

Compute-optimal scaling laws for data-constrained antibody language models

Jinglin Jian · Karenn Ng · Sarah M. Burbach · Shelby Ferrier · Mahdi Shafiei Neyestanak · Praneeth Gangavarapu · Benjamin Nemoz · Jonathan Hurtado · Stefano Forli · Bryan Briney

Scripps Research, La Jolla, CA, USA

Antibody MLMs must repeat their data. What is the recipe?

Scaling laws are the organizing principle for training large models: loss falls as a predictable power law in model size, data and compute⁴, and the Chinchilla law³ turns that surface into a recipe that grows the model and the data together (below left). Masked language models (MLMs) have the same law with a different objective: **Cheng et al.**¹ fitted compute-optimal rules for masked protein language models. But every token those models are fed is fresh.

Antibody masked language models sit in the more serve data-constrained setting. The corpus is small, fixed, and expensive to collect, so more tokens mean more epochs over the same data⁸. Following **Muennighoff et al.**⁵ we count effective data, where a repeated epoch counts for less than a fresh one, and we measure how much less.

Causal Language Model

The cat sat on the ?

given the first $n-1$ tokens as context, the model learns to predict the next one

What corpus does the model learn?

UNLIMITED DATA: every batch is new text

Chinchilla's recipe
 $L(N, D) = 1.69 + \frac{406.4}{N^{0.34}} + \frac{410.7}{D^{0.28}}$
 $N_{\text{opt}} \propto C^{0.51}$
so model and data grow together

Masked Language Model

Q V Q ? V E S

15% of the positions are hidden at random, and filled in from the residues on both sides

What corpus does the model learn?

REPETITIVE: the same training set, pass after pass

So what is the recipe here?

Hypothesis: repeated data still teaches an MLM something

A masked model hides a random 15% of the residues and predicts them from the rest, re-drawing **which** residues to hide on every pass. So we expect a masked model to tolerate far more passes over one corpus than a causal model can, and to reach the same error with a smaller model.

Pass 1 Q V Q V E S G G ? L V Q P G G S ? R L S C A A S G ? T F S S Y A M ? V W R Q A P G K G ? E W V ? A
Pass 2 Q ? Q L V E S ? G G L V Q P ? G S L R L S C ? A S G F T F S ? Y A M S W V R O ? P G K G L E ? V S A
Pass 3 Q V Q L V ? S G G G L V ? P G G S L R ? S C A A S G F T ? S S Y A M S W ? R Q A P ? K G L E W V S ?

same sequence, three different questions

The fitting law

Following **Muennighoff et al.**⁵ we ran the experiments in the next two columns to fit this surface. **L is the eval loss**, cross-entropy on held-out antibodies the model never trained on, so it scores the test task and lower is better.

$$L(N, U, R) = E + \frac{A}{N^\alpha} + \frac{B}{D_{\text{eff}}^\beta}$$

$$D_{\text{eff}} = U \left[1 + R_d^* \left(1 - \exp \left(-\frac{R-1}{R_d^*} \right) \right) \right]$$

L = eval loss, lower is better
 N = how big the model is
 U = unique data we hold, fixed
 R = how many passes over it
 E = the floor no model beats
 A, α = how fast error falls with size
 B, β = how fast it falls with data
 R_d^* = how fast a repeat loses value

