

NOBITA'S REASONING REVOLUTION

Provable Sample Efficiency of Curriculum Post-Training

BU DAKE

Doraemon! It's impossible!
I try to solve these reasoning
problems, but I never get
the reward! It's like finding a
needle in a galaxy!

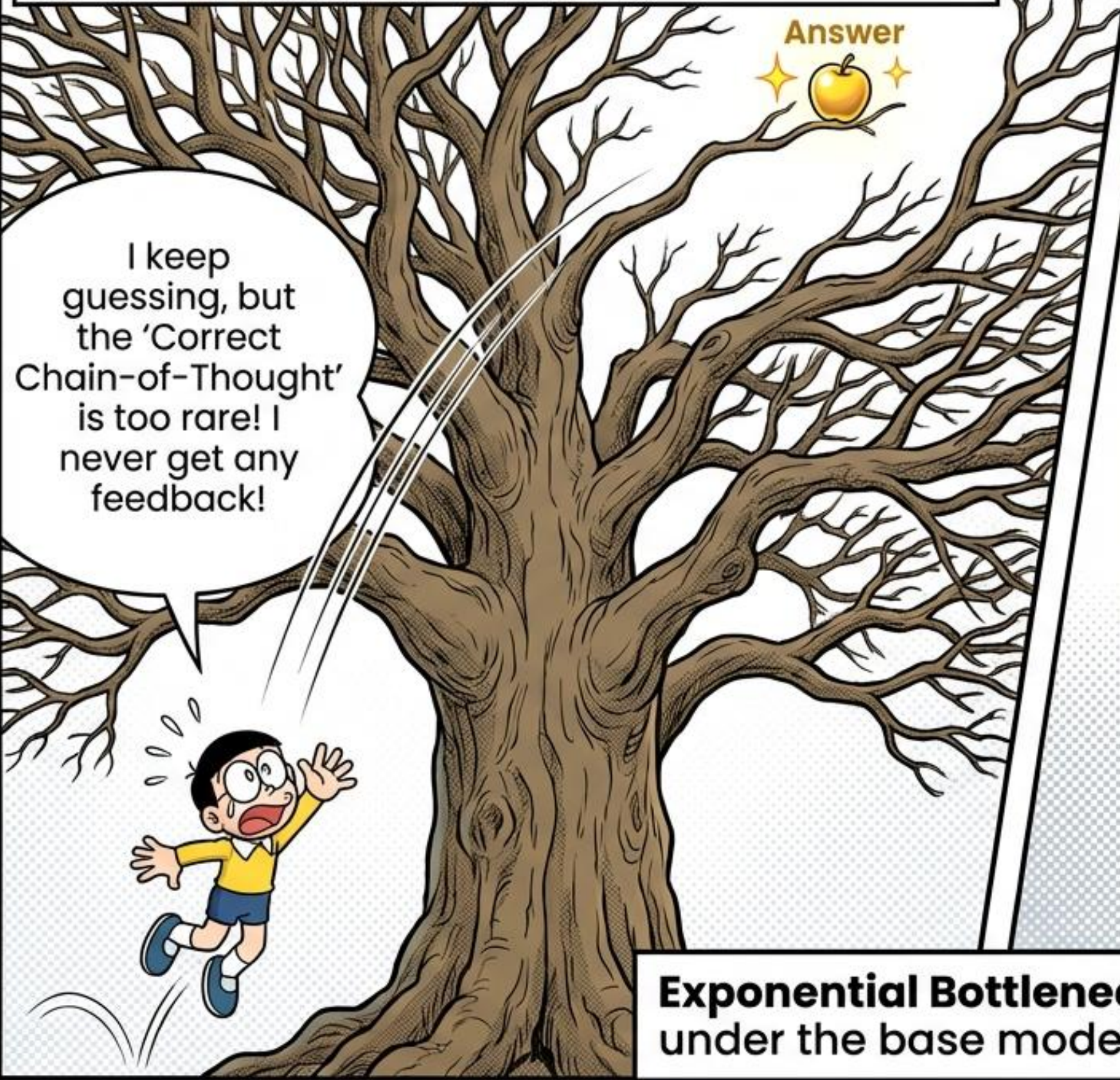


Standard Reinforcement
Learning again?
Nobita, you're hitting the

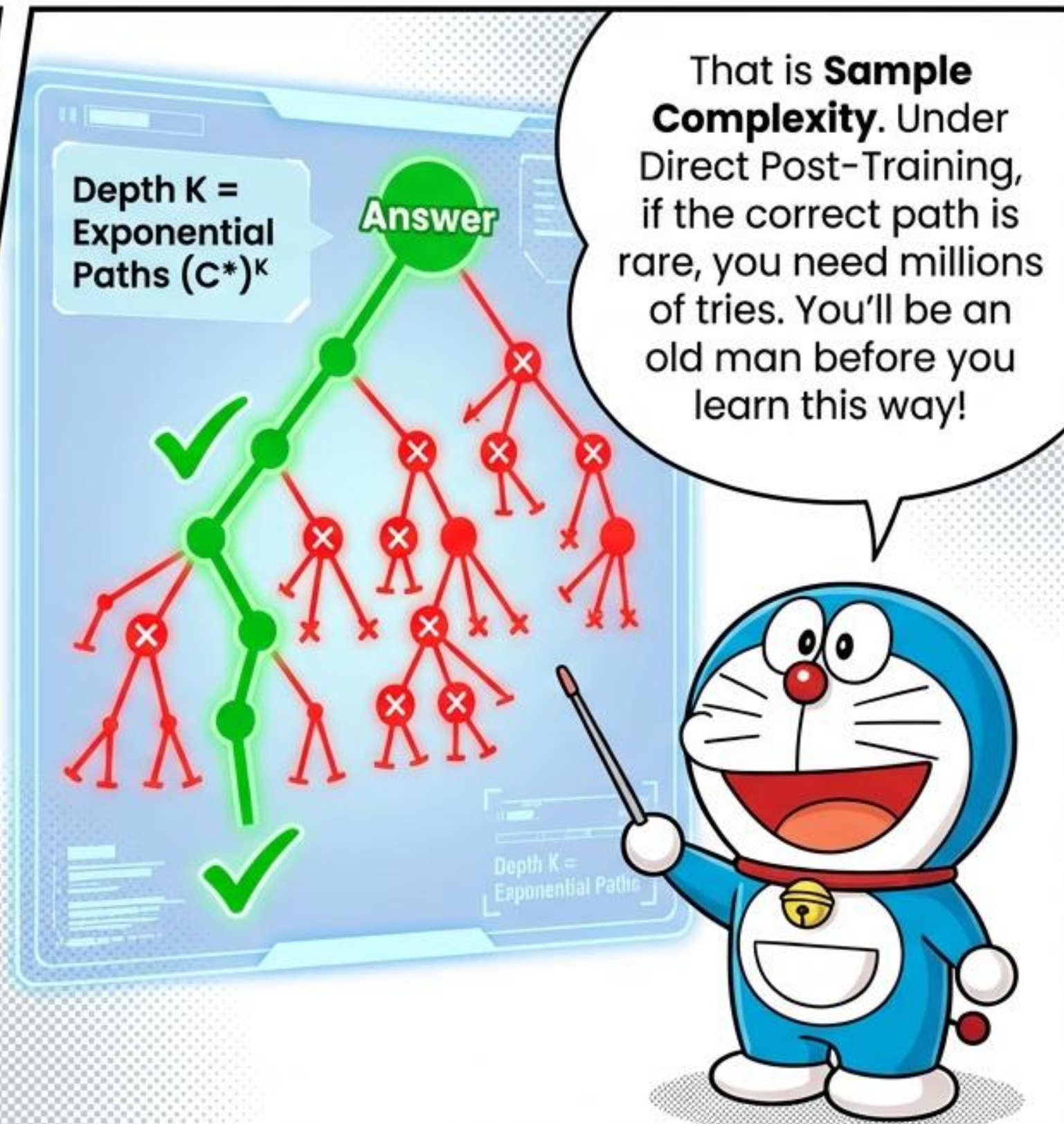
**Exponential
Bottleneck.**

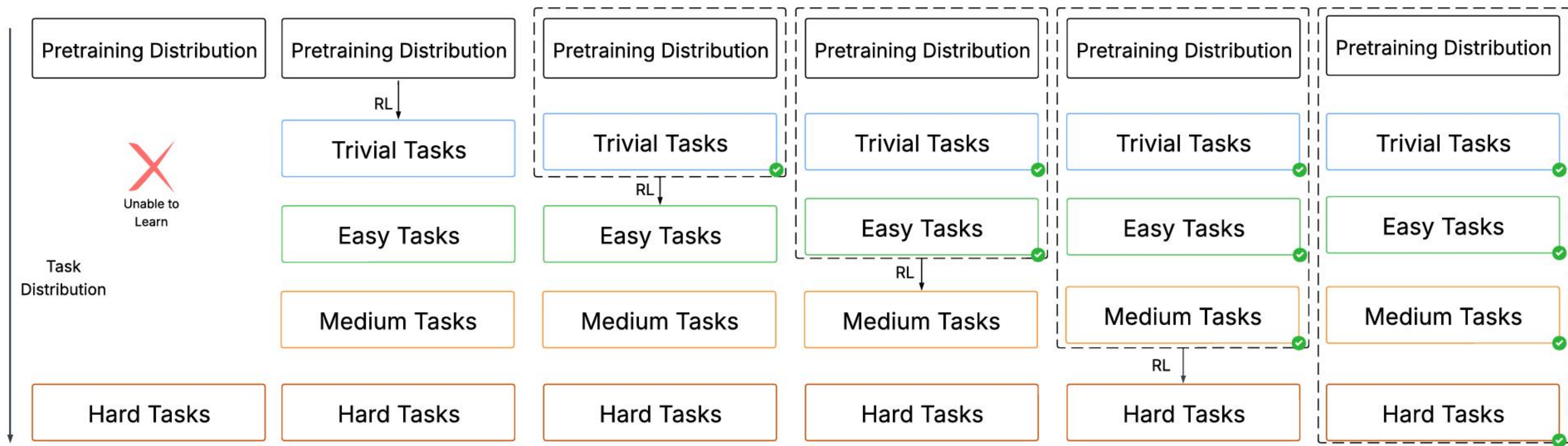


The Problem: Direct Reinforcement Learning

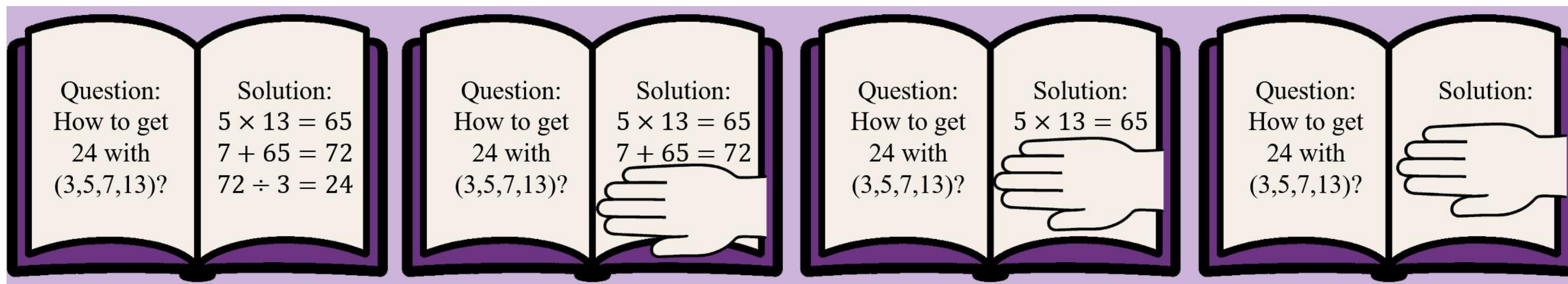


Exponential Bottleneck: The rarity of correct reasoning trajectories under the base model drives the sampling cost.





[1] Parashar et al (2025) Curriculum Reinforcement Learning from Easy to Hard Tasks Improves LLM Reasoning



[2] Liu et al. (2025) UFT: Unifying Supervised and Reinforcement Fine-Tuning

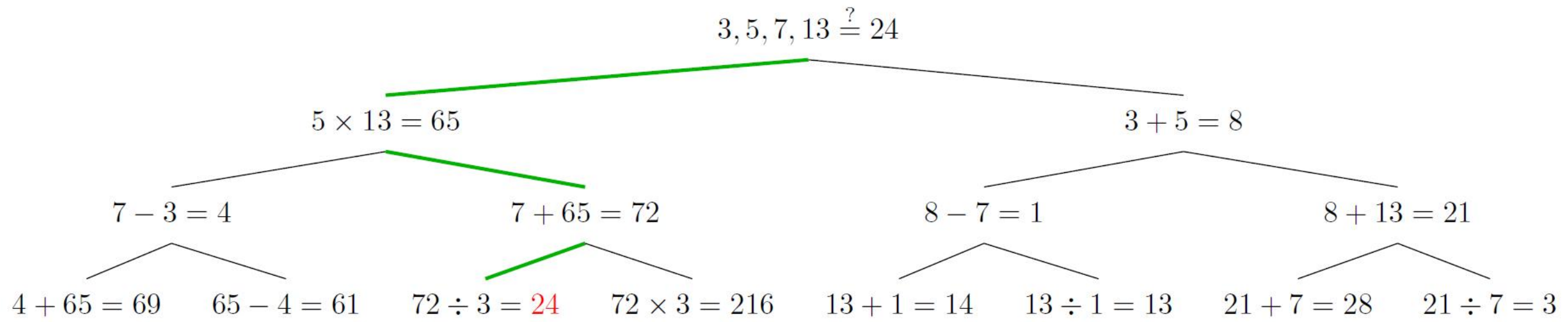


Figure 1. An illustration in Liu et al. (2025b) of the Chain-of-Thought for Countdown game, where the goal is to obtain 24 by applying basic arithmetic operations $(+, -, \times, \div)$ ($\Phi_l(\cdot, \cdot)$ in our Def. 2) between the current step’s number (e.g., 13, 65 or 72) and some unused number (e.g., 5, 7 or 3) in $\{3, 5, 7, 13\}$, targeting 24 as the final outcome. Per (Parashar et al., 2025), the difficulty measure of Countdown is the number of arithmetic operations required to solve an instance.

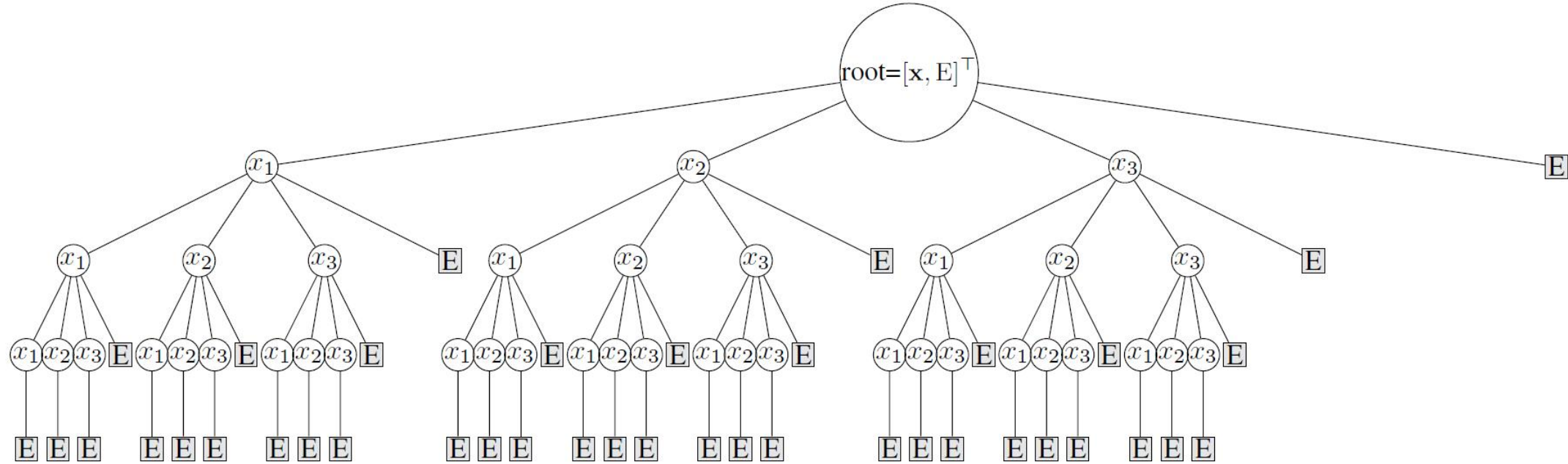
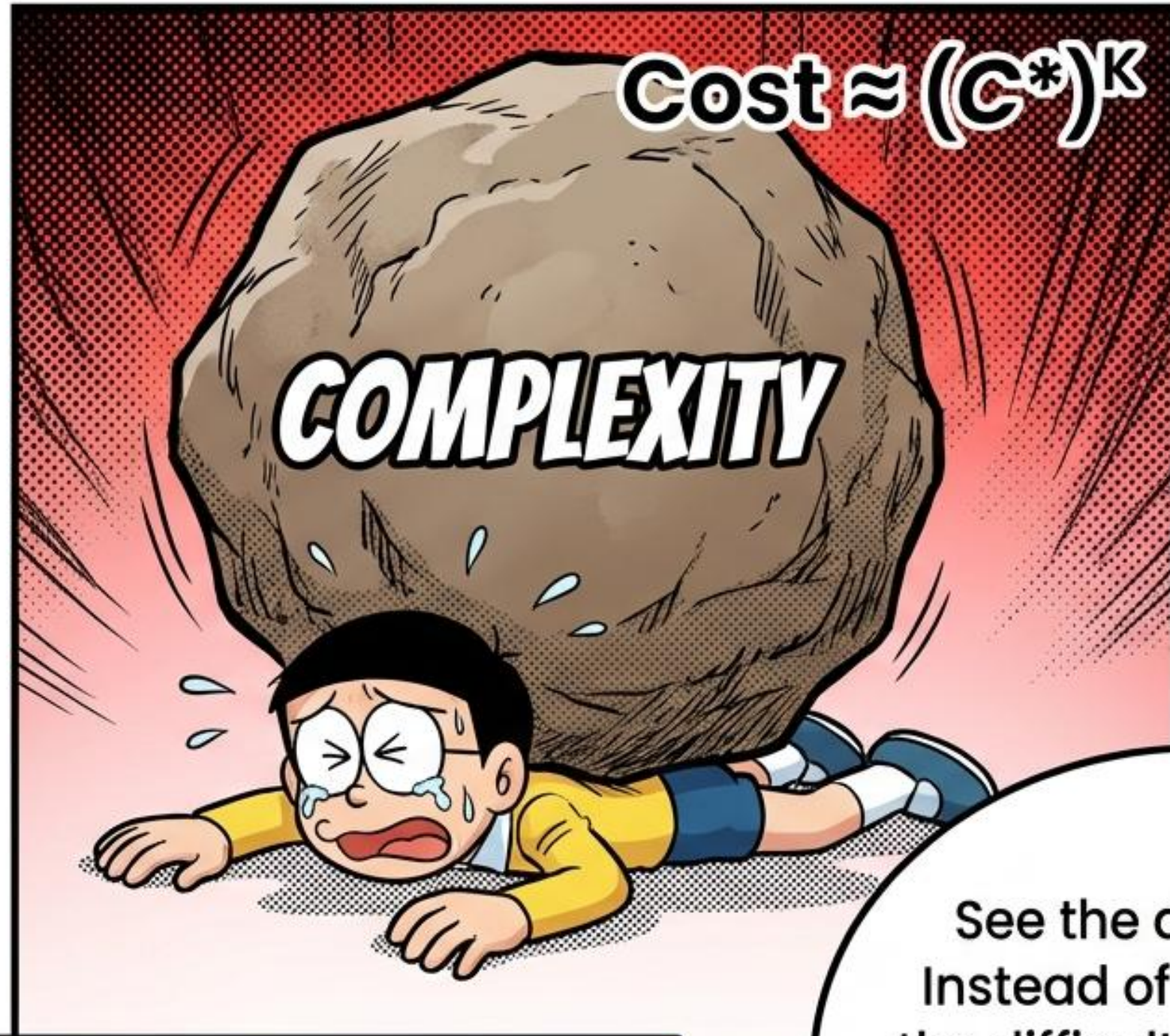
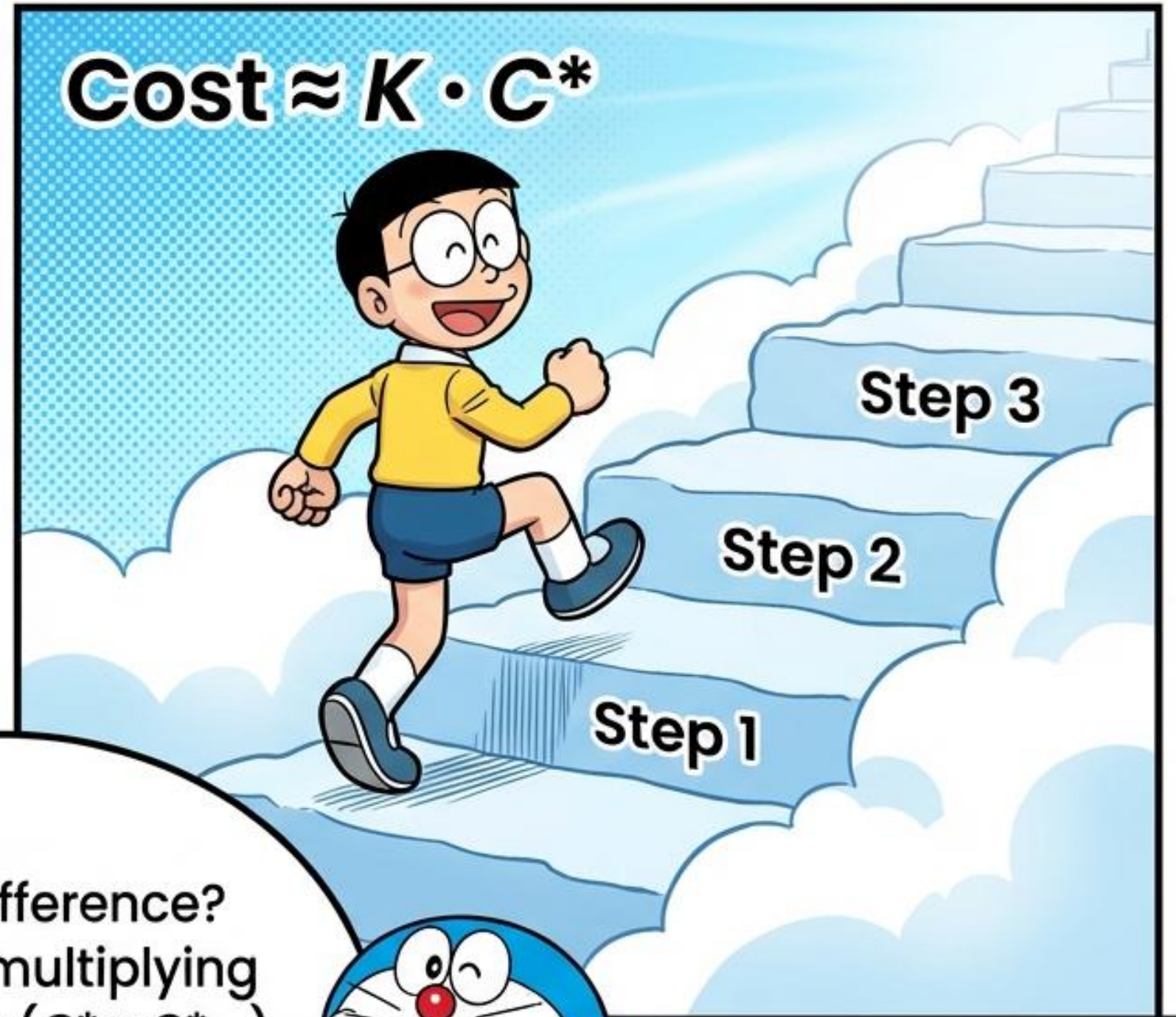


Figure 2. Reasoning tree for parity problems with $d = K = 3$ and input $x_1, x_2, x_3, \text{EOS}$. The nodes on the 2nd–4th levels denote hypotheses about which index is the current secret index, corresponding to x_{i_1}, x_{i_2} , and x_{i_3} , respectively. In our parity CoT class $f \in \mathcal{F}_{2\text{S-ART}}^{\text{parity}}$, each step actually consists of two actions: (i) choose the next secret index i_t ; (ii) after the choice, apply an XOR over z_{t-1} and x_{i_t} as $z_t = \Phi_l(z_{t-1}, x_{i_t}) := z_{t-1} \oplus x_{i_t}$, as formalized in Eq. (6). For visual clarity, the tree only displays the index-selection branches and omits the explicit XOR updates. E denotes the EOS token. For parity tasks, there are illegal children $\notin \mathcal{I}_l$ for each parent that violate the “legal” criteria: non-repeating indices and strictly increasing order ($i_1 < i_2 < i_3 < \dots$); thus, for any parity problem, CoTs with duplicate variables or decreasing index order are illegal.

