
CurLL: Curriculum Learning of Language Models

Anonymous Author(s)

Affiliation

Address

email

Abstract

We introduce a comprehensive continual learning dataset and benchmark (CURLL) grounded in human developmental trajectories from ages 5–10, enabling fine-grained assessment of models’ ability to acquire new skills. CURLL spans five developmental stages (0–4) covering ages 5–10, supported by a skill graph that breaks down broad skills into smaller abilities, concrete goals, and measurable indicators, while also capturing which abilities build on others. We generate a 23.4B-token synthetic dataset with controlled skill progression, vocabulary complexity, and format diversity, comprising paragraphs, comprehension-based QA (CQA), skill-testing QA (CSQA), and instruction–response (IR) pairs. Stage-wise token counts range from 2.12B to 6.78B tokens, supporting precise analysis of forgetting, forward and backward transfer. Using a 135M-parameter transformer trained under independent, joint, and sequential setups, we show trade-offs in skill retention and transfer efficiency.

1 Introduction

The capacity for lifelong learning in humans is not just a practical advantage but a fundamental aspect of intelligence itself [Kudithipudi et al., 2022, Yan et al., 2024, Schmidgall et al., 2023]. The continual learning (CL) problem thus is one of the grand challenges for achieving human-like artificial intelligence. It addresses the core problem of how computational systems can progressively acquire, integrate, and refine knowledge over extended periods without compromising earlier capabilities. For language models (LMs), this challenge is particularly interesting: despite their impressive performance across various tasks, these models face a fundamental limitation in that their skill-set and knowledge of the world become static after training, frozen at the point of deployment [Shi et al., 2024, Wu et al., 2024, Bell et al., 2025]. Despite the importance of the CL problem for LMs, current evaluation methodologies suffer from significant limitations: 1) Poor skill control: Existing benchmarks often lack precise control over the specific skills being tested, making it difficult to isolate the effects of learning new capabilities [Liu et al., 2025, Rivera et al., 2022]. 2) Unclear knowledge dependencies: The relationships between skills are rarely explicitly modeled, missing out on important transfer effects [Zheng et al., 2025, Nekoei et al., 2021]. 3) Inadequate forgetting metrics: Many evaluations fail to properly measure catastrophic forgetting across sequential learning tasks [Chen et al., 2023a, Huang et al., 2023].

To address these gaps, we introduce a dataset (CURLL) to train and evaluate continual learning algorithms for language models. Coming up with a set of skills with a rich structure and dependencies is a challenge in the construction of such a dataset. We find such a source of skills in human education. CURLL is grounded in the curriculum for human education from ages 5–10, divided into five developmental stages (0–4). Each of these stages represent one human-year. Our framework incorporates 1,300+ fine-grained skills with dependencies codified in a skill graph having skills as nodes with the edges capturing a prerequisite relationship. The edges are weighted on a scale of (1–5) to capture dependency strength. Starting from this set of skills, we generate a synthetic

dataset of 23.4B tokens, with controlled vocabulary complexity (stage-specific word sampling from Age-of-Acquisition data as seed) and multiple formats (paragraphs, comprehension QA, skill-testing QA, instruction-response). Each stage’s dataset ranges from 2.12B to 6.78B tokens, enabling fine-grained evaluation at indicator, skill, and stage levels. Our code, dataset (stages 0–4), and skill graph will be publicly released. Our contributions include: a) The idea of grounding skills in human education curriculum in the context of CL. b) A synthetic data generation pipeline spanning 5 developmental stages with stage-specific vocabulary and explicit skill dependencies. This pipeline gives us a benchmark with fine-grained control over measuring skill transfer, forgetting and sample efficiency c) A skill graph-based dependency model that explicitly captures prerequisite relationships between learning objectives, enabling nuanced analysis of skill transfer and forgetting.

2 Related Work

Many datasets and benchmarks exist for continual learning of LMs [Jang et al., 2021, Li et al., 2025]. TRACE [Wang et al., 2023] highlights that existing benchmarks are too simple or are already included in instruction-tuning sets. MMLM-CL [Zhao et al., 2025] notes the limited real world applicability in benchmarks. OCKL [Wu et al., 2023] proposes new metrics for measuring knowledge acquisition rate and knowledge gap but concentrates on knowledge-intensive tasks as compared to procedural tasks. TemporalWiki [Jang et al., 2022] is for updating factual information in LMs based on temporal data. SuperNI contains a variety of traditional NLP tasks and serves as a practical benchmark for continual learning of large language models [He et al., 2024]. Despite these developments, these benchmarks are often considered unsuitable for evaluating state-of-the-art LMs [Wang et al., 2023, Razdaibiedina et al., 2023, Scialom et al., 2022, Zhang et al., 2015]. These benchmarks often emphasize artificial task boundaries He et al. [2024], lack temporal and distributional complexity. Moreover, these datasets do not offer precise control over skills or information to validate the effectiveness of existing solutions for continual learning. Skill-it [Chen et al., 2023b] emphasizes the problem but only introduces a data sampling algorithm for continual pretraining by arranging the skills in a increasing order of complexity. Other existing works [Khetarpal et al., 2020, Greco et al., 2019, Xu et al., 2024] discuss the importance of skill distinction and its effect on evaluating continual learning. In contrast, our work is grounded in human developmental curricula and enables fine-grained evaluation of transfer, forgetting, and sample efficiency beyond what existing benchmarks support.

3 Dataset Setup

Our framework is grounded in human learning curriculum, with the dataset designed to mimic the developmental stages from age 5-10. We use two established educational frameworks to develop our skill taxonomy: the Early learning Outcomes framework (ELOF) for children aged 5¹ and the Cambridge curriculum for children aged 5-10². These frameworks help us define fine grained notion of skills as specified by a skill-tuple that consists of four components: 1) Skills³: High-level domains or subjects (e.g. Mathematics, Science). 2) Sub-skills: Specific components within a skill (e.g., Counting and Cardinality). 3) Goals: Broad statement of learning expectations within a sub-skill. 4) Indicators: Specific, observable behaviors that demonstrate mastery of a goal.

The ELOF framework has five broad areas: Approaches to Learning, Social and Emotional Development, Language and Literacy, Cognition, and Perceptual, Motor, and Physical Development. Cambridge Primary Curriculum covers subjects including English, Mathematics, Science, Computing, and Global Perspectives. The curriculum structure flows from subjects (renamed as skills in our framework) to domains/strands (renamed as subskills), then to substrands (goals), each with specific learning objectives (indicators). We also adopt the notion of stages from the Cambridge curriculum in our framework, where each stage corresponds to one year starting from age 5. Therefore, we have 5 stages in our framework, where stage 0 denotes ages up to 5, stage 1 denotes age 5-6 and so on. The number of skill-tuples in our framework is the same as the number of indicators present in stages 0-4, statistics of which are mentioned in Table 1. We construct a skill graph, which is a directed graph that has indicators as nodes, with edges representing prerequisite relationships weighted from 1-5 to indicate dependency strength. These edges model how skills are built on each other in developmental

¹U.S. Department of Health & Human Services, Administration for Children & Families [2024]

²Cambridge Assessment International Education [2025]

³"Skill" here has a specific meaning, which is different from the general notion of skill used before

Table 1: Dataset statistics across developmental stages (0–4), including total tokens

Stage	Skills & Goals				Instances			# Tokens in Bn
	# Skills	# Sub-skills	# Goals	# Indicators	# CQA	# CSQA	# IR Pairs	
0	7	24	59	182	1.0M	3.01M	3.30M	2.12
1	7	29	86	292	20.2M	4.04M	4.10M	3.47
2	6	26	67	249	23.5M	4.70M	4.78M	4.56
3	6	26	68	271	31.2M	6.24M	6.29M	6.47
4	6	23	70	349	27.4M	5.49M	5.52M	6.78

stages⁴. While the skill graph isn’t directly used for skill data generation, it provides insights for analyzing continual learning patterns and interpreting evaluation results (see Appendix A).

3.1 Synthetic Data Generation

Our synthetic data consists of instances, each mimicking a situation a child might encounter. Instances are of three types: (1) IR: an instruction-response pair, where the instruction is about some general world knowledge, (2) CQA: context-based question-answers for testing comprehension, (3) CSQA: context-based question-answers for testing skills. A context is a short piece of text which forms the basis of the corresponding question-answer pairs in the instance. Contexts can be of multiple types as specified by a template: e.g., a simple narrative, or a dialogue. IR-pairs can also have different types specified by templates, e.g., mimic action or follow simple direction (Appendix A for examples).

Instances are generated by prompting an LLM with a *seed*. A seed consists of a skill-tuple, vocabulary seed, instance type, template. This choice is crucial for ensuring diversity and coverage of our data. This tuple is also our way to ground the generations in the skill graph. To generate one instance of the data, we first construct a seed: each skill-tuple is combined with a vocabulary seed for that stage, an instance type and a template for that instance type. If the instance type is CQA or CSQA, then we first generate the context, and then using the context, we generate the corresponding question-answers. If the instance type is IR then we directly generate the instruction-response pairs. The prompts for all the generations and details of the generation process are presented in Appendix A. In our dataset, each instance includes the seed used to generate it as part of its metadata. We generated data for stages 0-4, containing a total of 23.4B tokens (Table 1).

We measure diversity of generated data using: 1) Diversity as reciprocal of compression ratio using gzip Gailly and Adler [1992]. 2) The intra- and inter-text deduplication rate as calculated by semantic deduplication. Cross-stage analysis shows higher diversity and lower deduplication rate (<5%) between stages compared to intra-stage results, confirming that content evolves meaningfully across developmental progression while maintaining stage-specific uniqueness. See Appendix B for more details. We also measure progression in the difficulty of the skills as the stage number increases. We sample 500K instances from each stage for each data type and run statistical readability tests⁵. Means across multiple readability metrics are reported in Appendix B.3. The readability tests show that as stages progress, the texts also become increasingly challenging. At least 50 random instances from each dataset per stage were manually analysed, revealing that CQA data for all stages was found to be accurate. IR and CSQA data had certain patterns like excessive use of discourse markers for early stages and verbose response to instructions. We choose 25 instances per indicator for test set resulting in 5k-7k samples per stage. Since the data is synthetically generated at scale, we reserve the highest quality samples for the test set. 100 instances per indicator instance type are sampled randomly and rated by LLM on a scale 1-5. Top 25 instances are selected for test, next 25 for validation set.

