
The Impact of Post-training on Data Contamination

Anonymous Author(s)

Affiliation

Address

email

Abstract

1 We present a controlled study of how dataset contamination interacts with the
2 post-training stages in large language models. Starting from clean checkpoints
3 of Qwen2.5 (0.5B/1.5B) and Gemma3 (1B/4B), we inject five copies of GSM8k
4 and MBPP test items into the first 2B tokens of an otherwise 25B token extended
5 pre-training dataset. We then compare the contaminated and clean models both
6 immediately after pre-training and again after supervised fine-tuning (SFT) or
7 reinforcement learning (RL). The post-training steps do not have any contami-
8 nation. Across math and coding benchmarks, we find two consistent patterns:
9 (i) Contamination causes performance spikes that are gradually diminished with
10 continued pre-training. After even 25B tokens, the apparent performance inflation
11 of contamination can become close to zero. (ii) Both SFT and RL resurface the
12 leaked information, but with different patterns: SFT inflates scores only on the
13 contaminated tasks (GSM8k, MBPP), whereas RL also improves performance
14 on uncontaminated counterparts (GSMPlus, HumanEval). Our results underscore
15 the need for contamination audits *after* post-training and suggest that RL-based
16 post-training, although not immune, can help alleviate overestimation problems.

17 1 Introduction

18 Recent work has shown the prevalence of data contamination (Singh et al., 2024; Sainz et al., 2024a).
19 While this overlap is concerning, how much this overlap impacts model performance and affects
20 our evaluations is a separate question. In this regard, most existing contamination analyses focus
21 exclusively on the pre-training stage and the impact of contamination right after pre-training (Kocyigit
22 et al., 2025; Jiang et al., 2024). However, state-of-the-art LLMs are almost always subjected to
23 one or more post-training procedures (Wei et al., 2022; Chung et al., 2022; Ouyang et al., 2022;
24 Zhang et al., 2024b). These procedures can materially reshape the model’s internal representation
25 consequently, contamination that appears dormant or innocuous after the pre-training stage may be
26 amplified, systematically exploited, or conversely attenuated once the model is steered by a different
27 optimization objective.

28 There is growing evidence that the type of post-training schema applied can also impact how much
29 models can generalize. Previous work suggests that SFT is more prone to causing memorization
30 while RL is shown to introduce generalization capabilities not direct memorization (Chu et al., 2025).

31 In this work, we studied this problem by deliberately injecting contamination from well-studied
32 mathematics and coding benchmarks and performing extended pre-training on models of up to 4B
33 parameters, Qwen2.5, 0.5B and 1.5B and Gemma3 1B and 4B. Following the completion of clean and
34 contaminated pre-training, we apply two widely adopted post-training paradigms, SFT and RL, on the
35 corresponding training splits and quantify how contamination influences downstream performance by
36 comparing contaminated models to contamination-free baselines.

37 With this experimental setup, we aim to answer the following questions: Does post-training alleviate
38 or intensify the performance overestimation caused by contamination? Do results change depending
39 on the type of post-training method used?

40 2 Related Work

41 Early warnings about evaluation-set leakage in LLMs emphasized that even minimal overlap between
42 training corpora and test datasets can inflate evaluation scores (Singh et al., 2024). Position papers
43 and surveys such as Cheng et al. (2025); Sainz et al. (2024b) catalog a broad range of contamination
44 pathways and call for community norms such as encrypted benchmarks, one-shot test releases, and
45 data audits to preserve the validity of leaderboards (Sainz et al., 2024a).

46 Empirical studies also aim to measure the impact of contamination for pre-training more precisely
47 by injecting controlled contamination into the pre-training mix (Jiang et al., 2024; Kocyigit et al.,
48 2025). These papers show that contamination yields large performance jumps that, more critically,
49 scale with model size (e.g. ~ 30 BLEU on MT).

50 While these studies evaluate contamination after pre-training, users rarely interact with raw pre-trained
51 models because most LLMs undergo post-training before deployment. This makes post-training a
52 relevant point of contact for real-world applications and, consequently, a critical stage for evaluating
53 the effects of contamination.

54 To the best of our knowledge, only Magar & Schwartz (2022) study a similar problem by separat-
55 ing pre-training and fine-tuning stages and introduce two metrics: memorization and exploitation.
56 However, their experiments are limited to small models (BERT-base/large) and standard SFT classifi-
57 cation benchmarks (SST, SNLI), where contamination dynamics may differ significantly from those
58 observed in more complex generative tasks such as mathematics or coding.

59 Nonetheless, no prior work jointly studies contamination across models exceeding one billion
60 parameters, along with variations in post-training methods such as SFT and RL. Our study fills this
61 gap by systematically comparing SFT and RL on contaminated versus clean continuations of the
62 same pre-trained checkpoints, allowing us to disentangle how contamination actually impacts model
63 performance after modern post-training.

64 3 Method

65 Very simply, we train the same model with and without injected contamination and compare their
66 performance after pre-training and post-training via SFT or RL. Below, we detail the components of
67 this experimental setup. A visual overview of our method is given in Figure 5 (Appendix).

68 **Data:** Ideally, full pre-training runs would allow for more realistic experiments. However, due to
69 compute constraints, we ran our experiments as extended pre-training runs. To avoid *overstating*
70 the impact of contamination, we used a relatively large extended pre-training mixture comprising
71 25B tokens, based on the findings of Kocyigit et al. (2025). This mixture includes web text from
72 FineWeb-Edu (Penedo et al., 2024), code data from CodeParrots (von Werra et al., 2021), and
73 mathematical content from OpenMath-Instruct (Toshniwal et al., 2024).