EXPONENTIAL VS. POLYNOMIAL



EXPONENTIAL NIGHTMARE



POLYNOMIAL WALK IN THE PARK

See the difference?
Instead of multiplying
the difficulty ($C^* \times C^* \dots$),
we just add it ($C^* + C^* \dots$)!



Contribution: A High-level Sketch

➤ High-level Intuition of the cost under direct and curriculum scenario.

To sample π_k from π_{ref} with prob $1 - \delta$, standard rejection sampling needs $\Theta(\|\frac{\pi_k^*}{\pi_{\text{ref}}}\|_\infty \log(\delta^{-1}))$ trials.

Assumptions:

(A1: Realizability): $\|\frac{\pi_{k+1}^*}{\pi_k^*}\|_\infty, \|\frac{\pi_K^*}{\pi_{\text{ref}}}\|_\infty < \infty$

(A2: Suitable Step-wise Difficulty): $\|\frac{\pi_{k+1}^*}{\pi_k^*}\|_\infty \leq \Theta(C^p), \forall k$ for some $p \ll K$ and constant $C > 1$

(A3: Intensive Direct Difficulty): $\|\frac{\pi_K^*}{\pi_{\text{ref}}}\|_\infty \geq \Theta(C^{bK-c})$ for the same constant $C > 1$ $1 \leq b \ll K, c \ll bK$

Conclusions: Direct Learning requires $\Theta(C^{bK-c})$ trials Curriculum Learning requires $\Theta(KC^p)$ trials

For any $\varepsilon > 0$, let $N_\varepsilon(\pi_2 \mid \pi_1)$ denote the ε -precision sample complexity¹ required to reach the optimal policy π^2 starting from a policy π_1 . Consider learning the optimal policy π^* from a base policy π_{ref} via a step-wise curriculum consisting of K intermediate subtask policies

$$\pi_0^* = \pi_{\text{ref}}, \quad \pi_1^*, \quad \dots, \quad \pi_K^* := \pi^*.$$

By definition, curriculum-style post-training is more sample-efficient than direct post-training if and only if the cumulative sample complexity of the intermediate stages is smaller than that of directly learning π^* from π_{ref} , namely,

$$\sum_{k \in [K]} N_\varepsilon(\pi_k^* \mid \pi_{k-1}^*) < N_\varepsilon(\pi^* \mid \pi_{\text{ref}}).$$



Corollary 1 (Sufficient Condition for Curriculum Exponential Improvement). *For any $\varepsilon > 0$, consider a base policy $\pi_{\text{ref}} = \pi_0^*$ and a target optimal policy $\pi^* \neq \pi_{\text{ref}}$. Suppose there exists a curriculum of K subtask optimal policies $\pi_1^*, \dots, \pi_K^* := \pi^*$ such that*

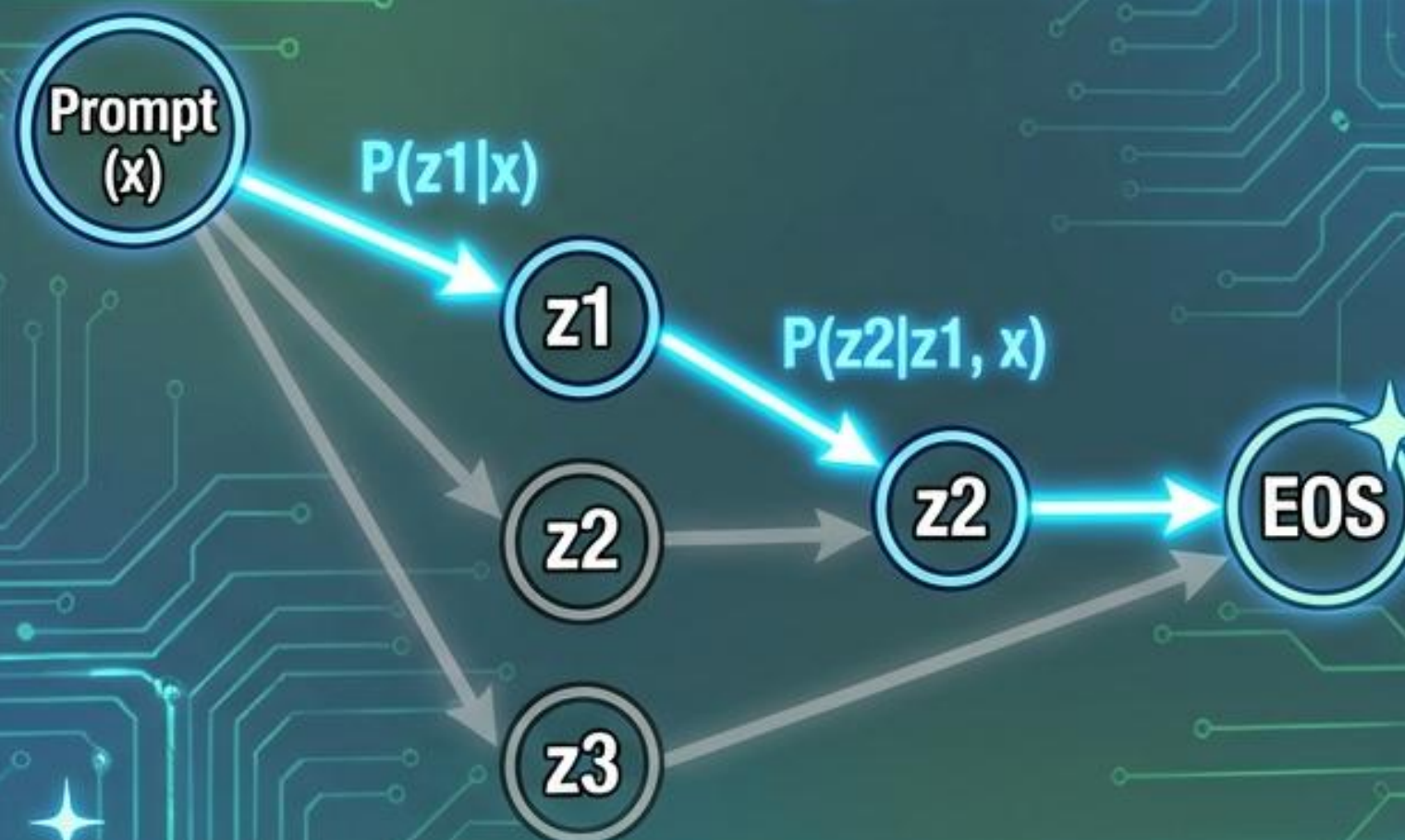
$$N_\varepsilon(\pi_k^* \mid \pi_{k-1}^*) = \Theta(\sqrt[K]{N_\varepsilon(\pi^* \mid \pi_{\text{ref}})}), \quad (1)$$

for all $k \in [K]$. Then, the ratio between the sample complexity of direct estimation $N_\varepsilon^{\text{direct}}$ and that of step-wise curriculum estimation $N_\varepsilon^{\text{curriculum}}$ satisfies

$$\frac{N_\varepsilon^{\text{direct}}}{N_\varepsilon^{\text{curriculum}}} = \frac{N_\varepsilon(\pi^* \mid \pi_{\text{ref}})}{\sum_{k \in [K]} N_\varepsilon(\pi_k^* \mid \pi_{k-1}^*)} = \Theta\left(\frac{(C^*)^K}{KC^*}\right),$$

where $C^ := \sqrt[K]{N_\varepsilon(\pi^* \mid \pi_{\text{ref}})} > 1$.*

2S-ART (Two-Stage Autoregressive Reasoning Tree)



We model your thinking as a **State-Conditioned Autoregressive Reasoning Tree (2S-ART)**. Every step is a choice.

In "Outcome-Only" tasks like math, you only know if you won at the very end (**\$EOS**). One wrong turn early on, and the whole path fails!

Def. 2: Formalizing Chain-of-Thought as a probabilistic path through a decision tree. Sparse rewards mean signal is lost in deep trees.

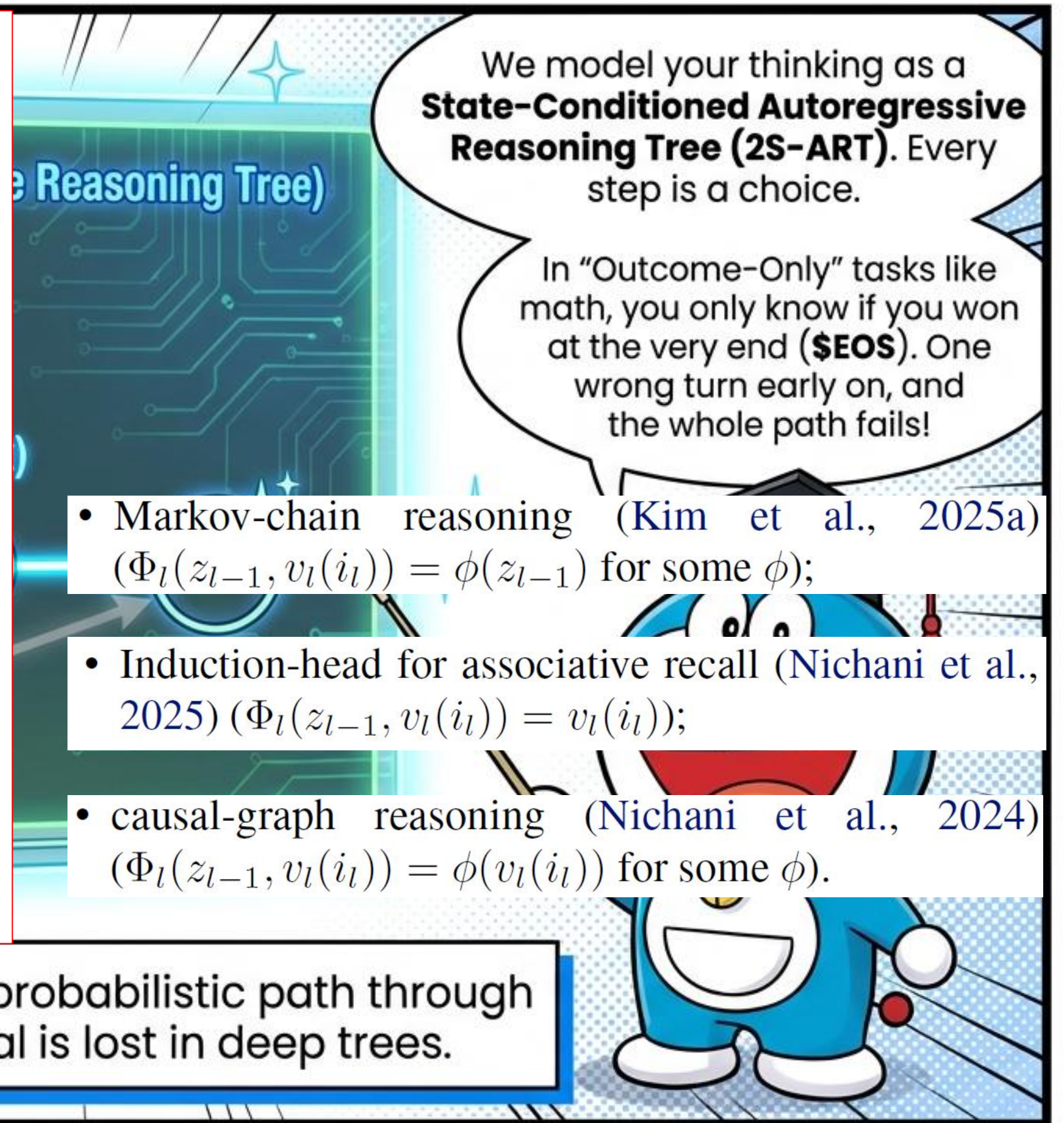
Definition 1 (2-States Conditioned Autoregressive Reasoning Tree (2S-ART)). A 2S-ART, denoted as $\mathcal{F}_{2S-ART}$ ($\{\Phi_l\}_{l \leq L}, \{\mathcal{I}_l\}_{l \leq L}$) is a class of autoregressive tasks. Given input $(\mathbf{x}, \text{EOS})$ where $\mathbf{x} = (x_1, \dots, x_d) \in [K]^d$ is followed by a EOS token, a task $f_{S^k} \in \mathcal{F}_{2S-ART}$ generates a CoT $(z_1, \dots, z_k, \text{EOS})$ deterministically according to its $(k+1)$ -size index path $S^k = (i_1, \dots, i_k, d+1)$ where $k < L$: at step l , f_{S^k} chooses an index $i_l \leq d+1$ from a legal set $\mathcal{I}_l$ with size $\Theta(d)$, reads the corresponding element v_{i_l} from the current sequence, and updates its reasoning state by a two-states map $z_l = \Phi_l(z_{l-1}, v_{i_l})$, with $z_0 := \text{EOS}$, $i_{k+1} = d+1$ and $\Phi_{k+1}(z_k, \text{EOS}) = \text{EOS}$.

The **function value** $f_{S^k}(\mathbf{x})$ of task f_{S^k} is defined by the pre-EOS token, i.e.

$$f_{S^k}(\mathbf{x}) := \Phi_k(z_{k-1}, v_{i_k}).$$

The associated curriculum subtask family is then characterized by the EOS-ended prefix paths $S^l = (i_1, \dots, i_l, d+1)$ with function values $f_{S^l}(\mathbf{x}) := \Phi_k(z_{l-1}, v_{i_l}), l < k$.

Def. 2: Formalizing Chain-of-Thought as a probabilistic path through a decision tree. Sparse rewards mean signal is lost in deep trees.



Definition 1 (2-States Conditioned Autoregressive Reasoning Tree (2S-ART)). A 2S-ART, denoted as $\mathcal{F}_{2S-ART}(\{\Phi_l\}_{l \leq L}, \{\mathcal{I}_l\}_{l \leq L})$ is a class of autoregressive tasks. Given input $(\mathbf{x}, \text{EOS})$ where $\mathbf{x} = (x_1, \dots, x_d) \in [K]^d$ is followed by a EOS token, a task $f_{S^k} \in \mathcal{F}_{2S-ART}$ generates a CoT $(z_1, \dots, z_k, \text{EOS})$ deterministically according to its $(k+1)$ -size index path $S^k = (i_1, \dots, i_k, d+1)$ where $k < L$: at step l , f_{S^k} chooses an index $i_l \leq d+1$ from a legal set $\mathcal{I}_l$ with size $\Theta(d)$, reads the corresponding element v_{i_l} from the current sequence, and updates its reasoning state by a two-states map $z_l = \Phi_l(z_{l-1}, v_{i_l})$, with $z_0 := \text{EOS}$, $i_{k+1} = d+1$ and $\Phi_{k+1}(z_k, \text{EOS}) = \text{EOS}$.

The **function value** $f_{S^k}(\mathbf{x})$ of task f_{S^k} is defined by the pre-EOS token, i.e.

$$f_{S^k}(\mathbf{x}) := \Phi_k(z_{k-1}, v_{i_k}).$$

The associated curriculum subtask family is then characterized by the EOS-ended prefix paths $S^l = (i_1, \dots, i_l, d+1)$ with function values $f_{S^l}(\mathbf{x}) := \Phi_k(z_{l-1}, v_{i_l}), l < k$.

Def. 2: Formalizing Chain-of-Thought as a problem on a decision tree. Sparse rewards mean signal is

Transformer (TF($\cdot$; $\mathbf{W}$)) as Learning Model. Let $\mathbf{U} = [\mu_1, \dots, \mu_K, \mu_{\text{EOS}}] \in \mathbb{R}^{d_x \times (K+1)}$ denote the embedding vocabulary, where the k -th column corresponds to the embedding of token k , and μ_{EOS} denotes the embedding of EOS. Following Li et al. (2025), we adopt mutually orthogonal positional encodings $\mathbf{P} := [\mathbf{p}_1, \dots, \mathbf{p}_{d+1}]$ satisfying $\mathbf{p}_i \perp \mathbf{U}$. Token embeddings and positional encodings are concatenated as $\mathbf{E}[x_m] = [\mu^{x_m}^\top, \mathbf{p}_m^\top]^\top$, $\mathbf{E}[\text{EOS}] = [\mu_{\text{EOS}}^\top, \mathbf{p}_{d+1}^\top]^\top$, and $\mathbf{E}[z_l] = [\mu^{z_l}^\top, \mathbf{p}_{i_l}^\top]^\top$, for all $m \in [d]$, $l \in [L+1]$, and $i_l \in [d+1]$. Given the embedded inputs $\mathbf{E}[x_1], \dots, \mathbf{E}[x_d], \mathbf{E}[\text{EOS}] \in \mathbb{R}^{d_E}$. For notational convenience, we define $\mathbf{E}[z_{-d}] := \mathbf{E}[x_1], \dots, \mathbf{E}[z_{-1}] := \mathbf{E}[x_d]$, and $\mathbf{E}[z_0] := \mathbf{E}[\text{EOS}]$. The transformer performs autoregressive next-token prediction via

$$\text{TF}(\mathbf{E}[z_{-d}], \dots, \mathbf{E}[z_{l-1}]; \mathbf{W}) = \mathbf{E}[z_l], \quad (4)$$

where $\mathbf{W}$ denotes the model parameters. Throughout the reasoning process, the original input embeddings remain fixed, while the CoT states $\mathbf{E}[z_l]$ are generated autoregressively. Specifically, the forward computation at step l is

$$\begin{aligned} \hat{\mathbf{p}}_l &= \sum_{j=-d}^{l-2} \mathbf{V}\left(\mathbf{E}[z_j] \text{softmax}(\mathbf{E}[z_j]^\top \mathbf{K}^\top \mathbf{Q} \mathbf{E}[z_{l-1}])\right) \in \mathbb{R}^{d_x}, \\ \mathbf{p}_{i_l} &\sim \text{softmax}(\hat{\mathbf{p}}_l^\top \mathbf{P} / \beta), \\ \hat{\mu}^{z_l} &= \text{FFN}_l(\mu^{x_{i_l}}, \hat{\mu}^{z_{l-1}}), \quad \mathbf{E}[z_l] = [\hat{\mu}^{z_l}, \mathbf{p}_{i_l}]^\top \in \mathbb{R}^{d_E}. \end{aligned} \quad (5)$$

Definition 1 (2-States Conditioned Autoregressive Reasoning Tree (2S-ART)). A 2S-ART, denoted as $\mathcal{F}_{2S-ART}$ ($\{\Phi_l\}_{l \leq L}, \{\mathcal{I}_l\}_{l \leq L}$) is a class of autoregressive tasks. Given input $(\mathbf{x}, \text{EOS})$ where $\mathbf{x} = (x_1, \dots, x_d) \in [K]^d$ is followed by a EOS token, a task $f_{S^k} \in \mathcal{F}_{2S-ART}$ generates a CoT $(z_1, \dots, z_k, \text{EOS})$ deterministically according to its $(k+1)$ -size index path $S^k = (i_1, \dots, i_k, d+1)$ where $k < L$: at step l , f_{S^k} chooses an index $i_l \leq d+1$ from a legal set $\mathcal{I}_l$ with size $\Theta(d)$, reads the corresponding element v_{i_l} from the current sequence, and updates its reasoning state by a two-states map $z_l = \Phi_l(z_{l-1}, v_{i_l})$, with $z_0 := \text{EOS}$, $i_{k+1} = d+1$ and $\Phi_{k+1}(z_k, \text{EOS}) = \text{EOS}$.

The **function value** $f_{S^k}(\mathbf{x})$ of task f_{S^k} is defined by the pre-EOS token, i.e.

$$f_{S^k}(\mathbf{x}) := \Phi_k(z_{k-1}, v_{i_k}).$$

The associated curriculum subtask family is then characterized by the EOS-ended prefix paths $S^l = (i_1, \dots, i_l, d+1)$ with function values $f_{S^l}(\mathbf{x}) := \Phi_k(z_{l-1}, v_{i_l}), l < k$.

Def. 2: Formalizing Chain-of-Thought as a problem as a decision tree. Sparse rewards mean signal is

CoT and Curriculum Subtask Family for Parity. Per Def. 1, the resulting XOR-based CoT (Wen et al., 2025a; Abedsoltan et al., 2025) is:

$$z_1 = x_{i_1}, z_2 = x_1 \oplus x_{i_2}, \dots, z_k = z_{k-1} \oplus x_{i_k} = f_S(\mathbf{x}), \quad (6)$$

where $i_1 < i_2 < \dots < i_k$. In this case, the curriculum subtask family is the prefix-index family $\mathcal{F}_S = \{f_{S^1}, \dots, f_{S^{k-1}}\}$ with $S^{k'} = (i_1, \dots, i_{k'}, d+1)$ and $f_{S^{k'}}(\mathbf{x}) = z_{k'}$ for $k' \in [k]$; early termination by EOS follows Def. 1. In particular, the legal sets $\{\mathcal{I}_l\}$ is defined to satisfy $i_1 < i_2 < \dots < i_k$.

Transformer (TF($\cdot$; $\mathbf{W}$)) as Learning Model. We follow Eq. (4), Eq. (5), set the vocabulary as $\mathbf{U} = [\mu^0, \mu^1, \mu_{\text{EOS}}] = [e_1, e_2, e_3]$ for 0, 1, EOS, and the FFN_l is instantiated for XOR operation:

$$\text{FFN}_l = \mathbf{W}_2 \text{ReLU}[\mathbf{W}_1(\hat{\mu}^{x_{i_l}} + \mathbf{E}[z_{l-1}]_{:d_X})] \quad (7)$$

where

$$\mathbf{W}_1 = \begin{bmatrix} \frac{1}{2} & \frac{1}{2} & \frac{1}{2} \\ 0 & 1 & 0 \\ -\frac{1}{2} & \frac{1}{2} & -\frac{1}{2} \end{bmatrix}, \quad \mathbf{W}_2 = \begin{bmatrix} 1 & -1 & 2 \\ 0 & 1 & -2 \\ 0 & 0 & 0 \end{bmatrix}. \quad (8)$$

It can be checked directly that this ensure our FFN_l replicates the $\Phi_l(z_{l-1}, v_l(i_l)) = z_{l-1} \oplus v_l(i_l)$.