4 Experiments and Results

We conduct preliminary experiments to validate that the dataset exposes meaningful challenges: whether models can retain earlier-learned skills, how sequential training affects generalization, and to what extent transfer across related skills occurs. By analyzing these at the granularity of skills, we demonstrate that CURLL enables insights that are not visible in existing benchmarks. Unlike

⁴an LLM (Gemma3-27B-IT is used for all LLM inferences throughout this work) is used to predict the edges

⁵Uses pre-defined words to predict the grade of a text (<https://github.com/cdimascio/py-readability-metrics>)

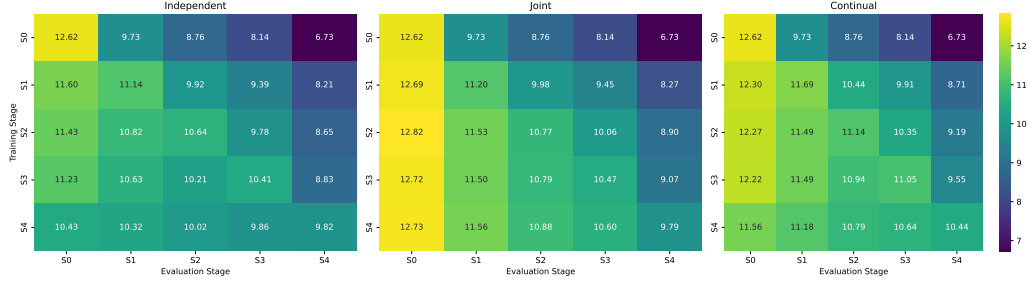


Figure 1: Stage-wise results for all training setups. Independent corresponds to models trained on a single stage, joint to models trained on mixtures of data up to a stage, and continual to sequential up to a stage. Heatmaps report summed correctness scores across all test formats (IR, CQA, CSQA)

traditional language model training that includes two stages: pretraining and then finetuning, we do a single phase training. All instance types i.e. CQA, CSQA, and IR are included in the same phase. Since all of them are question-answers, with and without context paragraphs, we use a standard chat template to train the language models from scratch. Smollm2-135M parameter model is used as the base architecture. All training runs are performed on one full epoch of the data. Learning rate of $5e-3$ and effective batch size of 1536 instances remain unchanged across experiments. We use a context length of 1024. Other training and inference related hyper-parameters are mentioned in the Appendix C. We perform three types of training: 1) Independent (M_i): The model is trained from scratch on data of each stage independently 2) Joint (M_{ij}): Jointly trained on a mixture of stages. The data from different stages is combined and shuffled randomly. 3) Continual (M_{i-j}): The model is first trained on stage i , then stage j , then stage k and so on. To evaluate the trained models, the instances from test set are passed through the chat template and the model is asked to complete the generation post instruction. These inferences along with the prompt is passed to an LLM to rate on a scale of 1-5. This is followed for all three types of test sets. Each model is evaluated on test sets of all stages. The main objective of the rating is to evaluate the correctness of the model inference with some weightage to the stage on which the model is being evaluated. The summation of scores across test set types (IR, CQA, CSQA) is presented in Figure 1. The individual scores, prompts and rubrics for evaluation are available in the Appendix E.

Joint models (M_{ij}) generalize better to later stages and maintain strong performance on trained stages compared to independent models (M_i). Continual models (M_{i-j}), however, achieve the best performance on later stages but suffer degradation on earlier ones. Sequential (continual) ordering improves generalization but also induces forgetting of earlier skills, which is counter-intuitive since later skills depend on foundational ones. The skill graph helps explain this. The largest performance gaps between joint and continual training occur for “Perceptual, Motor, and Physical Development” and “Digital Literacy”. Both have very few outgoing edges in the skill graph (Appendix D), meaning their indicators rarely serve as prerequisites for later skills.

5 Conclusion

We introduced (CURLL), a novel continual learning evaluation framework for language models grounded in human developmental curricula. (CURLL) combines a directed, weighted skill graph of over 1,300 fine-grained skills with a 23.4B-token synthetic dataset that controls stage-wise vocabulary, difficulty, and format. The skill graph serves as a diagnostic tool: its metadata enables fine-grained control over the number of instances and skills seen during training, supports evaluation of sample efficiency, and allows targeted testing of transfer effects (e.g., whether learning Skill A improves Skill B). Forgetting, forward transfer, backward transfer, and data efficiency can all be measured at the levels of skills, sub-skills, and indicators. This enables richer analysis than stage- or task-level metrics in existing benchmarks, which typically report only overall accuracy on entire tasks (e.g., classification or QA) without revealing which underlying abilities are gained or lost. Our experiments with independent, joint, and sequential training demonstrate that simply changing the order of data presentation affects both generalization and forgetting. Finally, the scalable data generation pipeline enables exploring continual pretraining in a controlled yet realistic setting.

References

- Jack Bell, Luigi Quarantiello, Eric Nuerthey Coleman, Lanpei Li, Malio Li, Mauro Madeddu, Elia Piccoli, and Vincenzo Lomonaco. The future of continual learning in the era of foundation models: Three key directions. *ArXiv*, abs/2506.03320, 2025. URL <https://api.semanticscholar.org/CorpusId:279155222>.
- Cambridge Assessment International Education. International curriculum. <https://www.cambridgeinternational.org/why-choose-us/benefits-of-a-cambridge-education/international-curriculum/>, 2025. URL <https://www.cambridgeinternational.org/why-choose-us/benefits-of-a-cambridge-education/international-curriculum/>. Accessed: 2025-08-20.
- Ernie Chang, Matteo Paltenghi, Yang Li, Pin-Jie Lin, Changsheng Zhao, Patrick Huber, Zechun Liu, Rastislav Rabatin, Yangyang Shi, and Vikas Chandra. Scaling parameter-constrained language models with quality data. *ArXiv*, abs/2410.03083, 2024. URL <https://api.semanticscholar.org/CorpusId:273162494>.
- Jiefeng Chen, Timothy Nguyen, Dilan Gorur, and Arslan Chaudhry. Is forgetting less a good inductive bias for forward transfer? *ArXiv*, abs/2303.08207, 2023a. URL <https://api.semanticscholar.org/CorpusId:257532608>.
- Mayee F. Chen, Nicholas Roberts, K. Bhatia, Jue Wang, Ce Zhang, Frederic Sala, and Christopher Ré. Skill-it! a data-driven skills framework for understanding and training language models. *ArXiv*, abs/2307.14430, 2023b. URL <https://api.semanticscholar.org/CorpusId:260203057>.
- Ronen Eldan and Yuanzhi Li. Tinstories: How small can language models be and still speak coherent english?, 2023. URL <https://arxiv.org/abs/2305.07759>.
- Jean Gailly and Mark Adler. GNU gzip. GNU Operating System, 1992.
- Claudio Greco, Barbara Plank, R. Fernández, and R. Bernardi. Psycholinguistics meets continual learning: Measuring catastrophic forgetting in visual question answering. In *Annual Meeting of the Association for Computational Linguistics*, 2019. URL <https://api.semanticscholar.org/CorpusId:184488333>.
- Jinghan He, Haiyun Guo, Kuan Zhu, Zihan Zhao, Ming Tang, and Jinqiao Wang. Seekr: Selective attention-guided knowledge retention for continual learning of large language models. In *Conference on Empirical Methods in Natural Language Processing*, 2024. URL <https://api.semanticscholar.org/CorpusId:273901289>.
- Heng Huang, Li Shen, Enneng Yang, and Zhenyi Wang. A comprehensive survey of forgetting in deep learning beyond continual learning. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 47:1464–1483, 2023. URL <https://api.semanticscholar.org/CorpusId:259951356>.
- Joel Jang, Seonghyeon Ye, Sohee Yang, Joongbo Shin, Janghoon Han, Gyeonghun Kim, Stanley Jungkyu Choi, and Minjoon Seo. Towards continual knowledge learning of language models. *ArXiv*, abs/2110.03215, 2021. URL <https://api.semanticscholar.org/CorpusId:238419458>.
- Joel Jang, Seonghyeon Ye, Changho Lee, Sohee Yang, Joongbo Shin, Janghoon Han, Gyeonghun Kim, and Minjoon Seo. Temporalwiki: A lifelong benchmark for training and evaluating ever-evolving language models. In *Conference on Empirical Methods in Natural Language Processing*, 2022. URL <https://www.aclanthology.org/2022.emnlp-main.418.pdf>.
- Khimya Khetarpal, M. Riemer, I. Rish, and Doina Precup. Towards continual reinforcement learning: A review and perspectives. *J. Artif. Intell. Res.*, 75:1401–1476, 2020. URL <https://api.semanticscholar.org/CorpusId:229679944>.

216 D. Kudithipudi, Mario Aguilar-Simon, Jonathan Babb, M. Bazhenov, Douglas Blackiston, J. Bongard,
217 Andrew P. Brna, Suraj Chakravarthi Raja, Nick Cheney, J. Clune, A. Daram, Stefano Fusi, Peter
218 Helfer, Leslie M. Kay, Nicholas A. Ketz, Z. Kira, Soheil Kolouri, J. Krichmar, Sam Kriegman,
219 Michael Levin, Sandeep Madireddy, Santosh Manicka, Ali Marjaninejad, Bruce L. McNaughton,
220 R. Miikkulainen, Zaneta Navratilova, Tej Pandit, Alice Parker, Praveen K. Pilly, S. Risi, T. Se-
221 jnowski, Andrea Soltoggio, Nicholas Soures, A. Talias, Darío Urbina-Meléndez, F. Valero-Cuevas,
222 Gido M. van de Ven, J. Vogelstein, Felix Wang, Ron Weiss, A. Yanguas-Gil, Xinyun Zou, and
223 H. Siegelmann. Biological underpinnings for lifelong learning machines. *Nature Machine Intelli-*
224 *gence*, 4:196 – 210, 2022. URL <https://doi.org/10.1038/s42256-022-00452-0>.

225 Jeffrey Li, Mohammadreza Armandpour, Iman Mirzadeh, Sachin Mehta, Vaishaal Shankar, Raviteja
226 Vemulapalli, Samy Bengio, Oncel Tuzel, Mehrdad Farajtabar, Hadi Pouransari, and Fartash Faghri.
227 Tic-lm: A web-scale benchmark for time-continual llm pretraining. *ArXiv*, abs/2504.02107, 2025.
228 URL <https://api.semanticscholar.org/CorpusId:277510618>.

229 Jia Liu, Jinguo Cheng, Xiangming Fang, Zhenyuan Ma, and Yuankai Wu. Evaluating temporal
230 plasticity in foundation time series models for incremental fine-tuning. *ArXiv*, abs/2504.14677,
231 2025. URL <https://api.semanticscholar.org/CorpusId:277954770>.

232 Hadi Nekoei, Akilesh Badrinaaraayanan, Aaron C. Courville, and Sarath Chandar. Continuous
233 coordination as a realistic scenario for lifelong learning. *ArXiv*, abs/2103.03216, 2021. URL
234 <https://api.semanticscholar.org/CorpusId:232110854>.

235 Anastasia Razdaibiedina, Yuning Mao, Rui Hou, Madian Khabsa, Mike Lewis, and Amjad Almahairi.
236 Progressive prompts: Continual learning for language models. *ArXiv*, abs/2301.12314, 2023. URL
237 <https://api.semanticscholar.org/CorpusId:256390383>.

238 Corban G. Rivera, C. Ashcraft, Alexander New, J. Schmidt, and Gautam K. Vallabha. Latent
239 properties of lifelong learning systems. *ArXiv*, abs/2207.14378, 2022. URL <https://api.semanticscholar.org/CorpusId:251196582>.

241 Samuel Schmidgall, Jascha Achterberg, Thomas Miconi, Louis Kirsch, Rojin Ziaei, S. P. Hajiseye-
242 draz, and Jason Eshraghian. Brain-inspired learning in artificial neural networks: a review. *ArXiv*,
243 abs/2305.11252, 2023. URL <https://api.semanticscholar.org/CorpusId:258823273>.

244 Thomas Scialom, Tuhin Chakrabarty, and Smaranda Muresan. Fine-tuned language models are
245 continual learners. In *Conference on Empirical Methods in Natural Language Processing*, 2022.
246 URL <https://api.semanticscholar.org/CorpusId:252815378>.

