



How We Learn Structured Knowledge with Limited Resources



Hanqi Zhou

hanqi.zhou@uni-tuebingen.de

*Cluster of Excellence “Machine Learning for Science”, Tübingen AI Center,
University of Tübingen*

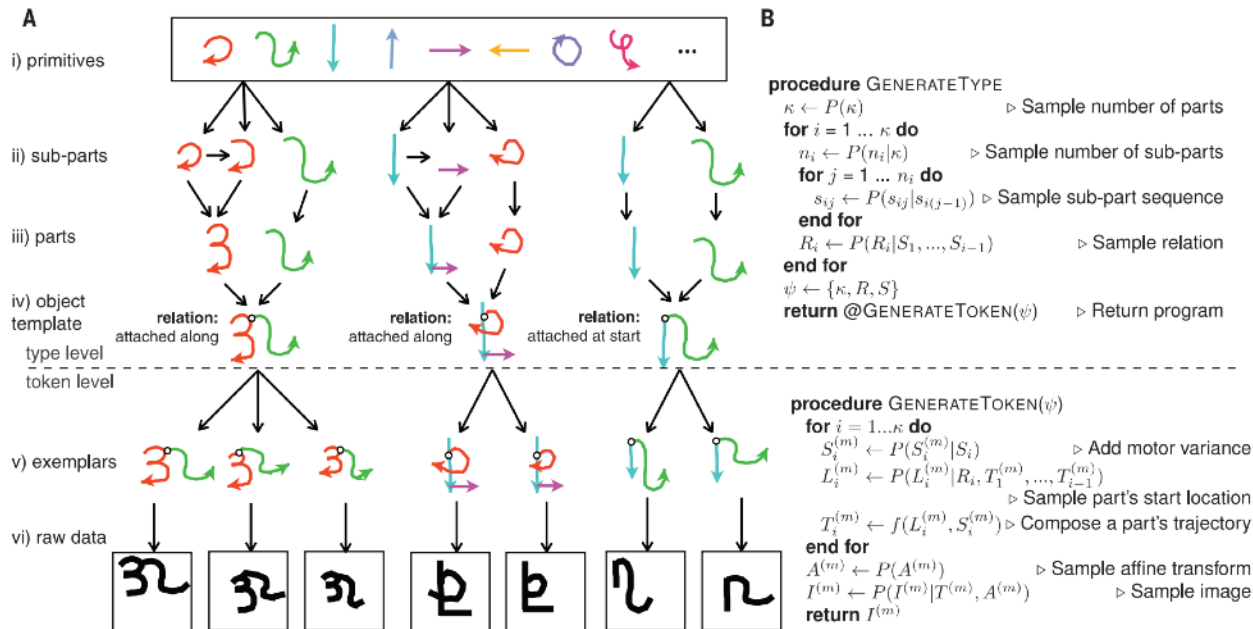


Fig. 3. A generative model of handwritten characters. (A) New types are generated by choosing primitive actions (color coded) from a library (i), combining these subparts (ii) to make parts (iii), and combining parts with relations to define simple programs (iv). New tokens are generated by running these programs (v), which are then rendered as raw data (vi). (B) Pseudocode for generating new types ψ and new token images $I^{(m)}$ for $m = 1, \dots, M$. The function $f(\cdot, \cdot)$ transforms a subpart sequence and start location into a trajectory.

Lake, B. M., Salakhutdinov, R., & Tenenbaum, J. B. (2015). Human-level concept learning through probabilistic program induction. *Science*, 350(6266), 1332-1338.

- **Resource:** we are resource constrained, e.g., attention, memory, ...
- **Rationality:** we act to maximize our expected utility
- **Trade-off:** we make use of limited resources in an optimal or rational manner

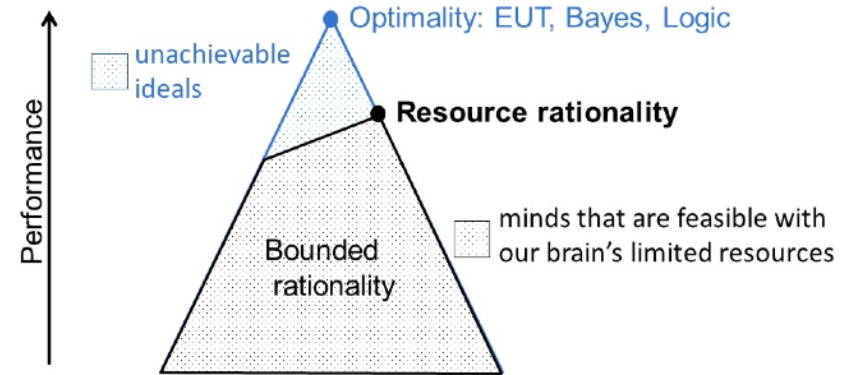


Figure: the best biologically feasible mind out of the infinite set of bounded-rational minds

Lieder, F., & Griffiths, T. L. (2020). Resource-rational analysis: Understanding human cognition as the optimal use of limited computational resources. *Behavioral and brain sciences*, 43, e1.

Rate-distortion theory (RDT) is used as a computational model

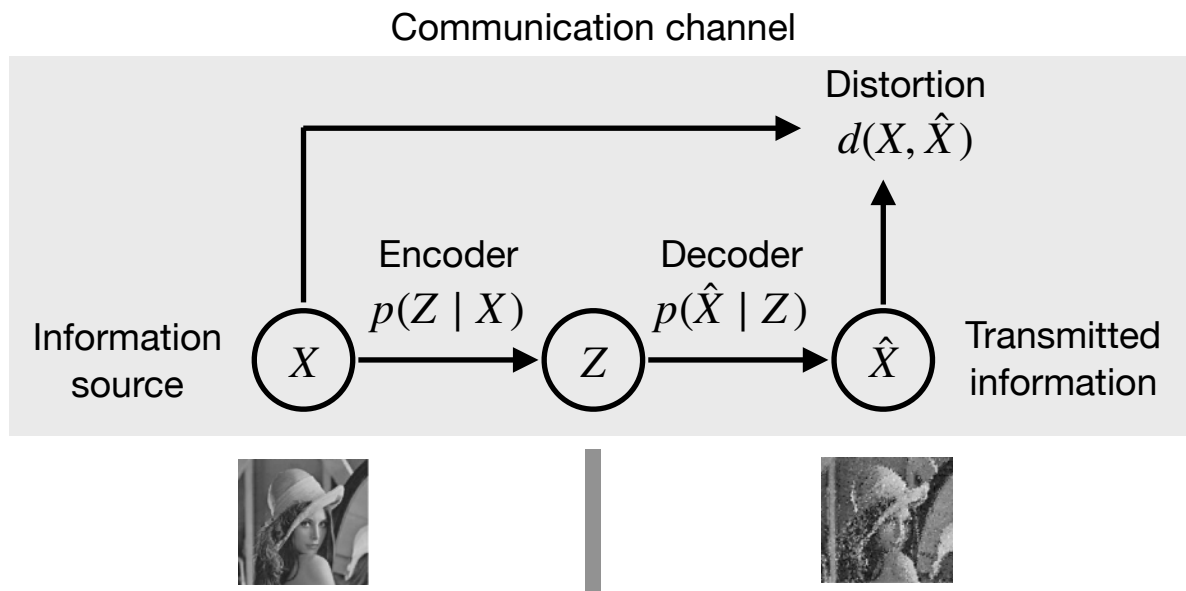
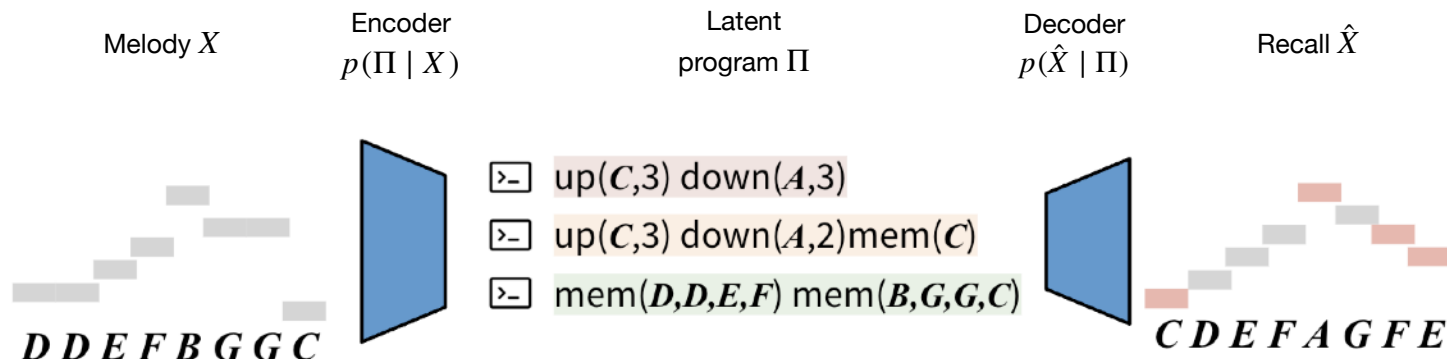
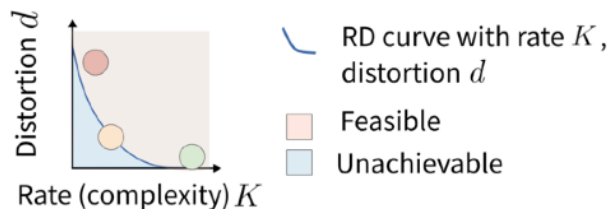
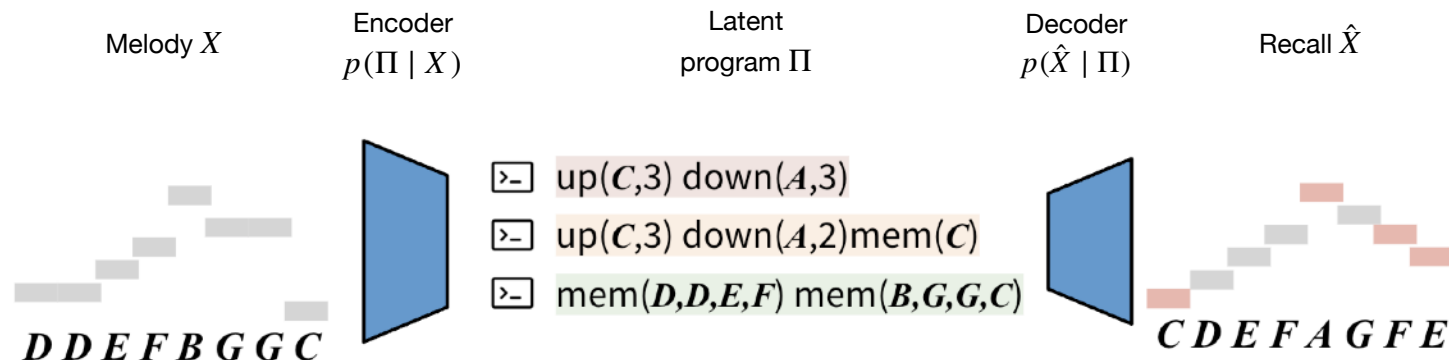


Figure adapted from Gershman, S. J. (2021). The rational analysis of memory. Oxford handbook of human memory.