74 **Models:** Our experiments use Qwen2.5(0.5B, 1.5B) (Qwen et al., 2025) and Gemma3 (1B, 4B)
75 (Team et al., 2025) as baseline models. These models were selected based on their demonstrated
76 capabilities in math and coding tasks, as well as the availability of pre-trained checkpoints without any
77 post-training. Since our experimental design involves extended pre-training, access to checkpoints
78 without intermediate fine-tuning steps is crucial to avoid confounding effects.

79 **Post-training and Evaluation:** The SFT step we fine-tune the model on reasoning steps and final-
80 answer tokens, details are shared in Appendix 5. For RL, we use Group Relative Policy Optimization
81 (GRPO) (Shao et al., 2024) with rule-based reward functions, detailed in Appendix 5. We choose
82 GRPO because it is a simple, effective method for math/code and has been shown to help smaller
83 models. Both the SFT and GRPO phases are capped at approximately the same number of update
84 steps to ensure comparability.

85 We evaluate contamination effects using GSM8k (Cobbe et al., 2021) and MBPP (Austin et al., 2021)
86 as the contaminated tasks. We chose these benchmarks for their widespread use and the availability

of a training set. Since we wanted to run post-training, we needed a benchmark that had a disjoint training set as well.

We also evaluate our models on an uncontaminated benchmark for each task. The objective of this is to understand the generalization gap generated by contamination. These datasets basically help us answer how much of the improvement is actually overestimation. For the math benchmark, we chose GSMPlus (Li et al., 2024), for coding, we chose HumanEval (Chen et al., 2021) since it is a high-quality coding benchmark that aims to measure Python programming performance. While the levels of difficulty are not directly comparable, we format HumanEval to mirror MBPP for comparability.

4 Results

In Figure 1, the blue bars show that, for the model families and benchmarks that we consider in our experiments, five copies of the test set shows no measurable and consistent performance inflation after pre-training.

On the other hand, we also present the performance difference between the contaminated and clean models after the specific post-training step, Figure 1 the green and red bars. Here we show that post-training can resurface the measured performance gap. After post-training the performance gap is above 2% for almost all models for both the math and coding benchmarks. The performance gap reaches 4% for the smallest Qwen2.5-0.5B model.

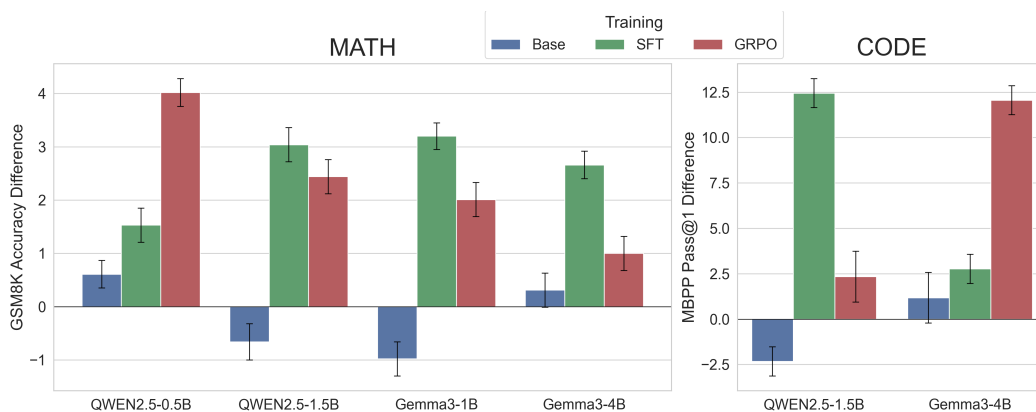


Figure 1: Performance Difference on Math and Code: Accuracy difference between Contaminated and Clean models right after pre-training (base) and after the SFT and GRPO steps. We observe that while the Base differences show little to no impact from contamination post-training can actually uncover the information acquired by the model in pre-training even after additional training seem to have suppress it.

This suggests that while continued pre-training after contamination masks the advantage attained by the contaminated model, the information is not forgotten and can be uncovered with task-specific fine-tuning of the model. We also observe that with a few exceptions, SFT causes a larger performance gap for the contaminated model compared to GRPO. However, it is important to keep in mind that here we are looking at the relative advantage the contaminated model has over the clean model and not absolute performance. Notwithstanding, we can draw the conclusion that SFT uncovers the impact of contamination in the pre-training comparatively better compared to GRPO for most models we experiment with.

4.1 Are performance overestimations actual overestimations?

In this context, it is reasonable to question whether injecting the high-quality test set into the pre-training mixture prompted with the same evaluation prompt could simply improve the model, meaning the gap might not reflect overestimation. To test this, we compare the models' performance on parallel, uncontaminated benchmarks that aim to measure the same underlying ability.

Results reveal another interesting pattern between SFT and GRPO. When comparing the base-blue (just pre-trained) markers with the SFT-green markers, we observe that the average movement is horizontal. This means that while SFT introduces larger performance gaps due on the contaminated benchmark, the impact of contamination on an external benchmark remains constant. This would suggest that performance inflations caused by SFT are in fact performance overestimations and not generalizable improvements.

On the other hand, when comparing the base-blue markers with the GRPO-red markers we observe a diagonal movement, meaning that the performance gap between the contaminated and clean model grow for both the contaminated and the uncontaminated test sets. This suggests that GRPO can extract generalizable improvements from the contamination. We suspect that this is a combination of higher quality data and the usage of the evaluation prompt in the contamination.