247 Haizhou Shi, Zihao Xu, Hengyi Wang, Weiyi Qin, Wenyuan Wang, Yibin Wang, and Hao Wang.
248 Continual learning of large language models: A comprehensive survey. *ACM Computing Surveys*,
249 2024. URL <https://api.semanticscholar.org/CorpusId:269362836>.

250 U.S. Department of Health & Human Services, Administration for Children & Families. Head start
251 early learning outcomes framework: Ages birth to five, 2024. URL <https://headstart.gov/school-readiness/article/head-start-early-learning-outcomes-framework>. Last
252 updated: December 23, 2024; Accessed: 2025-08-20.

254 Xiao Wang, Yuan Zhang, Tianze Chen, Songyang Gao, Senjie Jin, Xianjun Yang, Zhiheng Xi,
255 Rui Zheng, Yicheng Zou, Tao Gui, Qi Zhang, and Xuanjing Huang. Trace: A comprehensive
256 benchmark for continual learning in large language models. *ArXiv*, abs/2310.06762, 2023. URL
257 <https://api.semanticscholar.org/CorpusId:263830425>.

258 Tongtong Wu, Linhao Luo, Yuan-Fang Li, Shirui Pan, Thuy-Trang Vu, and Gholamreza Haffari.
259 Continual learning for large language models: A survey. *ArXiv*, abs/2402.01364, 2024. URL
260 <https://api.semanticscholar.org/CorpusId:267406164>.

261 Yuhao Wu, Tongjun Shi, Karthick Sharma, Chun Seah, and Shuhao Zhang. Online continual
262 knowledge learning for language models. *ArXiv*, abs/2311.09632, 2023. URL <https://api.semanticscholar.org/CorpusId:265221422>.

264 Yongxin Xu, Philip S. Yu, Zexin Lu, Xu Chu, Yujie Feng, Bo Liu, and Xiao-Ming Wu. Klf:
265 Knowledge localization and fusion for language model continual learning. 2024. URL <https://api.semanticscholar.org/CorpusId:271843361>.

- 267 Lixiang Yan, Samuel Greiff, Ziwen Teuber, and D. Gaević. Promises and challenges of generative
 268 artificial intelligence for human learning. *Nature human behaviour*, 8 10:1839–1850, 2024. URL
 269 <https://api.semanticscholar.org/CorpusId:271924303>.
- 270 Xiang Zhang, Junbo Jake Zhao, and Yann LeCun. Character-level convolutional networks for
 271 text classification. In *Neural Information Processing Systems*, 2015. URL <https://api.semanticscholar.org/CorpusId:368182>.
- 273 Hongbo Zhao, Fei Zhu, Rundong Wang, Gaofeng Meng, and Zhaoxiang Zhang. Mllm-cl: Continual
 274 learning for multimodal large language models. *ArXiv*, abs/2506.05453, 2025. URL <https://api.semanticscholar.org/CorpusId:279243888>.
- 276 Junhao Zheng, Xidi Cai, Qiuke Li, Duzhen Zhang, Zhongzhi Li, Yingying Zhang, Le Song, and
 277 Qianli Ma. Lifelongagentbench: Evaluating llm agents as lifelong learners. *ArXiv*, abs/2505.11942,
 278 2025. URL <https://api.semanticscholar.org/CorpusId:278739762>.

279 A Dataset Construction

280 Figure 2 gives an overview of our dataset, including examples of skills, subskills, goal and indicator.
 281 We also present an example of an edge from the skill graph. Figure 3 shows the number of incoming
 282 and outgoing edges from each stage. Figure 4 explains the data construction process and gives
 283 examples from each stage of the data generation pipeline.

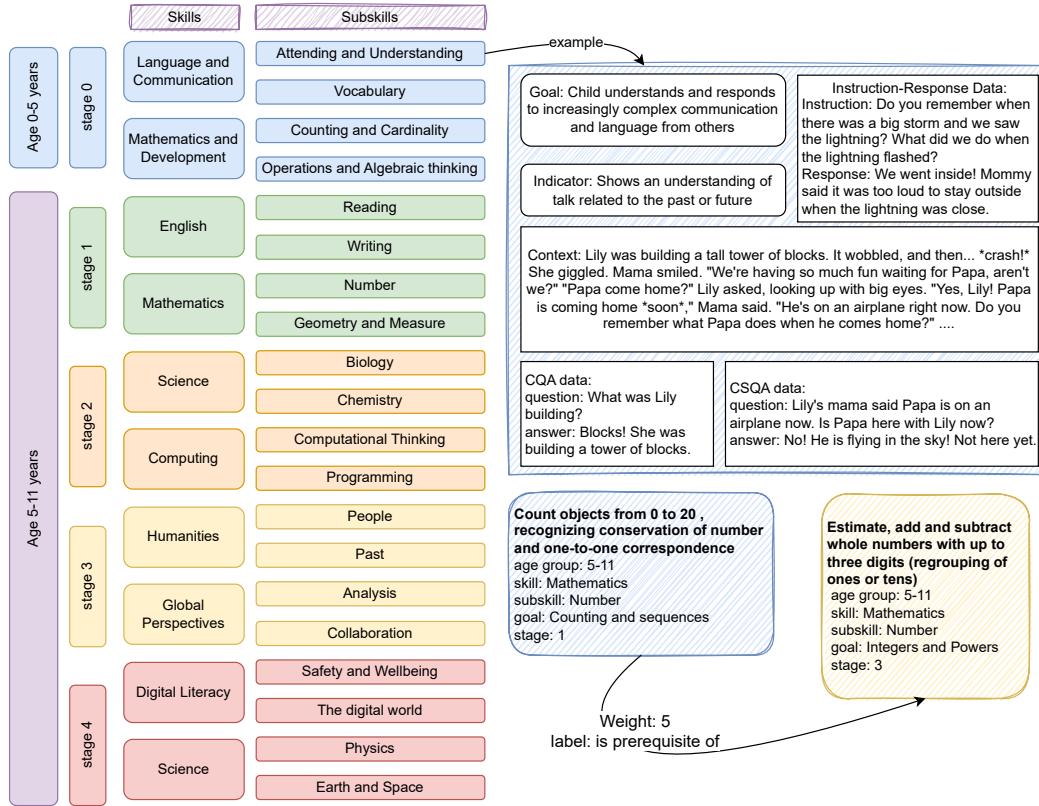


Figure 2: Developmental framework for children aged 0-11 years, categorized into stages (0-4). Only examples of skills and subskills are mentioned here. An example of how the data looks like is given in the top right. Two nodes and an edge from the skill graph is given in the bottom right.

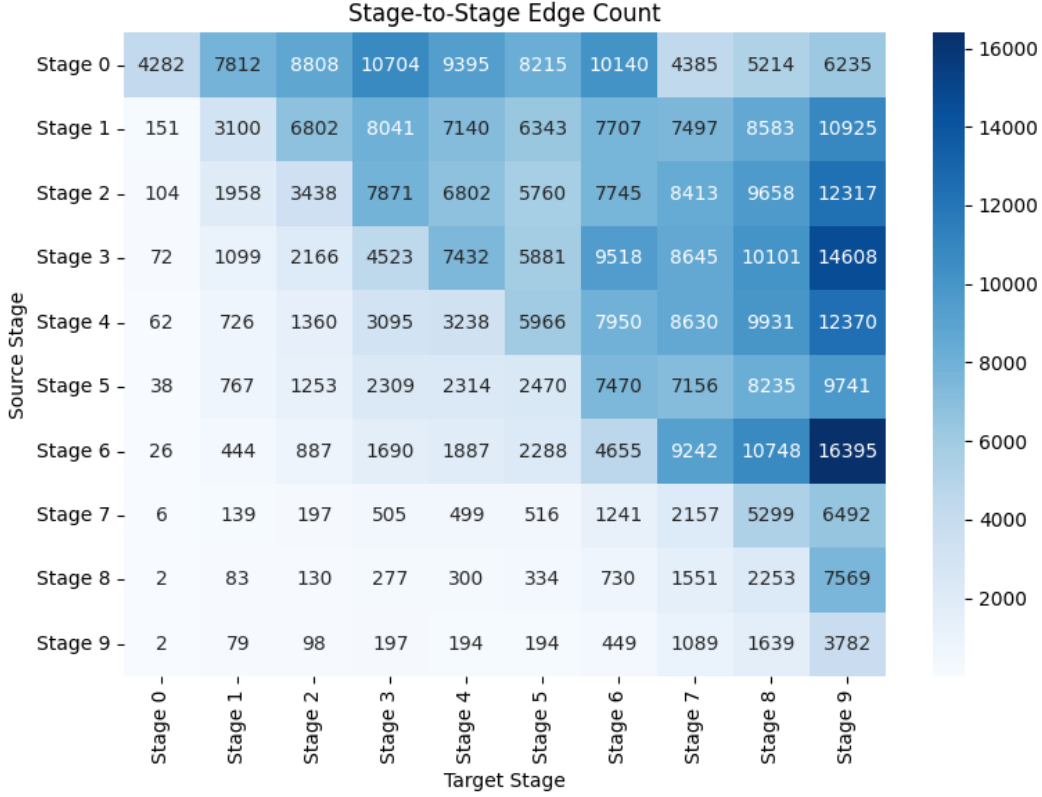


Figure 3: Heatmap showing the number of prerequisite edges between stages in the skill graph. Rows correspond to source stages, columns to target stages, and color intensity indicates the number of connections.

284 B Data Verification

285 For both the methods, 500K texts are sampled from each of the Paragraphs and Instruction-response
286 pairs.

287 B.1 Diversity

288 For the diversity of the text, we follow Chang et al. [2024] and calculate the compression ratio of the
289 text as

$$CR(D) = \frac{Originalsizeof D \text{ (bytes)}}{Compressedsizeof D \text{ (bytes)}},$$

290 and define diversity by

$$Dr(D) = 1/CR(D).$$

291 A higher compression ratio $CR(D)$ indicates greater redundancy, meaning lower diversity in the
292 text. Thus, diversity $Dr(D)$ increases when redundancy decreases. We see diversity ranging between
293 30.77% and 35.60%, which is similar to other work. As a comparison, we also calculated the diversity
294 of 500K samples from the validation set of TinyStories, a paper exploring synthetic data generation
295 to train a small language model. Their text diversity ranges from 31.04% to 32.66% within the
296 pretraining and instruct data, respectively Eldan and Li [2023].