Melody program induction under constraints

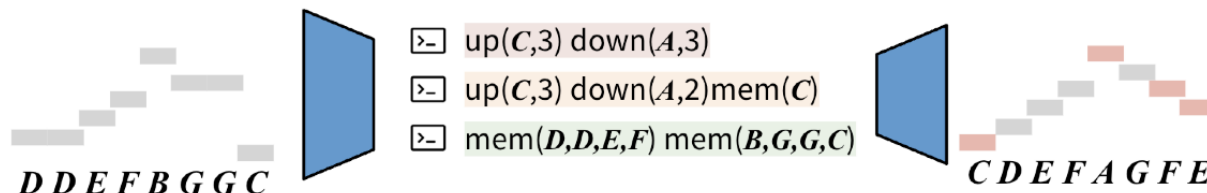


Melody program induction under constraints

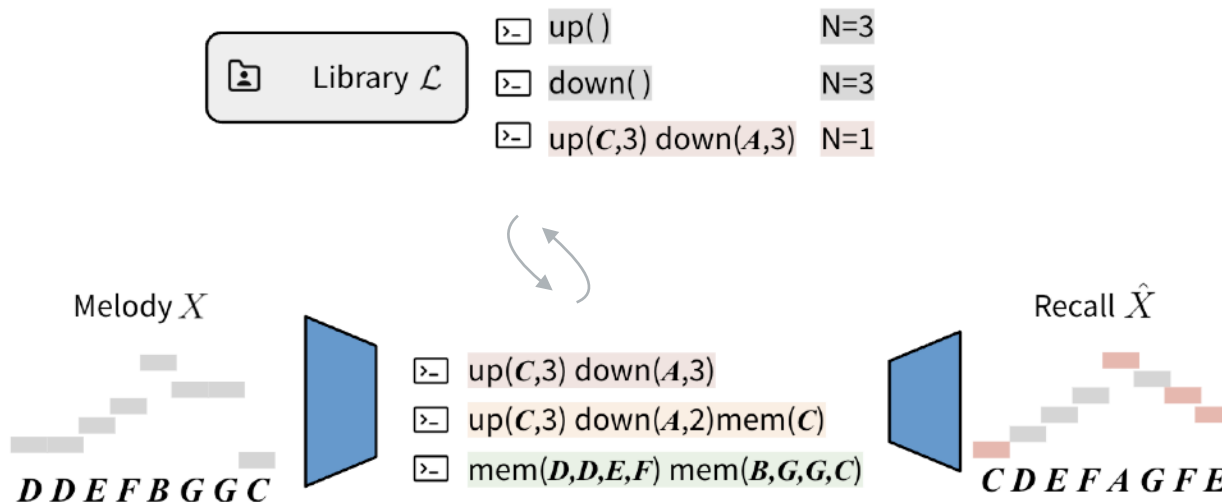


What is missing when modeling human learning with RDT?

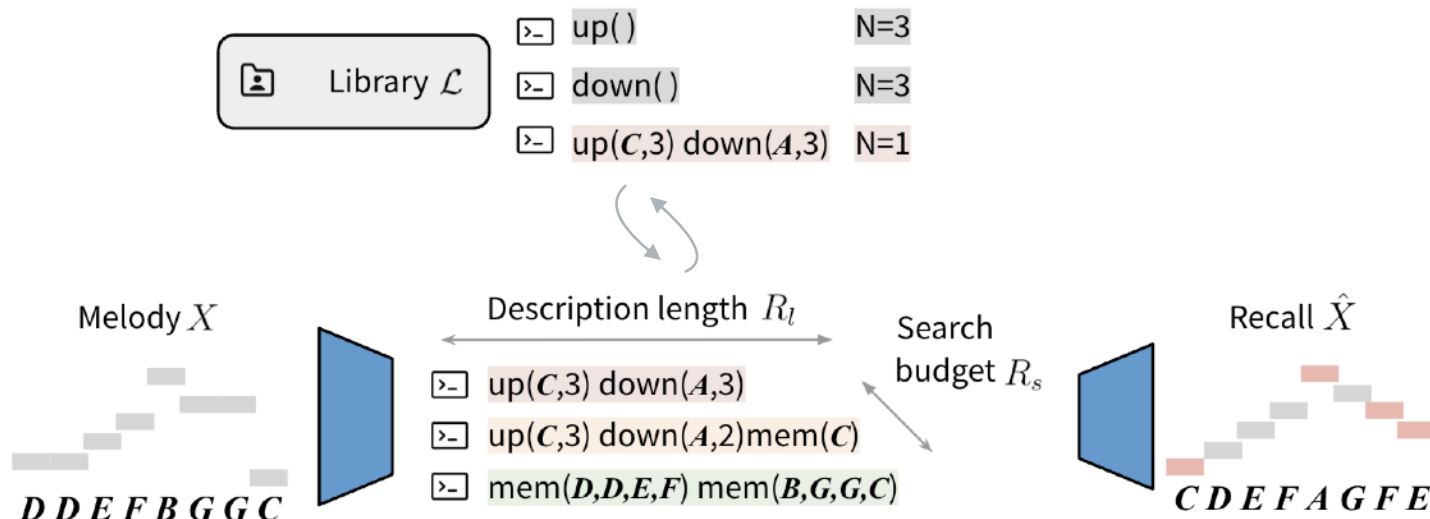
- RDT: How can we transmit a message as efficiently as possible while minimizing information loss?
- Human: How do our internal compression models evolve as we learn? / How do we determine what constitutes the “optimal” encoding over time?

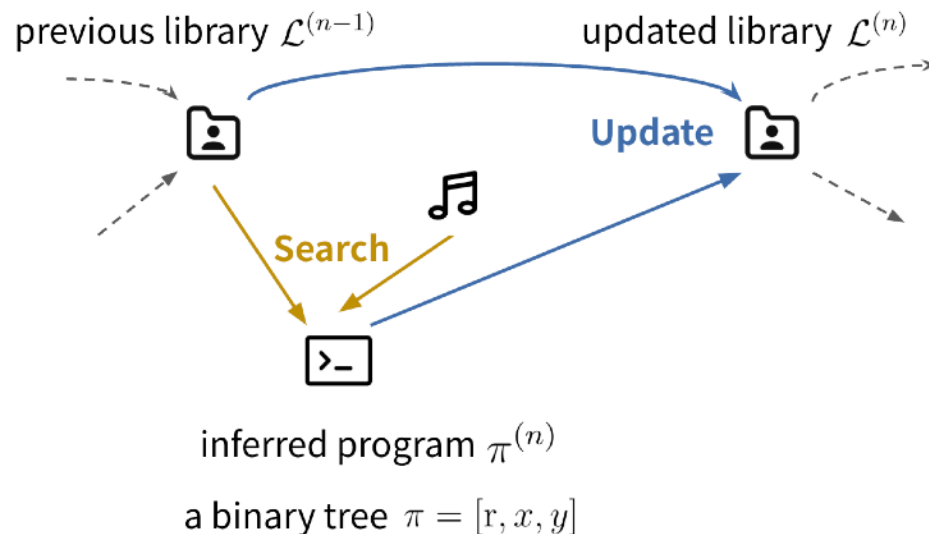


Melody program induction under constraints



Melody program induction under constraints

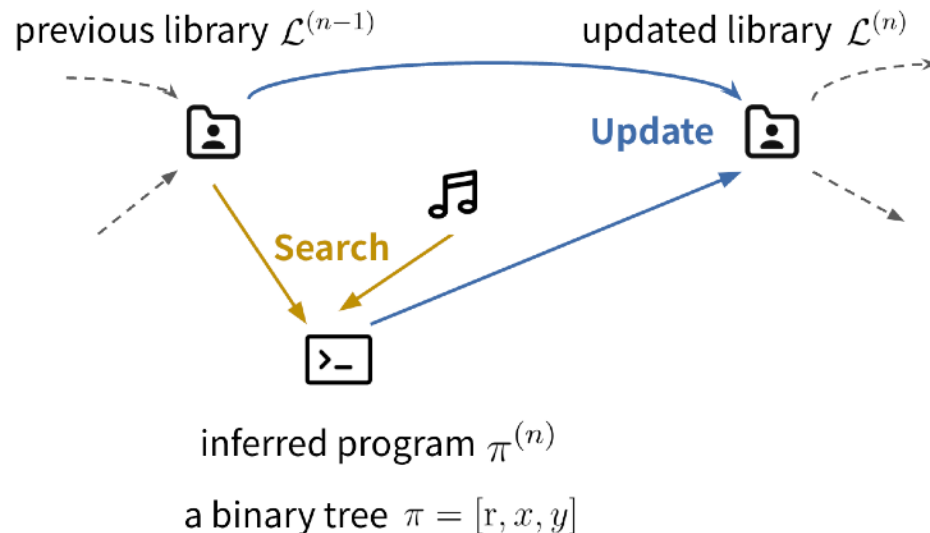




Search for programs (within a melody):

- Description length R_l : An upper limit on the length of the programs
- Search budget R_s : An upper limit on how many programs are considered

$$p(\pi \mid X^{(n)}) \propto p(X^{(n)} \mid \pi) p(\pi)$$

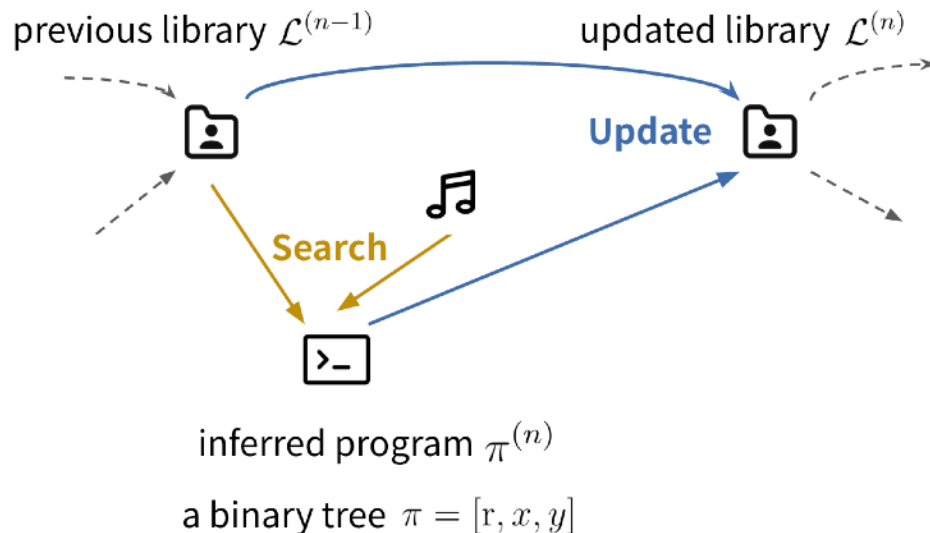


Search for programs (within a melody):