THE CURRICULUM CONTROLLER

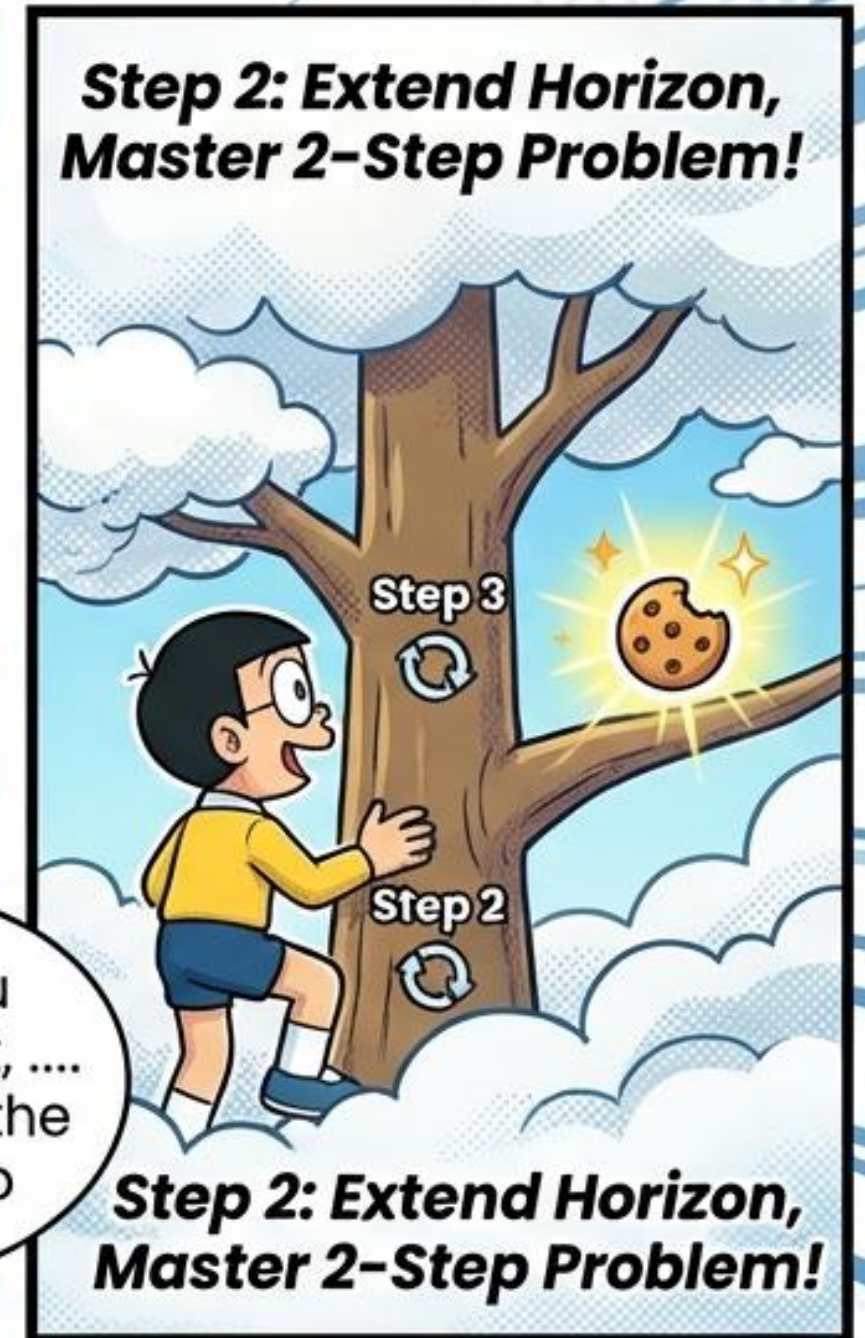
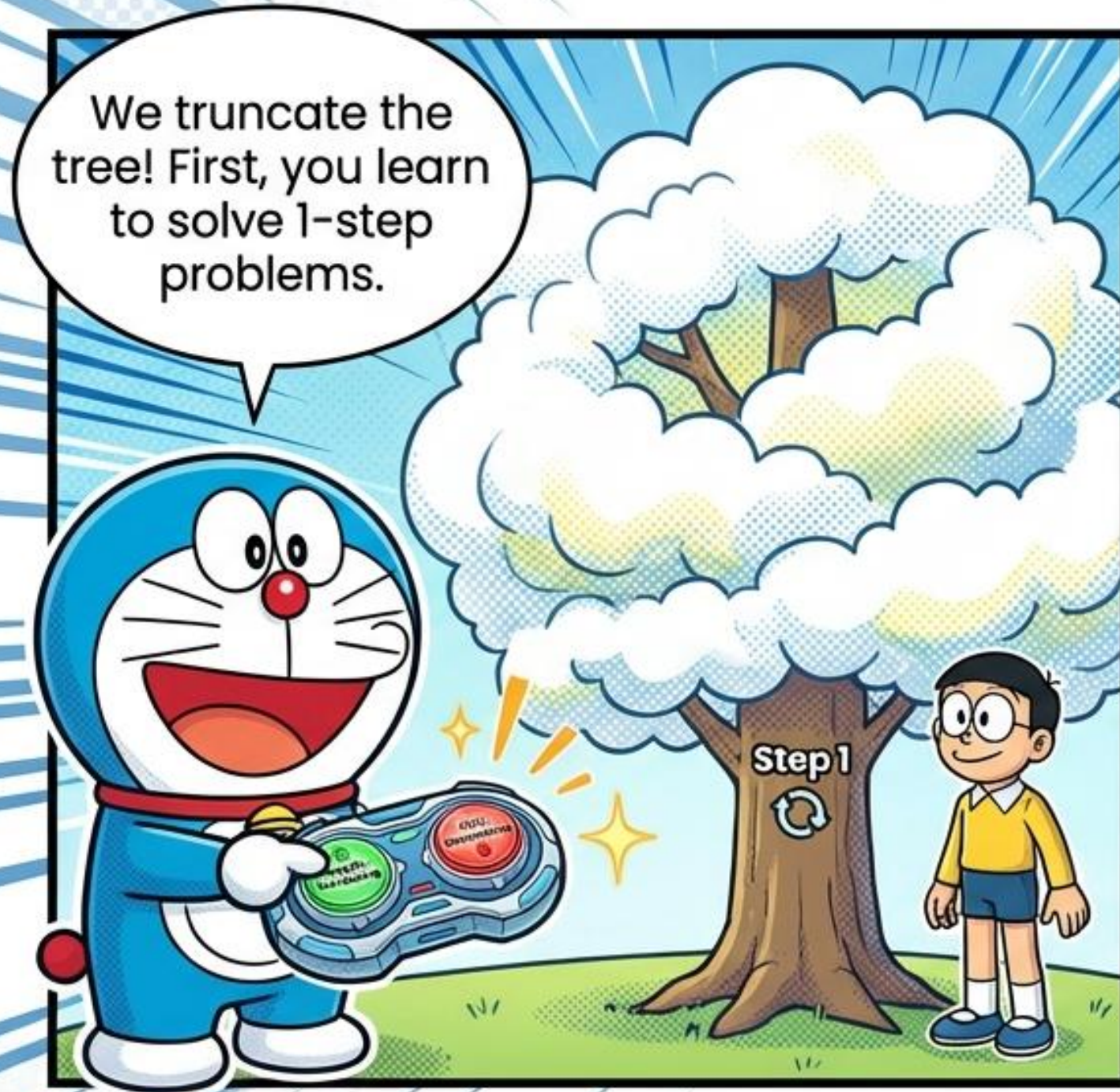
A gadget?
Will it do the
homework for me?

Don't work harder,
Nobita, work smarter!
We need **Curriculum
Post-Training!**

No! It breaks the
"Hard" task into a
sequence of "Easy"
subtasks. We turn the
vertical wall into a
gentle staircase.

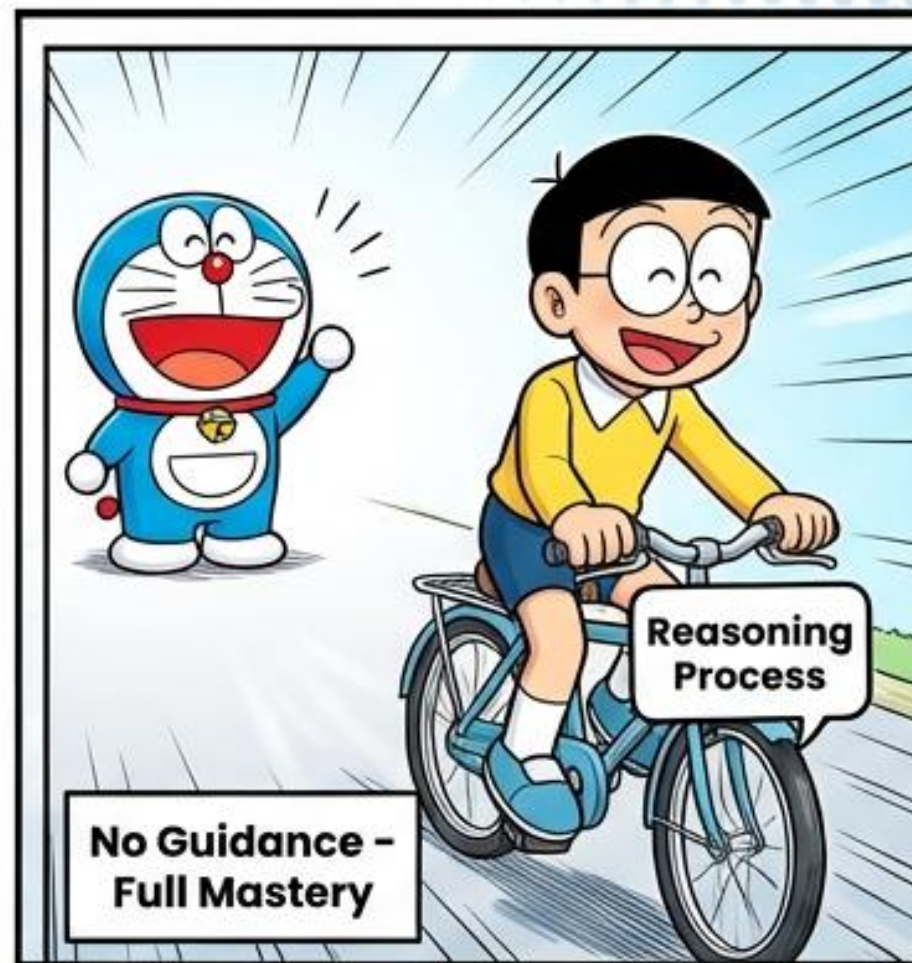
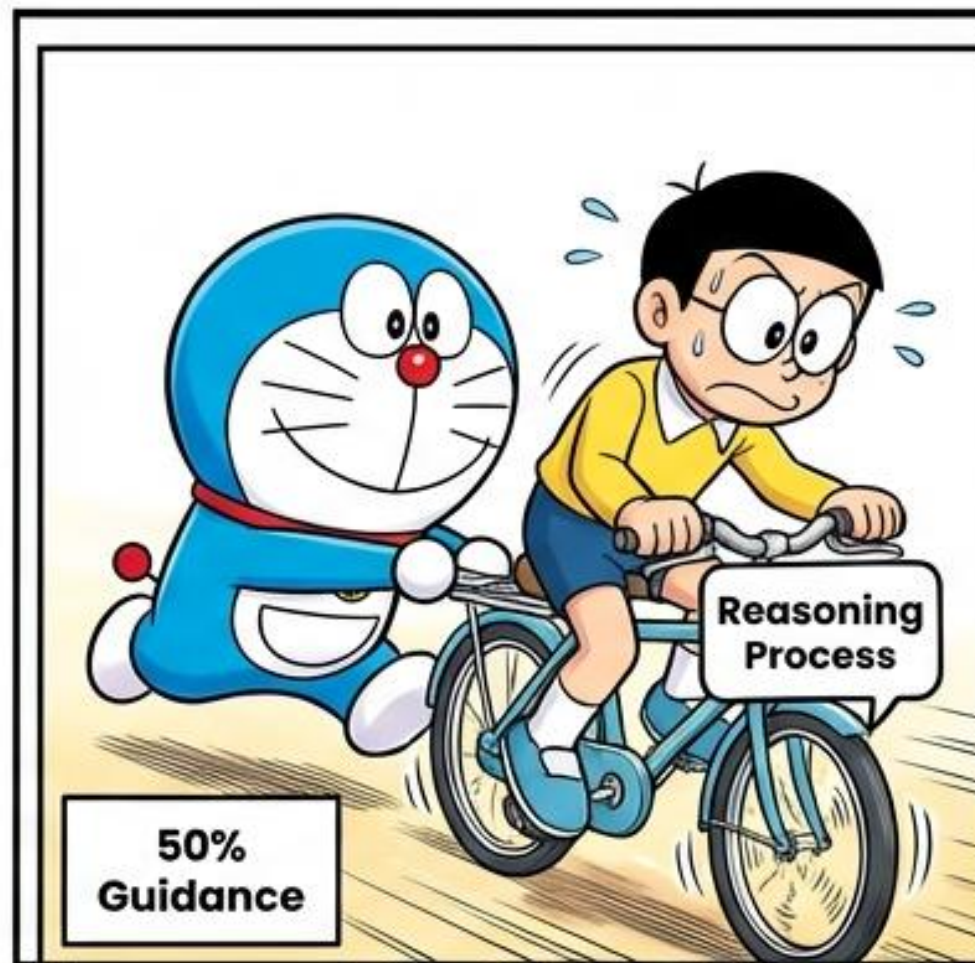
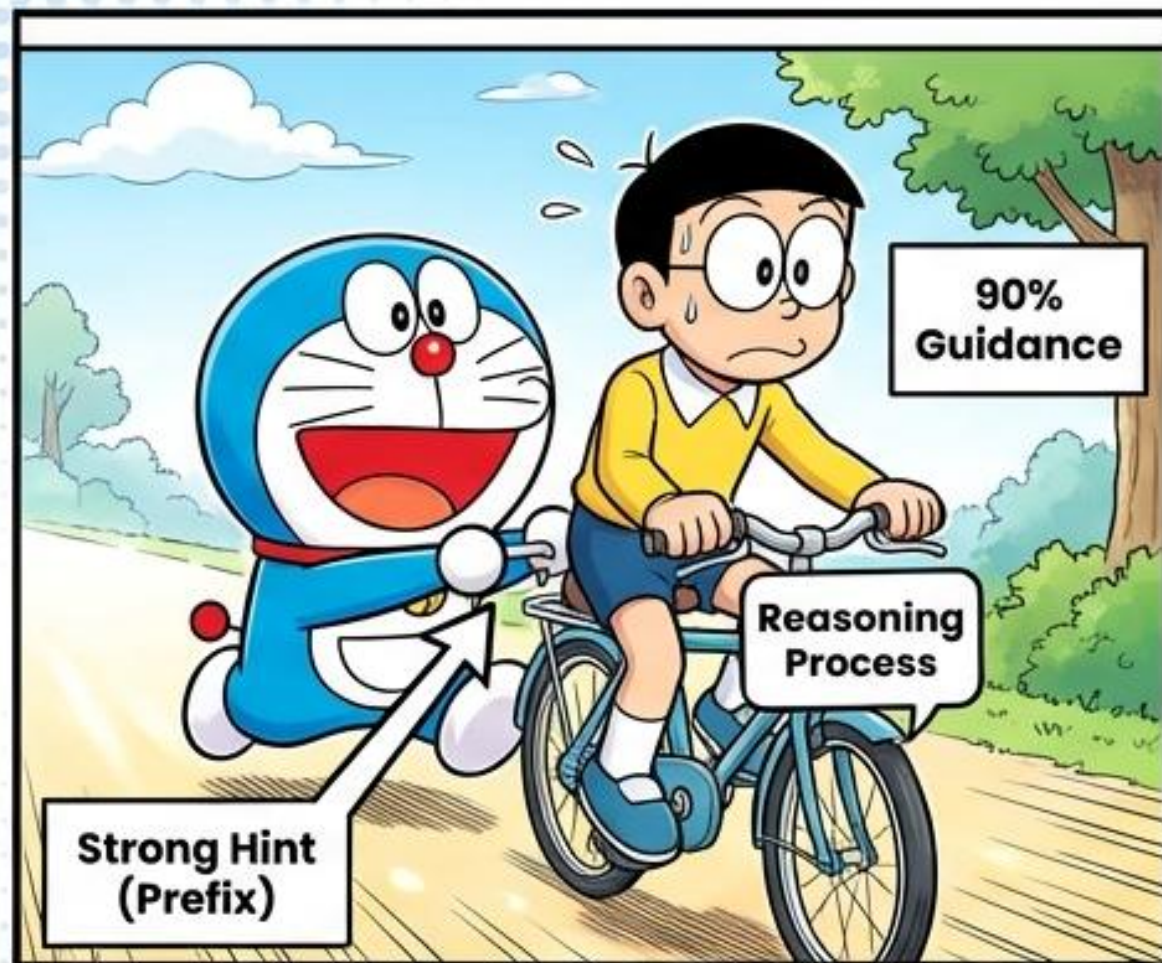
Concept: Transforming the difficult target distribution π^*
into a sequence of manageable steps $\pi_0^*, \pi_1^*, \dots, \pi_K^*$.

GADGET MODE 1: DEPTH-INCREASING



Depth-Increasing Curriculum: Progressively extending the reasoning horizon. The model learns short chains before attempting long ones.

GADGET MODE 2: HINT-DECREASING



Mode 2: **Hint-Decreasing!**
I give you the correct prefix (a hint)
and you just finish the last step.

I get it! It's like reverse-engineering!
You give me fewer hints until I can
do the whole chain myself!

Hint-Decreasing Curriculum: Completing the Chain-of-Thought
from increasingly shorter correct prefixes.



Algorithm 1 No-Curriculum REINFORCE Finetuning (A)

- 1: Inputs: input distribution $\mathbf{x} \sim \text{Unif}([K]^d)$, error budget $\varepsilon > 0$, target length k^* , terminal oracle $\mathbf{R}_{\mathbf{x}}^{f_{S^*}}(\cdot)$, sample complexity $n_0 = \tilde{\Omega}(d^{2k^*+4} (1 - \rho_{\text{spur}}^{\text{sup}, >1})^{-2})$, learning rate $\eta^A = \Theta(\beta \log(d) d^{k^*+1} (1 - \rho_{\text{spur}}^{\text{sup}, >1})^{-1})$.
 - 2: Initialize transformer parameters $\mathbf{W}^{(0)}$.
 - 3: Draw a single batch of n_0 i.i.d. samples $(\mathbf{x}^{(s)}, \tau^{(s)})_{s=1}^{n_0}$ where $\mathbf{x}^{(s)} \sim \text{Unif}([K]^d)$ and $\tau^{(s)} \sim p_{\mathbf{W}^{(0)}}(\cdot \mid \mathbf{x}^{(s)})$.
 - 4: Compute empirical REINFORCE gradient:
5: $\hat{\mathbf{g}} = \frac{1}{n_0} \sum_{s=1}^{n_0} \mathbf{R}_{\mathbf{x}^{(s)}}^{f_{S^*}}(\hat{\mathbf{p}}^{z_{k^*}}(\tau_{-2}^{(s)})) \sum_{\ell=1}^{k^*+1} \nabla_{\mathbf{W}} \log \pi_{\mathbf{W}^{(0)}}(i_{\ell}^{(s)} \mid \mathbf{x}^{(s)}, \hat{\mathbf{p}}^{z_{1:\ell}})$.
 - 6: Update parameters: $\mathbf{W}^{(1)} = \mathbf{W}^{(0)} + \eta^A \cdot \hat{\mathbf{g}}$.
 - 7: Return $\mathbf{W}^{(1)}$.
-

Algorithm 2 Depth-Increasing Curriculum REINFORCE Finetuning (B)

- 1: Inputs: input distribution $\mathbf{x} \sim \text{Unif}([K]^d)$, error budget $\varepsilon > 0$, target length k^* , family oracle $\mathbf{R}_{\mathbf{x}}^{\mathcal{F}_{S^*}}(\cdot)$, per-step sample complexity $n_{\ell} = \tilde{\Theta}(d^2 (1 - \rho_{\text{spur}}^{\text{sup}, \ell})^{-2})$, learning rates $\eta_{\ell}^B = \Omega(\beta \log(d) d^2 (1 - \rho_{\text{spur}}^{\text{sup}, \ell})^{-1})$.
 - 2: Initialize transformer parameters $\mathbf{W}^{(0)}$.
 - 3: **for** $\ell = 1$ to $k^* + 1$ **do**
 - 4: Draw fresh samples $(\mathbf{x}^{(s)}, \tau^{(s)})_{s=1}^{n_{\ell}}$ where $\mathbf{x}^{(s)} \sim \text{Unif}([K]^d)$ and $\tau^{(s)} \sim p_{\mathbf{W}^{(\ell-1)}}(\cdot \mid \mathbf{x}^{(s)})$.
 - 5: If $\ell < k^* + 1$, for each sample s , truncate trajectory at length $\ell + 1$:
 - 6: If $\tau^{(s)}$ contains EOS at position $\leq \ell$, discard sample.
 - 7: Otherwise, append EOS to create truncated sequence $\tau_{\text{trunc}}^{(s)}$.
 - 8: If $\ell = k^* + 1$, $\tau_{\text{trunc}}^{(s)} = \tau^{(s)}$.
 - 9: Compute empirical REINFORCE gradient for step ℓ :
10: $\hat{\mathbf{g}}_{\ell} = \frac{1}{n_{\ell}} \sum_{s=1}^{n_{\ell}} \mathbf{R}_{\mathbf{x}^{(s)}}^{\mathcal{F}_{S^*}}(\tau_{\text{trunc}}^{(s)}) \nabla_{\mathbf{W}} \log \pi_{\mathbf{W}^{(\ell-1)}}(i_{\ell}^{(s)} \mid \mathbf{x}^{(s)}, \hat{\mathbf{p}}^{z_{1:\ell}})$.
 - 11: Update parameters: $\mathbf{W}^{(\ell)} = \mathbf{W}^{(\ell-1)} + \eta_{\ell}^B \cdot \hat{\mathbf{g}}_{\ell}$.
 - 12: **end for**
 - 13: Return $\mathbf{W}^{(k^*+1)}$.
-

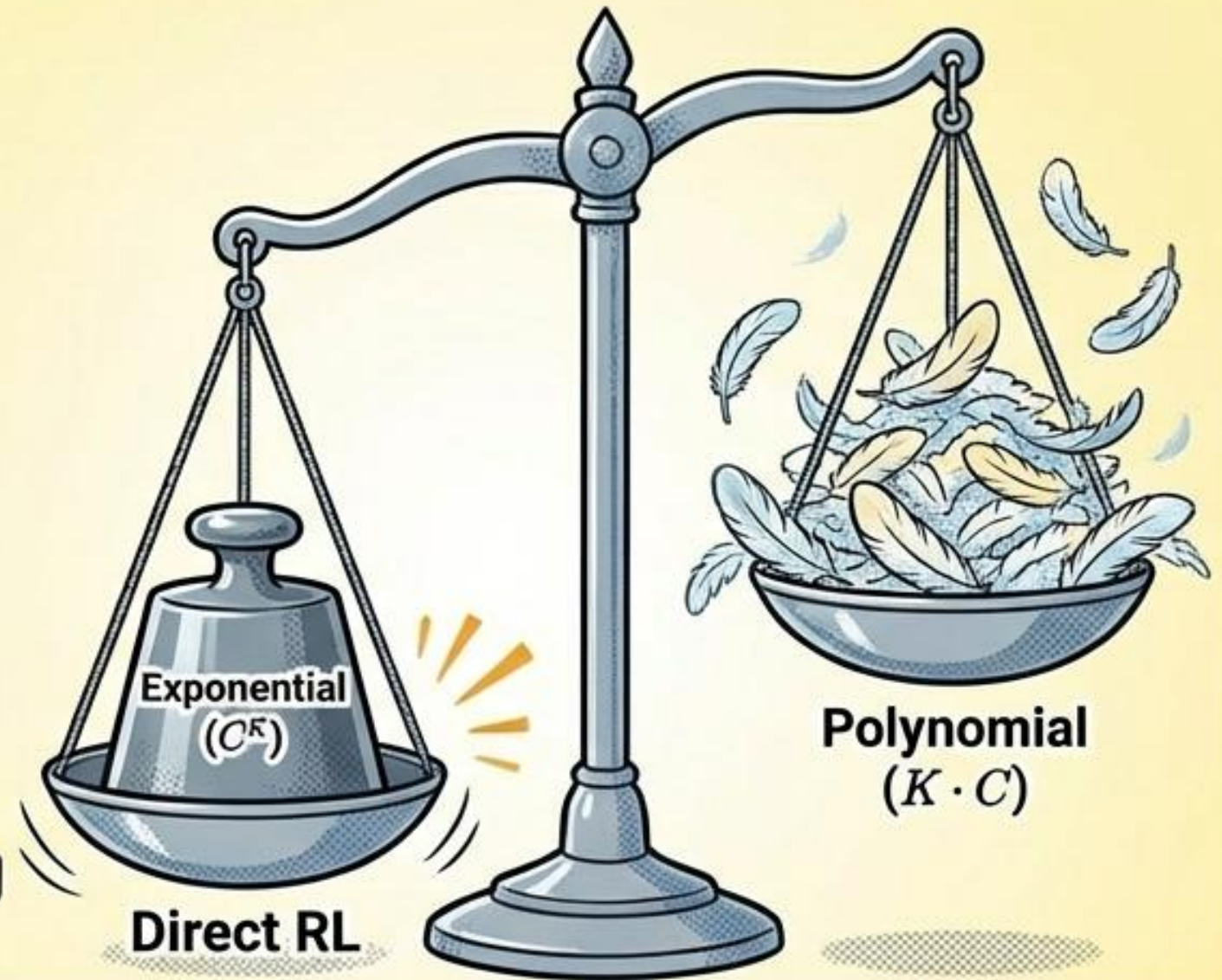
THE SECRET SAUCE: COVERAGE COEFFICIENT.