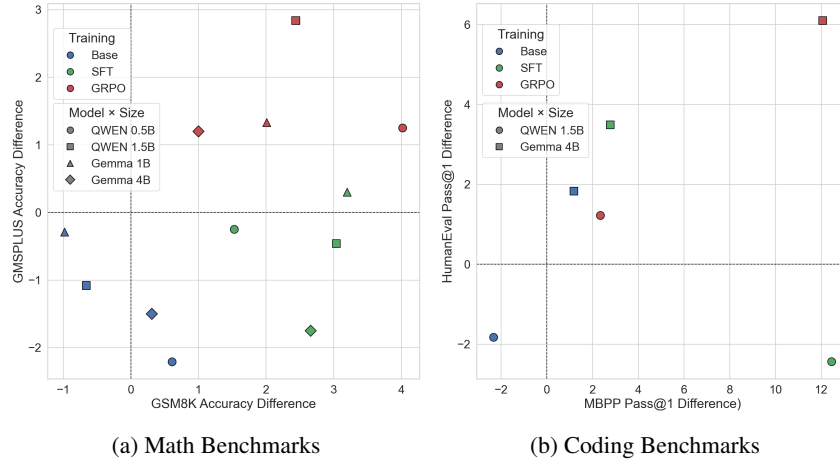


Figure 2: Comparison of Performance gap on contaminated and uncontaminated datasets. We observe that the Base models behave roughly the same on the contaminated and uncontaminated datasets for both Math and Coding. GRPO fine-tuned models have a positive gap on the contaminated dataset but also have a smaller but still positive gap on the uncontaminated dataset, suggesting the models learn some generalizable patterns. The SFT models on the other hand only have a larger gap in the contaminated dataset and show almost no improvement on the uncontaminated tests.

5 Discussion

Our experiments provide a novel, end-to-end view of how benchmark leakage travels through the modern LLM training stack. We experiment with two model families on two task types and four benchmarks. Overall we observe that our findings seems to be mostly consistent across the two model families (Qwen2.5 and Gemma3). Through our experiments two themes emerge:

When you measure matters. Figure 1 shows that the apparent gap between contaminated and clean models almost vanishes once pre-training resumes on clean data, a result consistent with Kocyigit et al. (2025). Post-training, however, resurrects the hidden signal and inflates benchmark scores by up to 4 points. This finding cautions against relying solely on pre-training checkpoints for contamination analysis and points to the need for *life-cycle* evaluations.

Post-training type can affect whether leakage overfits or generalizes. Both SFT and GRPO widen the gap on the contaminated task, yet they diverge sharply on uncontaminated benchmarks (Figure 2). SFT’s gains are almost purely local: the contaminated model answers GSM8k or MBPP questions better compared to the base model, but exhibits no extra competence on GSMPlus or HumanEval. GRPO, in contrast, creates a gap between these two models on both the contaminated and the uncontaminated datasets. This potentially suggests that it learns more generally useful reasoning patterns, partially mitigating the impact of contamination.

References

- Austin, J., Odena, A., Nye, M., Bosma, M., Michalewski, H., Dohan, D., Jiang, E., Cai, C., Terry, M., Le, Q., and Sutton, C. Program synthesis with large language models, 2021. URL <https://arxiv.org/abs/2108.07732>.
- Chen, M., Tworek, J., Jun, H., Yuan, Q., de Oliveira Pinto, H. P., Kaplan, J., Edwards, H., Burda, Y., Joseph, N., Brockman, G., Ray, A., Puri, R., Krueger, G., Petrov, M., Khlaaf, H., Sastry, G., Mishkin, P., Chan, B., Gray, S., Ryder, N., Pavlov, M., Power, A., Kaiser, L., Bavarian, M., Winter, C., Tillet, P., Such, F. P., Cummings, D., Plappert, M., Chantzis, F., Barnes, E., Herbert-Voss, A., Guss, W. H., Nichol, A., Paino, A., Tezak, N., Tang, J., Babuschkin, I., Balaji, S., Jain, S., Saunders, W., Hesse, C., Carr, A. N., Leike, J., Achiam, J., Misra, V., Morikawa, E., Radford, A., Knight, M., Brundage, M., Murati, M., Mayer, K., Welinder, P., McGrew, B., Amodei, D., McCandlish, S., Sutskever, I., and Zaremba, W. Evaluating large language models trained on code, 2021. URL <https://arxiv.org/abs/2107.03374>.
- Cheng, Y., Chang, Y., and Wu, Y. A survey on data contamination for large language models, 2025. URL <https://arxiv.org/abs/2502.14425>.
- Chu, T., Zhai, Y., Yang, J., Tong, S., Xie, S., Schuurmans, D., Le, Q. V., Levine, S., and Ma, Y. Sft memorizes, rl generalizes: A comparative study of foundation model post-training, 2025. URL <https://arxiv.org/abs/2501.17161>.
- Chung, H. W., Hou, L., Longpre, S., Zoph, B., Tay, Y., Fedus, W., Li, Y., Wang, X., Dehghani, M., Brahma, S., Webson, A., Gu, S. S., Dai, Z., Suzgun, M., Chen, X., Chowdhery, A., Castro-Ros, A., Pellat, M., Robinson, K., Valter, D., Narang, S., Mishra, G., Yu, A., Zhao, V., Huang, Y., Dai, A., Yu, H., Petrov, S., Chi, E. H., Dean, J., Devlin, J., Roberts, A., Zhou, D., Le, Q. V., and Wei, J. Scaling instruction-finetuned language models, 2022. URL <https://arxiv.org/abs/2210.11416>.
- Cobbe, K., Kosaraju, V., Bavarian, M., Chen, M., Jun, H., Kaiser, L., Plappert, M., Tworek, J., Hilton, J., Nakano, R., Hesse, C., and Schulman, J. Training verifiers to solve math word problems. *arXiv preprint arXiv:2110.14168*, 2021.
- Gao, L., Tow, J., Abbasi, B., Biderman, S., Black, S., DiPofi, A., Foster, C., Golding, L., Hsu, J., Le Noac’h, A., Li, H., McDonell, K., Muennighoff, N., Ociepa, C., Phang, J., Reynolds, L., Schoelkopf, H., Skowron, A., Sutawika, L., Tang, E., Thite, A., Wang, B., Wang, K., and Zou, A. The language model evaluation harness, 07 2024. URL <https://zenodo.org/records/12608602>.
- Jiang, M., Liu, K. Z., Zhong, M., Schaeffer, R., Ouyang, S., Han, J., and Koyejo, S. Investigating data contamination for pre-training language models, 2024. URL <https://arxiv.org/abs/2401.06059>.
- Kocigit, M. Y., Briakou, E., Deutsch, D., Luo, J., Cherry, C., and Freitag, M. Overestimation in llm evaluation: A controlled large-scale study on data contamination’s impact on machine translation, 2025. URL <https://arxiv.org/abs/2501.18771>.
- Kydlicek, H., Lozovskaya, A., Habib, N., and Fourier, C. Fixing open llm leaderboard with math-verify, 2025. URL https://huggingface.co/blog/math_verify_leaderboard. Blog post.
- Li, Q., Cui, L., Zhao, X., Kong, L., and Bi, W. Gsm-plus: A comprehensive benchmark for evaluating the robustness of llms as mathematical problem solvers, 2024. URL <https://arxiv.org/abs/2402.19255>.
- Magar, I. and Schwartz, R. Data contamination: From memorization to exploitation, 2022. URL <https://arxiv.org/abs/2203.08242>.
- Ouyang, L., Wu, J., Jiang, X., Almeida, D., Wainwright, C. L., Mishkin, P., Zhang, C., Agarwal, S., Slama, K., Ray, A., Schulman, J., Hilton, J., Kelton, F., Miller, L., Simens, M., Askeel, A., Welinder, P., Christiano, P., Leike, J., and Lowe, R. Training language models to follow instructions with human feedback, 2022. URL <https://arxiv.org/abs/2203.02155>.