- Description length R_l : An upper limit on the length of the programs
- Search budget R_s : An upper limit on how many programs are considered

$$p(\pi \mid X^{(n)}) \propto p(X^{(n)} \mid \pi) p(\pi)$$

Prior
 $\propto \text{Description length}^{-1}$

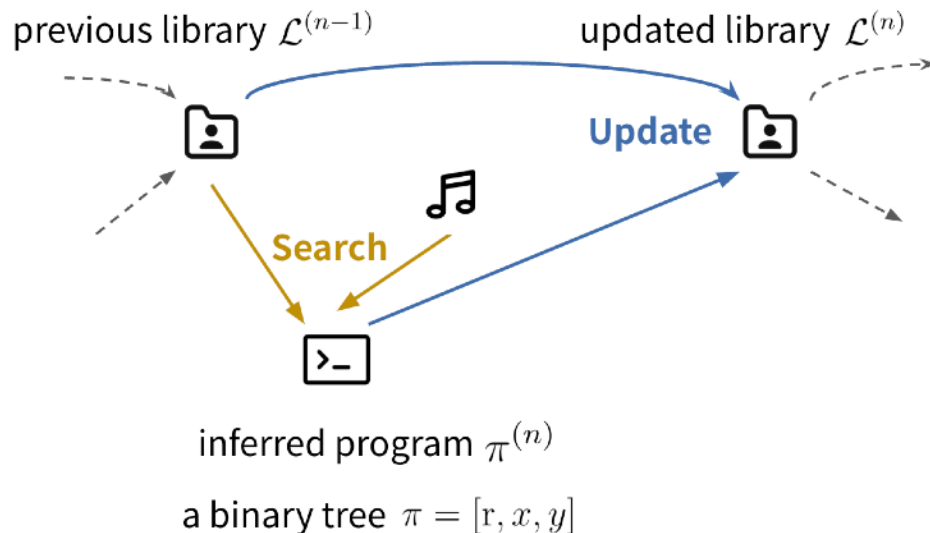


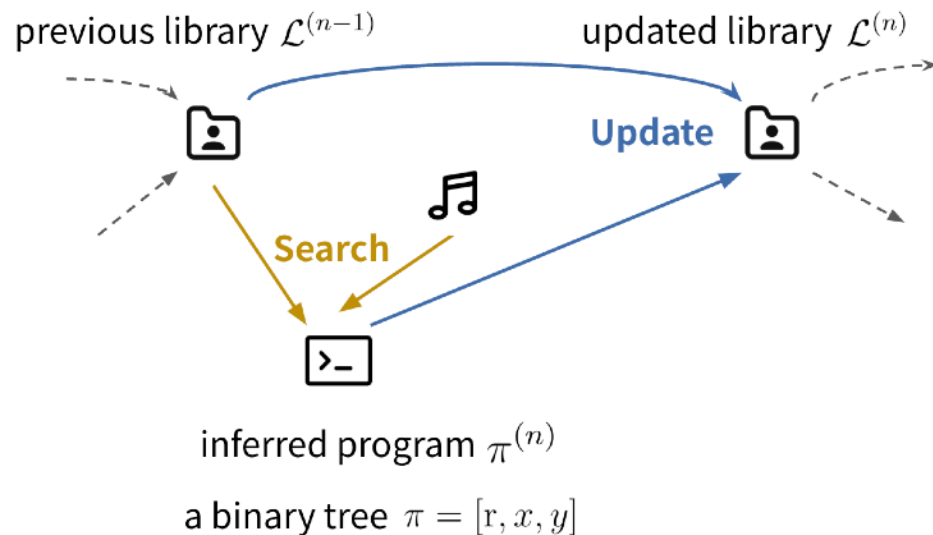
Search for programs (within a melody):

- Description length R_l : An upper limit on the length of the programs
- Search budget R_s : An upper limit on how many programs are considered

$$p(\pi \mid X^{(n)}) \propto p(X^{(n)} \mid \pi) p(\pi)$$

$$\begin{aligned} &\text{Likelihood} \\ &\propto \text{Distortion}^{-1} \end{aligned}$$





Update the library to prefer useful

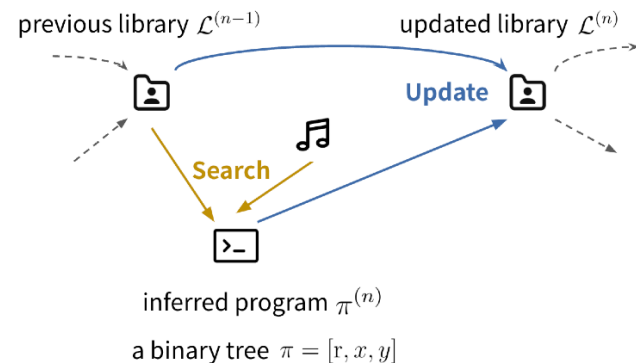
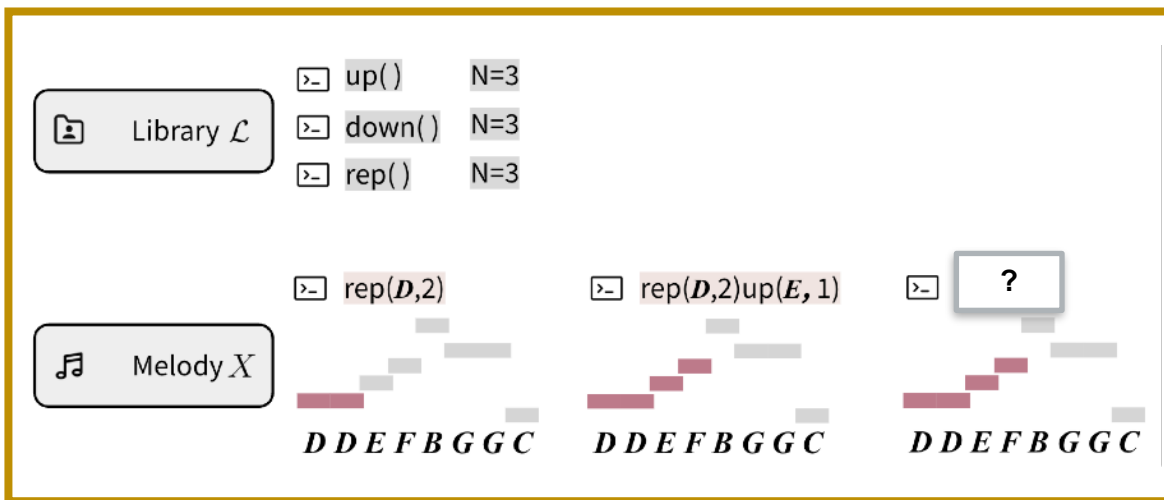
programs (across melodies):

- Initial prior: proportional to simplicity
- Updated prior: proportional to usage (entropy coding)

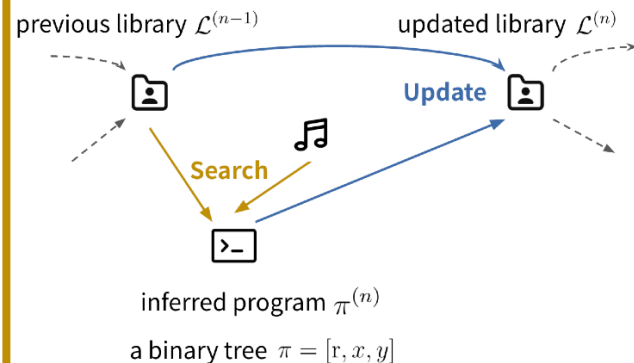
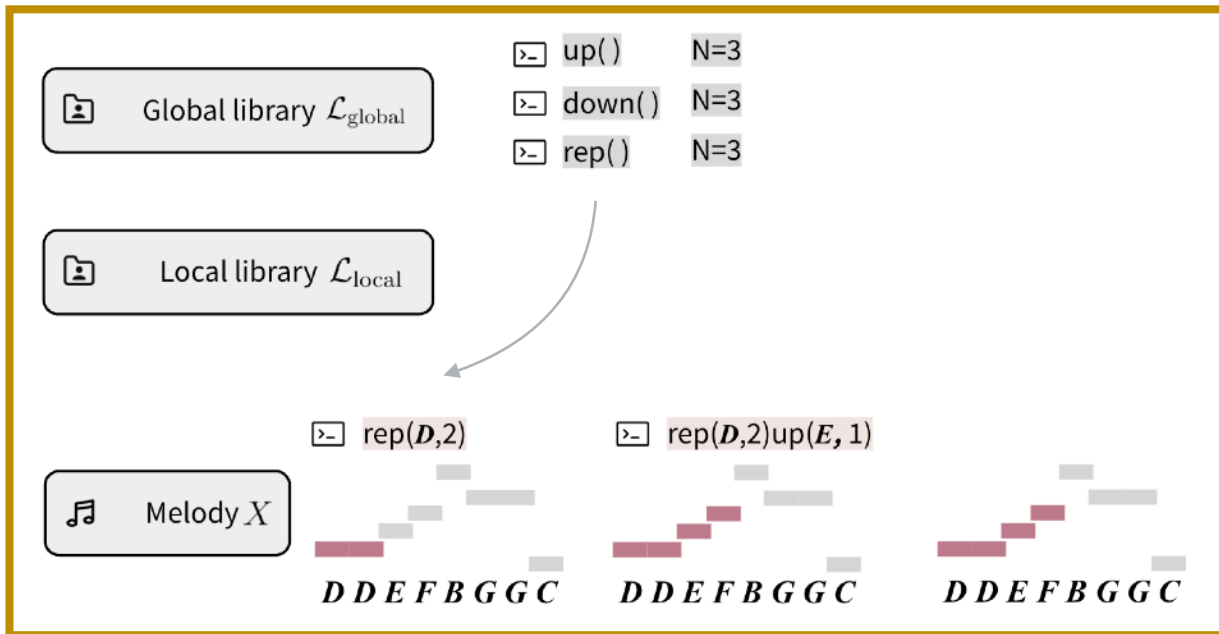
$$p(\pi | X^{(n)}) \propto p(X^{(n)} | \pi) p(\pi | X^{(1:n-1)})$$