Here is the secret sauce:
The **Coverage Coefficient** (C_{cov}).

$$N_{\text{curriculum}} \ll N_{\text{direct}}$$



It measures how 'surprised' you are by the answer. Direct learning has infinite surprise. But with the Curriculum Gadget, the 'surprise' at each step is small and bounded.



Theorem 2: Curriculum strategies achieve high accuracy with **Polynomial** sample complexity, avoiding the exponential bottleneck of direct training.

THE SECRET SAUCE: COVERAGE COEFFICIENT.

Theorem 3 (Curriculum RL Finetuning Avoid Exponential Bottleneck). *Let $\mathbf{U}$ be an orthogonal matrix with $d_X = \Theta(K)$, $d_E = \Theta(d + L)$, $\text{TF}(\cdot; \mathbf{W}_{base})$ per in Thm. 2 with trainable $\mathbf{W}$, and a target $f_{S_*} \in \mathcal{F}_{2S-ART}$ with $|S_*| = k^* + 1$. For any $\varepsilon \in (0, 1)$, using the RL objective in Eq. (9) and one gradient step per stage (a single step for no-curriculum; $k^* + 1$ online steps for curricula), with probability no less than $1 - \delta$, there exist learning-rate choices η such that*

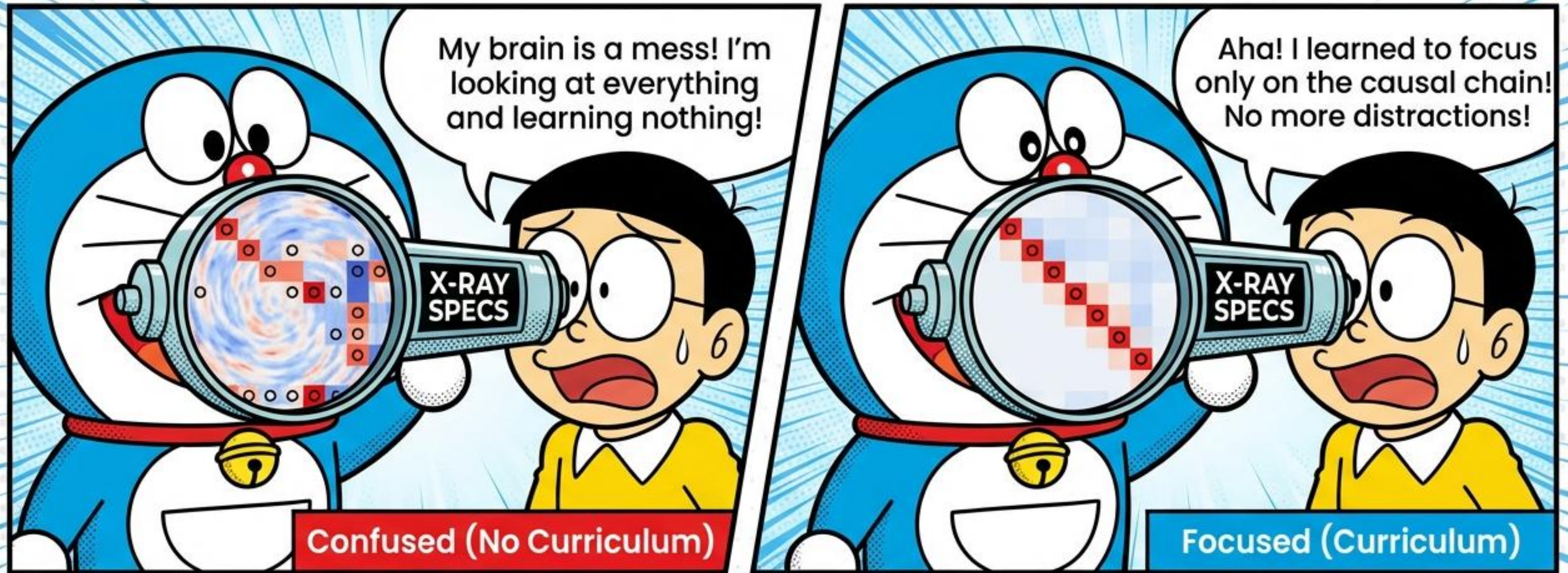
1. *No-curriculum ($R_{\mathbf{x}}^{f_{S_*}}(\cdot)$ as oracle, one-shot update with $\eta = \tilde{\Theta}(\beta d^{k^*+1})$): the sample complexity to achieve $\mathbb{E}_{\mathbf{x} \sim \text{Unif}([K]^d)} [R_{\mathbf{x}}^{f_{S_*}}(\text{TF}(\mathbf{x}; \mathbf{W}))] \geq 1 - \varepsilon$ is at least $n \geq \tilde{\Omega}(d^{2k^*+2})$.*
2. *Depth-increasing Curriculum and Hint-decreasing Curriculum ($\mathbf{R}_{\mathbf{x}}^{\mathcal{F}_{S_*}}(\cdot)$ as oracle, $k^* + 1$ online updates with $\eta = \tilde{\Theta}(\beta d^2)$): the sample complexity to achieve $\mathbb{E}_{\mathbf{x} \sim \text{Unif}([K]^d)} [R_{\mathbf{x}}^{f_{S_*}}(\text{TF}(\mathbf{x}; \mathbf{W}))] \geq 1 - \varepsilon$ is at most $n \leq \tilde{O}((k^* + 1)d^2)$.*

Here $\tilde{\Omega}(\cdot)$ and $\tilde{O}(\cdot)$ hide polylogarithmic factors in $d, 1/\delta$ and $1/\varepsilon$, and absolute constants (e.g., temperature β) as well as task-dependent spurious-acceptance parameters.

surprise at each step is small and bounded.

complexity, avoiding the exponential bottleneck of direct training.

SIMULATION: INSIDE THE ELECTRONIC BRAIN



Empirical Evidence: Curriculum training leads to clean attention heads that track the causal chain, while direct training gets lost in noise.

SIMULATION: INSIDE THE ELECTRONIC BRAIN

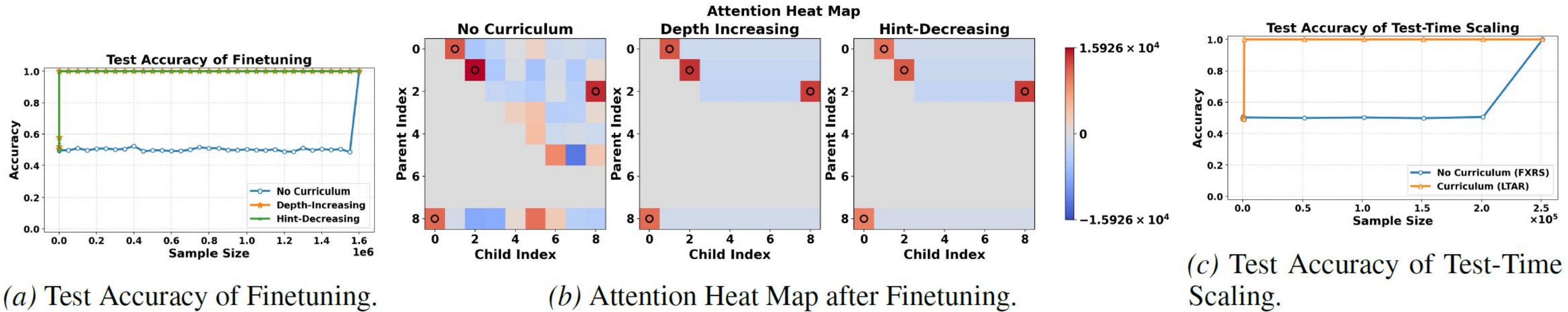
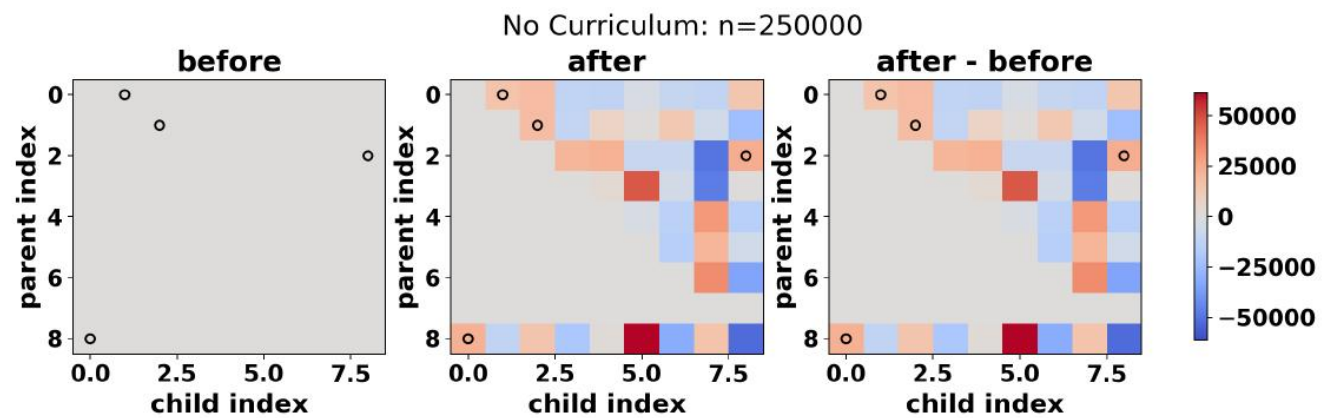
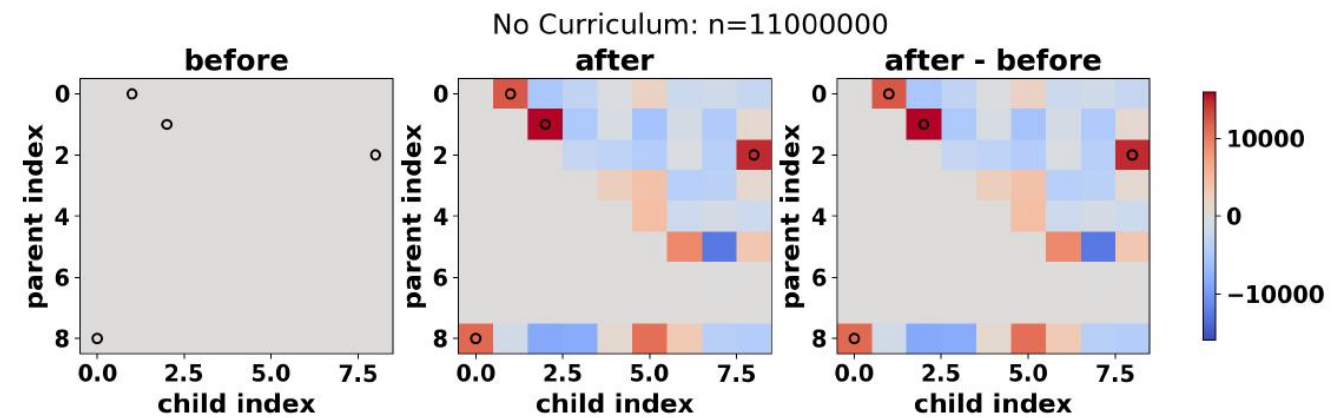


Figure 3. Dynamics of Finetuning and Test-time Scaling on the parity task ($d = 8$, $k^* = 3$, $S^* = [0, 1, 2]$). (a) Test accuracy evolution of finetuning along sample size; (b) attention heatmap after finetuning; (c) test accuracy evolution of test-time scaling along reward oracle query complexity. For finetuning, both curriculum algorithms (Alg. 1 & Alg. 2) converge at sample size $n = 76$, where no-curriculum counterpart (Alg. 1) converges at $n = 1600076$; Attention learning of curriculum algorithms focus on crucial parent-children maps, while no curriculum counterpart is contaminated with spurious correlations; For test-time scaling, curriculum algorithm (LTAR in Alg. 6) converges at $n = 1410$ while its no curriculum counterpart (Alg. 4) converges at $n = 251410$.

Empirical Evidence: Curriculum training leads to clean attention heads that track the causal chain, while direct training gets lost in noise.

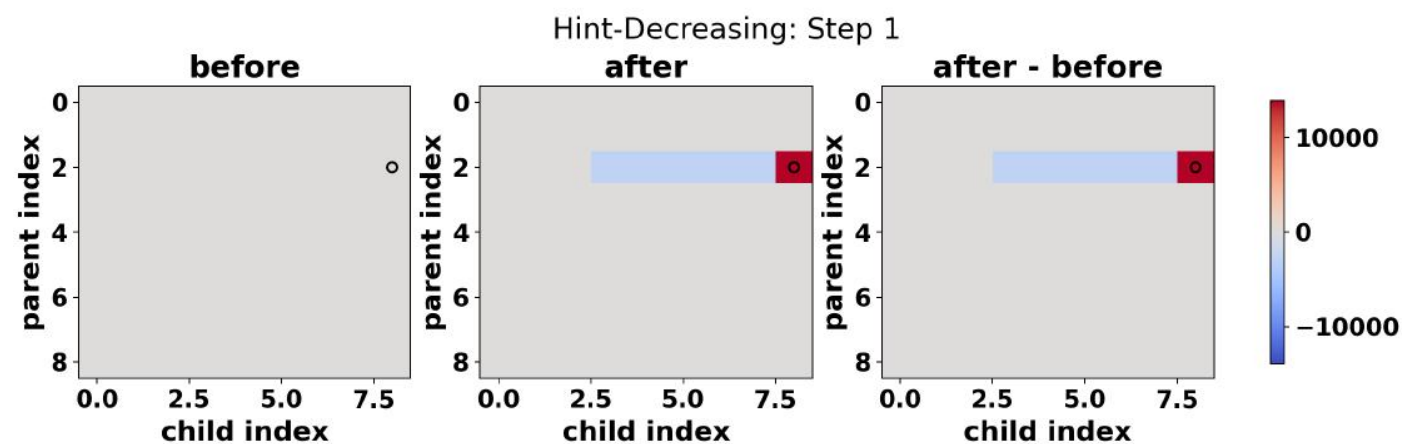


(a) $n = 250000$.

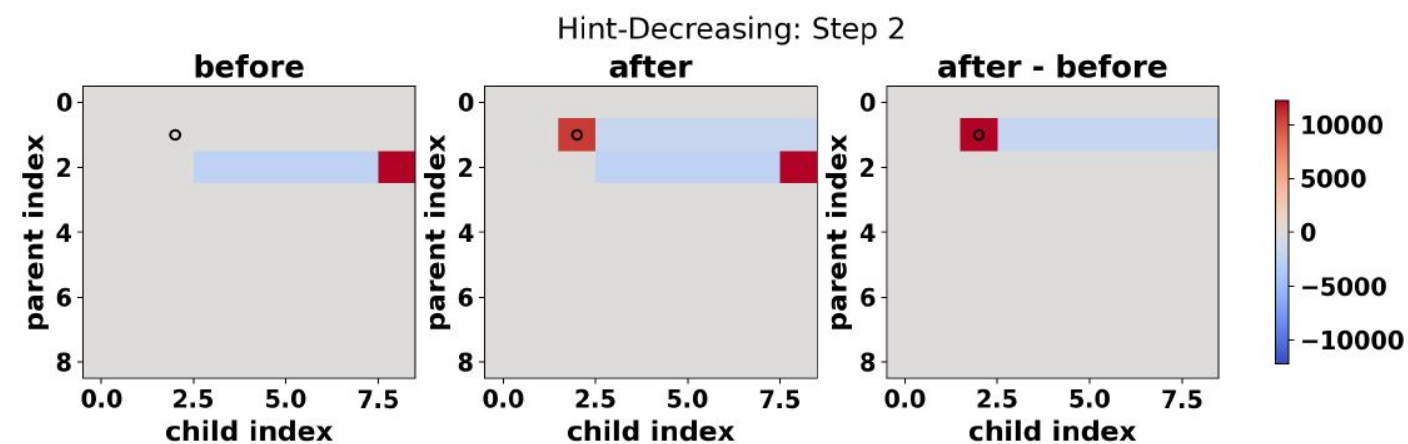


(b) $n = 11000000$.

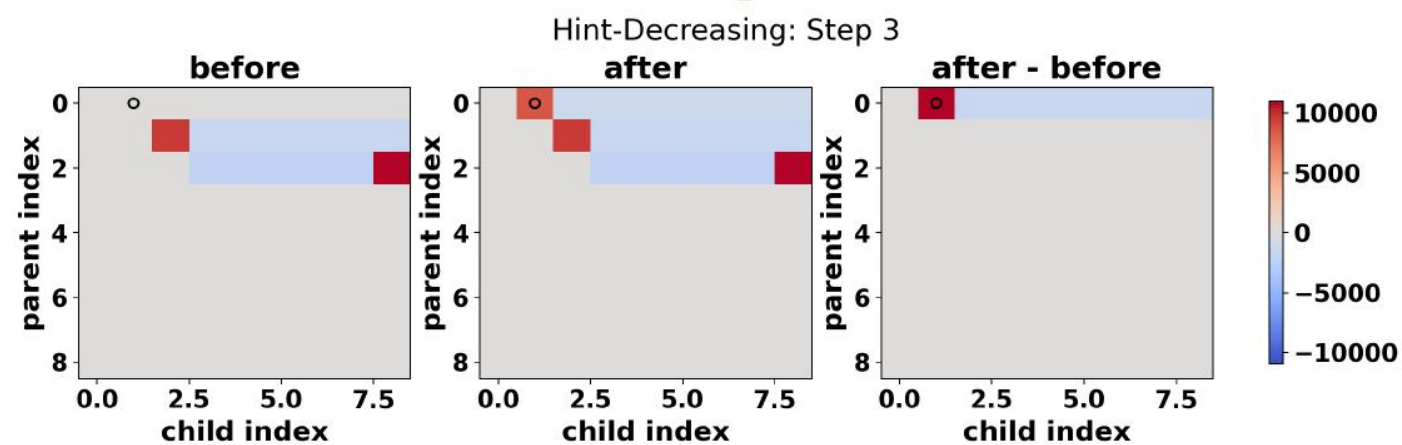
Figure 4. Attention Heatmaps of No-Curriculum (Direct Learning) with different sample complexity.



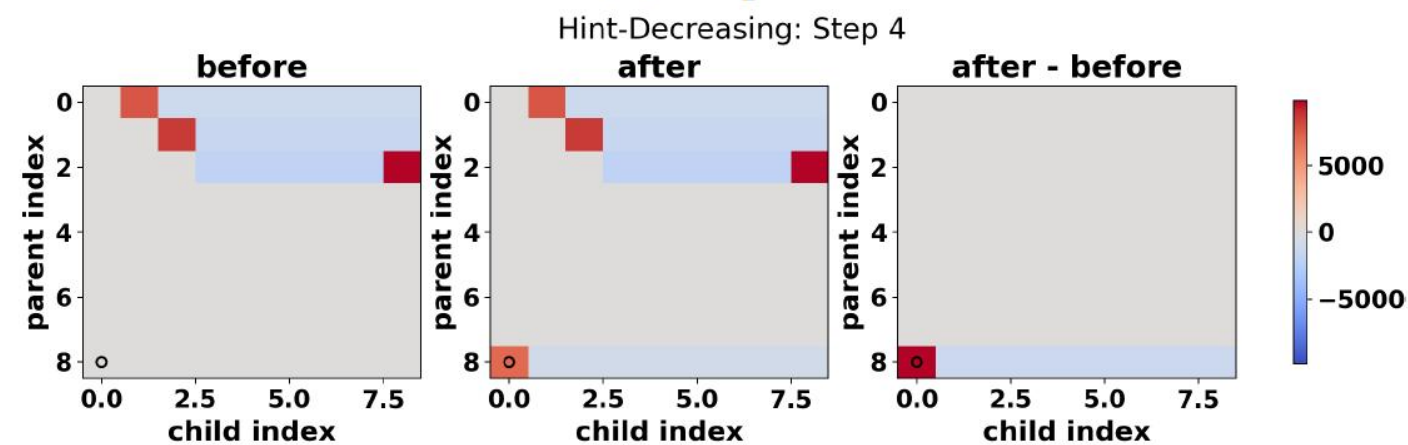
(a) Step 1.



(b) Step 2.

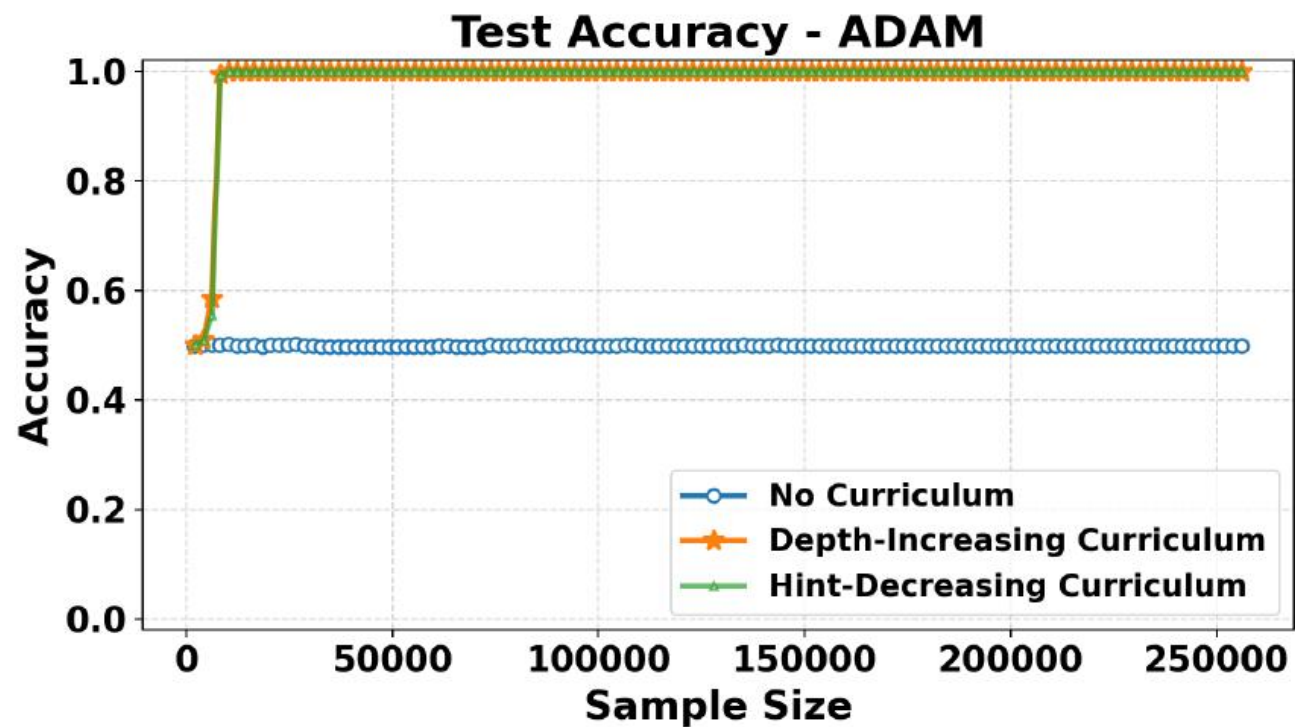


(c) Step 3.