194 Penedo, G., Kydlíček, H., allal, L. B., Lozhkov, A., Mitchell, M., Raffel, C., Werra, L. V., and
195 Wolf, T. The fineweb datasets: Decanting the web for the finest text data at scale, 2024. URL
196 <https://arxiv.org/abs/2406.17557>.

197 Qwen, :, Yang, A., Yang, B., Zhang, B., Hui, B., Zheng, B., Yu, B., Li, C., Liu, D., Huang, F., Wei,
198 H., Lin, H., Yang, J., Tu, J., Zhang, J., Yang, J., Yang, J., Zhou, J., Lin, J., Dang, K., Lu, K., Bao,
199 K., Yang, K., Yu, L., Li, M., Xue, M., Zhang, P., Zhu, Q., Men, R., Lin, R., Li, T., Tang, T., Xia,
200 T., Ren, X., Ren, X., Fan, Y., Su, Y., Zhang, Y., Wan, Y., Liu, Y., Cui, Z., Zhang, Z., and Qiu, Z.
201 Qwen2.5 technical report, 2025. URL <https://arxiv.org/abs/2412.15115>.

202 Sainz, O., García Ferrero, I., Agirre, E., Ander Campos, J., Jacovi, A., Elazar, Y., and Goldberg, Y.
203 (eds.). *Proceedings of the 1st Workshop on Data Contamination (CONDA)*, Bangkok, Thailand,
204 August 2024a. Association for Computational Linguistics. URL [https://aclanthology.org/](https://aclanthology.org/2024.conda-1.0/)
205 2024.conda-1.0/.

206 Sainz, O., García-Ferrero, I., Jacovi, A., Campos, J. A., Elazar, Y., Agirre, E., Goldberg, Y., Chen,
207 W.-L., Chim, J., Choshen, L., D’Amico-Wong, L., Dell, M., Fan, R.-Z., Golchin, S., Li, Y., Liu, P.,
208 Pahwa, B., Prabhu, A., Sharma, S., Silcock, E., Solonko, K., Stap, D., Surdeanu, M., Tseng, Y.-M.,
209 Udandaraao, V., Wang, Z., Xu, R., and Yang, J. Data contamination report from the 2024 conda
210 shared task, 2024b. URL <https://arxiv.org/abs/2407.21530>.

211 Shao, Z., Wang, P., Zhu, Q., Xu, R., Song, J., Bi, X., Zhang, H., Zhang, M., Li, Y. K., Wu, Y., and
212 Guo, D. Deepseekmath: Pushing the limits of mathematical reasoning in open language models,
213 2024. URL <https://arxiv.org/abs/2402.03300>.

214 Singh, A. K., Kocyigit, M. Y., Poulton, A., Esiobu, D., Lomeli, M., Szilvasy, G., and Hupkes, D.
215 Evaluation data contamination in llms: how do we measure it and (when) does it matter?, 2024.
216 URL <https://arxiv.org/abs/2411.03923>.