Updated prior
 $\propto$ Usage frequency

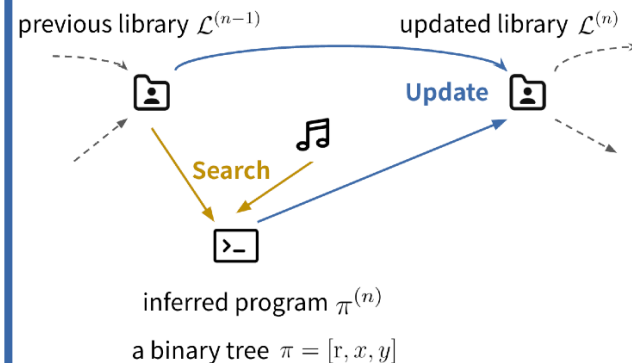
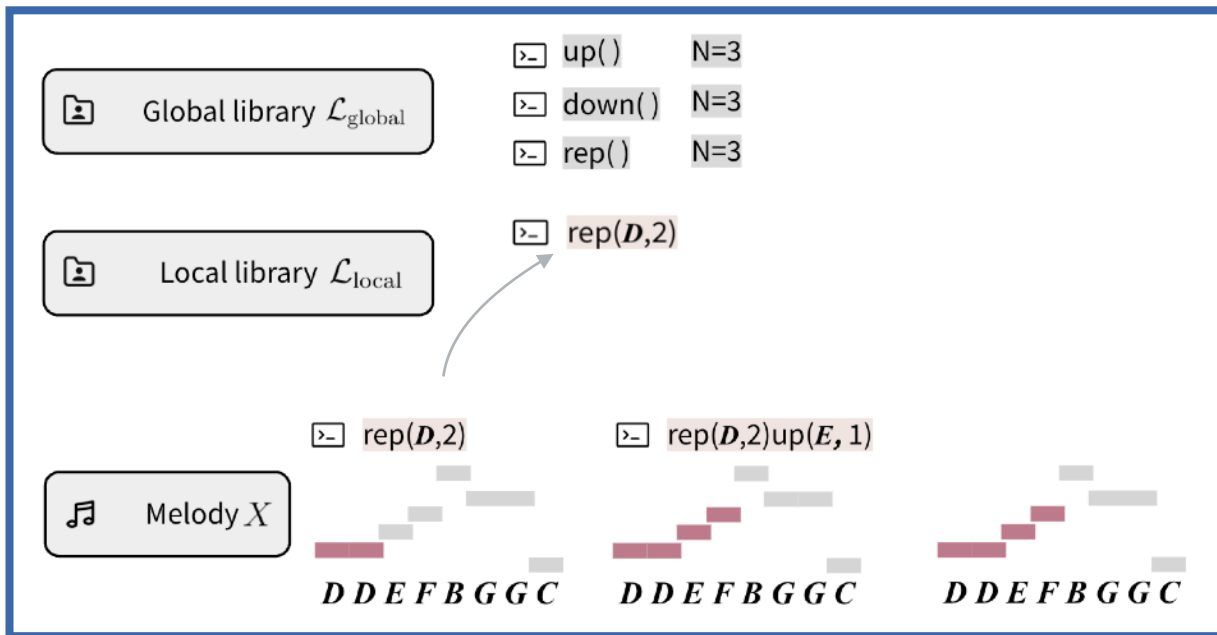
Augmented by 1) backtracking



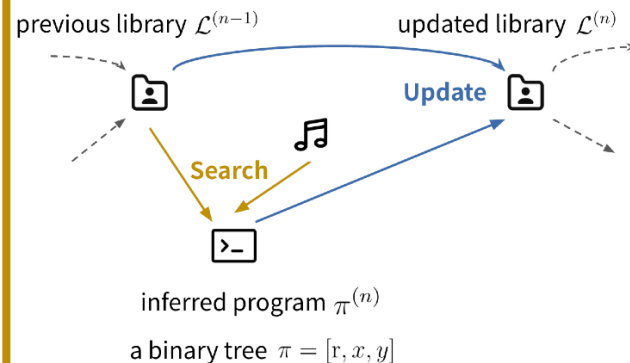
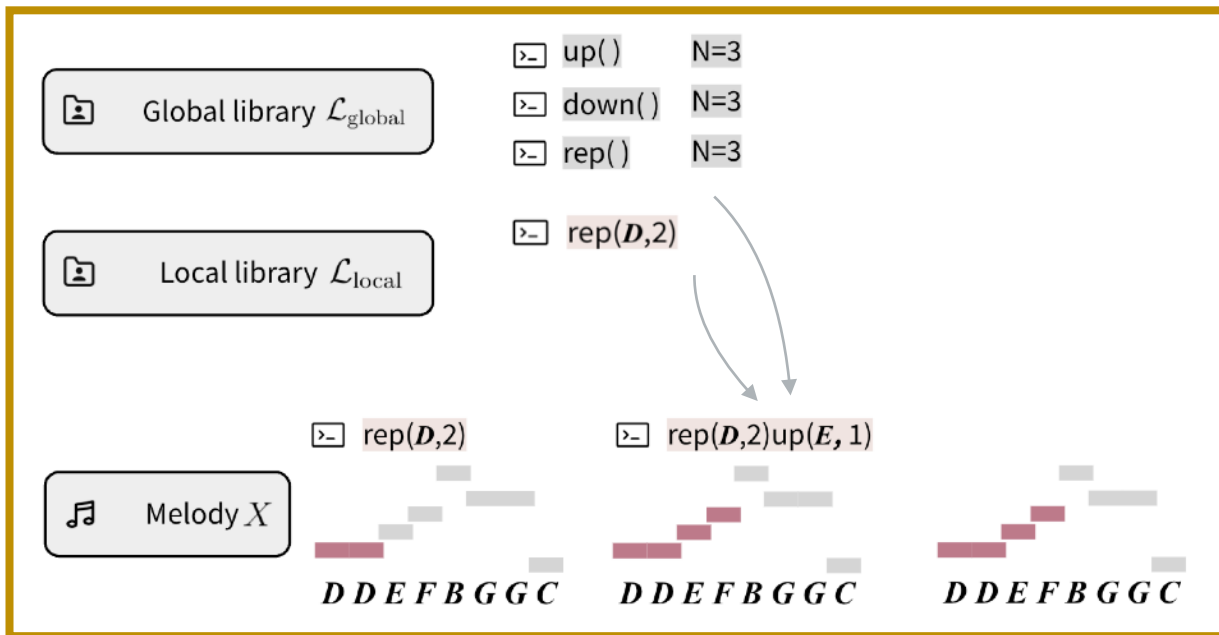
Augmented by 2) hierarchical library



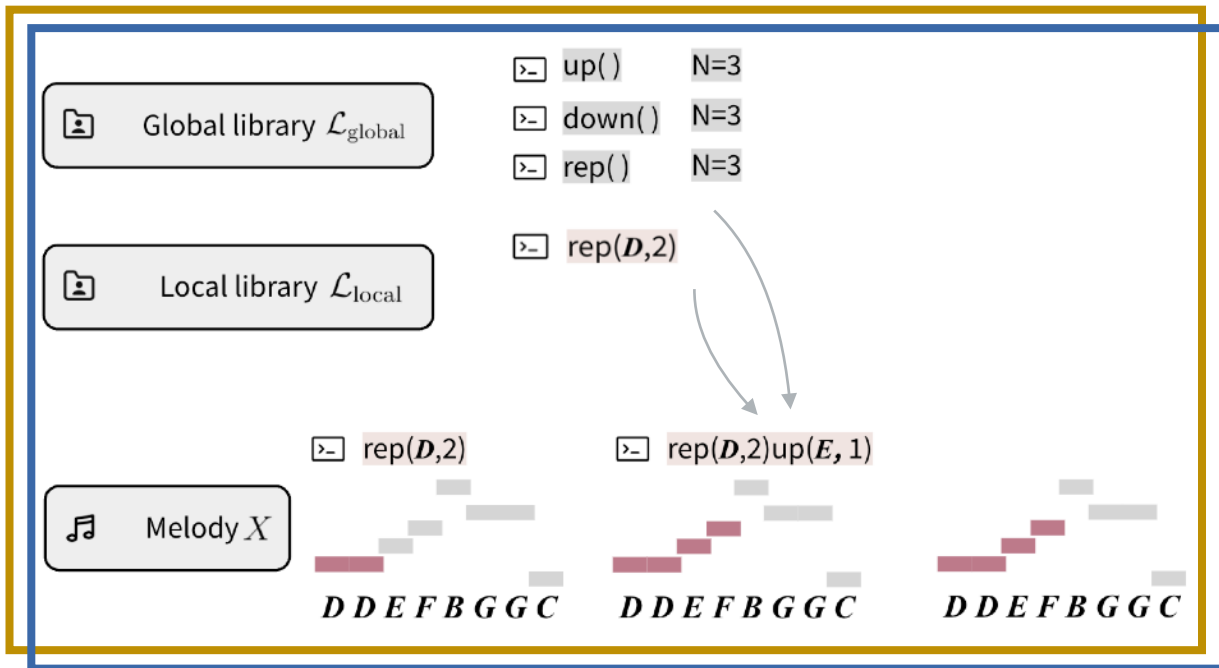
Augmented by 2) hierarchical library



Augmented by 2) hierarchical library



Augmented by 2) hierarchical library



All melodic
patterns

Pop, rock,
classical, jazz



All English
words

Fiction, history,
autobiography





Learning - recalling - predicting



Welcome to Chirp Champions, where every bird has a song and every song tells a story!

You are a young bird in a vibrant, musical forest, eager to learn the enchanting songs of your community. Your bird teachers will guide you in learning the songs of the various birds species in your area.

Experiment paradigm: ChirpChampion

Learning - recalling - predicting



Phase 1: Learning

Score: 23/100

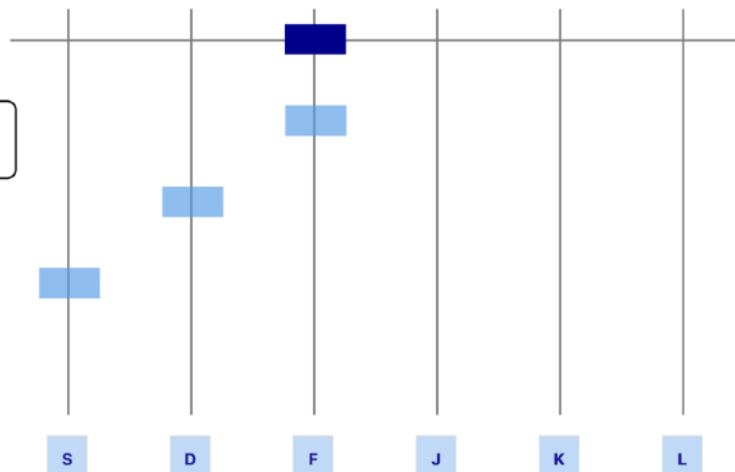
Notes: 3/13



Press the notes
highlighted in blue!

Your task:

- Repeat the note one by one.