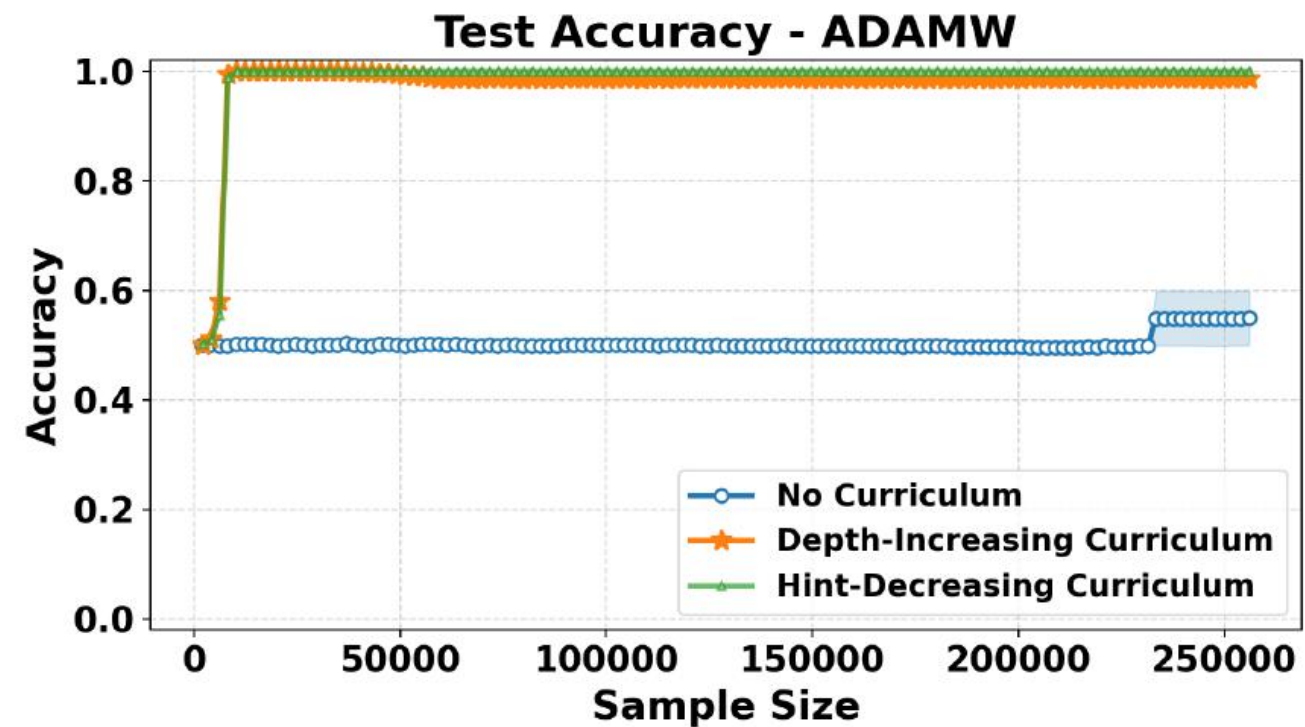


(d) Step 4.

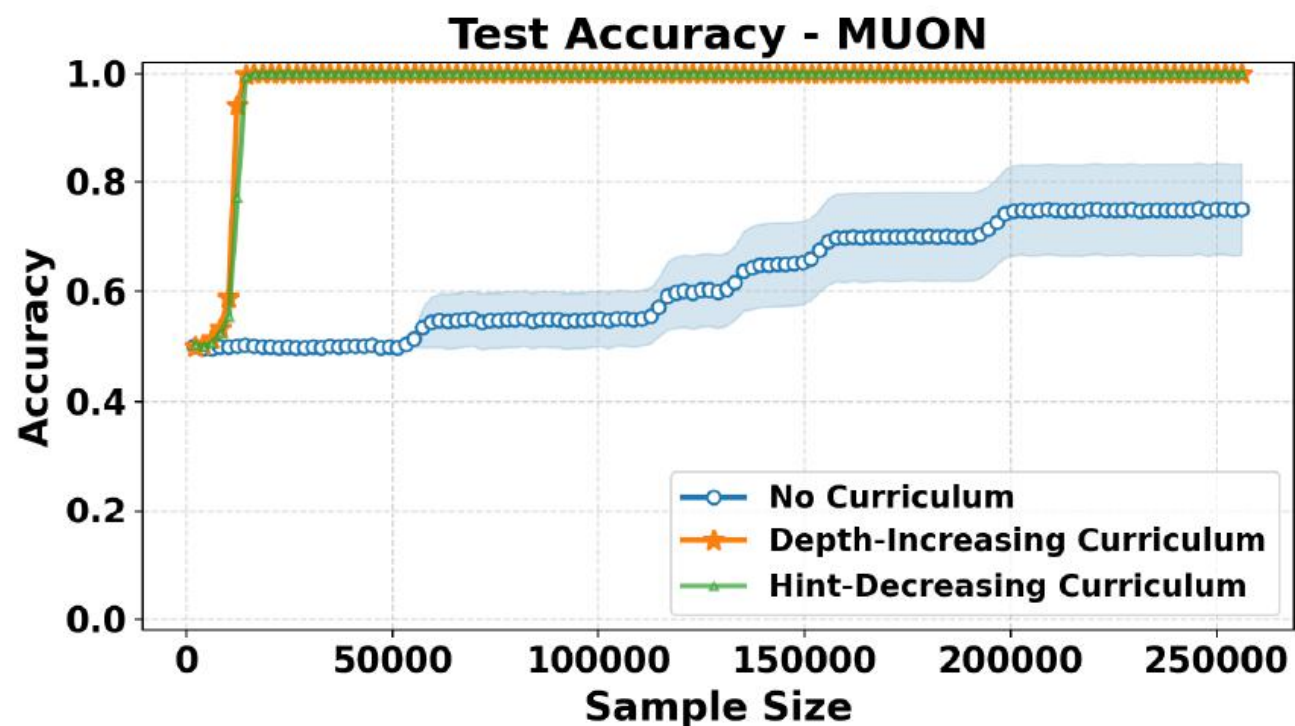
Figure 7. Step-wise Attention Heatmaps for Hint-Decreasing Curriculum ($d = 8, k^* = 3, S^* = [0, 1, 2], n = 100000$).



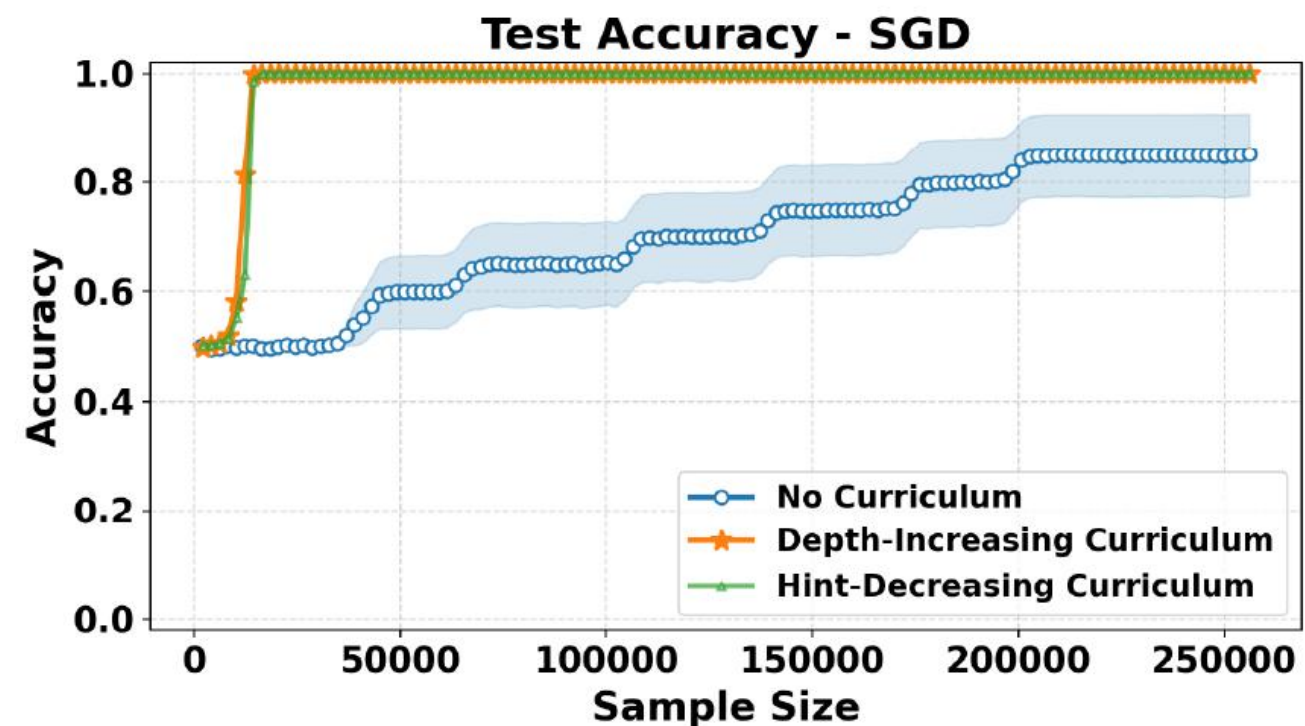
(a) Adam



(b) AdamW

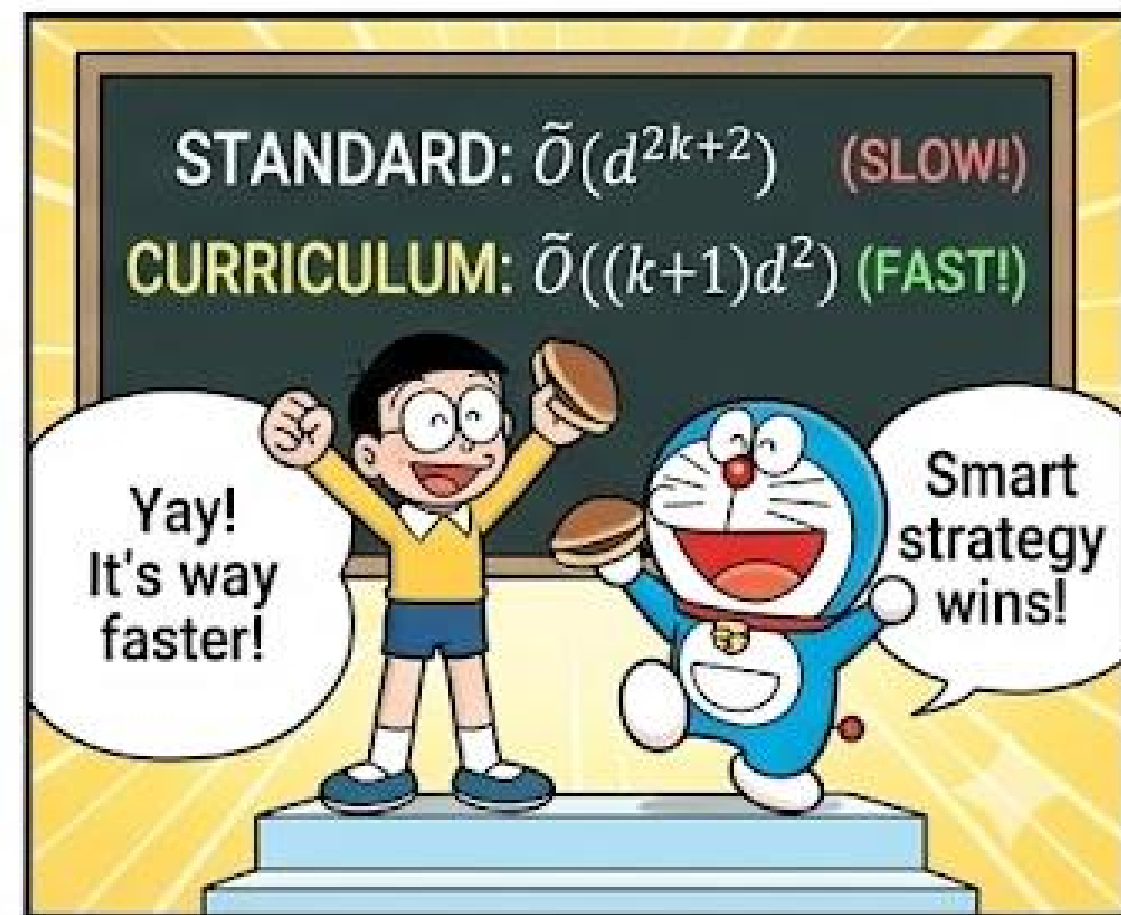
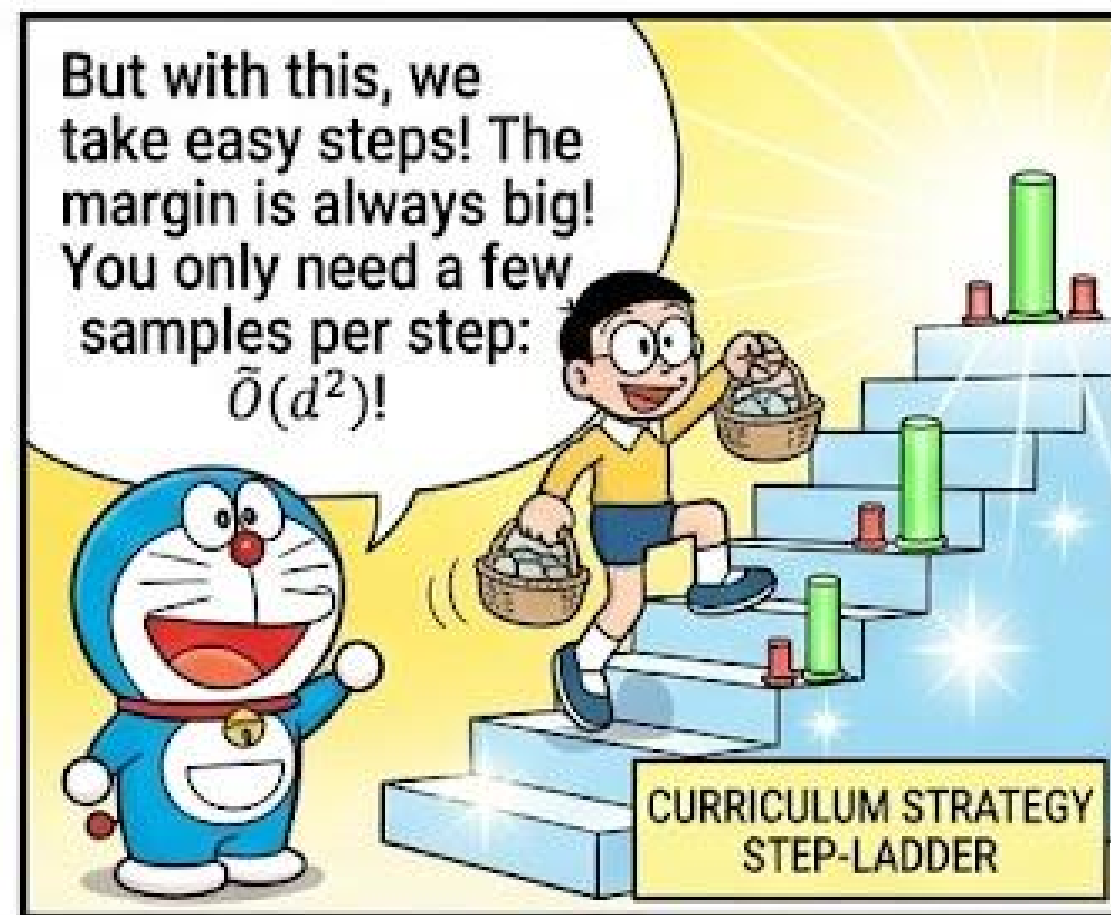
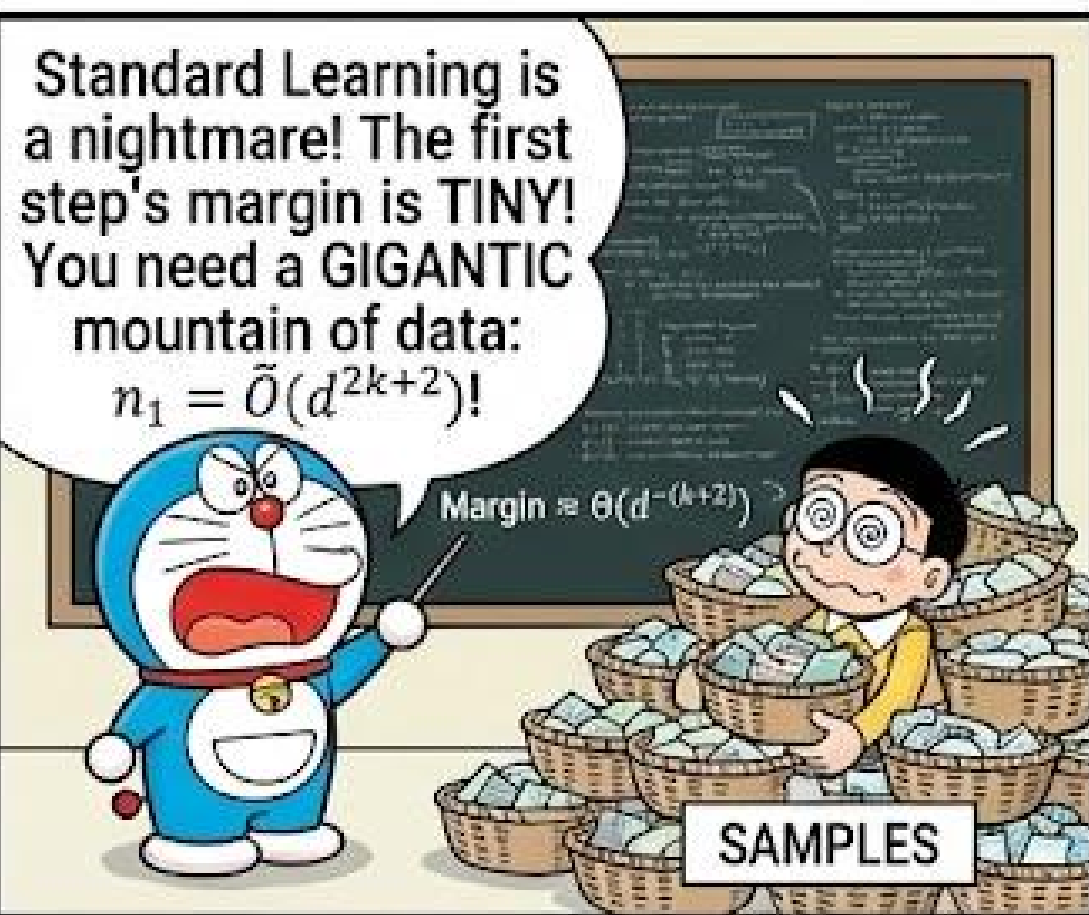
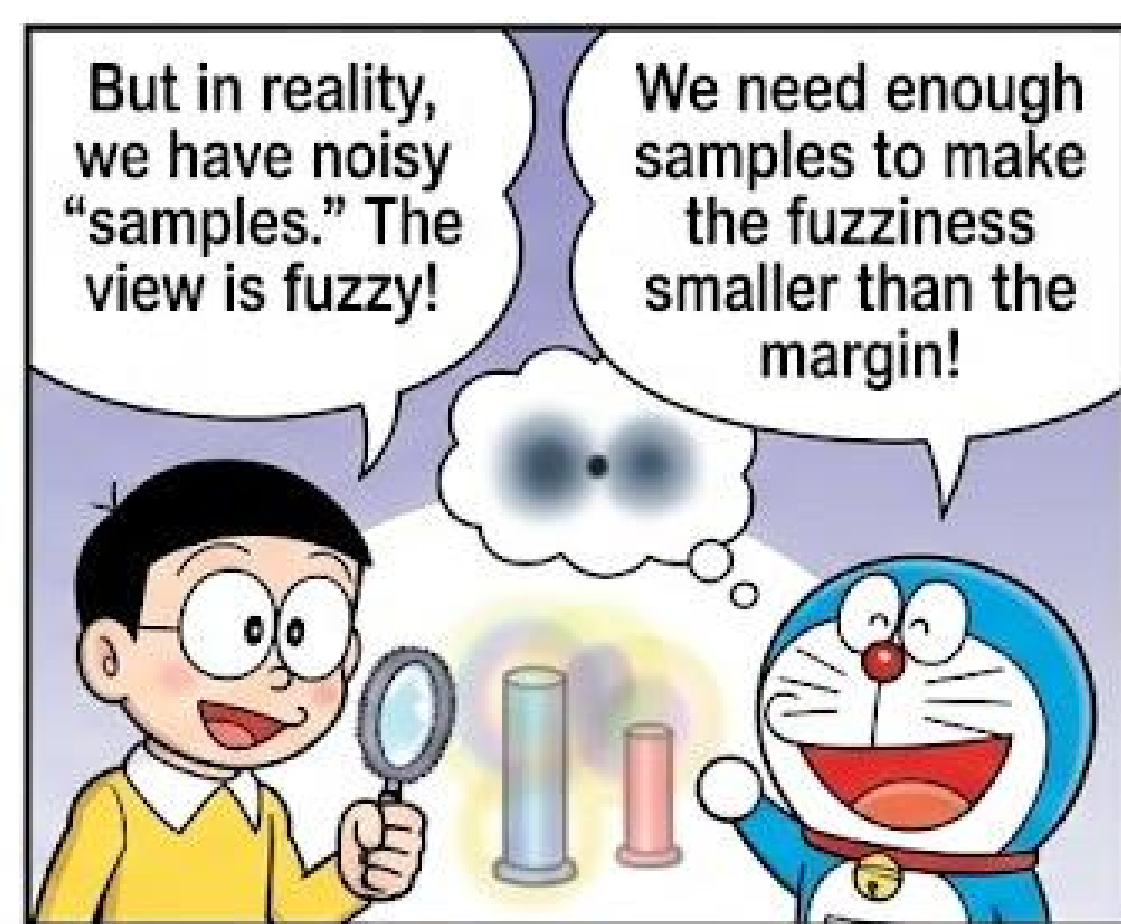
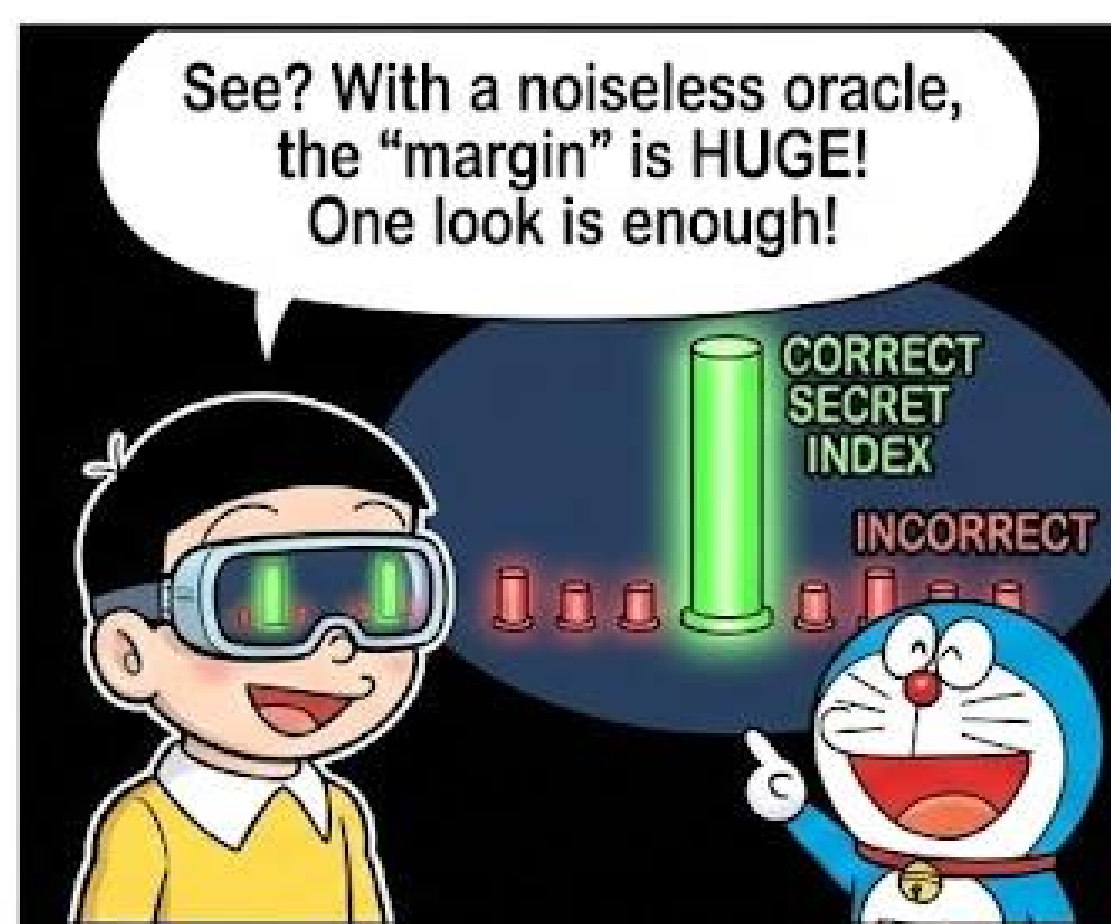
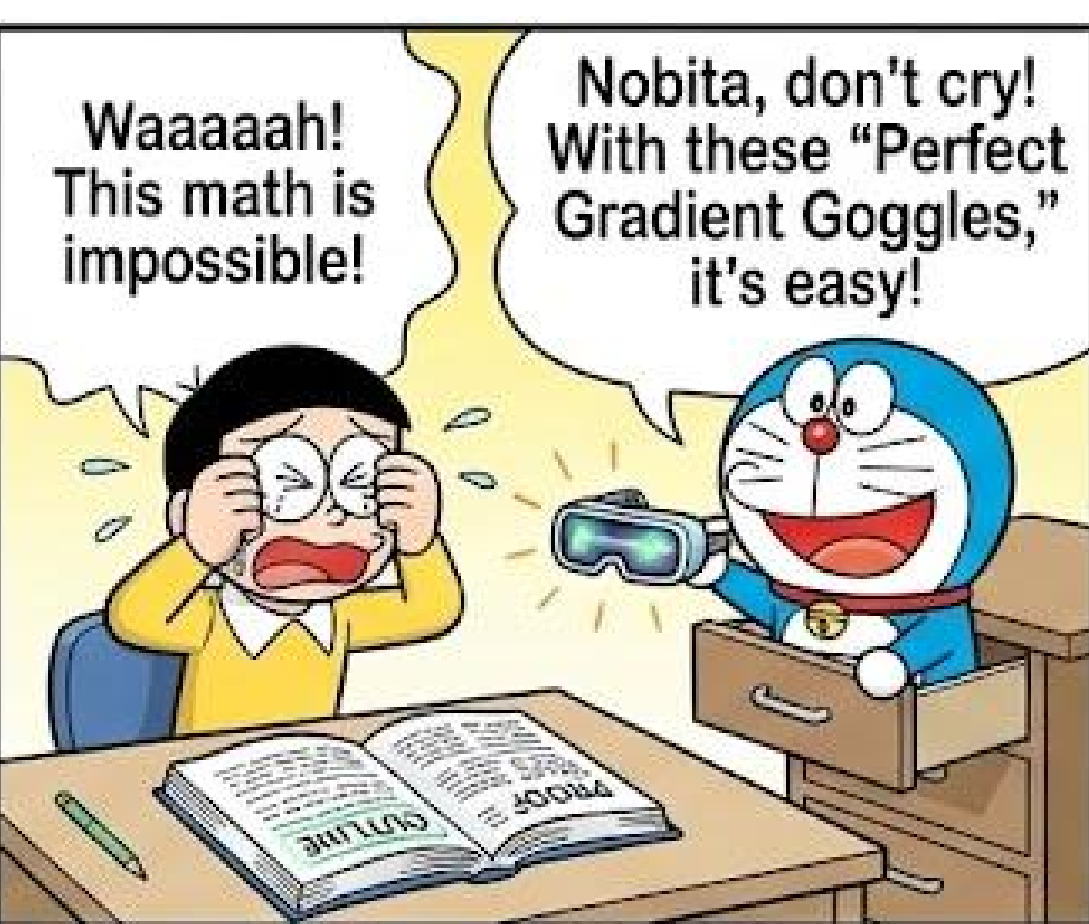