Step 1: muP + complete-P makes the sweep affordable

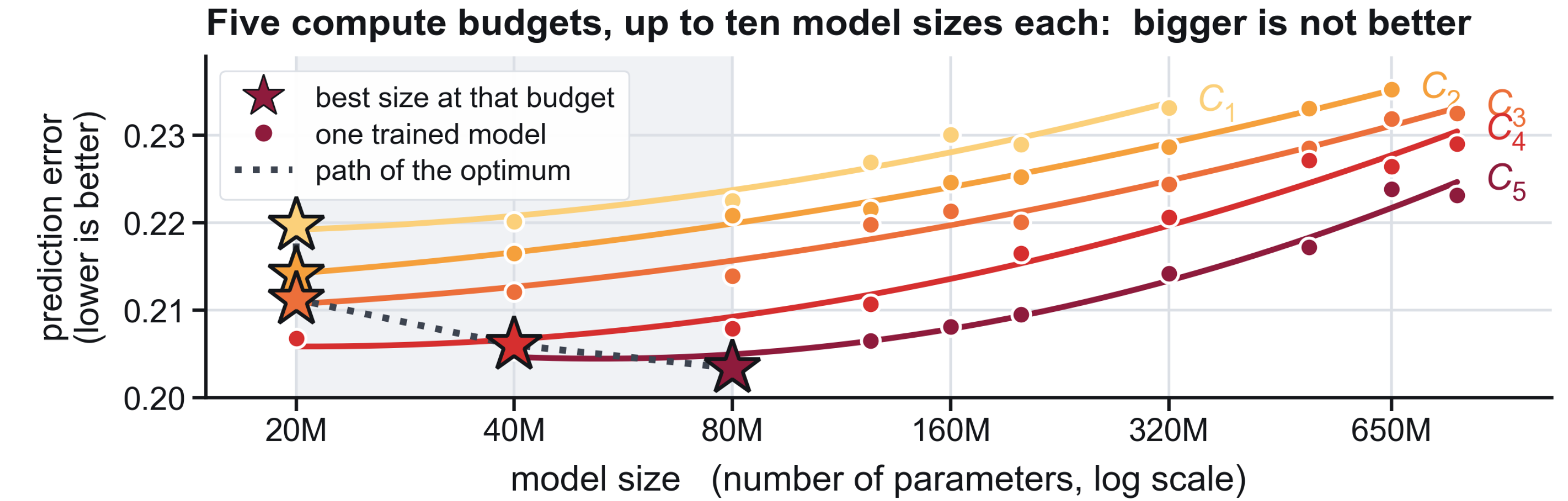
Setting. Tuning a learning rate at every size would make this ladder unaffordable. **muP⁹ + complete-P²** derives each tensor's scaling from width and depth, so one base rate serves every size: we swept base learning rates once on a 124M model and transferred it across the 20M to 800M ladder.

transfer score (lower is better)	$\alpha_d = 1.0$	$\alpha_d = 0.5$	no muP
across width	0.138	0.091	0.55
across depth	0.275	0.091	0.30
width + depth + heads at once	not run	0.075-0.120	1.024

Result. Coordinate checks pass at every cell of the width x depth matrix, **0.075 to 0.120** against a 0.20 bar. $\alpha_d = 1.0$ breaks depth transfer (0.275); $\alpha_d = 0.5$ holds it (0.091).