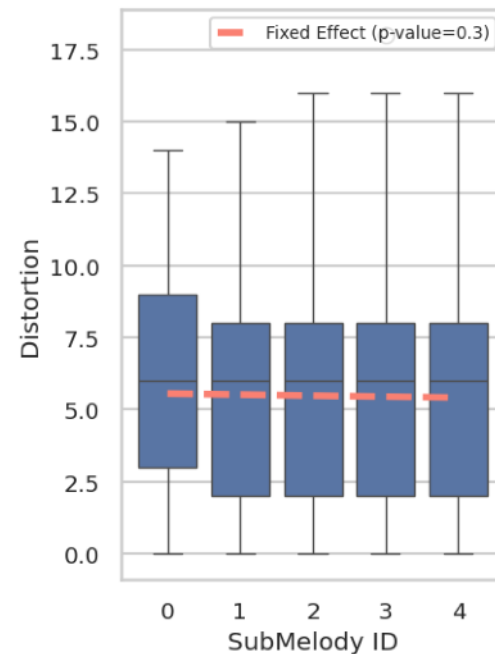
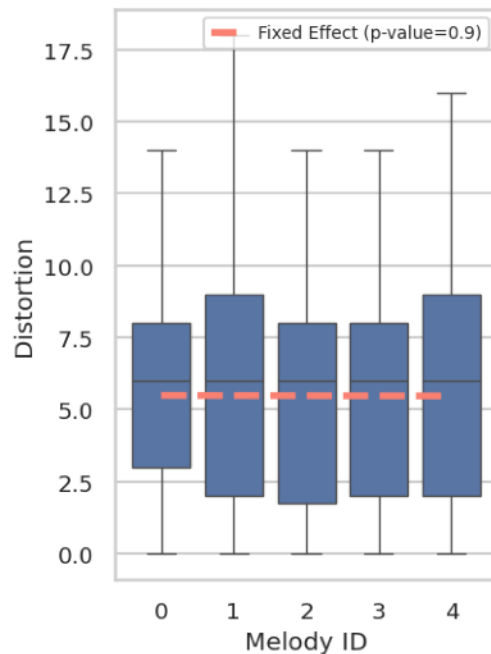


1. We do compression

- How can we explain performance (distortion in recall)?
- Hypothesis: Performance may improve with practice, but worsens as the melody becomes more ***complex***.

1. We do compression

- How can we explain performance (distortion in recall)?
- Hypothesis: Performance may improve with practice

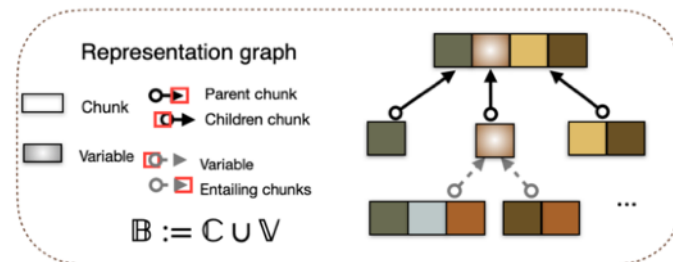


1. We do compression

- Hypothesis: Performance worsens as the melody becomes more **complex**.

distortion ~

program complexity
+ chunking complexity
+ entropy
+ length



order-1 transition probabilities

Planton, S., van Kerkoerle, T., Abbih, L., Maheu, M., Meyniel, F., Sigman, M., ... & Dehaene, S. (2021). A theory of memory for binary sequences: Evidence for a mental compression algorithm in humans. *PLoS computational biology*, 17(1), e1008598.

Wu, S., Thalmann, M., Dayan, P., Akata, Z., & Schulz, E. (2024). Building, Reusing, and Generalizing Abstract Representations from Concrete Sequences. *arXiv preprint arXiv:2410.21332*.

1. We do compression

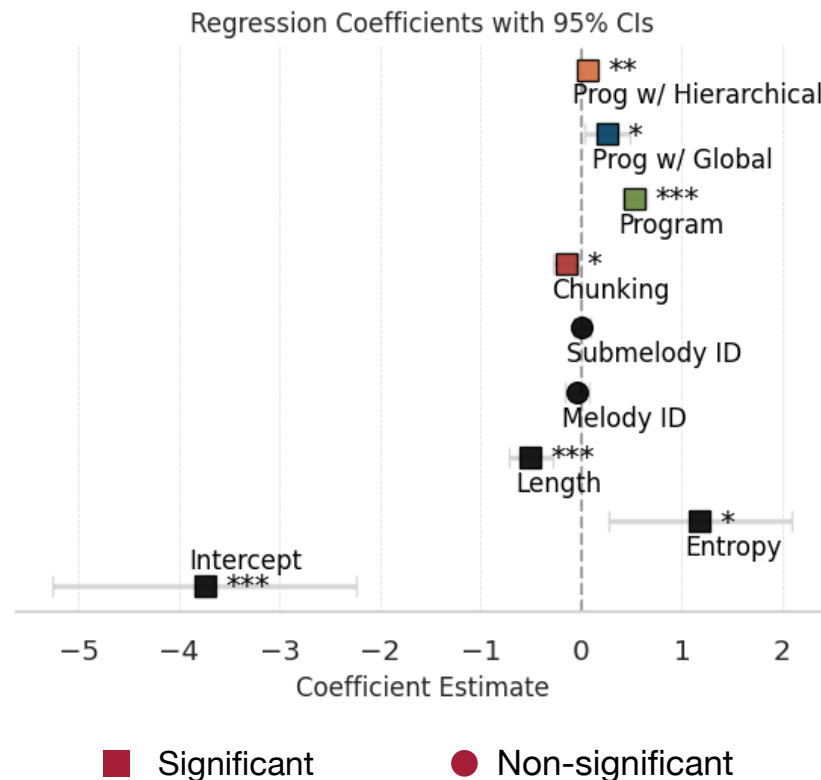
distortion ~

program complexity

+ chunking complexity

+ entropy

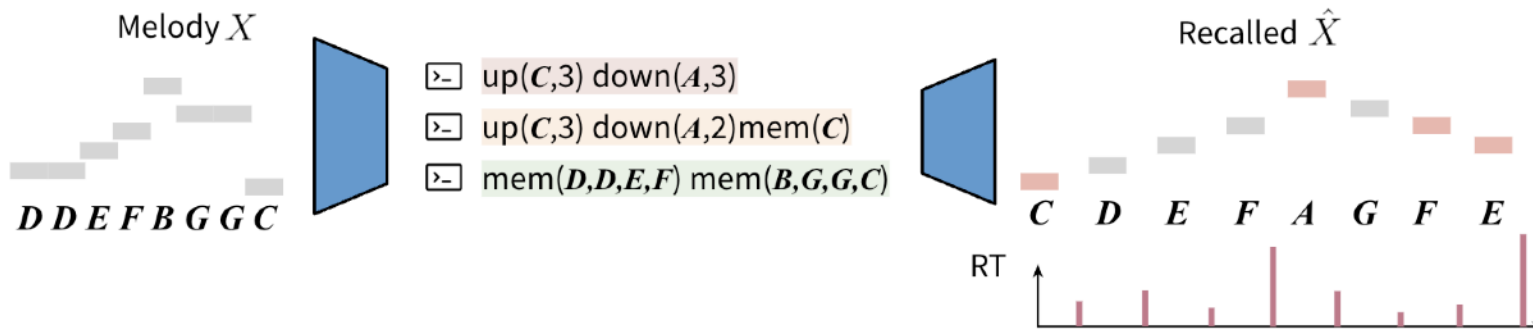
+ length



2. We do compression with programs

- 1) Reaction time: to identify 'chunking' boundaries, i.e., where participants naturally pause, indicating melody structure.
- 2) Error pattern: to identify whether participants make systematic errors.

2.1 Reaction time (human) vs. program boundary counts (model)



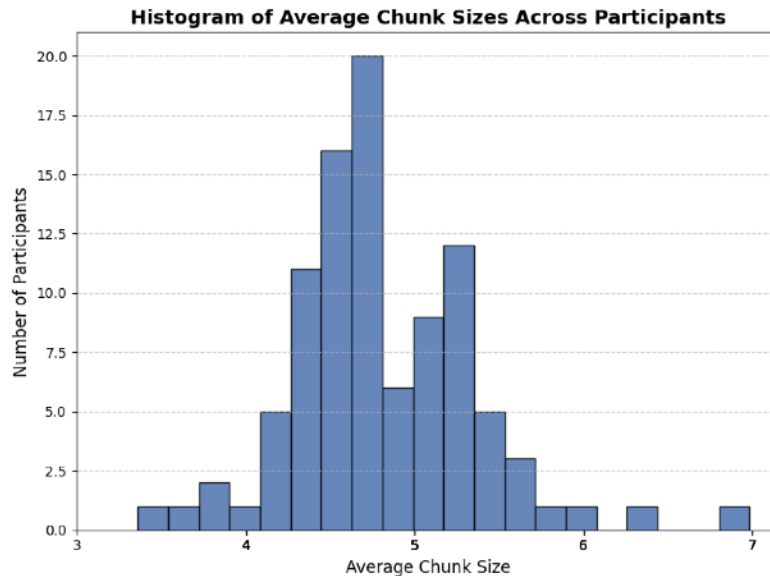


2.1 Reaction time (human) vs. program boundary counts (model)

Big changes in reaction time (RT) show when people are switching between mental “chunks”

- Compute first-order RT differences: $\Delta_i = RT_{i+1} - RT_i$
- Define a threshold to detect significant increases:

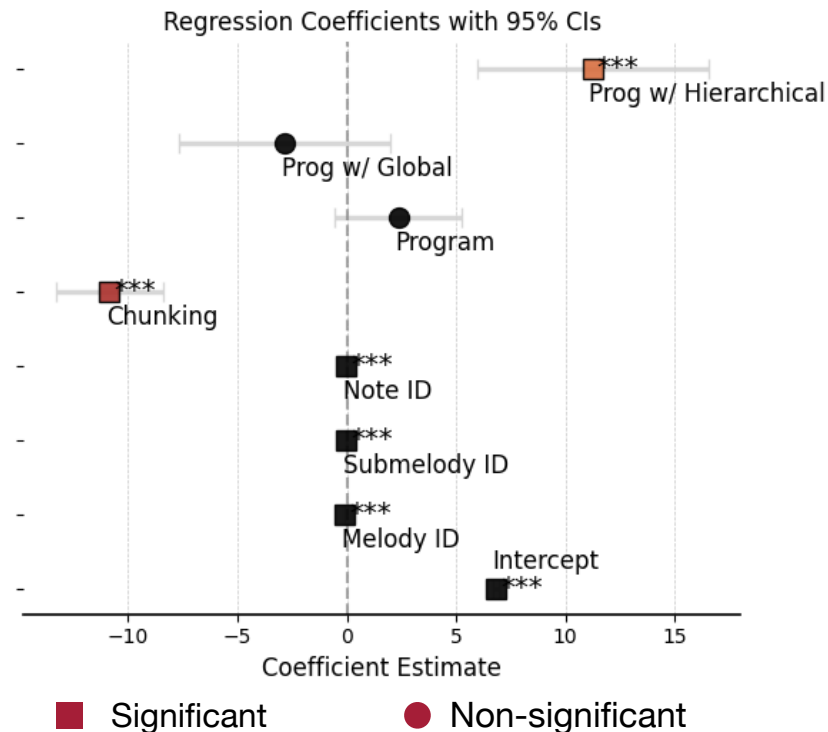
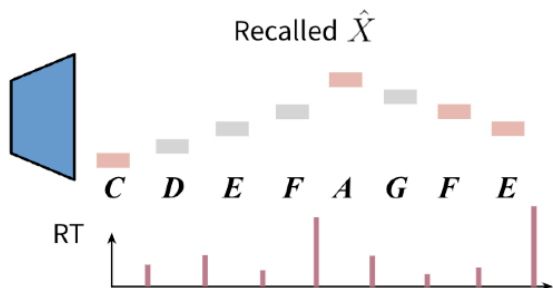
$$\text{Threshold} = \mu_{\Delta} + \lambda \cdot \sigma_{\Delta}$$