(c) Muon



(d) SGD

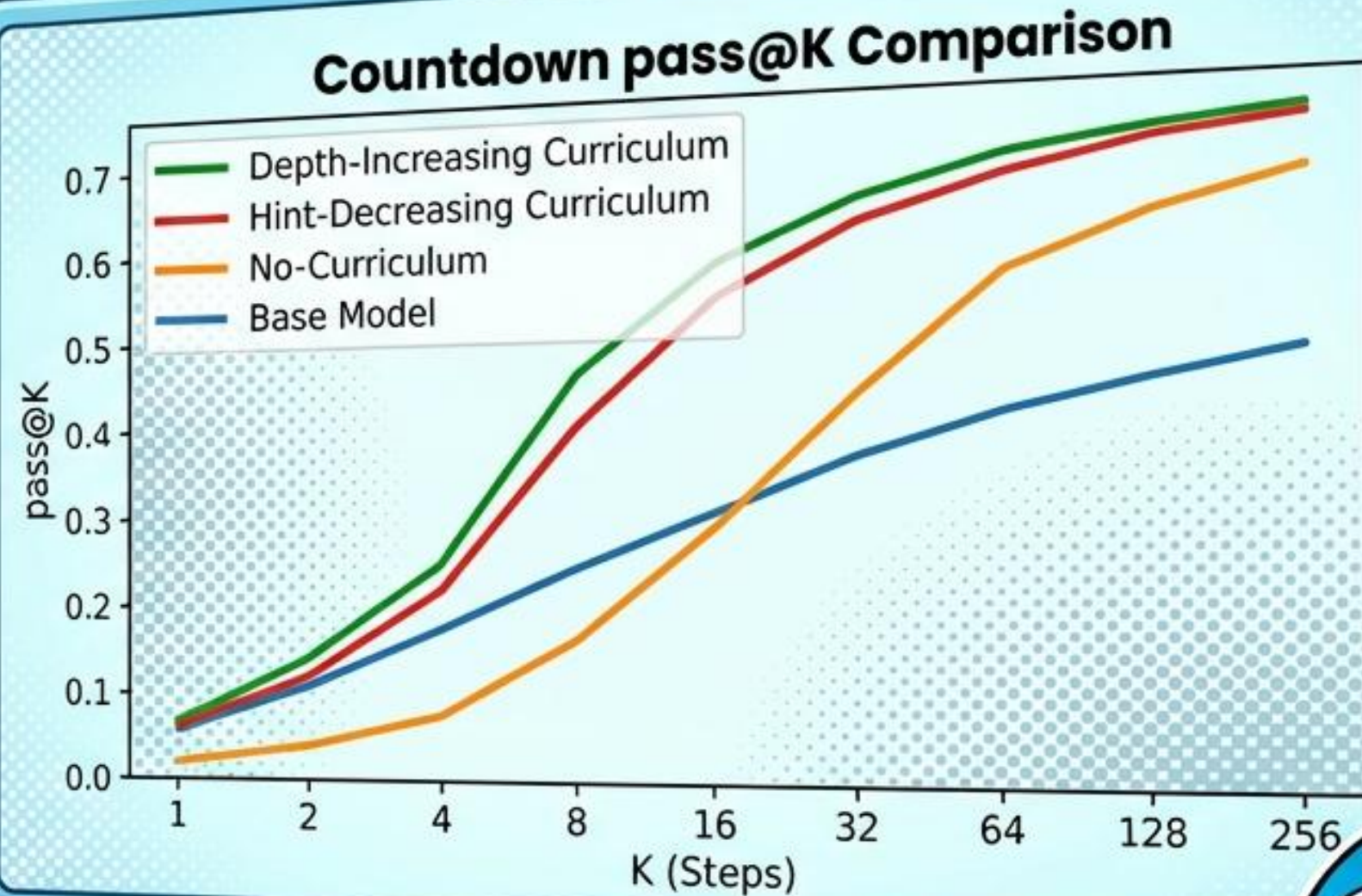
Figure 11. Small Learning rate regime ($d = 8$, $k^* = 3$, $S^* = [0, 1, 2]$, $\eta = 0.5$, $B = 256$).



****THE COUNTDOWN GAME CHALLENGE****

Countdown

Look at my score!
I can solve the
"Hard" level (5
operations) now!



Precisely. The No-Curriculum
version barely improved, but
the Curriculum methods
reached over 70% accuracy!

****THE COUNTDOWN GAME CHALLENGE****

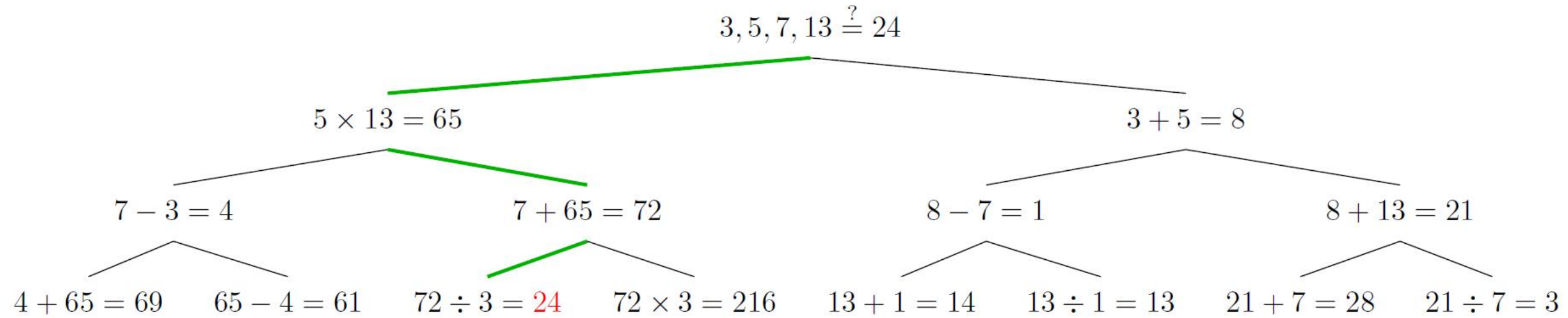
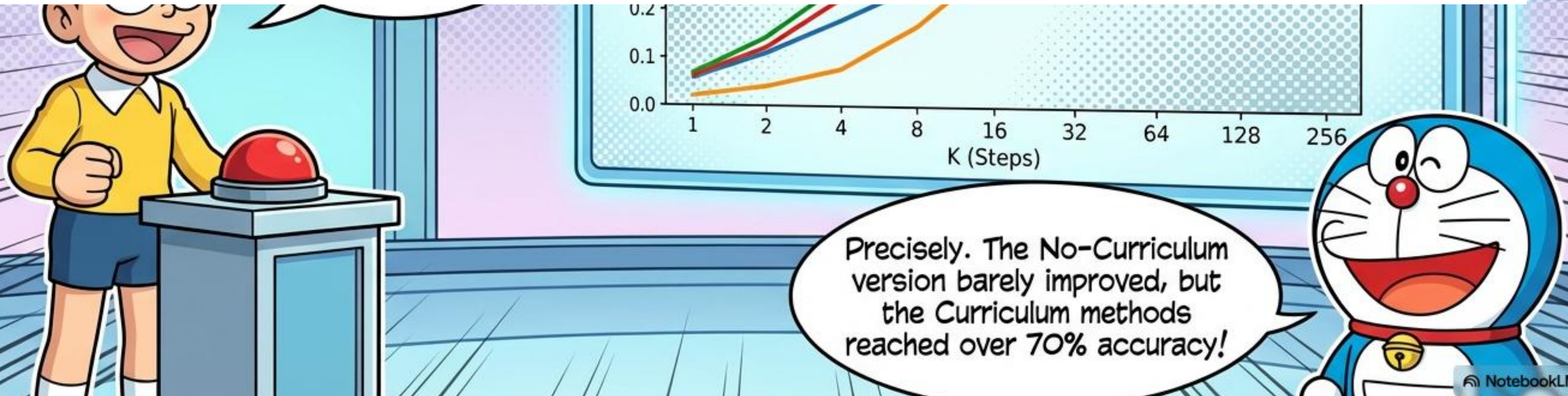
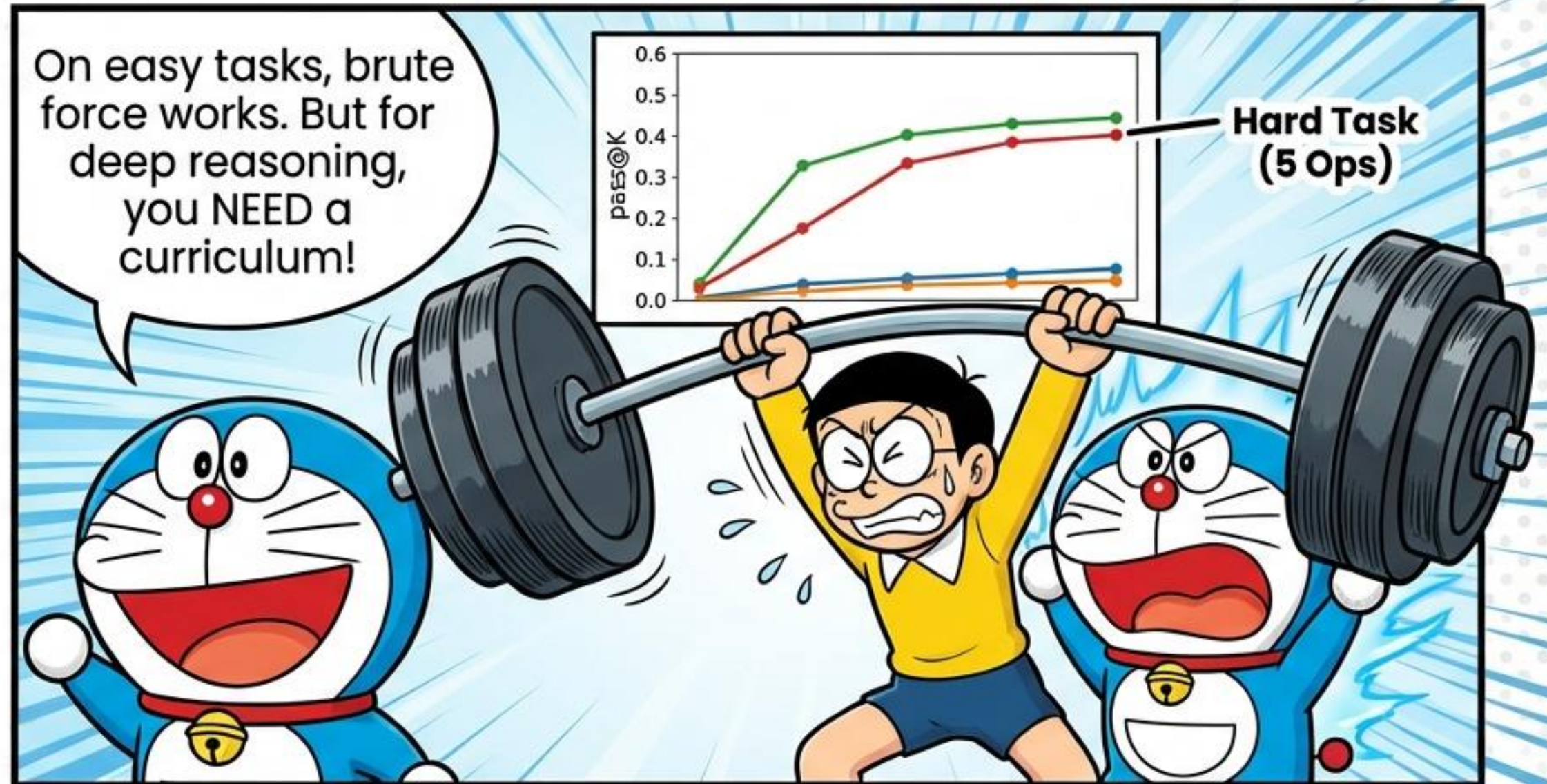
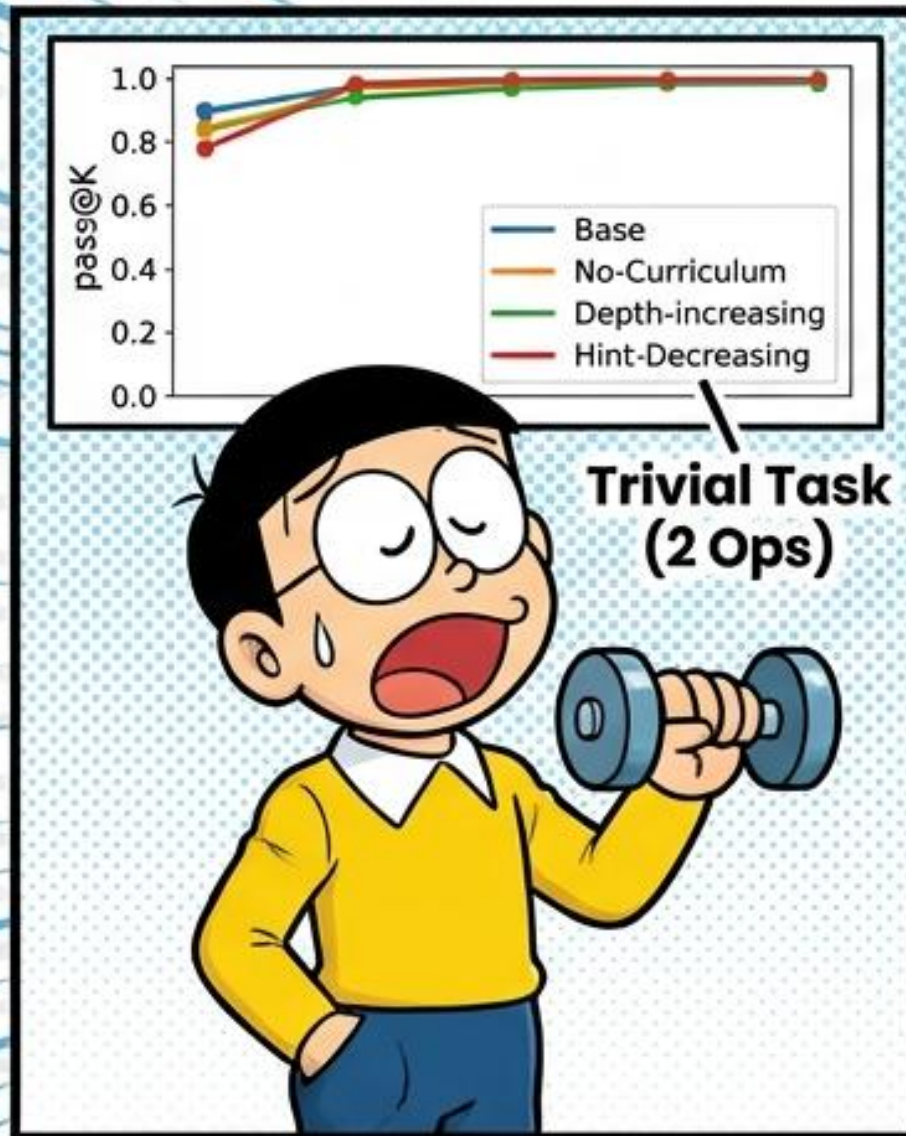


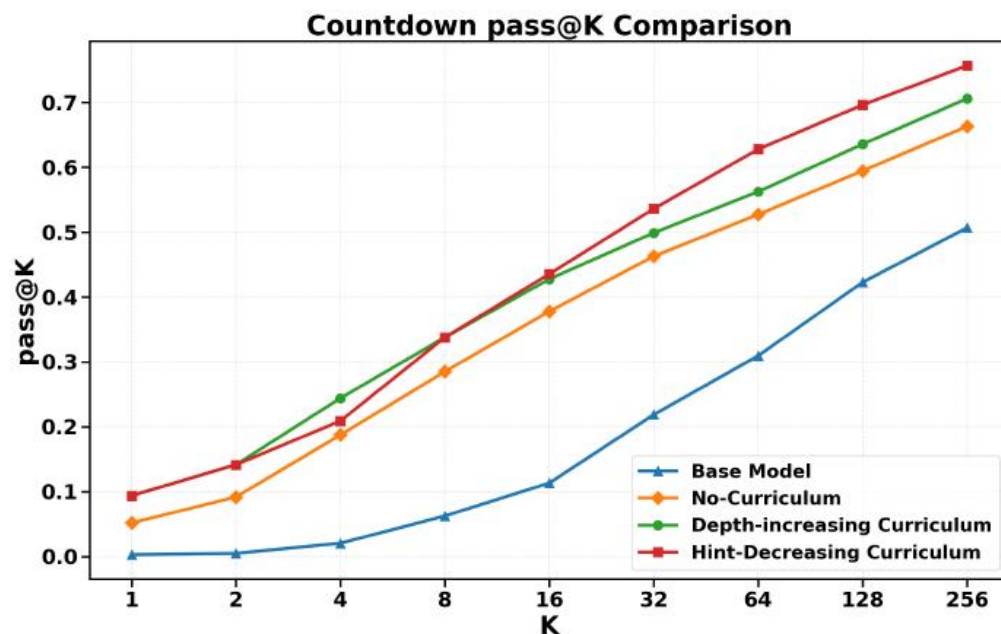
Figure 1. An illustration in Liu et al. (2025b) of the Chain-of-Thought for Countdown game, where the goal is to obtain 24 by applying basic arithmetic operations ($+$, $-$, $\times$, $\div$) ($\Phi_l(\cdot, \cdot)$ in our Def. 2) between the current step's number (e.g., 13, 65 or 72) and some unused number (e.g., 5, 7 or 3) in $\{3, 5, 7, 13\}$, targeting 24 as the final outcome. Per (Parashar et al., 2025), the difficulty measure of Countdown is the number of arithmetic operations required to solve an instance.