Step 2: the isoFLOP ladder locates the allocation

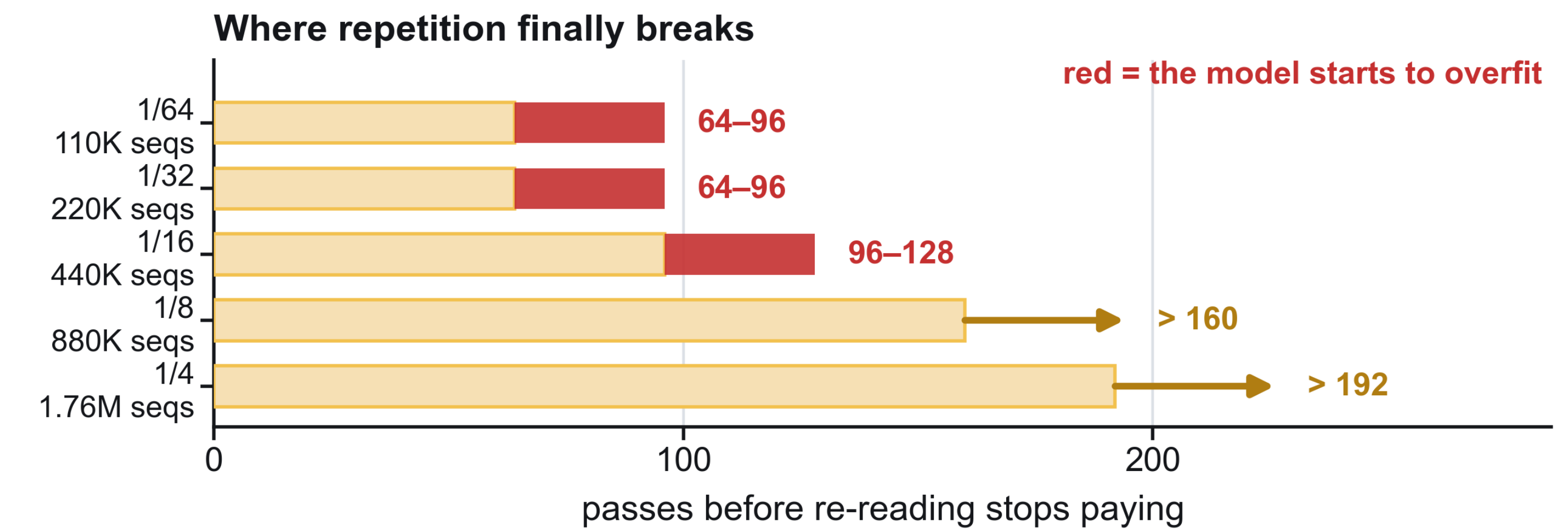
Setting. The corpus is held at $U_0 = 7,044,237$ sequences $\approx 2.254 \times 10^9$ tokens. Each budget C_k fixes the total tokens D , and we sweep N from 20M to 800M inside it. The corpus never grows, so the pass count $R = D / U_0$ is locked by the budget: the ladder is purely an allocation experiment.



Result. Every budget traces a bowl and the best model stays small, 20M to 80M over a 6.5× compute range. At the top of the ladder that is 80M, against 127M for the masked free-data rule¹ and 533M for the causal repetition rule⁵.

Step 3: the repetition arm prices a repeated pass

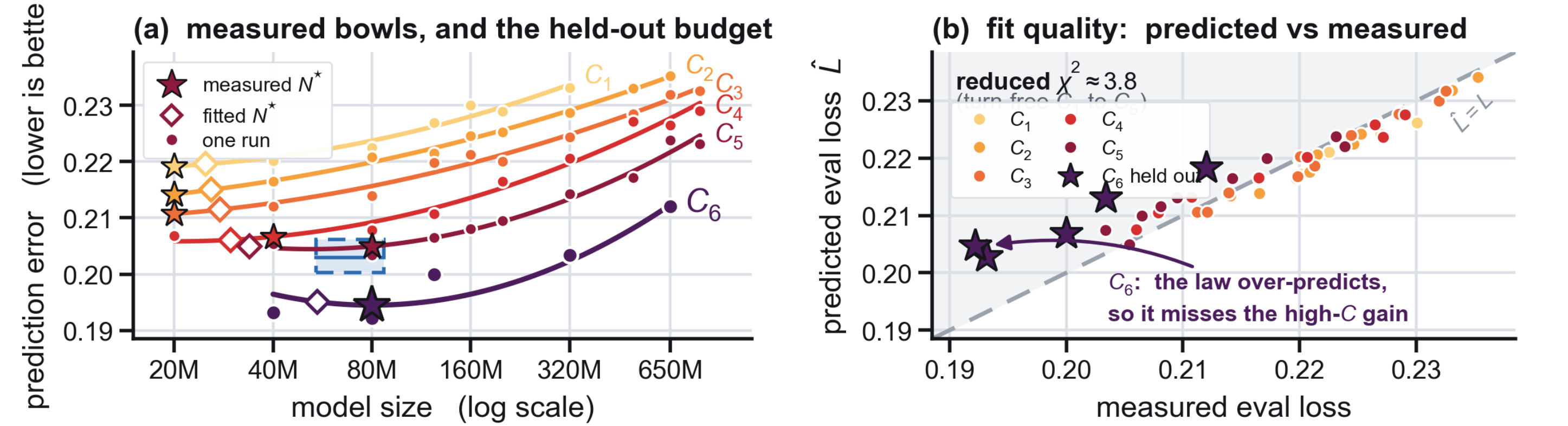
Setting. A fresh arm at $R = 1$ varies the unique subset U at 40M and 124M, which pins B and β . A **repetition grid** then sweeps the pass count R at five fractions of the corpus, pushing R as deep as 256 so that R_d^* separates from B .