2.1 Reaction time (human) vs. program boundary counts (model)

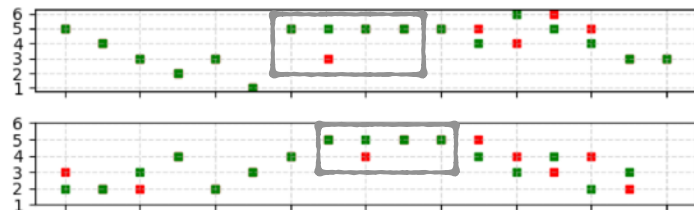
Reaction time (log space) $\sim$

- + Program boundary
- + Chunking boundary
- + Note index

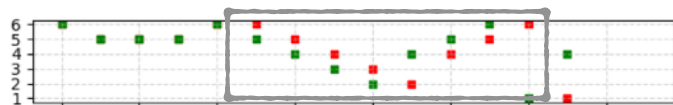


2.2 Error patterns

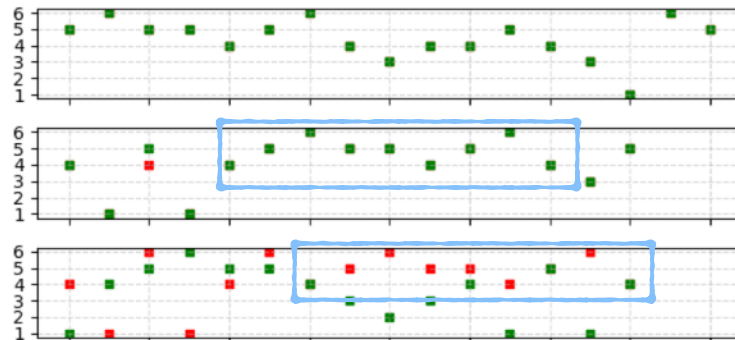
- Vertical shift



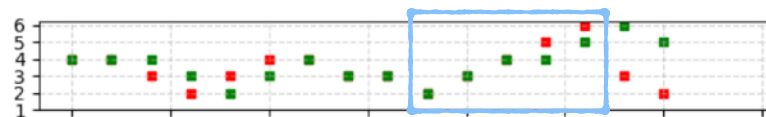
- Temporal shift



- Repeating



- Simple program



■ ground-truth

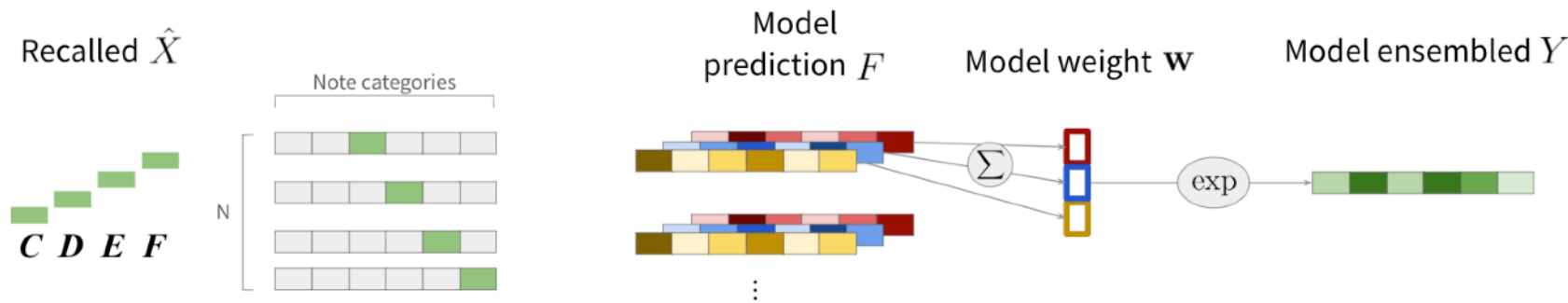
■ error

2.2 Recalled note sequence (human) vs. prediction (model)

Probability of pressing a key is based on a weighted sum of features:

$$\text{Choice}(x) \propto \exp(F \cdot \mathbf{w})$$

- F = feature matrix (what the model predicts)
- $\mathbf{w}$ = weights (learned from data)

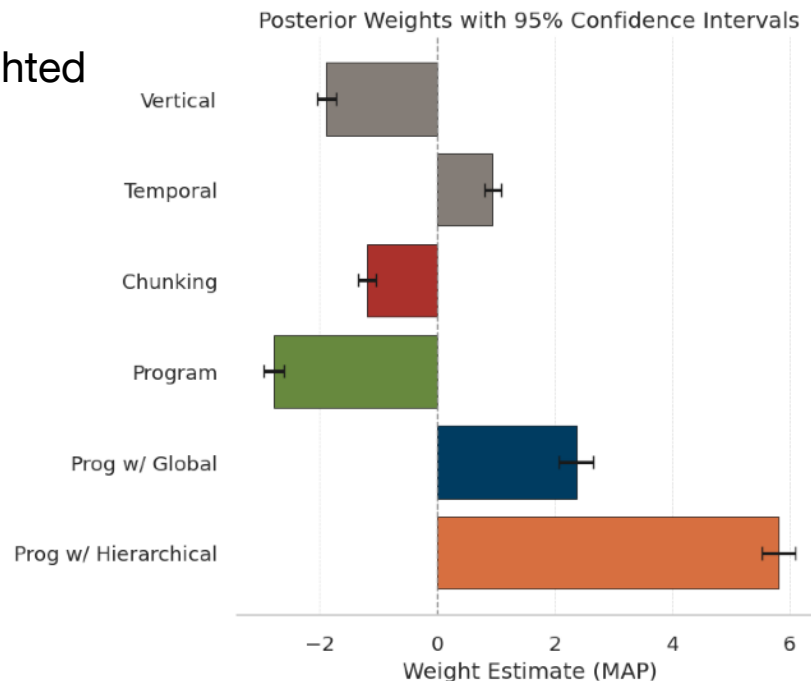


2.2 Recalled note sequence (human) vs. prediction (model)

Probability of pressing a key is based on a weighted sum of features:

$$\text{Choice}(x) \propto \exp(F \cdot \mathbf{w})$$

- F = feature matrix (what the model predicts)
- $\mathbf{w}$ = weights (learned from data)

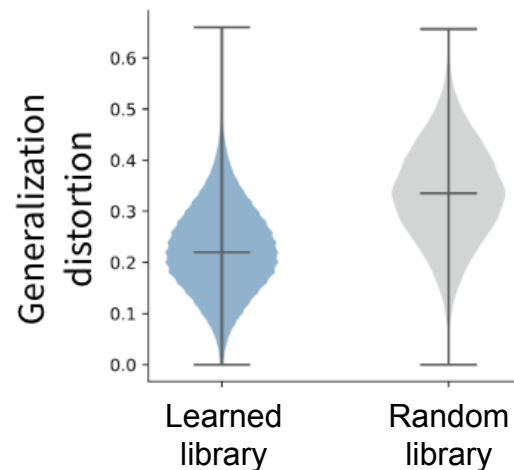




Sensitivity to curriculum

50 melodies under 100 different curricula:

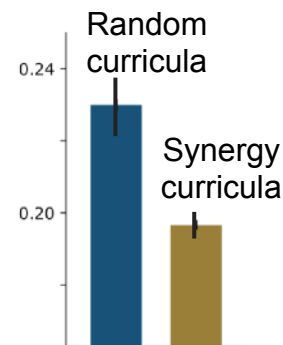
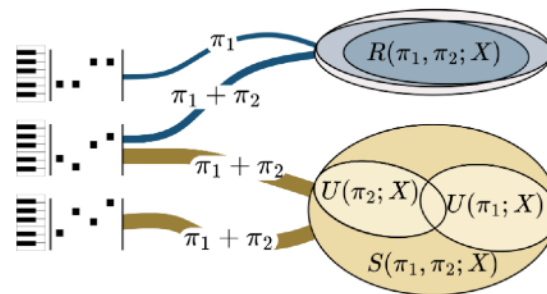
- Only ~3.19% of programs are the same under different curricula
- This cannot be accounted for by stochasticity of training alone
- Humans also show path-dependence

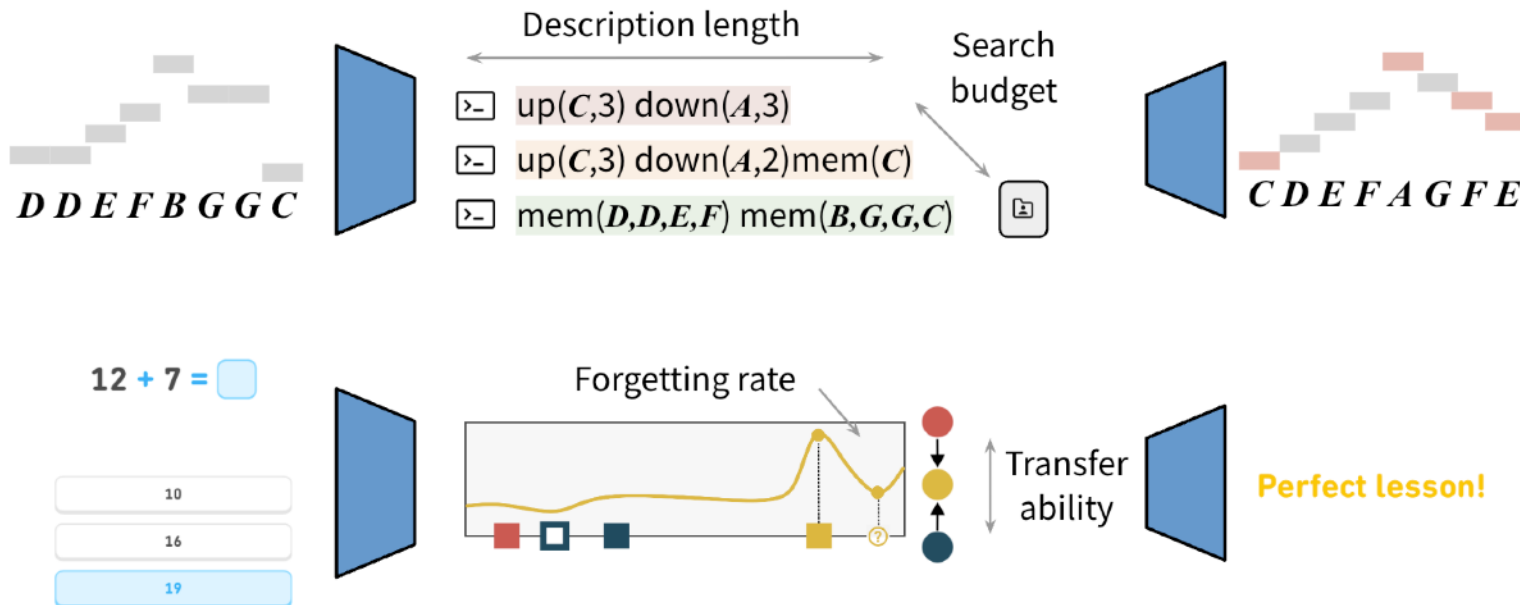


What is a good library?