217 Team, G., Kamath, A., Ferret, J., Pathak, S., Vieillard, N., Merhej, R., Perrin, S., Matejovicova,
218 T., Ramé, A., Rivière, M., Rouillard, L., Mesnard, T., Cideron, G., bastien Grill, J., Ramos, S.,
219 Yvinec, E., Casbon, M., Pot, E., Penchev, I., Liu, G., Visin, F., Kenealy, K., Beyer, L., Zhai, X.,
220 Tsitsulin, A., Busa-Fekete, R., Feng, A., Sachdeva, N., Coleman, B., Gao, Y., Mustafa, B., Barr,
221 I., Parisotto, E., Tian, D., Eyal, M., Cherry, C., Peter, J.-T., Sinopalnikov, D., Bhupatiraju, S.,
222 Agarwal, R., Kazemi, M., Malkin, D., Kumar, R., Vilar, D., Brusilovsky, I., Luo, J., Steiner, A.,
223 Friesen, A., Sharma, A., Sharma, A., Gilady, A. M., Goedeckemeyer, A., Saade, A., Feng, A.,
224 Kolesnikov, A., Bendebury, A., Abdagic, A., Vadi, A., György, A., Pinto, A. S., Das, A., Bapna,
225 A., Miech, A., Yang, A., Paterson, A., Shenoy, A., Chakrabarti, A., Piot, B., Wu, B., Shahriari,
226 B., Petrini, B., Chen, C., Lan, C. L., Choquette-Choo, C. A., Carey, C., Brick, C., Deutsch, D.,
227 Eisenbud, D., Cattle, D., Cheng, D., Paparas, D., Sreepathihalli, D. S., Reid, D., Tran, D., Zelle, D.,
228 Noland, E., Huizenga, E., Kharitonov, E., Liu, F., Amirkhanyan, G., Cameron, G., Hashemi, H.,
229 Klimczak-Plucińska, H., Singh, H., Mehta, H., Lehri, H. T., Hazimeh, H., Ballantyne, I., Szepektor,
230 I., Nardini, I., Pouget-Abadie, J., Chan, J., Stanton, J., Wieting, J., Lai, J., Orbay, J., Fernandez,
231 J., Newlan, J., yeong Ji, J., Singh, J., Black, K., Yu, K., Hui, K., Vodrahalli, K., Greff, K., Qiu,
232 L., Valentine, M., Coelho, M., Ritter, M., Hoffman, M., Watson, M., Chaturvedi, M., Moynihan,
233 M., Ma, M., Babar, N., Noy, N., Byrd, N., Roy, N., Momchev, N., Chauhan, N., Sachdeva, N.,
234 Bunyan, O., Botarda, P., Caron, P., Rubenstein, P. K., Culliton, P., Schmid, P., Sessa, P. G., Xu,
235 P., Stanczyk, P., Tafti, P., Shivanna, R., Wu, R., Pan, R., Rokni, R., Willoughby, R., Vallu, R.,
236 Mullins, R., Jerome, S., Smoot, S., Girgin, S., Iqbal, S., Reddy, S., Sheth, S., Pöder, S., Bhatnagar,
237 S., Panyam, S. R., Eiger, S., Zhang, S., Liu, T., Yacovone, T., Liechty, T., Kalra, U., Evci, U.,
238 Misra, V., Roseberry, V., Feinberg, V., Kolesnikov, V., Han, W., Kwon, W., Chen, X., Chow, Y.,
239 Zhu, Y., Wei, Z., Egyed, Z., Cotruta, V., Giang, M., Kirk, P., Rao, A., Black, K., Babar, N., Lo, J.,
240 Moreira, E., Martins, L. G., Sanseviero, O., Gonzalez, L., Gleicher, Z., Warkentin, T., Mirrokni,
241 V., Senter, E., Collins, E., Barral, J., Ghahramani, Z., Hadsell, R., Matias, Y., Sculley, D., Petrov,
242 S., Fiedel, N., Shazeer, N., Vinyals, O., Dean, J., Hassabis, D., Kavukcuoglu, K., Farabet, C.,
243 Buchatskaya, E., Alayrac, J.-B., Anil, R., Dmitry, Lepikhin, Borgeaud, S., Bachem, O., Joulin, A.,
244 Andreev, A., Hardin, C., Dadashi, R., and Hussenot, L. Gemma 3 technical report, 2025. URL
245 <https://arxiv.org/abs/2503.19786>.

246 Toshniwal, S., Du, W., Moshkov, I., Kisacanin, B., Ayrapetyan, A., and Gitman, I. Openmathinstruct-
247 2: Accelerating ai for math with massive open-source instruction data. *arXiv preprint*
248 *arXiv:2410.01560*, 2024.

249 von Werra, L., Allal, L. B., and Wolf, T. Codeparrot train v2 near-dedup. <https://huggingface.co/datasets/codeparrot/codeparrot-train-v2-near-dedup>, 2021. Dataset on the Hugging Face Hub. See the accompanying blog post “Training CodeParrot from Scratch”.

252 Wei, J., Bosma, M., Zhao, V. Y., Guu, K., Yu, A. W., Lester, B., Du, N., Dai, A. M., and Le, Q. V. Finetuned language models are zero-shot learners, 2022. URL <https://arxiv.org/abs/2109.01652>.

255 Zhang, P., Zeng, G., Wang, T., and Lu, W. Tinyllama: An open-source small language model. *arXiv preprint arXiv:2401.02385*, 2024a.

257 Zhang, S., Dong, L., Li, X., Zhang, S., Sun, X., Wang, S., Li, J., Hu, R., Zhang, T., Wu, F., and Wang, G. Instruction tuning for large language models: A survey, 2024b. URL <https://arxiv.org/abs/2308.10792>.