Result. All six parameters are fitted jointly by a robust, σ -weighted least squares from a grid of starts, dropping runs whose loss rises with epochs, with 95% intervals from 300 bootstrap resamples.

Held-out test at C_6

Setting. At $C_6 = 8.7 \times 10^{19}$ FLOPs, 1.6× past the fitting range, the fitted surface predicts a compute-optimal size of **54M to 87M** at an eval loss of **0.2029** (interval 0.2003 to 0.2062). We fixed that call, then ran the ladder.

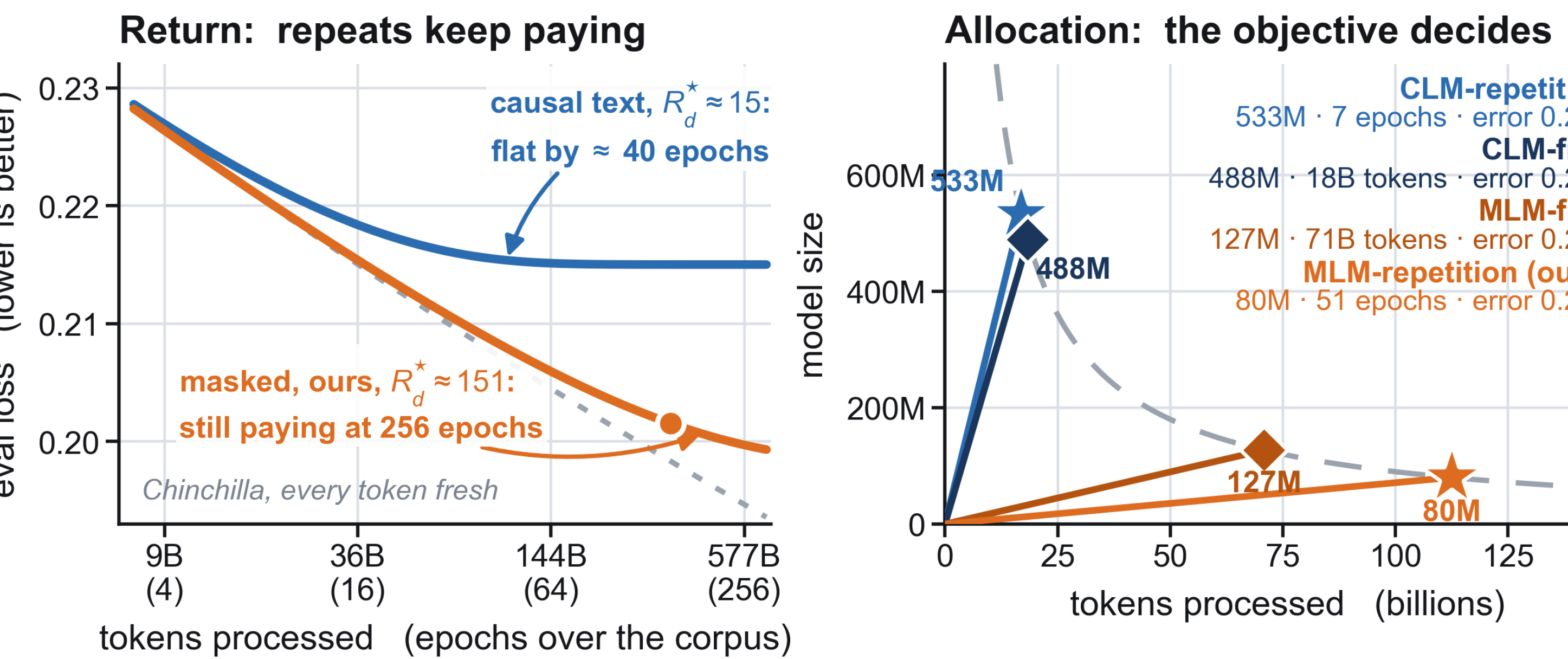


Results. The C_6 bowl bottoms at **80M** (eval loss **0.1922**), inside the predicted [54, 87]M interval and below the predicted loss, so the fit runs conservative at high compute; measured beating predicted is itself a robustness result. The loss lands 0.0107 under the point prediction, 0.0081 under the interval floor.

Conclusion

- **Infrastructure works.** muP + complete-P ($\alpha_d = 0.5$) transfers one learning rate across width and depth, over the whole 20M to 800M ladder.
- **Compute-optimal models are small.** Smaller models, more epochs. Extra compute buys epochs of a small model until they approach the repetition ceiling, and only then grows the model.
- **Repeated data is nearly free, until it isn't.** Fitted $R_d^* \approx 151$ for masked antibody LMs on our fixed 7.04M-sequence corpus, over budgets C_1 to C_5 inside the turn-free $R \leq 118$ domain.
- **Masking is why repeats keep paying.** Re-sampling the mask each epoch gives different supervised positions every pass, so a repeat is only partial: $R_d^* \approx 151$ for masked models against ≈ 15 for causal ones⁵.