A synergistic curriculum building method using principles of partial information decomposition

$$\begin{aligned} S(\mathcal{L}; X) &= I(\pi_1, \pi_2; X) \\ &\quad - R(\pi_1, \pi_2; X) \\ &\quad - U(\pi_1; X) \\ &\quad - U(\pi_2; X) \end{aligned}$$



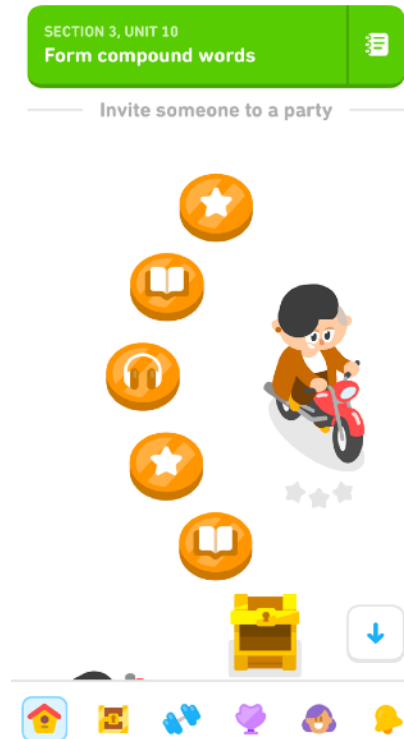
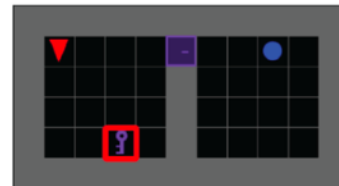
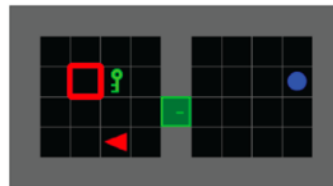




Curriculum in existing works



- Learning in the right order
- Given the target concept, learning with the right material



(Left) Campero, A., Raileanu, R., Küttler, H., Tenenbaum, J. B., Rocktäschel, T., & Grefenstette, E. (2020). Learning with amigo: Adversarially motivated intrinsic goals. arXiv preprint arXiv:2006.12122.

(Right) Figure source: (my own) Duolingo screenshot

Explanations in existing works

- What?

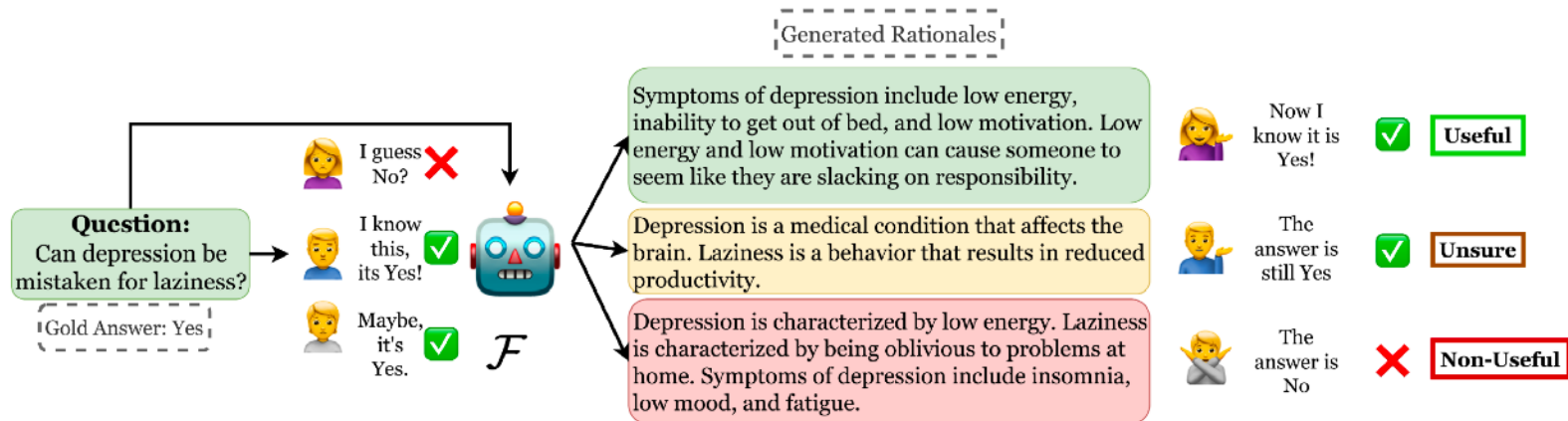


Figure 2: **An illustration of measuring human utility of machine rationales.** We evaluate whether a human's belief of the answer changes before and after seeing a rationale generated by an LM.

Explanations in existing works

- What & Why?



How...

- How to do it?
- How's my performance?
- How do I improve on it?
- How do you know all of this?
- ...

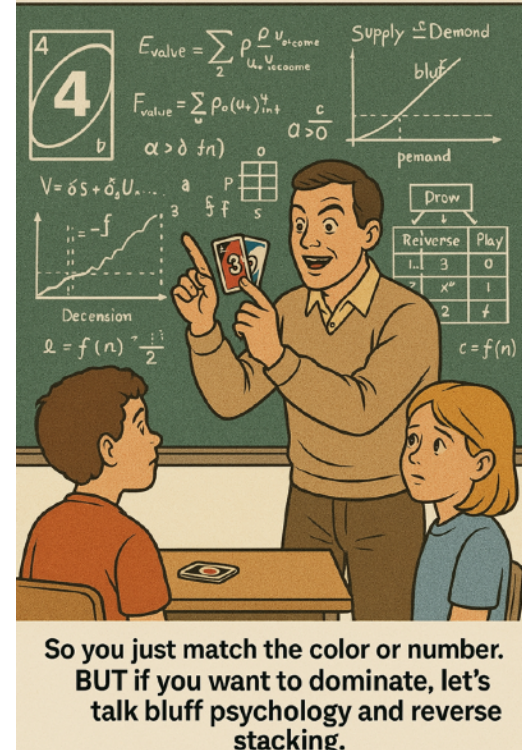
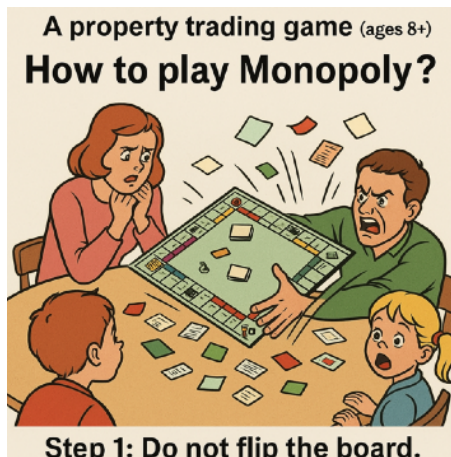


Figure source: GPT-4o (prompts are my own :)



Generate and evaluate explanations that are tuned to

- the learner's current knowledge and abilities, and
- pedagogical goals (engagement, curiosity, and learning transfer).

The big question is:

- What makes an explanation good, in context?
- And how can we design tools (LLMs) that don't just explain well, but explain like a great teacher would?





Has a task $T^{(n)}$ with a reward function R



- **Knowledge**

- Knows the task (and even the distribution $T \sim \mathcal{T}$)
- A (more “complete”) world model w_T
- Learner’s past experiences $T^{(1:n-1)}$

- **Goal**

- Extrinsic: help the learner finish the task with explanation e
- Intrinsic: pedagogy + unknown



- **Knowledge**

- A (limited) world model w_L

- **Goal**

- Extrinsic: finish the task and get the reward
- Intrinsic: unknown



Has a task $T^{(n)}$ with a reward function R



- What can be explained?
 - Policy π : here's exactly what you should do next.
 - World model w : here's how the terrain works -- pits block you, keys open gates.
 - Value function v : it's important to explore early, worth taking a detour for long-term gain.



- How much should be explained?
 - Memory load
 - Inference load

Has a task $T^{(n)}$ with a reward function R



- Inference

$$w^{(n)}, \pi^{(n)} \sim P_L(w, \pi \mid e, w_L^{(n-1)})$$

$$= P_L(\pi \mid w, e) P_L(w \mid e, w_L^{(n-1)})$$

- Utility function

$$U_L(e; T^{(n)}) = R(\pi^{(n)}) + R_L \quad \text{Reward}$$

$$- C(P_L) \quad \text{Learning cost}$$

	$e = \pi_T^{(n)}$	$e = w_T$	$e = v_T$
Compute	L	M	H
Memory	H	M	L

Has a task $T^{(n)}$ with a reward function R



- Inference

$$\hat{w}_L^{(n-1)} \sim P_T(w \mid T^{(1:n-1)})$$

$$\hat{w}_L^{(n)}, \hat{\pi}_L^{(n)} \sim P_T(w, \pi \mid e, \hat{w}_L^{(n-1)})$$

- Utility function

$$U_T(e; T^{(n)}) = R(\hat{\pi}^{(n)}) - C_T(P_L)$$

$$-C_T(e)$$

$$+\gamma \cdot \mathbb{E}_{T^{(n+1)}, \pi'} [R_{T^{(n+1)}}(\pi')] - C_T(P_L)$$

$$-C_T(P_T)$$

Learn's performance

Communication cost

Long-term pedagogy reward

Inference cost

- Inference

$$\begin{aligned} w^{(n)}, \pi^{(n)} &\sim P_L(w, \pi \mid e, w_L^{(n-1)}) \\ &= P_L(\pi \mid w, e) P_L(w \mid e, w_L^{(n-1)}) \end{aligned}$$

- Utility function

$$\begin{aligned} U_L(e; T^{(n)}) &= R(\pi^{(n)}) + R_L \\ &\quad - C(P_L) \end{aligned}$$





How to play the game?

How to get to the trader?