260 Appendix A - Details of training and evaluation

261 Table 1 provides the details of data composition. Preliminary experiments using only web text even
 262 with high-quality sources like books resulted in significant performance drops on math and coding
 263 tasks, rendering some experiments ineffective. Consequently, we opted for a more balanced mixture
 264 that includes task-specific data.

265 We perform extended pre-training with a short warm-up phase, followed by a fixed small learning
 266 rate typically used as the minimum learning rate in full pre-training schedules (Zhang et al., 2024a).
 267 This choice reflects the fact that the models have already been pre-trained on large corpora and the
 268 final learning rate in their training probably approached the minimum.

269 For evaluation, we employ the LM Evaluation Harness (Gao et al., 2024) and make necessary
 270 adjustments to prompts and tokenization to closely replicate baseline scores reported in prior work
 271 (Qwen et al., 2025; Team et al., 2025), minimizing variance from evaluation artifacts. We also use the
 272 math-verify library to parse responses for math tasks, as differences in output formatting especially
 273 in base models can significantly impact measured performance (Kydlicek et al., 2025).

274 For contamination, five copies of GSM8k and MBPP test sets prompted as shown in Appendix 5 are
 275 randomly inserted into the first 2B tokens of the contaminated training mixture. This way the model
 276 is trained on more than 23B tokens after it is exposed to the contamination, making our findings more
 277 realistic. Kocyyigit et al. (2025) show that late contamination, when set up in a correct way, does not
 278 necessarily yield higher performance inflation compared to uniformly distributing contamination
 279 across the entire training corpora.

Dataset	Token Count
OpenMath-Instruct	6,495,049,388
CodeParrots	3,647,426,066
FineWeb-Edu	14,869,906,469
Total	25B

Table 1: Token Counts for Components of the Pre-training data.

280 Appendix B - Change of performance with time

281 Initially, we examined how the contaminated and clean model performance changes over the training
 282 process. Specifically in Figure 3, we present Qwen2.5-1.5B’s contaminated and clean accuracies on
 283 the GSM8k benchmark. Similar to previous work (Kocyyigit et al., 2025), we observe that performance
 284 of the contaminated model spikes at the time of contamination then decreases back to the same level
 285 as the clean model. This observation is also supported by the base results shown i

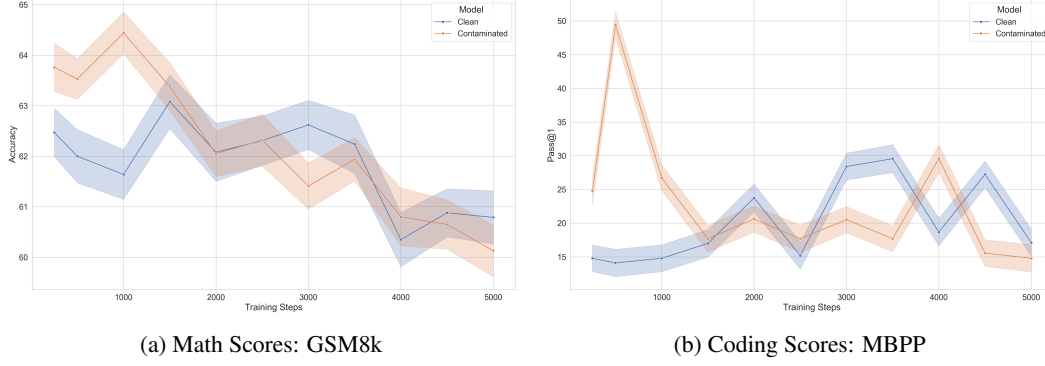


Figure 3: Performance Over Time: Accuracy and Pass@1 of the Clean and Contaminated Qwen2.5-1.5B models on the GSM8k Benchmark. The contamination is within the first 500 steps. We observe that the performance gap is much bigger at the exposure point but then closes as the model is trained on more data. For MBPP we observe that the peak is much higher meaning contamination of code is memorized better at sight however overtime the performance normalizes just like math.

Appendix C - How does model scale impact the generalization gap?

To analyze the impact of model scale on model generalization, we now track the **contamination-gap difference**, defined for each training recipe $t \in \{\text{Base, SFT, GRPO}\}$ and each model-size pair (m, s) as

$$d_{m,s,t} = \underbrace{\Delta M_1^t(m, s)}_{\text{GSM8k gain}} - \underbrace{\Delta M_2^t(m, s)}_{\text{GSMPlus gain}}, \quad (1)$$

where the per-metric contamination improvement is

$$\Delta M_k^t(m, s) = M_k^{\text{contam},t}(m, s) - M_k^{\text{clean},t}(m, s). \quad (2)$$

where $k \in \{1 \text{ (GSM8k), } 2 \text{ (GSMPlus)}\}$. A positive $d_{m,s,t}$ therefore means the contaminated model gains more on GSM8k than on GSMPlus after recipe t ; a negative value indicates the opposite. We present the results in Figure 4. Here we see that for the pure pre-trained Base model and the supervised fine-tuned model, the contamination-gap difference increases as model scale increases, which suggests more overestimation from contamination. However, for the model post-trained with GRPO, as the model scale increases, the model is actually able to learn more generalizable features and the contamination-gap difference is smaller.