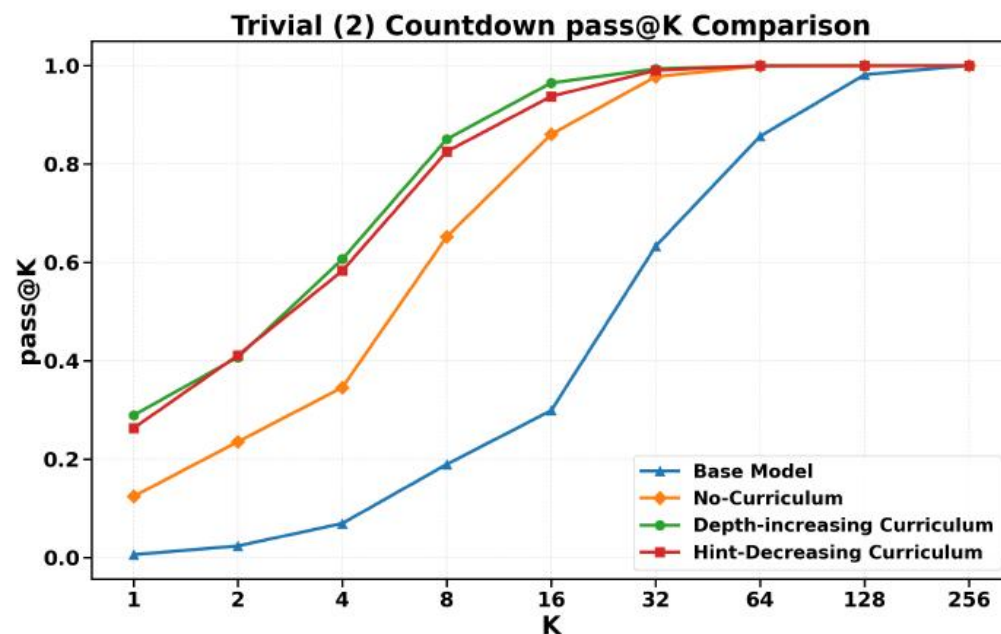
OUT-OF-DISTRIBUTION (OOD) GENERALIZATION



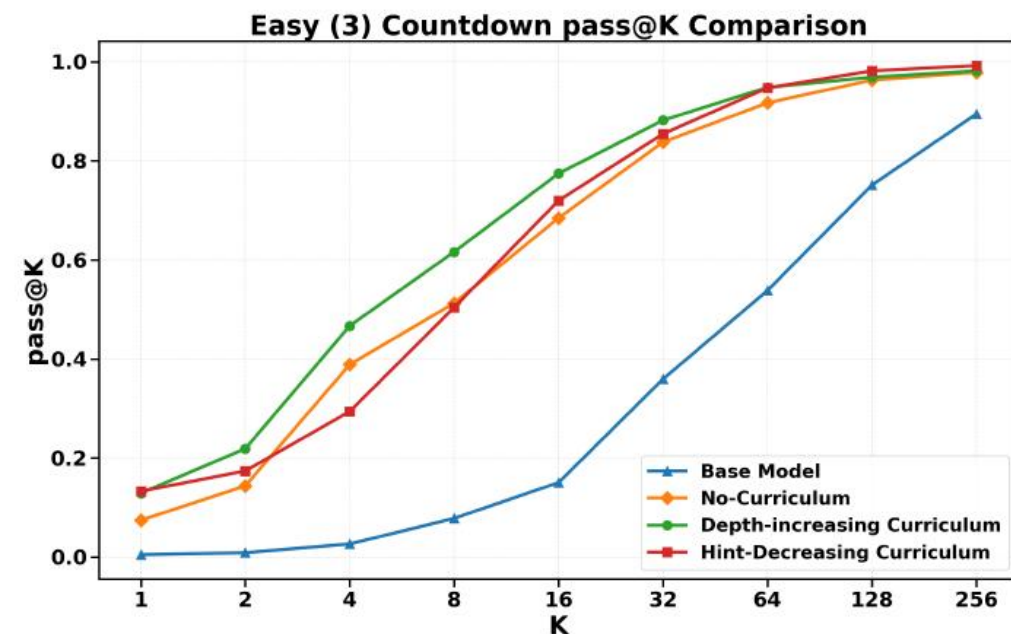
Insight: Curriculum methods show the largest relative gains on Out-Of-Distribution (deeper) tasks.



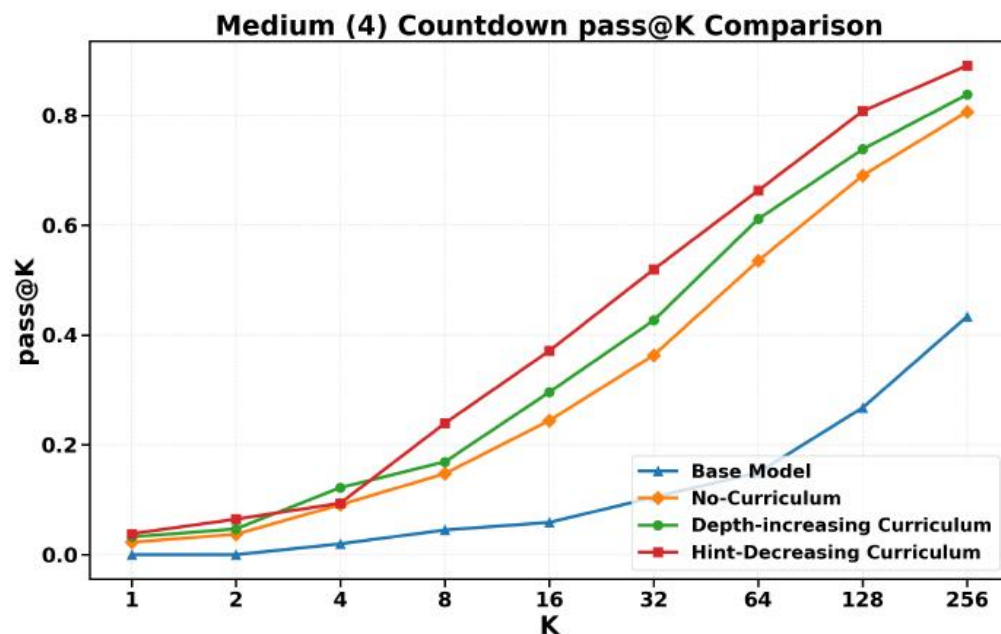
(a) Total validation set.



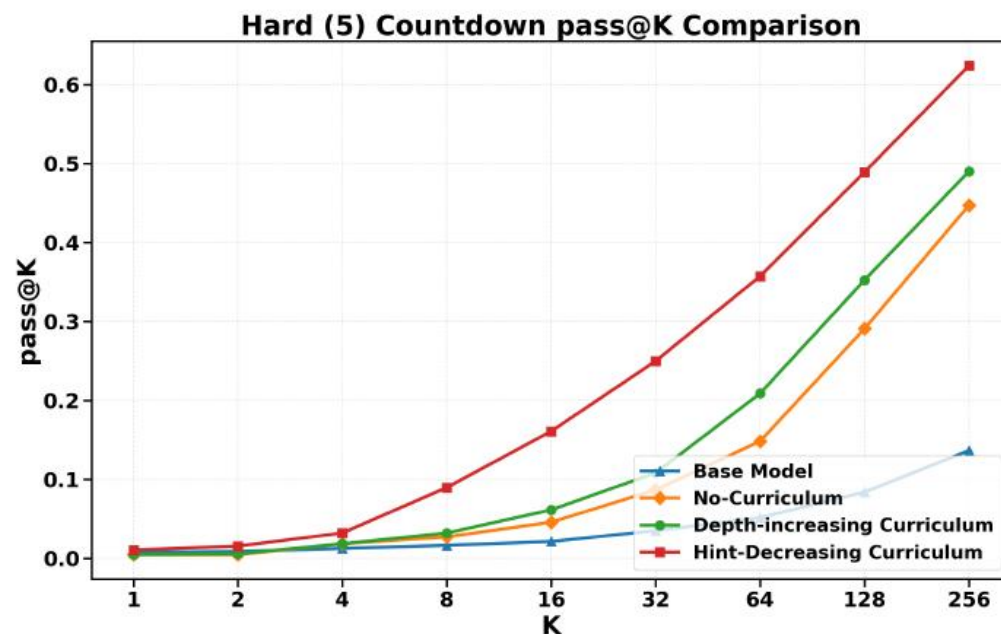
(b) Trivial (2 ops).



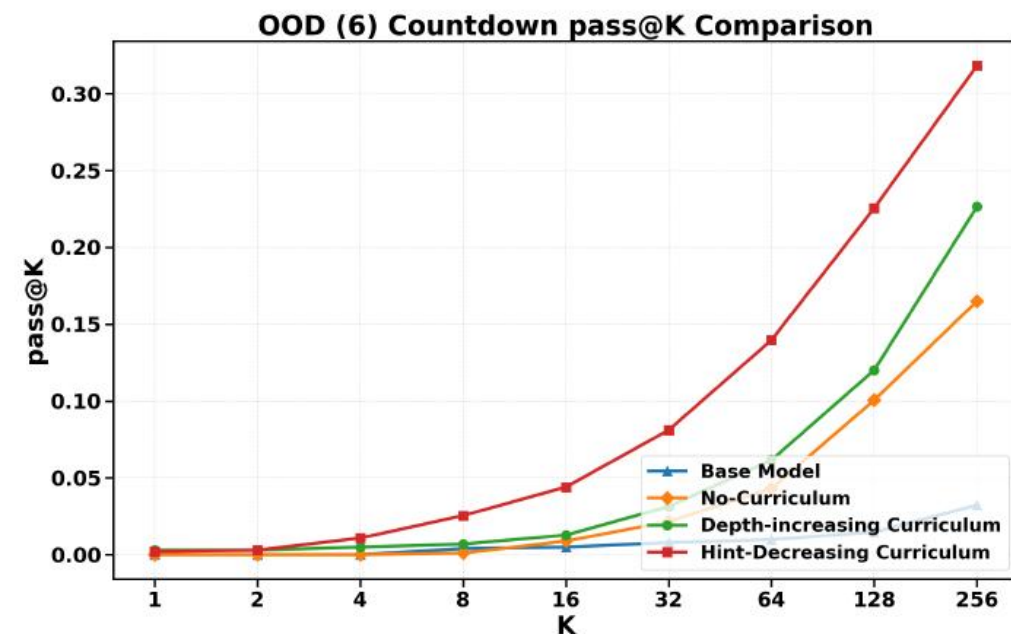
(c) Easy (3 ops).



(d) Medium (4 ops).



(e) Hard (5 ops).



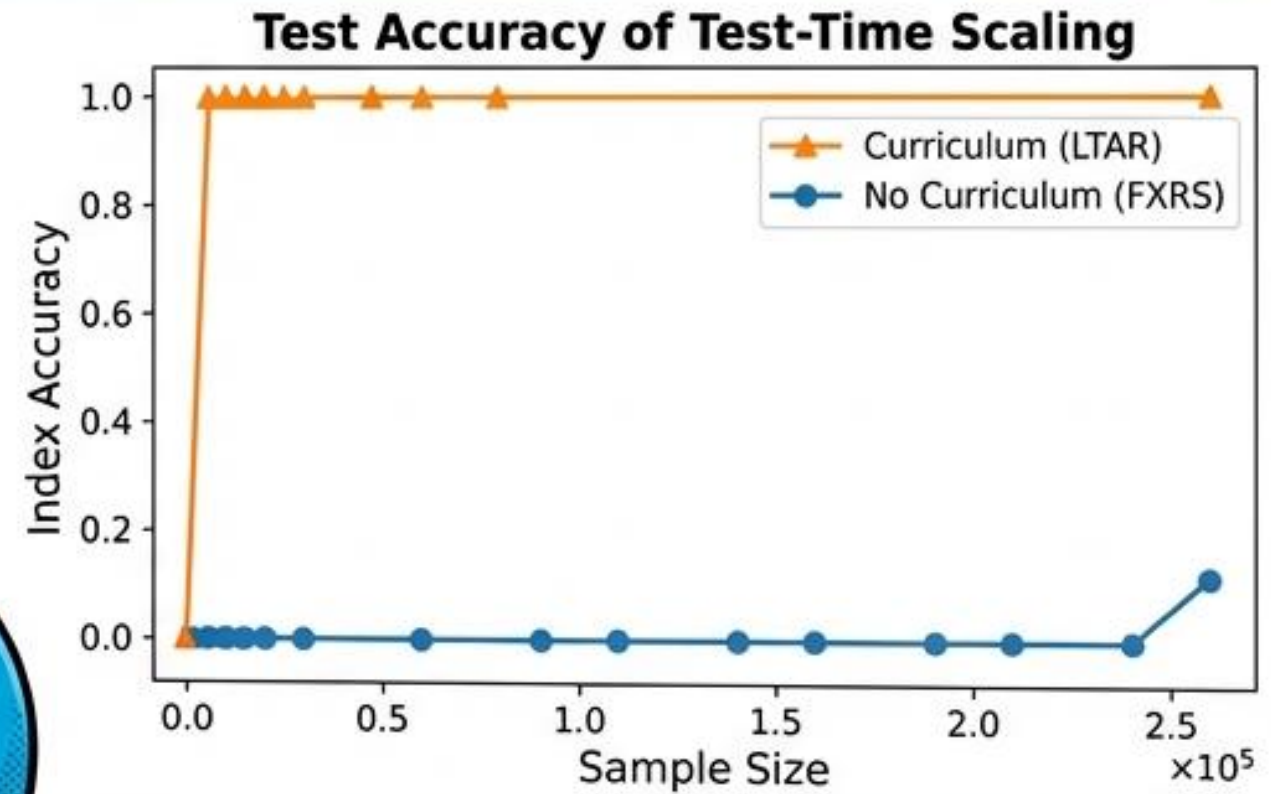
(f) OOD (6 ops).

Figure 12. Test performance on Countdown across difficulty strata. Both Depth-Increasing and Hint-Decreasing curricula outperform the No-Curriculum baseline as reasoning depth increases. Performance gaps are modest on Trivial and Easy subsets but widen substantially on Medium and Hard subsets. On the OOD (6-operation) subset, curriculum methods achieve the largest relative gains, indicating improved generalization to deeper reasoning chains.

TEST-TIME SCALING: THE ANYWHERE DOOR

This works for testing too!
We use a **Curriculum-Aware Oracle** to verify steps incrementally.

So I catch mistakes early? That saves so much time!



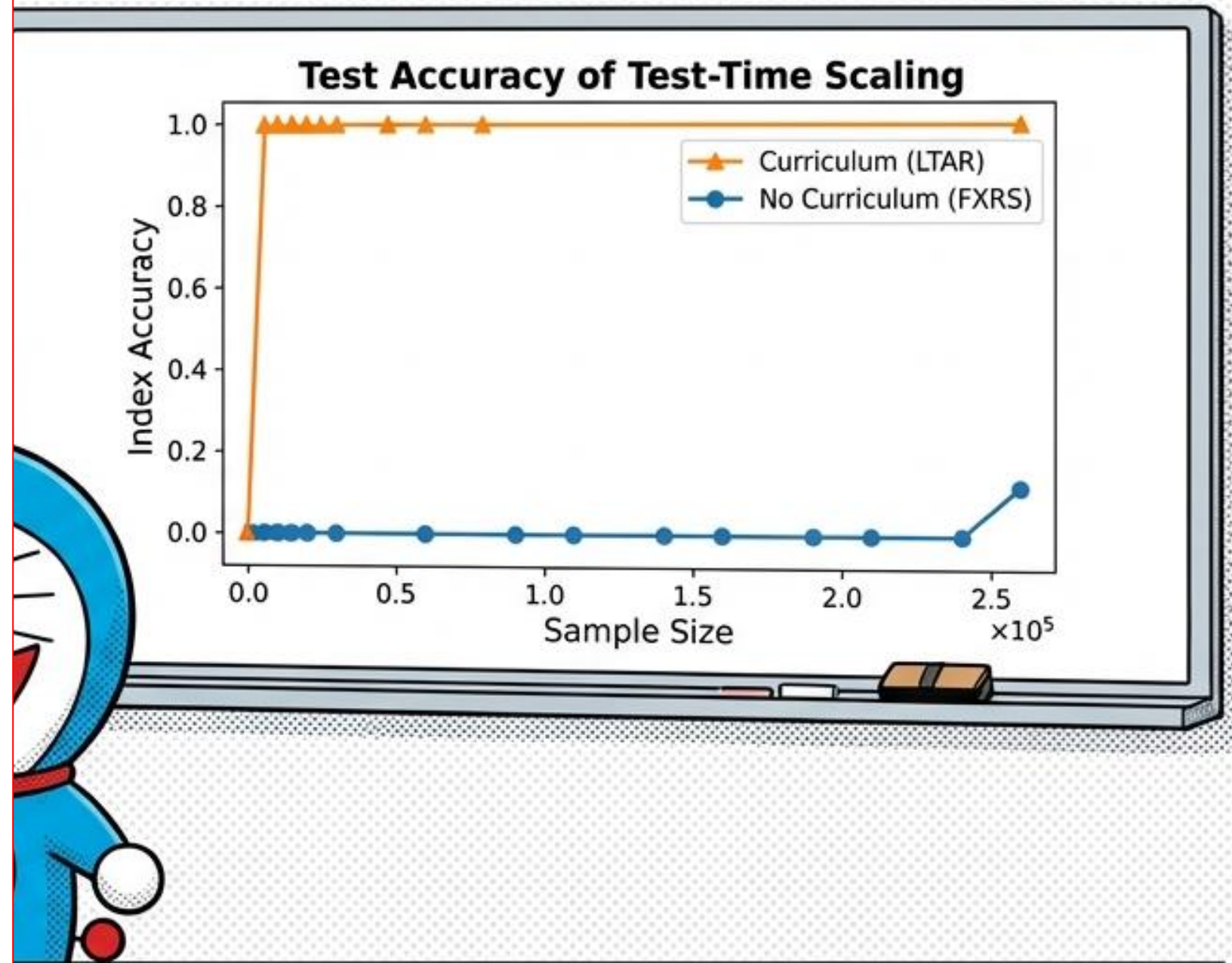
Theorem 3: Provable efficiency in Test-Time Scaling.
Reduces query cost from Exponential to Polynomial.

TEST-TIME SCALING: THE ANYWHERE DOOR

Theorem 3 (Curriculum Test-time Scaling Avoid Exponential Bottleneck). Let $\text{TF}_{\text{base}} := \text{TF}(\cdot; \mathbf{W}_{\text{base}})$ per defined be as in Thm. 1, and consider a target $f_{S_*} \in \mathcal{F}_{2S\text{-ART}}$ with $|S_*| = k^* + 1$. Then, for identifying the ground-truth path S_* with confidence $1 - \delta$:

1. Using only $\mathbf{R}_{\mathbf{x}}^{f_{S_*}}$, any adaptive procedure (including ones that force visible x -tokens by rejection sampling and then roll out) requires $T_{\text{data}}, T_{\text{comp}} \geq \tilde{\Omega}(d^{2k^*})$.
2. Using curriculum queries adaptively to $\mathbf{R}_{\mathbf{x}}^{\mathcal{F}_{S_*}}$, there exists a procedure with $T_{\text{data}} \leq \tilde{O}((k^* + 1)d^2)$ and $T_{\text{comp}} \leq \tilde{O}((k^* + 1)d^3)$.

Here $\tilde{\Omega}(\cdot)$ and $\tilde{O}(\cdot)$ hide polylogarithmic factors in d , $1/\delta$ and spurious-acceptance parameters.



Theorem 3: Provable efficiency in Test-Time Scaling. Reduces query cost from Exponential to Polynomial.

Algorithm 4 Forced X-token Rejection Sampling (FXRS) under token-only observability

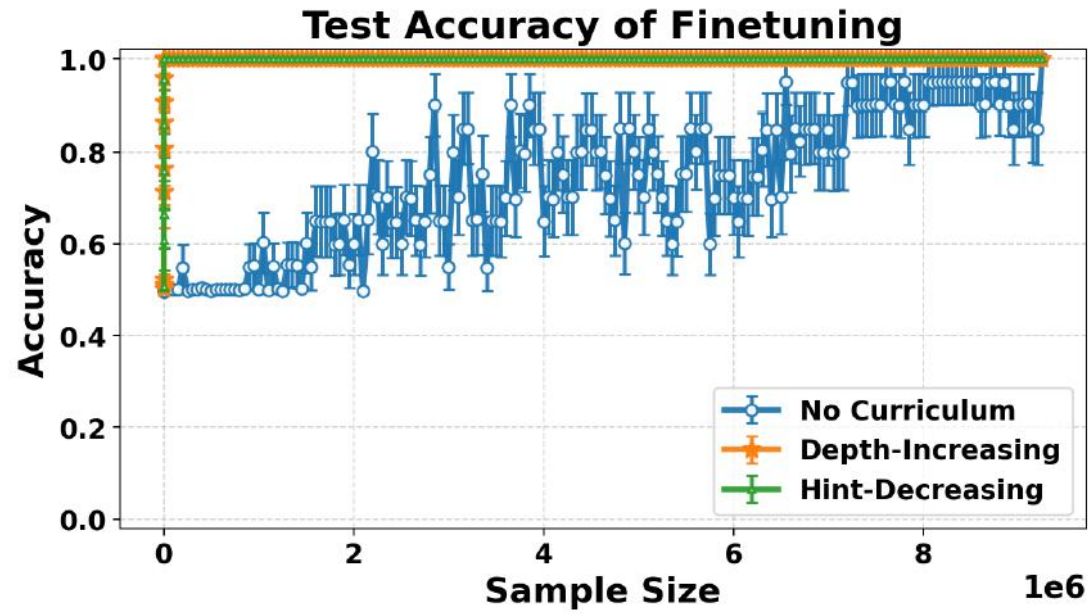
- 1: Inputs: depth ℓ , committed history $\mathbf{CoT}_{\ell-1}$, legal set $\mathcal{I}_\ell$, candidate $j \in \mathcal{I}_\ell$, model TF_{base} , oracle $\mathbf{R}_{\mathbf{x}}^{fs_\star}$, max trials $T_{\text{max}} = \Theta(d \log(\frac{d}{\delta}))$.
 - 2: **for** $t = 1$ to T_{max} **do**
 - 3: From context $\mathbf{CoT}_{\ell-1}$, sample one step with TF_{base} to emit $\hat{\mu}^{x_{i_\ell}}$.
 - 4: **if** $\hat{\mu}^{x_{i_\ell}} = \mu^{x_j}$ **then**
 - 5: Roll out the remaining suffix using TF_{base} until termination (EOS occurs at some depth by the copied PART behavior).
 - 6: Return the full rollout sequence and **stop**.
 - 7: **end if**
 - 8: **end for**
 - 9: If no match after T_{max} , return **fail** (increase T_{max} if needed).
-

Algorithm 5 BAI across depths with FXRS and $\mathbf{R}_{\mathbf{x}}^{fs_\star}$

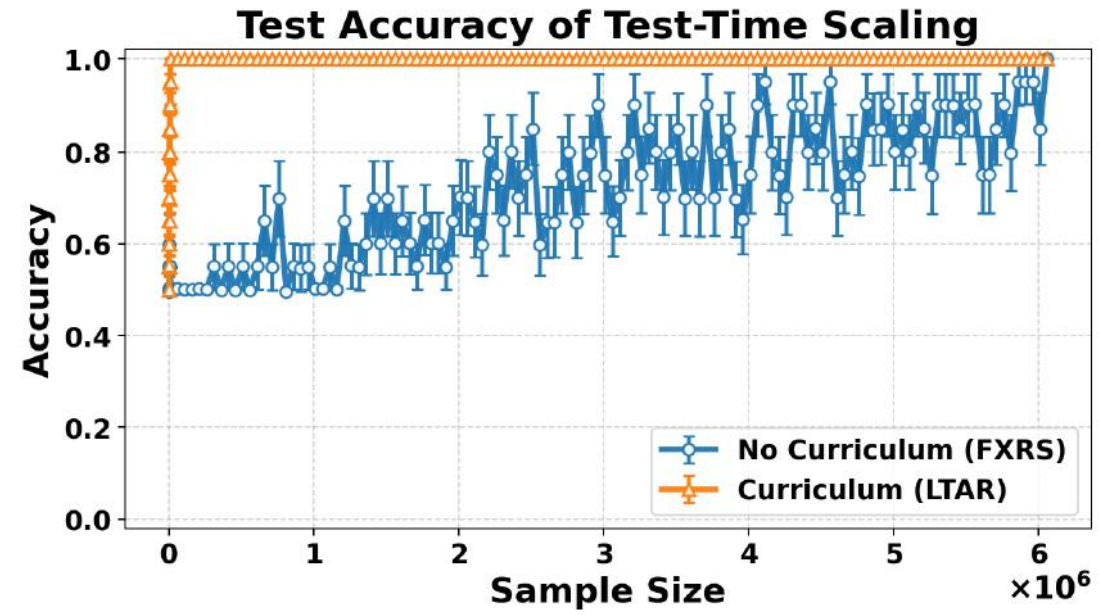
- 1: Inputs: legal sets $\{\mathcal{I}_\ell\}_{\ell=1}^{k^\star}$, model TF_{base} , oracle $\mathbf{R}_{\mathbf{x}}^{fs_\star}$, confidence δ .
 - 2: Initialize committed partial CoT $\mathcal{H} \leftarrow ()$.
 - 3: **for** $\ell = 1$ to $k^\star + 1$ **do**
 - 4: Set per-arm repetitions $m \leftarrow \Theta(d^{2(k^\star+1-\ell)} \log(d/\delta))$.
 - 5: **for** each $j \in \mathcal{I}_\ell$ **do**
 - 6: Initialize counter $c_j \leftarrow 0$.
 - 7: **for** $r = 1$ to m **do**
 - 8: Run FXRS (Alg. 4) with candidate j to obtain a full rollout (or retry until success). If FXRS fails, repeat until success.
 - 9: Query $\mathbf{R}_{\mathbf{x}}^{fs_\star}$ on the final pre-EOS token of the rollout; record one Bernoulli outcome $Y_{j,r} \in \{0, 1\}$.
 - 10: Update $c_j \leftarrow c_j + Y_{j,r}$.
 - 11: **end for**
 - 12: Set empirical acceptance $\hat{p}_j \leftarrow c_j/m$.
 - 13: **end for**
 - 14: Let $\hat{j}_\ell \leftarrow \arg \max_{j \in \mathcal{I}_\ell} \hat{p}_j$.
 - 15: Commit depth- ℓ : resample until $\hat{\mu}^{x_{i_\ell}} = \mu^{x_{\hat{j}_\ell}}$, then compute $\hat{\mu}^{z_\ell} = \text{FFN}_\ell(\hat{\mu}^{x_{i_\ell}}, \mathbf{E}[z_{\ell-1}]_{:d_{\mathbf{x}}})$; update $\mathcal{H} \leftarrow (\mathcal{H}, \hat{\mu}^{x_{\hat{j}_\ell}}, \hat{\mu}^{z_\ell})$.
 - 16: **end for**
 - 17: Return the committed path encoded by $\mathcal{H}$.
-