Walk

Locate

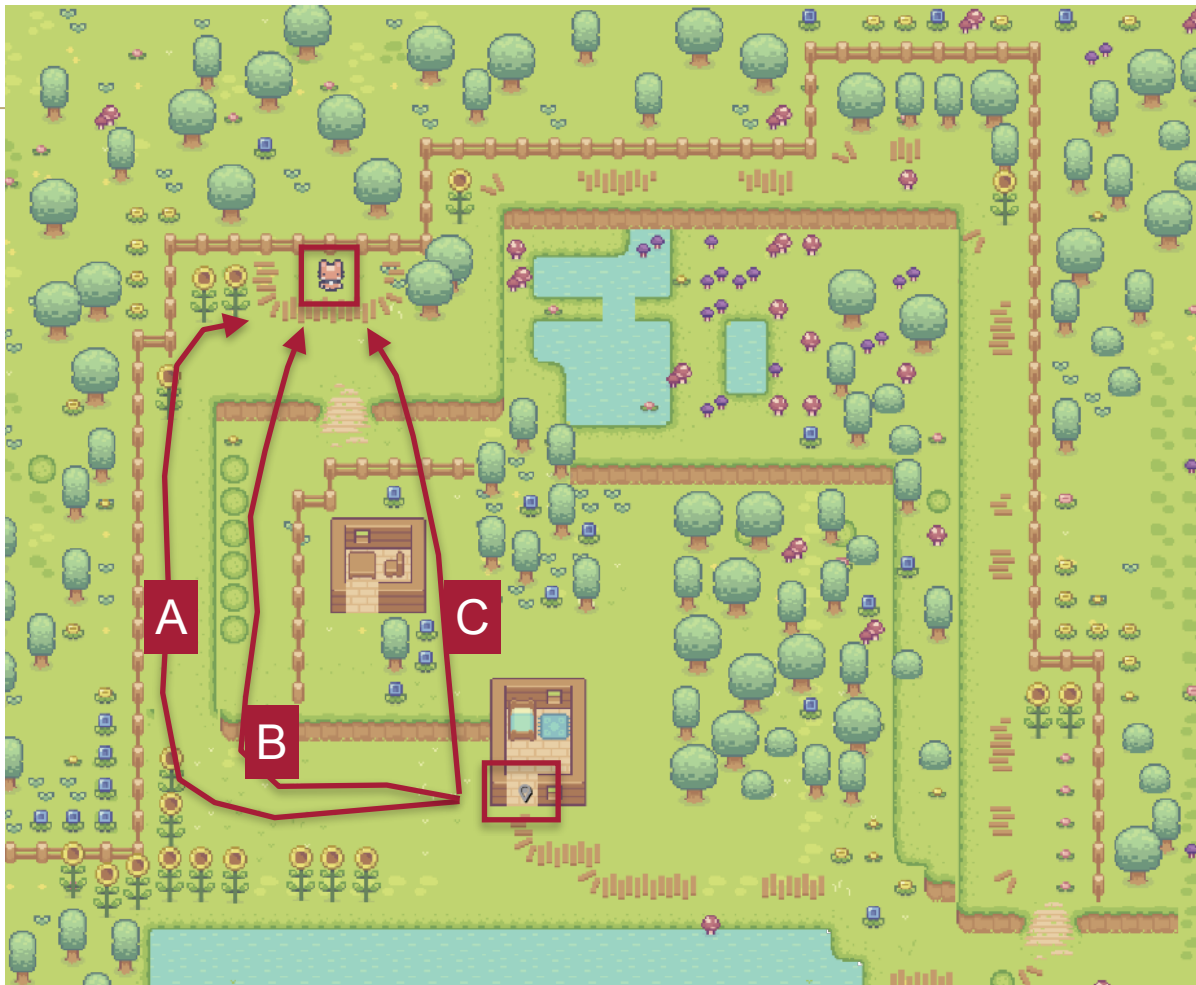
Clamber

Vault

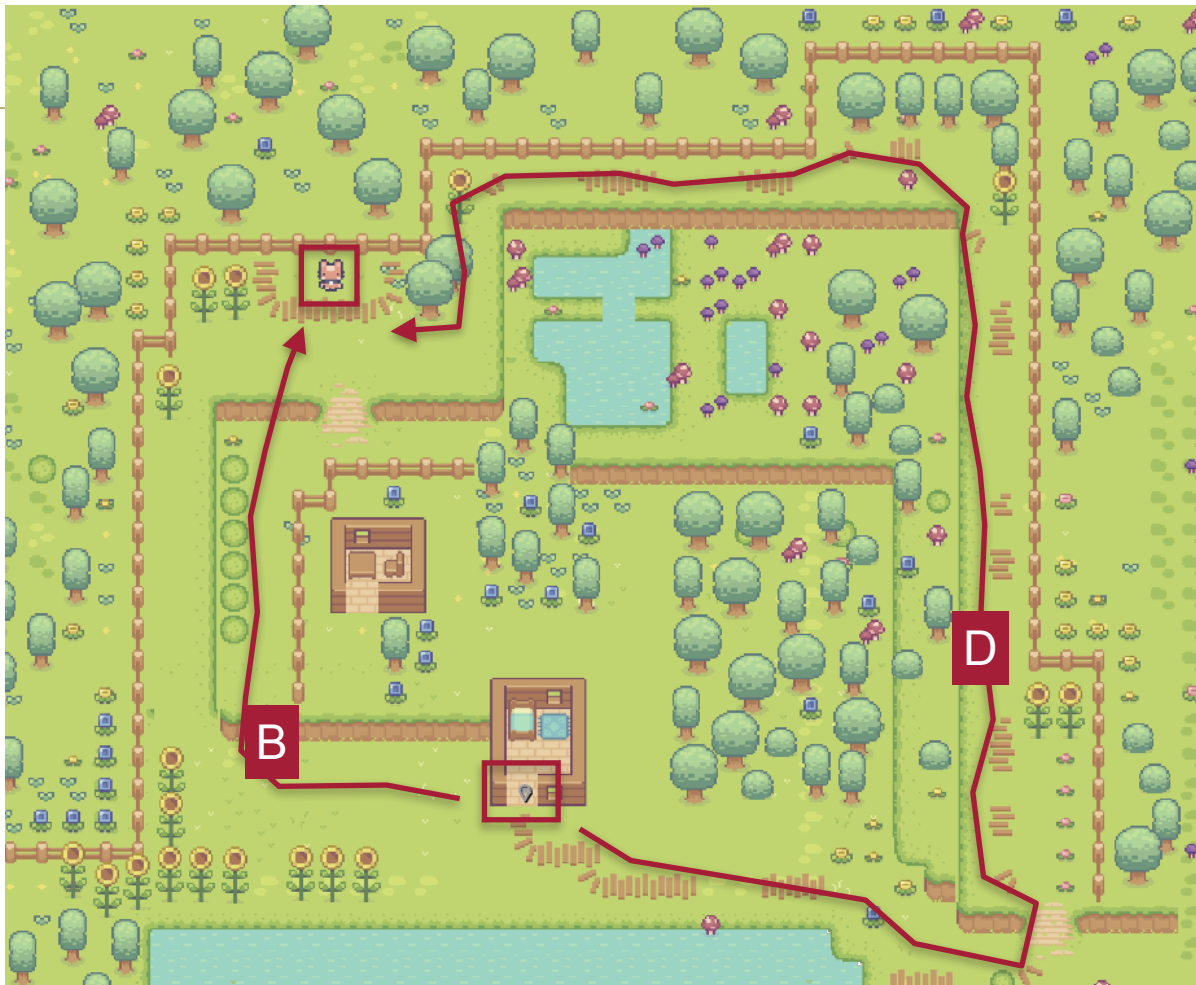
Hop

Forage

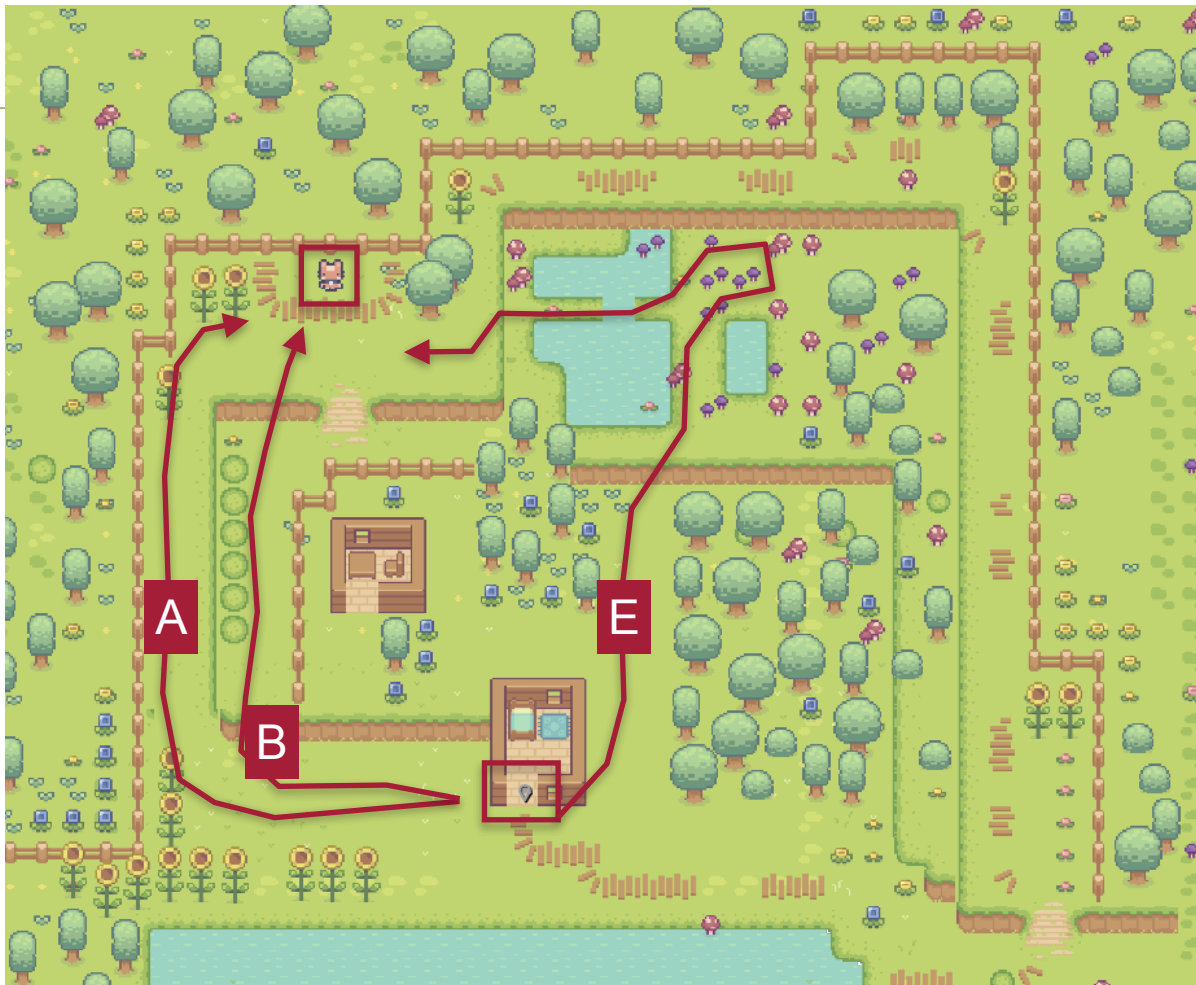




- Sensitive to learner's background
- A to B to C
 - Learner needs less effort (more reward)
 - Teacher needs more inference cost
 - Teacher needs more communication cost



- Sensitive to computation cost
- B and D:
 - Learner needs less effort on B (more reward)
 - Teacher needs more inference cost on B
 - Teacher needs more communication cost on B



- Sensitive to long-term goal
- A to B to E:
 - Learner needs more effort (more reward in the long run)
 - Teacher needs more inference cost
 - Teacher needs more communication cost



Thank you for listening!
Question?