Scale amplifies the contrast. Section 4 and Figure 4 reveal that larger SFT models extract *more* benefit from contamination that does not translate into a relative benefit on non-contaminated benchmarks. GRPO shows the opposite trend: bigger models channel additional capacity toward broad generalization and thus *dilute* relative overestimation. The interplay between scale and alignment methods, therefore, deserves careful attention, but our findings are in line with Chu et al. (2025)

Appendix D — GRPO Details

D.1 Dataset

Our experiments fine-tune on the *GSM8k* (openai/GSM8k, main split). Each prompt begins with a *system* message that prescribes the `<reasoning>` / `<answer>` XML schema, followed by a single worked example (“*What is 2 + 2?*” → 4), and finally the user question drawn from GSM8k. The gold integer answer is extracted from the dataset for the correctness reward.

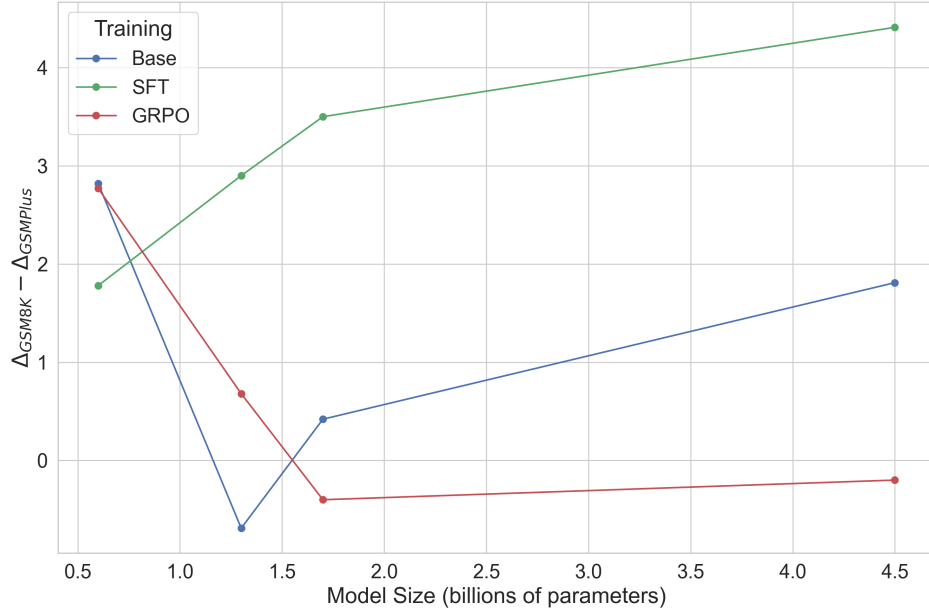


Figure 4: Contamination-gap Difference for Math: We measure the contamination-gap difference as in Eq 1. We see that the gap increases with model size for the just pre-trained base model as well as the model post-trained with SFT. However for the models post-trained with GRPO the gap actually gets smaller, even negative as the model size increases.

310 D.2 Model and Tokenizer

311 The script supports any causal-LM checkpoint that fits on GPU, and has been validated on Google’s
 312 *Gemma-3-1B-IT* and *Gemma-3-4B-IT* as well *Qwen2.5-0.5B-IT* and *Qwen2.5-1.5B-IT*. The tokenizer
 313 comes from the same directory; we set the padding token equal to `\eos` and enable a beginning-of-
 314 sequence token. This particularly impacts the performance of Gemma3 models.

315 D.3 GRPO Training Configuration

316 We train with GRPO (Group Regularised Policy Optimisation) for a single epoch over GSM8k train
 317 set. Optimisation uses 8-bit AdamW. The model produces sixteen candidate generations for each
 318 input example.

319 D.4 Reward Functions

320 Policy updates rely on lightweight heuristic rewards, each returning a scalar per generated sample;
 321 their contributions are summed without additional weighting. The complete Python implementation
 322 is reproduced below.

Listing 1: Reward functions used by GRPO

```

323
324 # -- Math Reward functions -----
325 def format_reward_func(completions, **_):
326     """1 pt for a perfectly formatted XML response."""
327     pattern = r"^(<reasoning>(?!</reasoning>).*</reasoning>\n"
328     \
329     r"<answer>(?!</answer>).*</answer> $"
330     responses = [c[0]["content"] for c in completions]
331     return [1.0 if re.match(pattern, r) else 0.0 for r in
332             responses]
333

```

```

334 def correctness_reward_func(prompts, completions, answer, **_):
335     """2 pt when <answer> matches the gold integer."""
336     responses = [c[0]["content"] for c in completions]
337     preds     = [extract_xml_answer(r) for r in responses]
338     return [2.0 if p == a else 0.0 for p, a in zip(preds, answer)]
339
340 def int_reward_func(completions, **_):
341     """0.5 pt if <answer> contains any integer."""
342     responses = [c[0]["content"] for c in completions]
343     preds     = [extract_xml_answer(r) for r in responses]
344     return [0.5 if p.isdigit() else 0.0 for p in preds]
345
346 def strict_format_reward_func(completions, **_):
347     """0.5 pt for newline-strict XML formatting."""
348     pat = r"^(<reasoning>\n.*?\n</reasoning>\n<answer>\n.*?\n</
349         answer>\n$"
350     responses = [c[0]["content"] for c in completions]
351     return [0.5 if re.match(pat, r) else 0.0 for r in responses]
352
353 def soft_format_reward_func(completions, **_):
354     """0.5 pt for a laxer XML pattern (tags may abut)."""
355     pat = r"<reasoning>.*?</reasoning>\s*<answer>.*?</answer>"
356     responses = [c[0]["content"] for c in completions]
357     return [0.5 if re.match(pat, r) else 0.0 for r in responses]
358
359 def xmlcount_reward_func(completions, **_):
360     """Fractional reward based on correct tag usage."""
361     responses = [c[0]["content"] for c in completions]
362     return [count_xml(r) for r in responses]