- **Data-constrained scaling is sharper than Chinchilla.** With the corpus dead at $U_0 \approx 2.25 \times 10^9$ tokens and repeat value saturating, effective data has a hard ceiling: $D_{\text{eff}}^{\text{max}} \approx U_0(1 + R_d^*) \approx$ **343 billion tokens**, about 1.07 billion fresh-equivalent sequences.

Main Takeaways



Left: what a repeated pass is worth. **Right:** where one fixed budget of 5.4×10^{19} FLOPs lands under four recipes.

WORKS CITED

1. Cheng, Xingyi, et al. "Training Compute-Optimal Protein Language Models." *Advances in Neural Information Processing Systems*, 2024.
2. Dey, Nolan, et al. "Don't Be Lazy: CompleteP Enables Compute-Efficient Deep Transformers." *Advances in Neural Information Processing Systems*, 2025.
3. Hoffmann, Jordan, et al. "Training Compute-Optimal Large Language Models." *arXiv*, arXiv:2203.15568, 2022.

4. Kaplan, Jared, et al. "Scaling Laws for Neural Language Models." *arXiv*, arXiv:2001.08361, 2020.
5. Muennighoff, Niklas, et al. "Scaling Data-Constrained Language Models." *arXiv*, arXiv:2305.16264, 2023.
6. Neyestanak, Mahdi Shafiei, et al. "Data-Optimal Scaling of Paired Antibody Language Models." *bioRxiv*, 2025.09.02.673765, 2025.

7. Olsen, Tobias H., et al. "Observed Antibody Space: A Diverse Database of Cleaned, Annotated, and Translated Unpaired and Paired Antibody Sequences." *Protein Science*, vol. 31, no. 1, 2022, pp. 141–46.
8. Qin, Tian, et al. "Bridging Compute- and Data-Optimal Pretraining." *arXiv*, arXiv:2607.25271, 2026.
9. Yang, Greg, et al. "Tensor Programs V: Tuning Large Neural Networks via Zero-Shot Hyperparameter Transfer." *arXiv*, arXiv:2203.03466, 2022.