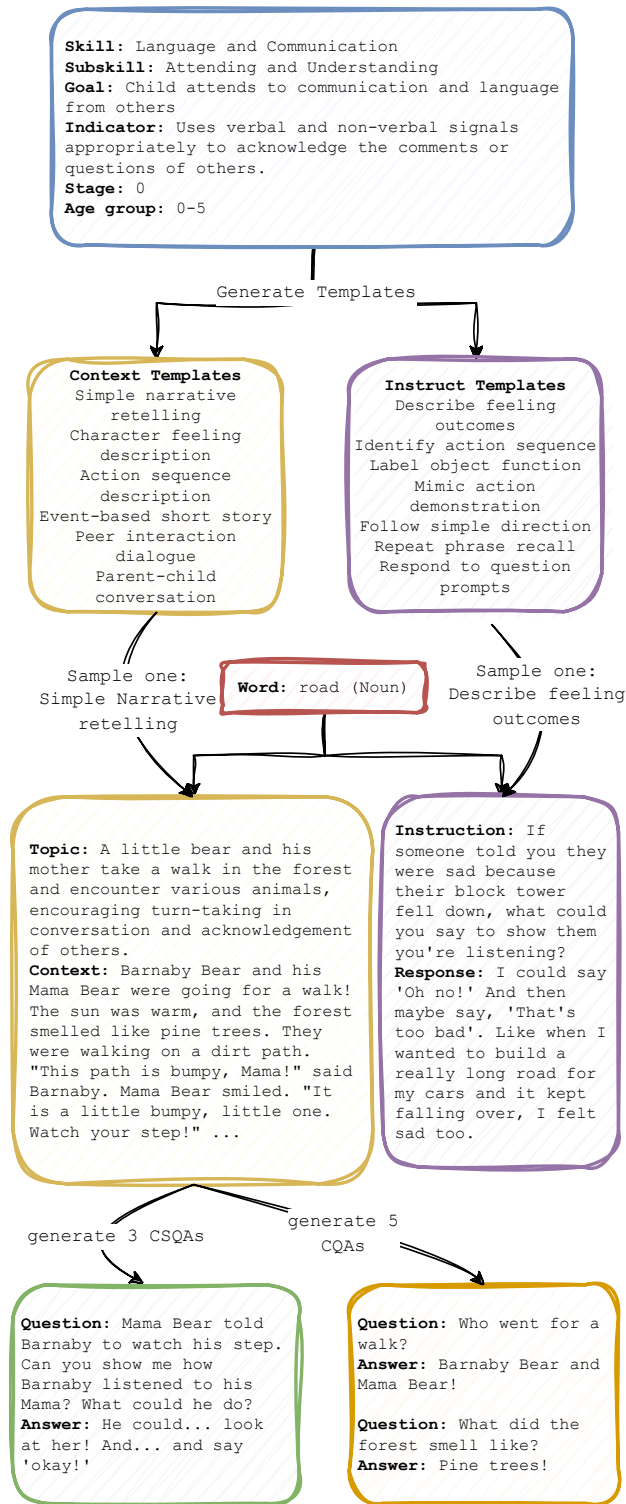


Figure 4: Synthetic data generation pipeline

Table 2: Diversity and Deduplication metrics for context and instruction–response data across stages

Stage	Context		IR	
	Div ↑	Dedup ↓	Div ↑	Dedup ↓
0	34.29%	11.83%	30.77%	3.50%
1	35.60%	5.36%	31.73%	3.85%
2	34.17%	15.47%	32.64%	2.54%
3	34.68%	14.86%	32.97%	2.09%
4	35.45%	13.41%	33.14%	1.93%

B.2 Deduplication

For semantic deduplication⁶, we pass the texts through a sentence encoder and find the deduplication rate as the percentage of sentences that have cosine similarity of at least 0.95 with another sentence in the same stage.

Table 3: Diversity and Deduplication Rates when Considering Pairwise Stages

Stage Pair	Context	
	Div ↑	Dedup ↓
0, 1	31.29%	0.3%
0, 2	31.96%	0.1%
0, 3	32.25%	0.0%
0, 4	32.50%	0.0%
1, 2	32.27%	0.3%
1, 3	32.52%	0.2%
1, 4	32.71%	0.1%
2, 3	32.82%	0.4%
2, 4	32.94%	0.2%
3, 4	33.07%	0.2%

B.3 Detailed Readability Metrics

Note that average grade of the data is slightly higher than the intended age of the data (especially for the first few stages). However, this is because not all skills we generate data for are, in real-life, text-based. Thus, demonstrating them in language ends up requiring complex words, which affects the readability score. For example, children can verbally reason about cause-and-effect in multi-turn conversations, but when written down, that same dialogue is rated at a much higher reading level than the child can actually read, leading to higher readability scores in our data.

Table 4: Average readability scores of generated data across stages, reported for context, comprehension QA (CQA), skill-testing QA (CSQA), and instruction–response (IR) data. Scores generally increase with stage, reflecting controlled growth in textual complexity aligned with developmental progression

Stage	Context	CQA	CSQA	IR
0	4.61 ^{1.87}	2.38 ^{2.88}	3.07 ^{2.26}	4.48 ^{1.52}
1	5.24 ^{1.72}	4.39 ^{1.81}	4.44 ^{1.62}	4.86 ^{1.41}
2	5.18 ^{1.93}	4.39 ^{1.80}	4.69 ^{1.54}	4.69 ^{1.59}
3	5.51 ^{1.85}	4.65 ^{1.70}	4.98 ^{1.46}	5.03 ^{1.50}
4	6.42 ^{1.79}	5.63 ^{1.44}	5.96 ^{1.30}	5.91 ^{1.34}

⁶We use the following repo for semantic deduplication: <https://github.com/MinishLab/semhash>

Table 5: Detailed Readability Metrics Across all 5 Stages and Datasets

Dataset	Stage	Flesch Kincaid	SMOG	Coleman Liau	Automated Readability	Dale Chall	Gunning Fog
Context	0	3.15 _{0.35}	6.90 _{0.76}	4.28 _{0.51}	1.68 _{0.47}	6.69 _{0.27}	4.94 _{0.34}
Context	1	3.68 _{0.35}	7.55 _{0.73}	5.18 _{0.50}	2.58 _{0.49}	6.70 _{0.29}	5.74 _{0.35}
Context	2	3.80 _{0.36}	7.54 _{0.75}	4.21 _{0.46}	2.25 _{0.48}	7.18 _{0.35}	6.12 _{0.38}
Context	3	4.16 _{0.36}	7.84 _{0.74}	4.58 _{0.48}	2.71 _{0.50}	7.27 _{0.36}	6.48 _{0.38}
Context	4	5.13 _{0.42}	8.76 _{0.79}	5.39 _{0.51}	3.77 _{0.56}	7.89 _{0.34}	7.58 _{0.45}
CQA	0	0.79 _{0.35}	5.06 _{0.59}	0.26 _{0.59}	-1.47 _{0.43}	6.89 _{0.30}	2.75 _{0.34}
CQA	1	2.73 _{0.37}	6.45 _{0.59}	4.10 _{0.54}	1.65 _{0.49}	6.47 _{0.26}	4.92 _{0.45}
CQA	2	2.74 _{0.38}	6.42 _{0.59}	4.00 _{0.53}	1.67 _{0.50}	6.44 _{0.28}	5.07 _{0.45}
CQA	3	3.04 _{0.37}	6.66 _{0.59}	4.37 _{0.52}	2.08 _{0.49}	6.41 _{0.27}	5.36 _{0.44}
CQA	4	4.08 _{0.38}	7.54 _{0.59}	5.59 _{0.48}	3.52 _{0.49}	6.50 _{0.25}	6.54 _{0.47}
CSQA	0	1.34 _{0.28}	5.37 _{0.70}	2.07 _{0.43}	-0.20 _{0.36}	6.21 _{0.20}	3.65 _{0.27}
CSQA	1	2.84 _{0.30}	6.36 _{0.71}	4.14 _{0.37}	2.04 _{0.40}	6.03 _{0.21}	5.24 _{0.33}
CSQA	2	3.16 _{0.29}	6.54 _{0.72}	4.33 _{0.37}	2.43 _{0.39}	6.08 _{0.23}	5.59 _{0.33}
CSQA	3	3.49 _{0.29}	6.81 _{0.70}	4.64 _{0.37}	2.87 _{0.39}	6.14 _{0.25}	5.96 _{0.32}
CSQA	4	4.62 _{0.33}	7.72 _{0.72}	5.56 _{0.41}	4.25 _{0.46}	6.50 _{0.27}	7.12 _{0.37}
IR	0	2.97 _{0.47}	6.32 _{0.64}	4.12 _{0.52}	2.25 _{0.62}	5.81 _{0.23}	5.43 _{0.48}
IR	1	3.40 _{0.45}	6.61 _{0.65}	4.51 _{0.50}	2.88 _{0.61}	5.76 _{0.25}	6.02 _{0.50}
IR	2	3.16 _{0.37}	6.62 _{0.72}	4.23 _{0.45}	2.33 _{0.50}	6.10 _{0.26}	5.68 _{0.43}
IR	3	3.55 _{0.37}	6.93 _{0.71}	4.62 _{0.46}	2.87 _{0.51}	6.13 _{0.27}	6.09 _{0.42}
IR	4	4.59 _{0.41}	7.66 _{0.73}	5.41 _{0.46}	4.13 _{0.56}	6.46 _{0.28}	7.20 _{0.47}

C Hyperparameters

All experiments were conducted with a consistent set of training hyperparameters to ensure comparability across runs. Models were initialized using the kaiming normal method unless otherwise specified, and trained with AdamW optimizer ($\beta_1 = 0.9$, $\beta_2 = 0.98$, $\epsilon = 1e - 8$) with weight decay of 0.01. We used a base learning rate of $5e - 3$, applied gradient clipping with a maximum norm of 1.0. We used gradient accumulation (8 steps with batch size 24 on 8 GPUs, yielding an effective batch size of 1536). Training was performed for one full epoch over each dataset split with a context length of 1024 tokens. Mixed precision was enabled with bfloat16 (bf16) for efficiency, while fp16 was disabled. All experiments were seeded with 42 for reproducibility. For inference, the model was loaded in bfloat16 precision with padding set to the EOS token and leftside padding for alignment. Prompts were tokenized with a maximum length of 512 tokens, and generation used a temperature of 0.7, top-p sampling of 0.95, and a maximum of 128 new tokens per prompt.

D Results

Table 6 gives the results of all experiments on IR test set. Table 7 gives the results of all experiments on CQA test set. Table 8 gives the results of all experiments on CSQA test set. Per-stage per-Indicator results can be found here: Results sheet. Forgetting analysis is shown in Figure 5. Relation of forgetting analysis to the skill graph can be drawn from Figure 6.

E Prompts

E.1 Edge Prediction

System prompt for Edge prediction

You are an expert in skill development and cognitive science. Your task is to analyze the relationship between two skill indicators and determine if there is a logical prerequisite dependency between them.

Each skill indicator is given with:

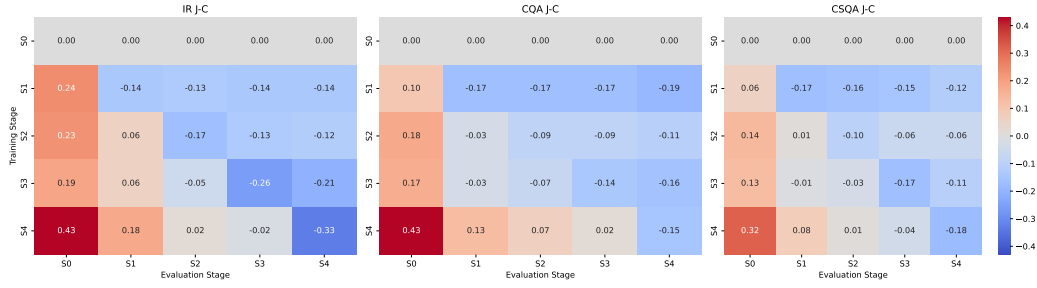


Figure 5: Forgetting analysis across training setups. The plots show performance differences between joint and continual training for IR, CQA, and CSQA test sets across stages 0–4. The Y-axis corresponds to models trained upto a stage. The X-axis corresponds to test set of mentioned stage.

Table 6: All results for IR test set. The column represents each stage on which a model is being evaluated.

