕಿದೆ. ಆದಾಗ್ಯೂ, ಈ ಆಧ್ಯಯನಗಳ ವ್ಯಾಖ್ಯಾನದ ಬಗ್ಗೆ ವ್ಯಾಪಕ ಚರ್ಚೆಗಳಿವೆ. ಅವು ಸರಳ ನಿರ್ಧಾರಗಳಿಗೆ ಮಾತ್ರ ಅನ್ವಯಿಸುತ್ತವೆ ಮತ್ತು ಸಂಕೀರ್ಣ, ಉದ್ದೇಶಪೂರ್ವಕ ಯೋಜನೆಗಳಿಗೆ ಅಲ್ಲ ಎಂದು ಕೆಲವರು ವಾದಿಸುತ್ತಾರೆ. ತಾಂತ್ರಿಕ ವಾದಗಳ ಬಲ ಮತ್ತು ದುರ್ಬಲತೆಗಳು: ಬಲ: ನರವಿಜ್ಞಾನದ ಆವಿಷ್ಕಾರಗಳು ನಮ್ಮ ನಿರ್ಧಾರ ಪ್ರಕ್ರಿಯೆಗಳ ಜೈವಿಕ ಆಧಾರದ ಬಗ್ಗೆ ಅಳವಡದ ಒಳನೋಟಗಳನ್ನು ನೀಡುತ್ತವೆ. ಇದು ಸ್ವಯಂ ಇಚ್ಛಾಶಕ್ತಿಯು ನಮ್ಮ ಗ್ರಹಿಕೆಯನ್ನು ಪುನರ್ನಿರ್ಮಿಸಲು ಪ್ರೇರೇಪಿಸುತ್ತದೆ. ಕರ್ಮ ಸಿದ್ಧಾಂತವು ವೈಯಕ್ತಿಕ ಜವಾಬ್ದಾರಿಗೆ ಪೂಜ್ಯ ಸೈತಿಕ ಭೌತಿಕವನ್ನು ನೀಡುತ್ತದೆ. ದುರ್ಬಲತೆ: ನರವಿಜ್ಞಾನದ ವಾದಗಳು ಸಂಪೂರ್ಣವಾಗಿ ನಿರ್ಣಾಯಕವಾಗಿಲ್ಲ ಮತ್ತು ಅವು ಪ್ರಜ್ಞೆ ಮತ್ತು ಆಯ್ಕೆ ಮಾಡುವ ಸಾಮರ್ಥ್ಯದ ಸಂಕೀರ್ಣತೆಯನ್ನು ಸಂಪೂರ್ಣವಾಗಿ ಗ್ರಹಿಸುವುದಿಲ್ಲ. ಕೇವಲ ಮೆದುಳಿನ ಚಟುವಟಿಕೆಯನ್ನು ಅಳಿಯುವ ಮೂಲಕ ಮಾನವ ಇಚ್ಛಾಶಕ್ತಿಯನ್ನು ಸಂಪೂರ್ಣವಾಗಿ ತಿಳಿದುಕೊಳ್ಳಲು ಸಾಧ್ಯವಿಲ್ಲ ಎಂಬ ವಾದವಿದೆ. ಭಾರತೀಯ ತತ್ವಶಾಸ್ತ್ರದಲ್ಲಿ, ಕರ್ಮ ಸಿದ್ಧಾಂತವು ಭವಿಷ್ಯವನ್ನು ನಿರ್ಧರಿಸುವ ಪೂರ್ವಜನ್ಯದ ಕರ್ಮಗಳ ಬಗ್ಗೆಯೂ ಮಾತನಾಡುತ್ತದೆ, ಇದು ಸಂಪೂರ್ಣ ಸ್ವಾತಂತ್ರ್ಯದ ಪರಿಕಲ್ಪನೆಯನ್ನು ಸಂಕೀರ್ಣಗೊಳಿಸುತ್ತದೆ. ಒಟ್ಟಾರೆ, ಸ್ವಯಂ ಇಚ್ಛಾಶಕ್ತಿಯ ಪ್ರಶ್ನೆಯು ಅಳವಡ ತತ್ವಶಾಸ್ತ್ರೀಯ ಮತ್ತು ವೈಜ್ಞಾನಿಕ ಪರಿಶೋಧನೆಯನ್ನು ಮುಂದುವರಿಸಿದೆ ಮತ್ತು ಭಾರತೀಯ ಚಿಂತನೆಯು ಈ ಜಾಗತಿಕ ಚರ್ಚೆಗೆ ಅನನ್ಯ ಮತ್ತು ಸಮೃದ್ಧ ದೃಷ್ಟಿಕೋನಗಳನ್ನು ಒದಗಿಸುತ್ತದೆ.

Figure 29: Pragmaan-IT data examples in category Information Seeking - (Hard, 1-Turn, Medium-IF, Safe, IC-3, Normal (No Thinking Trails)) settings configuration.

Figure 30: *Pragyaan-IT* data examples in categories Named Entity Recognition - (*Hard, 1-Turn, Complex-IF, Safe, IC-3, Normal (No Thinking Trails)*), Comprehension - (*Easy, 1-Turn, Simple-IF, Safe, IC-1, Normal (No Thinking Trails)*) and Inference - (*Hard, 1-Turn, Simple-IF, Safe, IC-1, Normal (No Thinking Trails)*) settings configuration.