```

363 For the code specific post-training the main difference is the reward models that are used

Listing 2: Reward functions used by GRPO

```

364
365 # -- Code Reward functions -----
366 def tests_reward_func(completions, *, tests: List[str], **_) ->
367     List[float]:
368     rewards = []
369     for c, t in zip(completions, tests):
370         code = clean_completion(c[0]["content"])
371         #print('--' * 20)
372         #print('Running code:', code)
373         try:
374             passed, _, _ = run_with_timeout(code, t, time_limit
375                 =1.0)
376             rewards.append(2.0 * passed / len(t))
377         except Exception:
378             rewards.append(0.0)
379     return rewards
380
381
382 def typehint_reward_func(completions, **_) -> List[float]:
383     codes = [c[0]["content"] for c in completions]
384     scores = []
385     for code in codes:
386         try:
387             tree = ast.parse(code)
388             hints = sum(bool(a.annotation) for a in ast.walk(tree))

```

```

389         if isinstance(a, (ast.arg, ast.FunctionDef
390             )))
391         scores.append(min(0.25, 0.05 * hints)) # cap at 0.25
392     except Exception:
393         scores.append(0.0)
394     return scores
395
396 def brevity_reward_func(completions, **_) -> List[float]:
397     codes = [c[0]["content"] for c in completions]
398     return [max(0.0, 0.25 - 0.002 * len(code.splitlines())) for
399         code in codes]

```

400 Appendix E — Prompt Templates

401 E.1 Training Prompt - GRPO

402 Each training sample is a sequence of four chat messages. The placeholder {QUESTION} is substituted
403 with the GSM8k problem text; the corresponding ground-truth integer answer is kept separately for
404 reward computation.

```

405 <system>
406 A conversation between User and Assistant. The user asks a question,
407 and the Assistant solves it. The assistant first thinks about the
408 reasoning process in the mind and then provides the user with the
409 answer. The reasoning process and answer are enclosed within
410 <reasoning> </reasoning> and <answer> </answer> tags, respectively,
411 i.e., <reasoning> reasoning process here </reasoning>
412 <answer> answer here </answer>. The answer must be a single integer.
413 Format example:
414 <reasoning>
415 ...
416 </reasoning>
417 <answer>
418 ...
419 </answer>
420 </system>
421
422 <user>
423 What is 2+2?
424 </user>
425
426 <assistant>
427 <reasoning>
428 To calculate 2+2, we simply add the numbers together: 2 + 2 = 4.
429 </reasoning>
430 <answer>
431 4
432 </answer>
433 </assistant>
434
435 <user>
436 {QUESTION}
437 </user>

```

438 E.2 Training Prompt - SFT

439 For SFT we take the training split of GSM8k and process the question, chain of thought and answer
440 separately. The model is trained on the prompt structure below but all the tokens before let's think

step-by-step are masked to not produce any loss. Thus the model is only trained on the chain of thought and answer given the input question.

Question: {QUESTION}\
 Let's think step by step.
 {CHAIN_OF_THOUGHT}.
 The answer is {ANSWER}

E.3 Evaluation Prompt

All evaluation follows the same format where we just prompt the model with the question followed by a "Let's think step by step." primer. We opted for zero shot evaluations as the few shot examples can impact how we have contaminated the data as well.

Question: {QUESTION}\
 Let's think step by step.

Appendix F - Method Overview Plot

We present an overview of our method visually for readers who prefer to take a quick look.

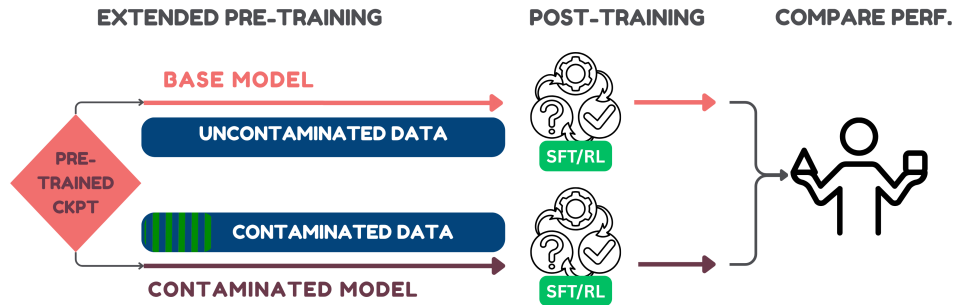


Figure 5: An Overview of our Method: We take existing pre-trained models and run them through extended pre-training with and without contamination. Afterwards we post-train them using SFT or RL methods and compare their performance. The pre-trained checkpoints here are from Qwen2.5 and Gemma3 non-instruction tuned models.

Appendix G - Limitations and Future Work

Our study purposefully isolates a *5-copy, late-injection* contamination scenario. Real-world leakage can be multi-pass, paraphrased, or distributed throughout pre-training, and its effects may differ. Moreover, we restrict model sizes to 4B parameters and focus on two open-weights families. That said, larger proprietary models, different architectures, or alternative RLHF algorithms could behave differently. Finally, our RL setting uses synthetic rule-based rewards; human-annotated preference signals might have different effects.

With these limitations in mind, an immediate first step for future research is to run our experiment on a larger scale. While we share some insights into how our results would scale, the analysis is still very local and there are potentially mixed signals that would benefit more from further exploration. Additionally, with more and more complex RL systems, the post-training stage can become just as complex as the pre-training stage, warranting further investigations regarding the compute that goes into the post-training stage and the impact it has on model behavior and performance.