Test type Stages	IR (rating out of 5)				
	0	1	2	3	4
M_0	4.16	3.29	2.97	2.83	2.49
M_1	3.70	3.70	3.21	3.08	2.80
M_2	3.71	3.55	3.56	3.27	3.00
M_3	3.64	3.45	3.35	3.57	3.07
M_4	3.38	3.35	3.32	3.34	3.55
M_{012}	4.22	3.81	3.55	3.34	3.07
M_{01}	4.19	3.73	3.25	3.12	2.84
M_{0-1}	3.94	3.87	3.38	3.26	2.98
M_{0123}	4.15	3.79	3.56	3.55	3.14
M_{0-1-2}	3.99	3.75	3.72	3.47	3.19
M_{01234}	4.16	3.80	3.60	3.60	3.46
$M_{0-1-2-3}$	3.97	3.73	3.61	3.82	3.34
$M_{0-1-2-3-4}$	3.73	3.63	3.58	3.62	3.78

Table 7: All results for CQA test set. The column represents each stage on which a model is being evaluated.

Test type Stages	CQA (rating out of 5)				
	0	1	2	3	4
M_0	4.16	3.29	2.97	2.83	2.49
M_1	3.70	3.70	3.21	3.08	2.80
M_2	3.71	3.55	3.56	3.27	3.00
M_3	3.64	3.45	3.35	3.57	3.07
M_4	3.38	3.35	3.32	3.34	3.55
M_{012}	4.22	3.81	3.55	3.34	3.07
M_{01}	4.19	3.73	3.25	3.12	2.84
M_{0-1}	3.94	3.87	3.38	3.26	2.98
M_{0123}	4.15	3.79	3.56	3.55	3.14
M_{0-1-2}	3.99	3.75	3.72	3.47	3.19
M_{01234}	4.61	4.27	4.05	3.87	3.45
$M_{0-1-2-3}$	4.42	4.27	4.09	3.97	3.45
$M_{0-1-2-3-4}$	4.17	4.14	3.97	3.85	3.60

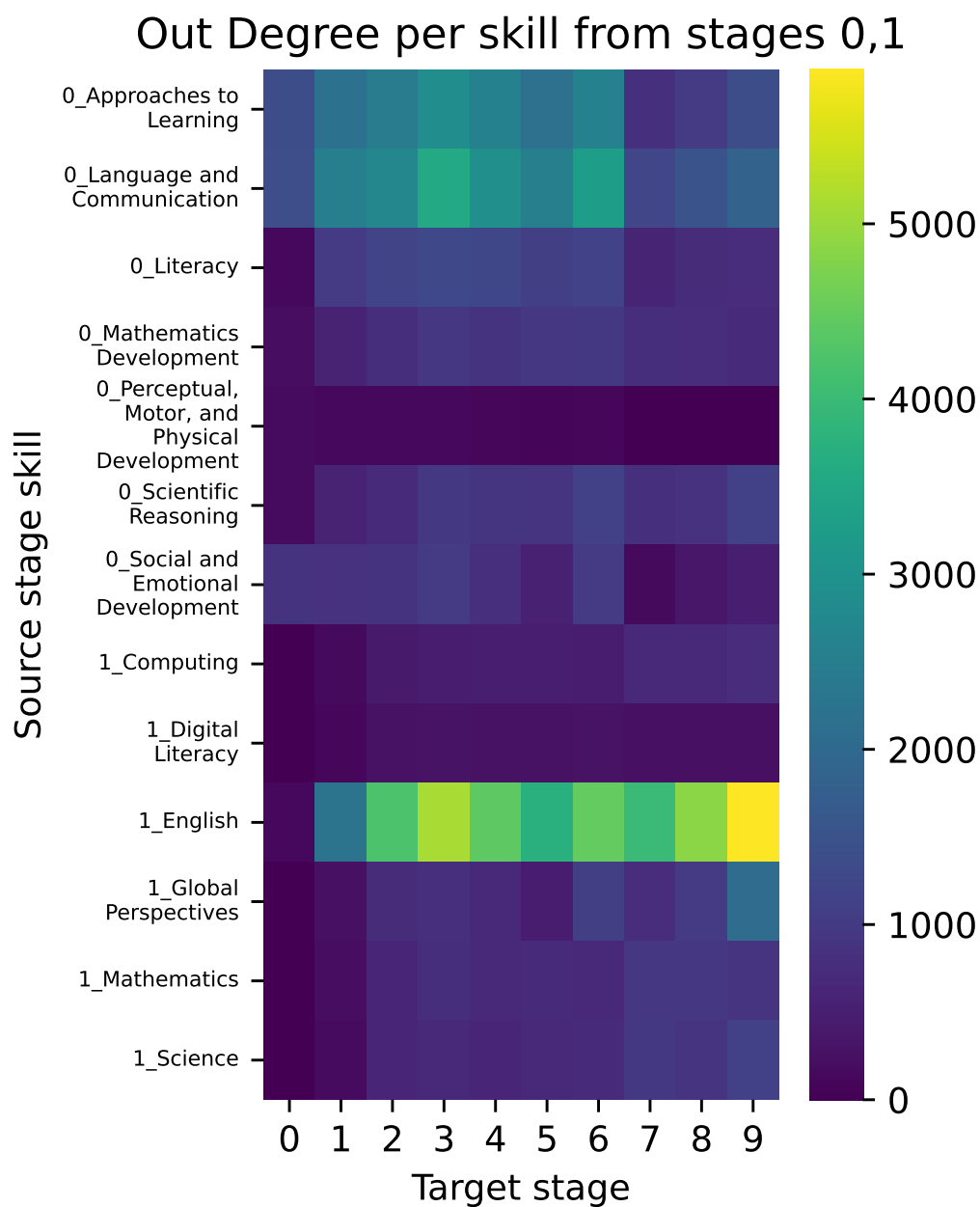


Figure 6: Out-degree distribution of skills from stages 0 and 1 in the skill graph. Skills with fewer outgoing prerequisite edges (e.g., Perceptual, Motor, and Physical Development; Digital Literacy) are less connected to later stages and are observed to be more vulnerable to forgetting in continual training.

Table 8: All results for CSQA test set. The column represents each stage on which a model is being evaluated.

Test type Stages	CSQA (rating out of 5)				
	0	1	2	3	4
M_0	3.89	2.85	2.52	2.33	1.95
M_1	3.63	3.35	2.92	2.75	2.39
M_2	3.53	3.25	3.15	2.87	2.53
M_3	3.51	3.22	3.03	3.10	2.61
M_4	3.29	3.13	3.00	2.93	2.89
M_{012}	3.97	3.48	3.21	2.96	2.61
M_{01}	3.93	3.37	2.93	2.76	2.40
M_{0-1}	3.87	3.55	3.09	2.91	2.51
M_{0123}	3.97	3.47	3.21	3.08	2.65
M_{0-1-2}	3.83	3.47	3.31	3.03	2.66
M_{01234}	3.97	3.49	3.24	3.13	2.88
$M_{0-1-2-3}$	3.83	3.48	3.24	3.26	2.76
$M_{0-1-2-3-4}$	3.65	3.41	3.23	3.17	3.05

```

334 - a_label and a_id
335 - b_label and b_id
336
337 These represent two distinct skill indicators. You must determine whether one is a
338 prerequisite for the other.
339
340 Instructions:
341 - A skill X is a prerequisite for skill Y if Y logically requires understanding or
342 demonstrating X beforehand.
343 - Compare the meaning of a_label and b_label to determine if:
344   - A depends on B edge from b_id to a_id
345   - B depends on A edge from a_id to b_id
346   - No clear dependency no edge
347
348 Output format:
349 Return a JSON object like:
350
351 ```json
352 {{
353   "edge": true or false,
354   "from": "source_id" or "NA",
355   "to": "target_id" or "NA",
356   "reason": "Brief explanation of the dependency or lack thereof"
357 }}
358 ```
359
360 - If there is a dependency, set edge: true, from as the prerequisite's ID, and to as
361 the dependent's ID.
362 - If there is no clear prerequisite relationship, set edge: false and "from": "NA",
363 "to": "NA" with a brief justification in reason.
364
365 Only base your answer on the textual meaning of the labels, and only report direct
366 dependencies (not transitive or indirect ones).
367

```

368 User prompt for Edge prediction

```

369 Given the following skill indicators:
370 - a_label: {label_1}
371 - a_id: {id_1}
372 - b_label: {label_2}
373 - b_id: {id_2}
374
375

```

```

376 Determine the dependency relationship and output the JSON:
377
378 ```json
379 {{
380   "edge": true or false,
381   "from": "source_id" or "NA",
382   "to": "target_id" or "NA",
383   "reason": "Brief explanation of the dependency or lack thereof"
384 }}
385 ```
386

```

387 E.2 Edge weight prediction

388 System prompt:

```

389 You are an expert in child development, skill acquisition, and cognitive science.
390 Your task is to rate the strength of a prerequisite relationship between two
391 skill indicators. Each input includes:
392   - from_label and to_label: the skill indicators (already determined to be in a
393     prerequisite relationship, where from_label is a prerequisite for to_label)
394   - Additional metadata: age groups, subskills, goals, developmental stages, and a
395     rationale for why the edge exists.
396
397 Instructions:
398 Rate the dependency strength on a scale from 1 to 5, where:
399   - 1 = Very weak dependency (minimal or contextual support, can often be developed
400     independently)
401   - 2 = Weak dependency (some support role, but not always required)
402   - 3 = Moderate dependency (often occurs first, but not strictly necessary)
403   - 4 = Strong dependency (usually needed before progressing)
404   - 5 = Very strong dependency (essential foundational step for the next)
405
406 Your response should consider:
407   1. The specific behaviors or understandings described in the two indicators.
408   2. Whether the earlier skill is conceptually or procedurally required to perform the
409     later one.
410   3. The closeness of developmental stages and subskills.
411
412 Output Format:
413 Return your decision as a JSON object:
414 ```json
415 {{
416   "weight": [an integer from 1 to 5],
417   "reason": "[a brief explanation of why this weight reflects the strength of the
418     dependency]"
419 }}
420 ```
421
422

```

423 User prompt:

```

424 Given the following information about a prerequisite relationship between two skill
425 indicators:
426
427   - from_label: {from_label}
428   - from_id: {from_id}
429     - age group: {from_age_group}
430     - skill: {from_skill}
431     - subskill: {from_subskill}
432     - goal: {from_goal}
433     - stage: {from_stage}
434
435 -----
436
437   - to_label: {to_label}
438

```

```

439 - to_id: {to_id}
440   - age_group: {to_age_group}
441   - skill: {to_skill}
442   - subskill: {to_subskill}
443   - goal: {to_goal}
444   - stage: {to_stage}
445
446 This relationship has already been labeled as a prerequisite edge (from_id to_id).
447
448 Rationale for this dependency:
449 "{reason}"
450
451 Rate the strength of this dependency on a scale from 1 to 5.
452
453 Output a JSON object:
454 ```json
455 {{
456   "weight": [an integer from 1 to 5],
457   "reason": "Brief explanation of why this weight reflects the strength of the
458             dependency"
459 }}
460 ```