Algorithm 6 Layer-wise Truncated Accept-Reject (LTAR) with $\mathbf{R}_x^{\mathcal{F}^{S^*}}$ (BAI with inline forcing and truncation)

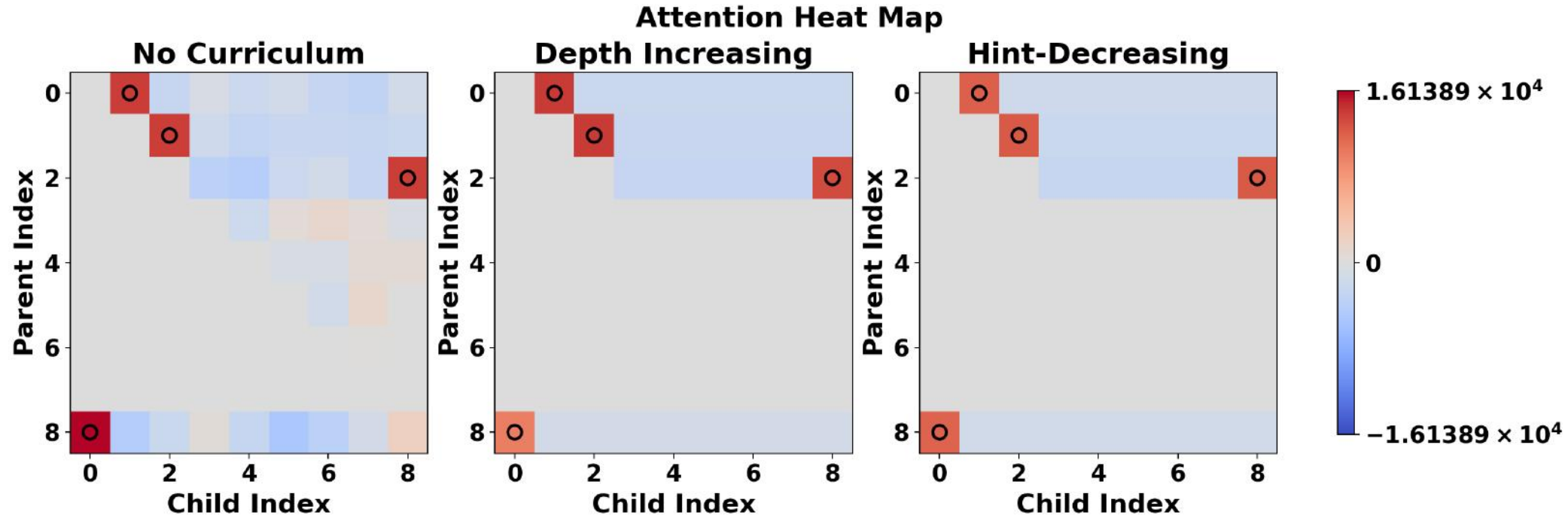
- 1: Inputs: legal sets $\{\mathcal{I}_\ell\}_{\ell=1}^{k^*}$, confidence level δ , per-depth budgets $m_\ell = \Theta(d^2(1 - \rho_{\text{spur}}^{\text{sup}, \ell})^{-2} \log(\frac{d}{\delta}))$, max trials $T_{\text{max}} = \Theta(d \log(\frac{d}{\delta}))$.
 - 2: Initialize committed partial CoT $\mathcal{H} \leftarrow ()$.
 - 3: **for** $\ell = 1$ to $k^* + 1$ **do**
 - 4: For each $j \in \mathcal{I}_\ell$, estimate acceptance by m_ℓ repetitions with inline forcing:
 - 5: Initialize $c_j \leftarrow 0$.
 - 6: **for** $r = 1$ to m_ℓ **do**
 - 7: **for** $t = 1$ to T_{max} **do**
 - 8: From context $\mathcal{H}$, sample one step with TF_{base} to emit $\hat{\mu}^{x_{i_\ell}}$.
 - 9: **if** $\hat{\mu}^{x_{i_\ell}} = \mu^{x_j}$ **then**
 - 10: Compute $\hat{\mu}^{z_\ell} \leftarrow \text{FFN}_\ell(\hat{\mu}^{x_{i_\ell}}, \mathbf{E}[z_{\ell-1}]_{:d_X})$.
 - 11: Externally append EOS if $\ell < k^* + 1$; query $\mathbf{R}_x^{\mathcal{F}^{S^*}}$ at depth ℓ ; record Bernoulli $Y_{j,r} \in \{0, 1\}$.
 - 12: Update $c_j \leftarrow c_j + Y_{j,r}$ and **break** the retry loop.
 - 13: **end if**
 - 14: **end for**
 - 15: If no match within T_{max} , optionally increase T_{max} and repeat this repetition.
 - 16: **end for**
 - 17: Set empirical acceptance $\hat{p}_j \leftarrow c_j/m_\ell$ for all $j \in \mathcal{I}_\ell$ and pick $\hat{j}_\ell = \arg \max_j \hat{p}_j$.
 - 18: Commit depth- ℓ : resample until $\hat{\mu}^{x_{i_\ell}} = \mu^{x_{\hat{j}_\ell}}$, then compute $\hat{\mu}^{z_\ell} = \text{FFN}_\ell(\hat{\mu}^{x_{i_\ell}}, \mathbf{E}[z_{\ell-1}]_{:d_X})$; update $\mathcal{H} \leftarrow (\mathcal{H}, \hat{\mu}^{x_{\hat{j}_\ell}}, \hat{\mu}^{z_\ell})$.
 - 19: **end for**
 - 20: Return the committed path encoded by $\mathcal{H}$.
-



(a) Test Accuracy of d.

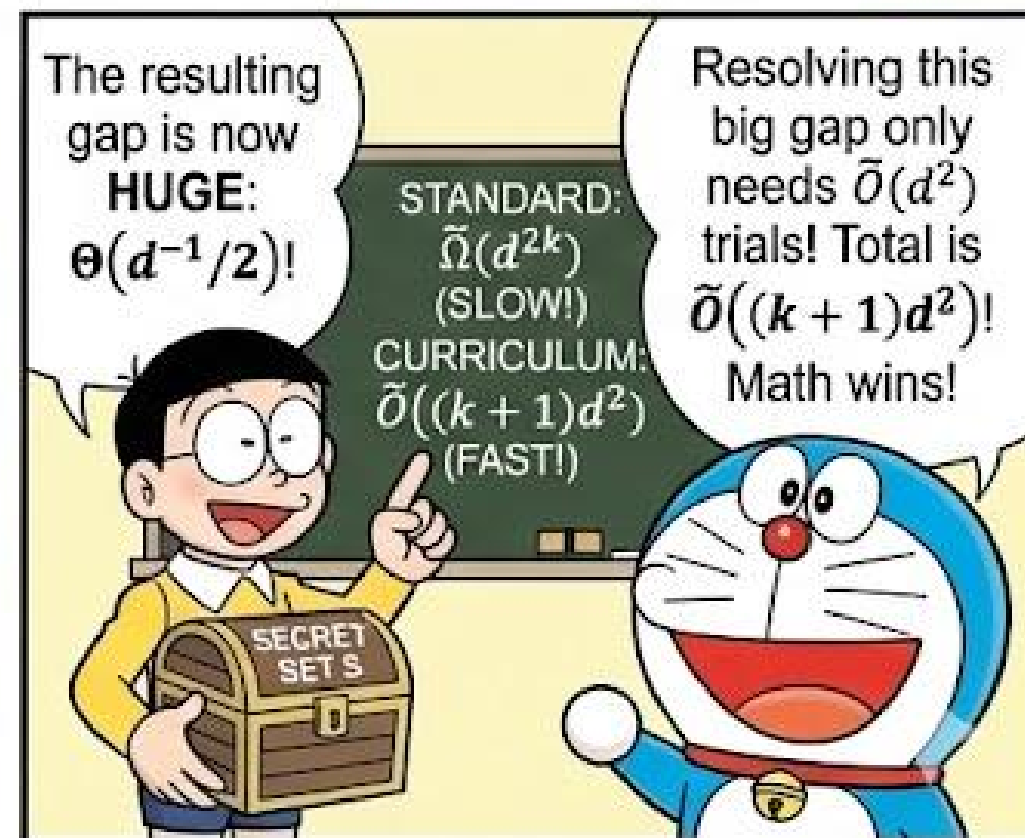
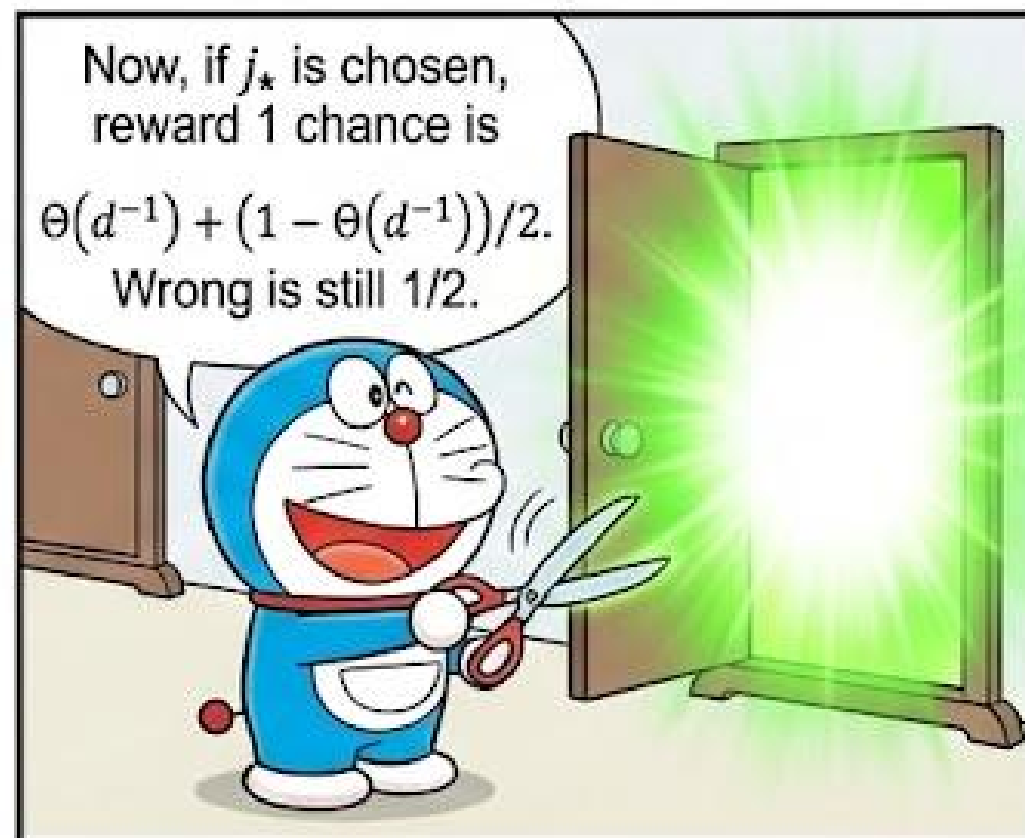
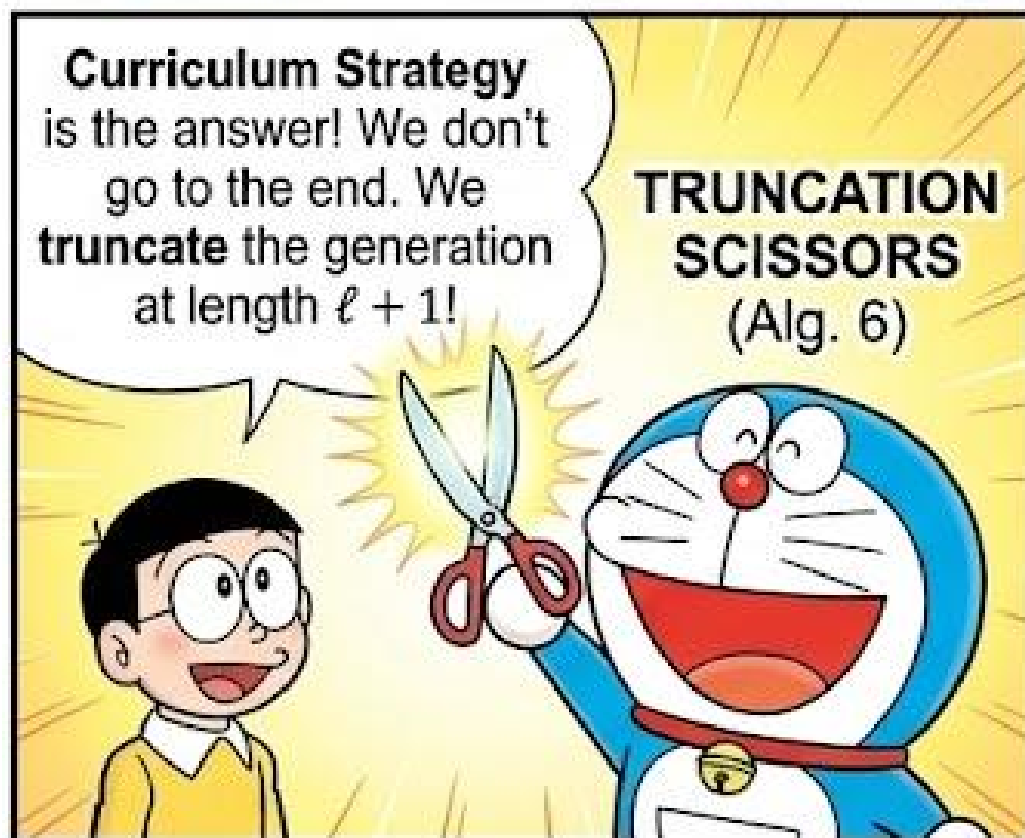
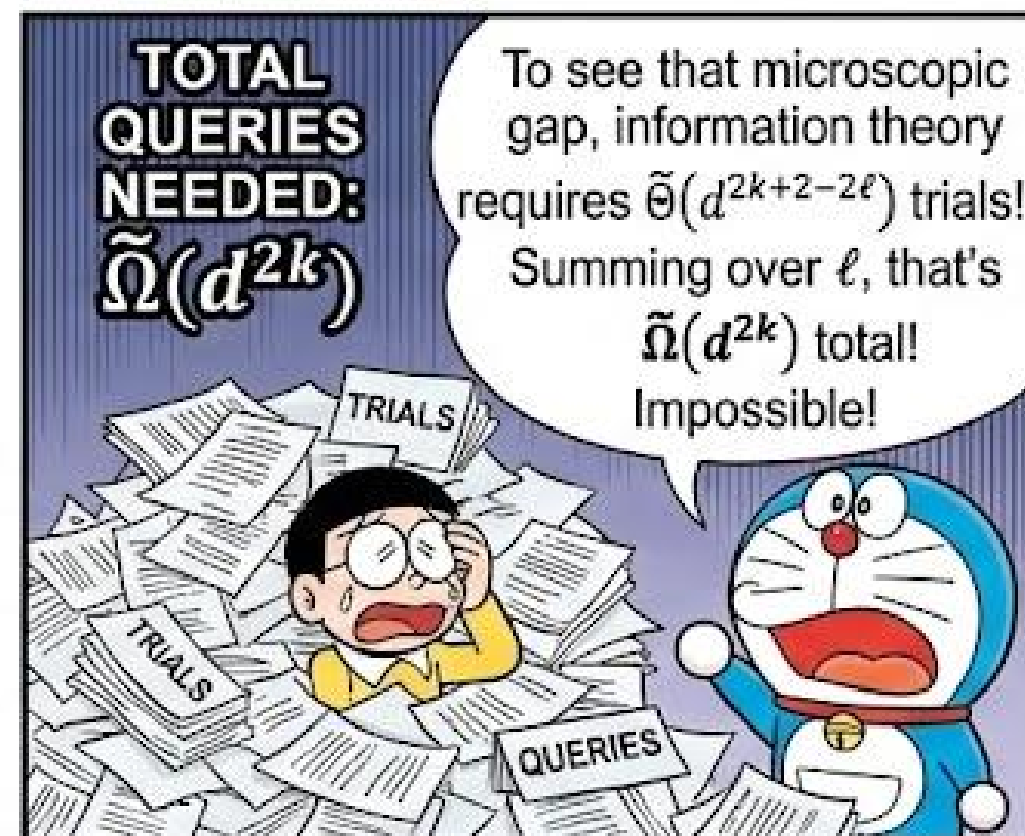
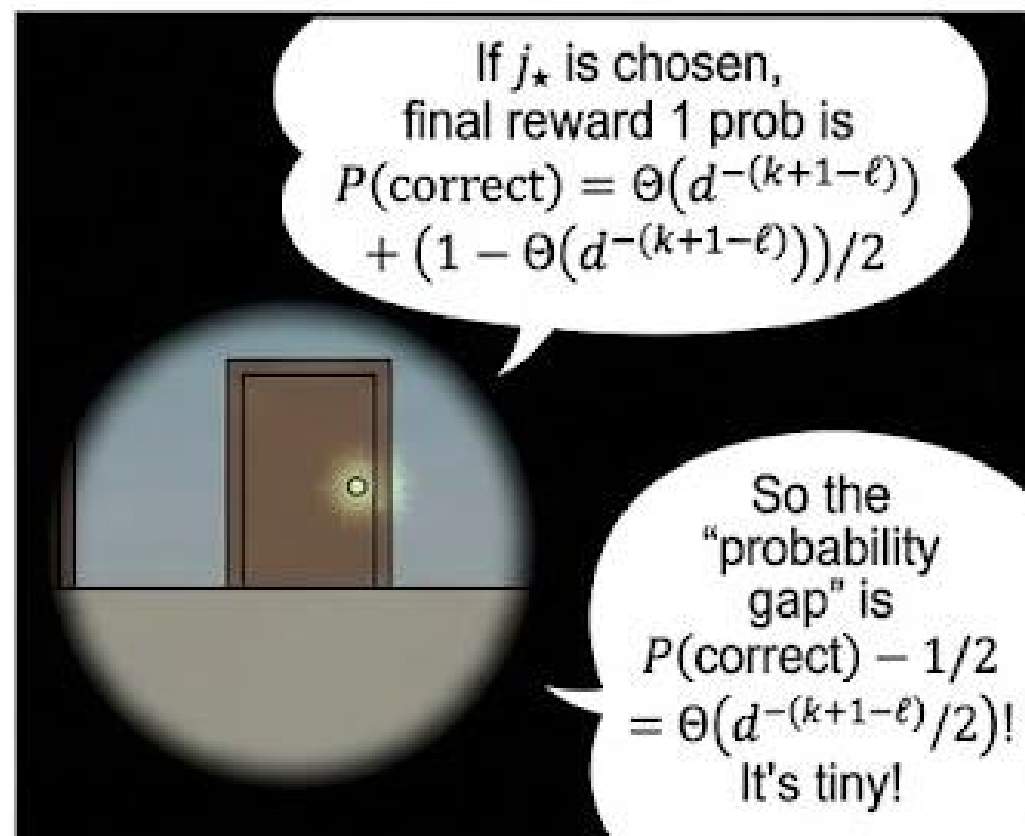
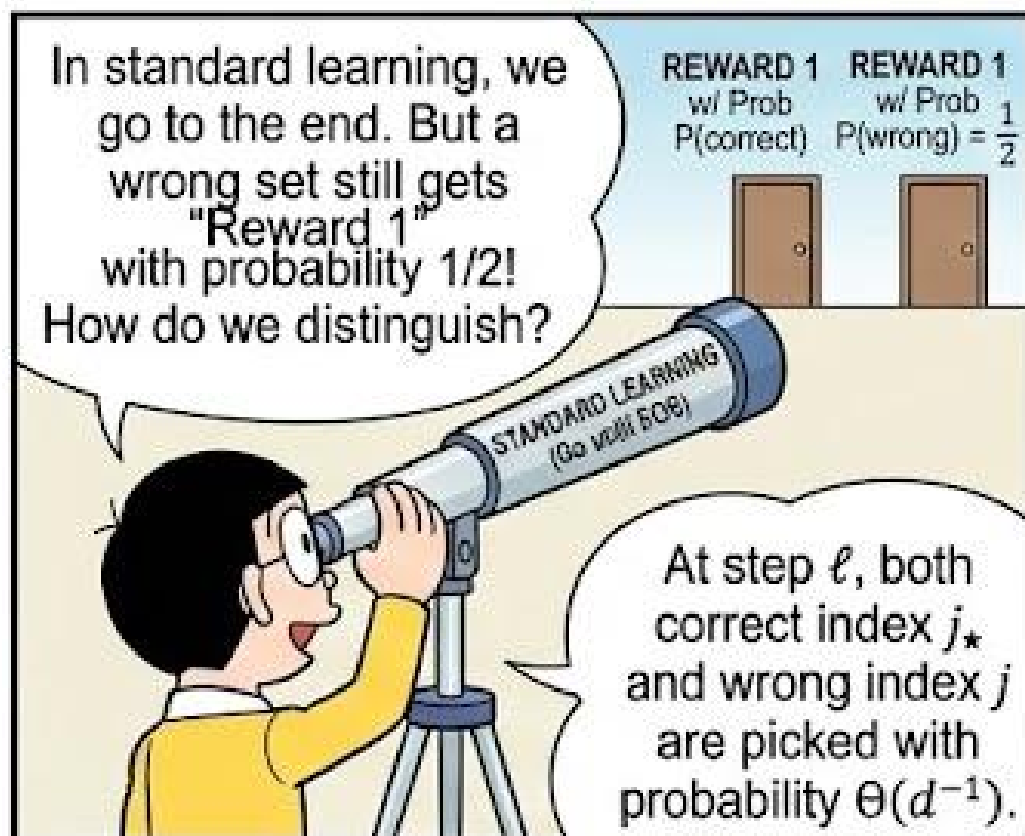


(b) Test Accuracy of Test-Time Scaling.



(c) Attention Heat Map of Finetuning.

Figure 8. Post-training Dynamics on the parity task with error bar ($d = 8$, $k^* = 8$, $S^* = [0, 1, 2]$, $\eta_A = 10^8$, $\eta_B = \eta_C = 10^5$). (a) Test accuracy evolution of finetuning along sample complexity; (b) Test accuracy evolution of Test Time-Scaling along sample complexity; (c) averaged attention heatmap after finetuning.



GADGET SUMMARY CHECKLIST



1. **The Problem:** Hard reasoning has sparse rewards → Exponential Bottleneck. **[CHECK]**



2. **The Solution:** Decompose the task! Use Depth-Increasing or Hint-Decreasing gadgets. **[CHECK]**



3. **The Guarantee:** The Coverage Coefficient proves Polynomial Sample Complexity. **[CHECK]**



4. **The Result:** Works for both Post-Training (RL) and Test-Time Scaling. **[CHECK]**

People always
felt that **curriculum**
learning worked.
But now, thanks to this
paper, we have the
theory to prove *why*.

Next time I have a
hard problem, I won't
panic. I'll just find a
good Curriculum!

THE END

Based on: Provable Sample Efficiency of Curriculum Post-Training
for Transformer Reasoning (ICML Submission)