```

462 E.3 Templates

463 System prompt for generating templates for IR data:

```

464 You are an expert in child development, skill acquisition, curriculum design, and
465 language model pretraining. Your task is to identify developmentally
466 appropriate and general non-instructional text types for synthetic
467 pretraining of a language model.
468
469 Each input includes:
470 - indicator: a natural language description of the learning objective or task
471 - age_group: developmental age (e.g., 05, 511, 1114)
472 - skill: broad academic or developmental domain (e.g., Mathematics, English,
473   Scientific Reasoning)
474 - subskill: a specific subdomain or area of focus (e.g., Listening, Measurement,
475   Problem-solving)
476 - goal: the purpose or nature of the learning (e.g., Application, Reflection,
477   Evaluation)
478 - stage: the curriculum stage (0 to 9, loosely corresponding to increasing age and
479   complexity)
480
481 Instructions:
482 Return a list of general non-instructional text types that:
483 - Are suitable for the learner's developmental stage
484 - Reflect naturalistic or structured formats that don't rely on explicit
485   instructionresponse pairs
486 - Can be used as abstract templates to generate content across many topics
487 - Are defined at a high level of abstraction (e.g., "peer dialogue", "narrative
488   description", "cause-effect explanation")
489
490 **CRITICALLY IMPORTANT**:
491 - Provide format categories, NOT specific content or scenarios
492 - Text types should be 2-5 words that describe a general format, not complete
493   sentences
494 - Each text type should be usable with ANY topic relevant to the age/skill
495   combination
496
497 **Examples of appropriate non-instructional text types**:
498 - "Narrative story with characters"
499 - "Peer conversation transcript"
500 - "Process description passage"
501

```



```

502 - "Personal reflection monologue"
503
504 **Examples of inappropriate text types** (too specific):
505 - "Story about a child going to the zoo"
506 - "Conversation between friends about toys"
507 - "Description of a butterfly's life cycle"
508
509 Output Format:
510 Return your result as a JSON object with the following structure:
511
512 ```json
513 {{
514   "text_types": ["...", "...", "..."]
515 }}
516 ```
517
518 Ensure the list is:
519 - 1520 items long
520 - Abstract enough to work across many topics
521 - Varied across narration, description, interaction, emotion, reasoning
522 - Appropriate in complexity for the given age group and learning goal
523
524 Only output the JSON object.
525

```

526 User prompt for generating templates for IR data:

```

527
528 Given the following information about a learning objective, return a list of general
529   , reusable non-instructional text formats that can serve as templates for
530   synthetic training data:
531
532 - indicator: {indicator}
533 - age_group: {age_group}
534 - skill: {skill}
535 - subskill: {subskill}
536 - goal: {goal}
537 - stage: {stage}
538
539 IMPORTANT: Provide ABSTRACT FORMAT CATEGORIES (2-5 words each), not specific content
540   or scenarios.
541
542 Examples of good non-instructional formats:
543 - "Peer dialogue transcript"
544 - "Sequential process description"
545 - "Character-driven narrative"
546 - "Emotional experience monologue"
547
548 Examples of unsuitable formats (too specific):
549 - "Conversation between friends about toys"
550 - "Description of a butterfly's life cycle"
551 - "Story about going to the beach"
552
553 Ensure your list contains:
554 - 15 to 20 developmentally appropriate text formats
555 - General templates that can be combined with ANY relevant topic
556 - Varied format types that don't rely on explicit instruction-response pairs
557
558 Return only a JSON object in the following format:
559
560 ```json
561 {{
562   "text_types": ["...", "...", "..."]
563 }}
564 ```
565

```

566 System prompt for generating templates for Context data:

```
567 You are an expert in child development, skill acquisition, curriculum design, and
568 language model pretraining. Your task is to identify developmentally
569 appropriate and general **instruction-response text types** for synthetic
570 pretraining of a language model.
571
572 Each input includes:
573 - indicator: a natural language description of the learning objective or task
574 - age_group: developmental age (e.g., 05, 511, 1114)
575 - skill: broad academic or developmental domain (e.g., Mathematics, English,
576 Scientific Reasoning)
577 - subskill: a specific subdomain or area of focus (e.g., Listening, Measurement,
578 Problem-solving)
579 - goal: the purpose or nature of the learning (e.g., Application, Reflection,
580 Evaluation)
581 - stage: the curriculum stage (0 to 9, loosely corresponding to increasing age and
582 complexity)
583
584 Instructions:
585 Return a list of **general instruction-response style text types** that:
586 - Are suitable for the learner's developmental stage
587 - Can be used in instruction tuning and task-based language modeling
588 - Involve a clearly defined instruction format that can be applied across many
589 topics
590 - Are defined at a high level of abstraction (e.g., "explain why X occurs", "compare
591 and contrast X and Y")
592
593 **CRITICALLY IMPORTANT**:
594 - Provide abstract instruction formats, NOT specific prompts or questions
595 - Text types should be 2-5 words describing a general instruction format
596 - Each text type should be usable with ANY topic relevant to the age/skill
597 combination
598
599 **Examples of appropriate instruction-response text types**:
600 - "Compare and contrast analysis"
601 - "Explain why reasoning"
602 - "Step-by-step instruction"
603 - "Open-ended reflection prompt"
604
605 **Examples of inappropriate text types** (too specific):
606 - "Explain why plants need water"
607 - "Compare dogs and cats"
608 - "Describe your favorite toy"
609
610 Output Format:
611 Return your result as a JSON object with the following structure:
612
613 ```json
614 {
615   "text_types": ["...", "...", "..."]
616 }
617 ```
618
619 Ensure the list is:
620 - 1520 items long
621 - Abstract enough to work across many topics
622 - Varied across explanation, reasoning, reflection, comparison, instruction,
623 imagination
624 - Appropriate in complexity for the given age group and learning goal
625
626 Only output the JSON object.
627
628
```

629 User prompt for generating templates for Context data:

```

630
631 Given the following information about a learning objective, return a list of general
632 , reusable instruction-response text formats that can serve as templates for
633 synthetic training data:
634
635 - indicator: {indicator}
636 - age_group: {age_group}
637 - skill: {skill}
638 - subskill: {subskill}
639 - goal: {goal}
640 - stage: {stage}
641
642 IMPORTANT: Provide ABSTRACT INSTRUCTION FORMATS (2-5 words each), not specific
643 questions or prompts.
644
645 Examples of good instruction formats:
646 - "Compare and contrast analysis"
647 - "Explain why reasoning"
648 - "Problem-solving walkthrough"
649 - "Open-ended reflection prompt"
650
651 Examples of unsuitable formats (too specific):
652 - "Explain why plants need water"
653 - "Compare dogs and cats"
654 - "Solve this math problem"
655
656 Ensure your list contains:
657 - 15 to 20 developmentally appropriate instruction formats
658 - General templates that can be combined with ANY relevant topic
659 - Varied instruction types that address different cognitive processes
660
661 Return only a JSON object in the following format:
662
663 ```json
664 {{
665   "text_types": ["...", "...", "..."]
666 }}
667 ```
668

```

669 E.4 Context

670 System prompt for generating context data:

```

671
672 You are an AI model generating training data to help language models simulate human
673 developmental skills at various stages from early childhood through early
674 adolescence.
675
676 Your task is to create engaging, developmentally appropriate texts based on provided
677 developmental indicators, skills, and a tuple of word and its part of speech.
678
679 Strictly follow these guidelines:
680
681 1. Developmental Appropriateness:
682   - Stage 0 (Age 5): Use simple sentences, concrete concepts, familiar experiences,
683     present tense focus
684   - Stages 1-3 (Ages 6-8): Introduce basic past/future concepts, simple cause-
685     effect, familiar settings
686   - Stages 4-6 (Ages 9-11): Include more complex reasoning, abstract thinking,
687     varied sentence structures
688   - Stages 7-9 (Ages 12-14): Incorporate hypothetical scenarios, multiple
689     perspectives, sophisticated vocabulary
690
691 2. Context Generation:

```

```

692 - Use the provided word and its part of speech to create a meaningful,
693     developmentally appropriate topic
694 - **Ensure the selected word and expanded topic fit the required Text Type
695     Template (context_template)**
696 - Expand the selected word into a more detailed, skill-aligned topic that
697     resonates with the target age group
698 - Generate a rich, complete, and engaging text matching the provided context
699     template
700 - The generated text must be **between 250 and 500 words regardless of
701     developmental stage**
702 - The text must clearly align with the skill, subskill, goal, and indicator
703 - The selected word does not need to explicitly appear in the final text
704
705 3. **Writing Style by Stage:**
706 - **Early Stages (0-3):** Simple vocabulary, short to medium sentences, concrete
707     experiences, repetitive patterns for reinforcement
708 - **Middle Stages (4-6):** More varied vocabulary, complex sentences,
709     introduction of abstract concepts, problem-solving scenarios
710 - **Later Stages (7-9):** Sophisticated vocabulary, complex sentence structures,
711     abstract reasoning, multiple viewpoints
712
713 4. **Content Enrichment:**
714 - Include age-appropriate actions, feelings, interactions, and sensory details
715 - Incorporate social situations relevant to the developmental stage
716 - Use scenarios that promote the specific skill being targeted
717 - Avoid overly abstract or culturally specific references unless appropriate for
718     the age group
719
720 5. **Output Format:** Strictly return the output in the following JSON structure:
721 ```json
722 {
723     "expanded_topic": "<expanded topic>",
724     "generated_text": "<generated text between 250 and 500 words>"
725 }
726 ```
727 Only output the JSON. No additional commentary.

```

729 User prompt for generating context data:

```

730 Generate a rich and engaging context text based on the following input:
731
732 - ID: {id}
733 - Indicator: {indicator}
734 - Skill: {skill}
735 - Sub-skill: {subskill}
736 - Goal: {goal}
737 - Age Group: {age_group}
738 - Stage: {stage}
739 - Text Type Template: {context_template}
740 - (Word, Part of speech): {word_list}
741
742 Instructions:
743 - Consider the developmental stage ({stage}) and age group ({age_group}) when
744     crafting vocabulary, sentence complexity, and content themes
745 - Expand the selected word into a skill-relevant topic **that fits the Text Type
746     Template**
747 - Generate a detailed text of **250-500 words** following the context template
748 - Enrich the text with developmentally appropriate actions, emotions, and
749     interactions
750 - Ensure the content promotes the specific skill and subskill being targeted
751
752 Output strictly in this format:
753 ```json
754 {
755     "expanded_topic": "<expanded topic>",
756

```

```

757 "generated_text": "<generated text between 250 and 500 words>"
758 }}
759 '''
760

```

761 E.5 CQA

762 System prompt for generating CQA data:

```

763 You are an AI model generating training data to help language models simulate human
764 reading comprehension skills at various stages from early childhood through
765 early adolescence.
766
767 Your task is to create 5 developmentally appropriate question-answer pairs based on
768 a provided text, ensuring all questions test understanding of the given
769 paragraph and can be answered directly from the text.
770
771 Strictly follow these guidelines:
772
773 1. Developmental Appropriateness by Stage:
774   - Stage 0 (Age 5): Simple "what/who/where" questions, literal comprehension,
775     single-step reasoning
776   - Stages 1-3 (Ages 6-8): Basic "why/how" questions, simple cause-effect, sequence
777     understanding, character feelings
778   - Stages 4-6 (Ages 9-11): Inference questions, comparing/contrasting, predicting
779     outcomes, understanding motivations
780   - Stages 7-9 (Ages 12-14): Complex analysis, multiple perspectives, abstract
781     concepts, theme identification
782
783 2. Question Creation Standards:
784   - All answers must be directly supported by information in the provided text
785   - No questions requiring outside knowledge or information not present in the text
786   - Questions should test different types of comprehension appropriate to the
787     developmental stage
788   - Vary question types to assess different reading skills (literal, inferential,
789     evaluative)
790   - Use vocabulary and sentence complexity appropriate to the age group
791   - Ensure questions are engaging and relevant to the child's interests and
792     experiences
793
794 3. Question Types by Stage:
795   - Early Stages (0-3): Literal recall, identifying main characters/objects,
796     simple sequence, basic emotions
797   - Middle Stages (4-6): Cause-effect relationships, character motivations,
798     comparing details, simple predictions
799   - Later Stages (7-9): Drawing conclusions, analyzing relationships,
800     evaluating actions, understanding themes
801
802 4. Answer Generation:
803   - Create authentic child responses that demonstrate comprehension at the target
804     developmental stage
805   - Use vocabulary and sentence structures appropriate to the age group
806   - Include natural speech patterns and expressions typical of the developmental
807     stage
808   - Ensure answers are complete but not overly elaborate for the age group
809   - Answers should sound conversational and natural, not textbook-like
810
811 5. Content Guidelines:
812   - Purely verbal exchanges - no references to physical gestures or non-verbal
813     actions
814   - No formatting (bold, italics, markdown)
815   - Questions should flow naturally and cover different aspects of the text
816   - Ensure logical progression from simpler to more complex questions when
817     appropriate
818

```

```

819 - Include a mix of question types (factual, inferential, personal connection when
820 text-supported)
821
822 6. Quality Standards:
823 - Every question must be answerable using only information provided in the text
824 - Questions should test genuine comprehension, not just memory of isolated facts
825 - Avoid questions with obvious or trivial answers
826 - Ensure questions are meaningful and help assess understanding of key text
827 elements
828 - Create questions that feel natural in an educational setting
829
830 7. Output Format: Strictly return the output in the following JSON structure:
831 ```json
832 {{
833   "question_answer_pairs": [
834     {{
835       "question": "<question 1>",
836       "answer": "<answer 1>"
837     }},
838     {{
839       "question": "<question 2>",
840       "answer": "<answer 2>"
841     }},
842     {{
843       "question": "<question 3>",
844       "answer": "<answer 3>"
845     }},
846     {{
847       "question": "<question 4>",
848       "answer": "<answer 4>"
849     }},
850     {{
851       "question": "<question 5>",
852       "answer": "<answer 5>"
853     }}
854   ]
855 }}
856 ```
857 Only output the JSON. No additional commentary or explanations.

```

859 User prompt for generating CQA data:

```

860 Generate 5 developmentally appropriate reading comprehension question-answer pairs
861 based on the following input:
862
863 - Text: {output}
864 - Age Group: {age_group}
865 - Stage: {stage}
866
867 Instructions:
868 - Consider the developmental stage ({stage}) and age group ({age_group}) when
869 crafting question complexity and answer expectations
870 - Create questions that test different types of comprehension appropriate to the
871 developmental level
872 - Ensure all questions can be answered directly from the provided text
873 - Generate authentic child responses that demonstrate comprehension at the target
874 stage
875 - Use vocabulary and sentence structures appropriate to the age group
876 - Create a mix of question types that genuinely assess understanding of the text
877
878 Output strictly in this format:
879 ```json
880 {{
881   "question_answer_pairs": [
882     {{

```

```

884         "question": "<question 1>",
885         "answer": "<answer 1>"
886     }},
887     {{
888         "question": "<question 2>",
889         "answer": "<answer 2>"
890     }},
891     {{
892         "question": "<question 3>",
893         "answer": "<answer 3>"
894     }},
895     {{
896         "question": "<question 4>",
897         "answer": "<answer 4>"
898     }},
899     {{
900         "question": "<question 5>",
901         "answer": "<answer 5>"
902     }}
903 ]
904 }}
905 '''
906

```

907 E.6 CSQA

908 System prompt for generating CSQA data:

```

909 You are an AI model generating training data to help language models simulate human
910 developmental skills at various stages from early childhood through early
911 adolescence.
912
913 Your task is to create 3 skill-based instruction-response pairs between an educator
914 and a child that use a provided text as context to test specific developmental
915 skills, rather than simple reading comprehension.
916
917 Strictly follow these guidelines:
918
919 1. Developmental Appropriateness by Stage:
920   - Stage 0 (Age 5): Simple vocabulary, short sentences, concrete thinking, present
921     -focused, immediate experiences
922   - Stages 1-3 (Ages 6-8): Basic past/future concepts, simple reasoning, familiar
923     contexts, beginning abstract thought
924   - Stages 4-6 (Ages 9-11): Complex reasoning, abstract thinking, varied sentence
925     structures, hypothetical scenarios
926   - Stages 7-9 (Ages 12-14): Sophisticated vocabulary, multiple perspectives,
927     advanced abstract reasoning, nuanced responses
928
929 2. Skill-Based Instruction Creation:
930   - Use the provided text as context, not as the primary focus
931   - Create instructions that test the specific skill, subskill, goal, and indicator
932     provided
933   - Instructions should prompt the child to demonstrate the target skill using
934     elements from the text
935   - Avoid simple recall questions - focus on skill application, analysis, synthesis
936     , or evaluation
937   - Vary instruction starters - avoid overusing "Imagine..." or "Tell me about..."
938   - Include necessary context within the instruction if recall is required
939   - Use developmentally appropriate language and concepts for the target stage
940   - Make instructions engaging and thought-provoking for the age group
941
942 3. Response Generation:
943   - Create authentic child responses that clearly demonstrate the target indicator
944   - Use vocabulary, sentence complexity, and reasoning appropriate to the
945     developmental stage
946

```

```

947 - Include natural speech patterns and expressions typical of the age group
948 - Ensure responses show genuine skill application, not just text recall
949 - Responses should be verifiable through either:
950   * Information provided in the instruction or text
951   * Common world knowledge appropriate for the child's developmental level
952   * Typical personal experiences for that age group
953 - Avoid arbitrary claims or purely imaginative details unless the skill
954   explicitly encourages creativity
955
956 4. **Context Integration:**
957   - Use the provided text as a springboard for skill demonstration
958   - Connect text elements to real-world applications of the skill
959   - Encourage children to apply their skills to analyze, extend, or relate to the
960     text content
961   - Ensure the skill being tested is meaningfully connected to the text context
962
963 5. **Content Guidelines:**
964   - **Purely verbal exchanges** - no references to physical objects, gestures, or
965     non-verbal actions
966   - No formatting (bold, italics, markdown)
967   - Instructions should feel natural and appropriate for educational settings
968   - Responses should sound natural and spontaneous, not rehearsed
969   - Include appropriate emotional expressions and personal connections when
970     relevant
971   - Ensure logical consistency between instruction and response
972   - Focus on the skill demonstration rather than text comprehension
973
974 6. **Quality Standards:**
975   - The exchange must demonstrate clear alignment with the skill, subskill, goal,
976     and indicator
977   - Each instruction must clearly target the specific developmental parameters
978     provided
979   - Instructions should be distinct from each other, testing different aspects of
980     the same skill
981   - Both instruction and response should feel authentic to a real classroom or
982     learning interaction
983   - Responses must demonstrate clear mastery or development of the target skill
984   - The text should serve as meaningful context, not just background information
985   - Avoid overly abstract concepts for younger stages or overly simple concepts for
986     older stages
987   - Ensure developmental appropriateness in both challenge level and expectations
988
989 7. **Output Format:** Strictly return the output in the following JSON structure:
990   ```json
991   {{
992     "skill_based_pairs": [
993       {{
994         "instruction": "<instruction 1>",
995         "response": "<response 1>"
996       }},
997       {{
998         "instruction": "<instruction 2>",
999         "response": "<response 2>"
1000     }},
1001     {{
1002       "instruction": "<instruction 3>",
1003       "response": "<response 3>"
1004     }}
1005   ]
1006 }}
1007   ```
1008 Only output the JSON. No additional commentary or explanations.
1009

```

1010 User prompt for generating CSQA data:


```

1011
1012 Generate 3 developmentally appropriate skill-based instruction-response pairs based
1013     on the following input:
1014
1015 - Text: {output}
1016 - Age Group: {age_group}
1017 - Stage: {stage}
1018 - Skill: {skill}
1019 - Sub-skill: {subskill}
1020 - Goal: {goal}
1021 - Indicator: {indicator}
1022
1023 Instructions:
1024 - Consider the developmental stage ({stage}) and age group ({age_group}) when
1025     crafting instruction complexity and response expectations
1026 - Use the provided text as context to create instructions that test the specific
1027     skill ({skill}) and subskill ({subskill})
1028 - Create instructions that elicit demonstration of the goal ({goal}) and indicator
1029     ({indicator})
1030 - Focus on skill application and demonstration, not text comprehension
1031 - Generate authentic child responses that show clear mastery of the target skill at
1032     the developmental stage
1033 - Use vocabulary and sentence structures appropriate to the age group
1034 - Create 3 distinct instructions that test different aspects of the same skill
1035
1036 Output strictly in this format:
1037 ```json
1038 {{
1039     "skill_based_pairs": [
1040         {{
1041             "instruction": "<instruction 1>",
1042             "response": "<response 1>"
1043         }},
1044         {{
1045             "instruction": "<instruction 2>",
1046             "response": "<response 2>"
1047         }},
1048         {{
1049             "instruction": "<instruction 3>",
1050             "response": "<response 3>"
1051         }}
1052     ]
1053 }}
1054 ```
1055

```

1056 E.7 IR

1057 System prompt for generating IR data:

```

1058 You are an AI model generating training data to help language models simulate human
1059     developmental skills at various stages from early childhood through early
1060     adolescence.
1061
1062 Your task is to create realistic instruction-response pairs between an educator and
1063     a child, based on developmental indicators, skills, and a tuple of word and its
1064     part of speech.
1065
1066 Strictly follow these guidelines:
1067
1068 1. Developmental Appropriateness by Stage:
1069     - Stage 0 (Age 5): Simple vocabulary, short sentences, concrete thinking, present
1070       -focused, immediate experiences
1071     - Stages 1-3 (Ages 6-8): Basic past/future concepts, simple reasoning, familiar
1072       contexts, beginning abstract thought
1073

```

```

1074 - Stages 4-6 (Ages 9-11): Complex reasoning, abstract thinking, varied sentence
1075 structures, hypothetical scenarios
1076 - Stages 7-9 (Ages 12-14): Sophisticated vocabulary, multiple perspectives,
1077 advanced abstract reasoning, nuanced responses
1078
1079 2. **Instruction Creation:**
1080 - Use the provided word and its part of speech to meaningfully inspire the
1081 interaction topic
1082 - **Ensure the topic aligns with the Text Type Template (instruct_template)**
1083 - Craft prompts that naturally elicit demonstration of the specific indicator and
1084 skill
1085 - Vary instruction starters - avoid overusing "Imagine..." or "Tell me about..."
1086 - Include necessary context within the instruction if recall is required
1087 - Use developmentally appropriate language and concepts for the target stage
1088 - Make instructions engaging and thought-provoking for the age group
1089
1090 3. **Response Generation:**
1091 - Create authentic child responses that clearly demonstrate the target indicator
1092 - Use vocabulary, sentence complexity, and reasoning appropriate to the
1093 developmental stage
1094 - Include natural speech patterns and expressions typical of the age group
1095 - Ensure responses are verifiable through either:
1096 * Information provided in the instruction
1097 * Common world knowledge appropriate for the child's developmental level
1098 * Typical personal experiences for that age group
1099 - Avoid arbitrary claims or purely imaginative details unless storytelling is
1100 explicitly encouraged
1101
1102 4. **Content Guidelines:**
1103 - **Purely verbal exchanges** - no references to physical objects, gestures, or
1104 non-verbal actions
1105 - No formatting (bold, italics, markdown)
1106 - Responses should sound natural and spontaneous, not rehearsed
1107 - Include appropriate emotional expressions and personal connections when
1108 relevant
1109 - Ensure logical consistency between instruction and response
1110
1111 5. **Quality Standards:**
1112 - The exchange must demonstrate clear alignment with the skill, subskill, goal,
1113 and indicator
1114 - Both instruction and response should feel authentic to a real classroom or
1115 learning interaction
1116 - Avoid overly abstract concepts for younger stages or overly simple concepts for
1117 older stages
1118 - Ensure the selected word meaningfully influences the dialogue topic
1119
1120 6. **Output Format:** Strictly return the output in the following JSON structure:
1121 ```json
1122 {{
1123     "instruction": "<instruction>",
1124     "response": "<response>"
1125 }}
1126 ```
1127 Only output the JSON. No additional commentary or explanations.

```

1129 User prompt for generating IR data:

```

1130
1131 Generate a developmentally appropriate instruction-response pair based on the
1132 following input:
1133
1134 - ID: {id}
1135 - Indicator: {indicator}
1136 - Skill: {skill}
1137 - Sub-skill: {subskill}
1138 - Goal: {goal}

```

```

1139 - Age Group: {age_group}
1140 - Stage: {stage}
1141 - Text Type Template: {instruct_template}
1142 - (Word, Part of speech): {word_list}
1143
1144 Instructions:
1145 - Consider the developmental stage ({stage}) and age group ({age_group}) when
1146   crafting language complexity and content themes
1147 - Use the selected word to meaningfully inspire the interaction topic that fits
1148   the Text Type Template
1149 - Create an engaging instruction that naturally elicits demonstration of the target
1150   indicator
1151 - Generate an authentic child response that clearly shows mastery of the skill and
1152   subskill
1153 - Ensure the exchange feels natural and appropriate for a real educational
1154   interaction
1155
1156 Output strictly in this format:
1157 ```json
1158 {{
1159   "instruction": "<instruction>",
1160   "response": "<response>"
1161 }}
1162 ```
1163

```

1164 E.8 Evaluating CQA

1165 System prompt for evaluating trained model's response for questions from CQA:

```

1166 You are a developmental expert evaluating how well a child's answer to a reading
1167 comprehension question reflects appropriate understanding and reasoning for a
1168 specific developmental stage.
1169
1170 You will receive:
1171 - The original context paragraph
1172 - A question based on the context
1173 - The child's answer to the question
1174 - The child's developmental stage (09)
1175 - The child's age group (e.g., '05', '511', '1114')
1176
1177 Your job is to:
1178 1. Rate the child's answer on a scale from 1 to 5, using the following criteria:
1179   - 5 Excellent: Fully correct, precise, and well-formed for the stage. Shows
1180     strong comprehension and reasoning.
1181   - 4 Strong: Mostly correct and appropriate; may have minor phrasing issues
1182     or slight gaps in reasoning.
1183   - 3 Adequate: Understands the gist but may be vague, partially incorrect, or
1184     simplistic for the stage.
1185   - 2 Limited: Misunderstands part of the question or context; reasoning is
1186     weak or off-track.
1187   - 1 Inadequate: Confused, incorrect, or clearly not appropriate for the
1188     stage.
1189
1190 2. Consider developmental expectations for language and reasoning:
1191   - Stage 0 (Age 5): Very basic phrases, literal recall, present-focused
1192     answers
1193   - Stages 13 (Ages 68): Simple reasoning, sequencing, basic cause-effect,
1194     clear answers
1195   - Stages 46 (Ages 911): Logical inference, comparative language, clear
1196     justification
1197   - Stages 79 (Ages 1214): Abstract reasoning, complex ideas, nuanced
1198     explanations
1199
1200 3. Evaluate:
1201

```

```

1202 - Does the child's answer meaningfully address the question using the provided
1203 context?
1204 - Is the reasoning and language appropriate for the stage?
1205 - Does it reflect comprehension of the text and question?
1206
1207 4. Output Format:
1208 Only return the following dictionary:
1209 ```json
1210 {
1211     "rating": <integer from 1 to 5>,
1212     "explanation": "<23 sentence rationale>"
1213 }
1214 ```
1215 Do not add any other text or formatting. Only return the JSON object.

```

1217 User prompt for evaluating trained model's response for questions from CQA:

```

1218 Evaluate the child's answer to a reading comprehension question. Consider the context
1219 and the developmental stage.
1220
1221 Context:
1222 {context}
1223
1224 Question:
1225 {question}
1226
1227 Answer:
1228 {answer}
1229
1230 Stage: {stage}
1231 Age group: {age_group}
1232 Index: {q_index}
1233
1234 Output Format:
1235 ```json
1236 {
1237     "rating": <integer from 1 to 5>,
1238     "explanation": "<23 sentence rationale>"
1239 }
1240 ```
1241

```

1243 E.9 Evaluating CSQA

1244 System prompt for evaluating trained model's response for questions from CSQA:

```

1245 You are a developmental expert evaluating how well a child's response demonstrates a
1246 specific developmental skill at a given stage, using a provided instruction
1247 and background text.
1248
1249 You will receive:
1250 - A short text (used as context for the instruction)
1251 - A skill-based instruction given to the child
1252 - The child's response
1253 - The child's developmental stage (09)
1254 - The child's age group (e.g., '05', '511', '1114')
1255 - The target skill, subskill, goal, and indicator that the
1256 instruction was designed to assess
1257
1258 Your job is to:
1259 1. Rate the child's response on a scale from 1 to 5, using these criteria:
1260 - 5 Excellent: Fully demonstrates the targeted skill/indicator with clarity
1261 and developmental appropriateness. Strong reasoning, appropriate expression,
1262 and alignment with instruction.
1263

```

```

1264 - **4 Strong:** Mostly appropriate and well-formed. Some minor gaps in
1265     completeness, precision, or phrasing, but shows the intended skill.
1266 - **3 Adequate:** Response attempts the skill but may be vague, simplistic, or
1267     only partially aligned with the goal/indicator.
1268 - **2 Limited:** Weak or unclear demonstration of the skill. Response is
1269     partially off-track, underdeveloped, or barely relevant.
1270 - **1 Inadequate:** Fails to demonstrate the intended skill. Response is
1271     irrelevant, confusing, or clearly inappropriate for the stage.
1272
1273 2. **Use stage-specific developmental expectations**:
1274     - **Stage 0 (Age 5):** Short, concrete, present-focused responses with simple
1275       vocabulary
1276     - **Stages 13 (Ages 68):** Clear expression of ideas, simple cause-effect,
1277       emotional awareness, basic reasoning
1278     - **Stages 46 (Ages 911):** Logical structure, hypothetical thinking, connections
1279       to personal experience, comparisons
1280     - **Stages 79 (Ages 1214):** Advanced abstraction, multiple perspectives,
1281       justification, nuanced expression
1282
1283 3. **Evaluate**:
1284     - Does the child's response meaningfully follow the instruction?
1285     - Does it demonstrate the **targeted skill and indicator**?
1286     - Is the language, reasoning, and expression developmentally appropriate for the
1287       stage?
1288     - Is the response authentic and logically consistent with the instruction and the
1289       context text?
1290
1291 4. **Output Format**:
1292     Return only the following dictionary:
1293     ```json
1294     {{
1295         "rating": <integer from 1 to 5>,
1296         "explanation": "<23 sentence rationale>"
1297     }}
1298     ```
1299     Do not add any other text or formatting. Only return the JSON object.
1300

```

1301 User prompt for evaluating trained model's response for questions from CSQA:

```

1302 Evaluate the child's response to a skill-based instruction using the provided text
1303     and developmental context. Focus on how well the response demonstrates the
1304     intended skill.
1305
1306 Context:
1307 {context}
1308
1309 Instruction:
1310 {instruction}
1311
1312 Response:
1313 {response}
1314
1315 Stage: {stage}
1316 Age group: {age_group}
1317 Skill: {skill}
1318 Subskill: {subskill}
1319 Goal: {goal}
1320 Indicator: {indicator}
1321 Index: {q_index}
1322
1323 Output format:
1324 ```json
1325 {{
1326     "rating": <integer from 1 to 5>,
1327     "explanation": "<23 sentence rationale>"
1328

```

```
1329 }}
1330 '''
1331
```

1332 E.10 Evaluating IR

1333 System prompt for evaluating trained model's response for questions from IR:

```
1334 You are a developmental expert rating how well a child's response to a prompt
1335 demonstrates age-appropriate reasoning and language for a given developmental
1336 stage.
1337
1338 You will receive:
1339 - An instruction given to the child
1340 - The child's response
1341 - The child's developmental stage (09)
1342 - The child's age group (e.g., '05', '511', '1114')
1343
1344 Your job is to:
1345 1. Rate the response on a scale from 1 to 5, using the following criteria:
1346   - 5 Excellent: The response fully addresses the instruction with clear,
1347     developmentally appropriate reasoning and language. It meets expectations for
1348     the stage with no major issues.
1349   - 4 Strong: Mostly appropriate and coherent; minor gaps in clarity, depth,
1350     or completeness.
1351   - 3 Adequate: A reasonable attempt that partially addresses the instruction;
1352     may be vague, brief, or contain small misunderstandings.
1353   - 2 Limited: Weak or underdeveloped response; minimal reasoning or limited
1354     relevance to the instruction.
1355   - 1 Inadequate: Response is off-topic, confusing, or clearly inappropriate
1356     for the stage.
1357
1358 2. Use stage-specific developmental expectations:
1359   - Stage 0 (Age 5): Very simple sentences, concrete ideas, focused on here and
1360     now
1361   - Stages 13 (Ages 68): Simple reasoning, some past/future thinking, familiar
1362     examples
1363   - Stages 46 (Ages 911): Logical structure, comparisons, abstract or
1364     hypothetical reasoning
1365   - Stages 79 (Ages 1214): Nuanced reasoning, multi-step thinking, advanced
1366     vocabulary
1367
1368 3. Evaluate:
1369   - Does the child's response meaningfully address the instruction?
1370   - Is the language and reasoning developmentally appropriate for the stage?
1371   - Is the response authentic and logically consistent?
1372
1373 4. Output Format:
1374 Only return the following dictionary:
1375 '''json
1376 {{
1377   "rating": <integer from 1 to 5>,
1378   "explanation": "<23 sentence rationale>"
1379 }}
1380 '''
1381 Do not add any other text or formatting. Only return the JSON object.
1382
1383
```

1384 User prompt for evaluating trained model's response for questions from IR:

```
1385 Evaluate the child's response to the instruction below based on the developmental
1386 stage and age group. Return a numerical rating (15) and a short explanation.
1387
1388 Instruction: {instruction}
1389 Response: {response}
1390 Stage: {stage}
1391
```

```

1392 Age group: {age_group}
1393 Index: {q_index}
1394
1395 **Output Format:**
1396 Only return the following dictionary:
1397 ```json
1398 {{
1399     "rating": <integer from 1 to 5>,
1400     "explanation": "<23 sentence rationale>"
1401 }}
1402 ```
1403

```