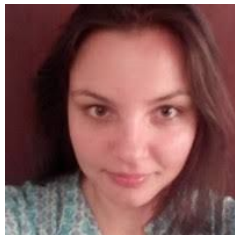


Reason to Rote: Rethinking Memorization in Reasoning

Yupei Du¹³, Philipp Mondorf², Silvia Casola², Yuekun Yao³, Robert Litschko²,
and Barbara Plank²

Utrecht University¹, LMU Munich², and Saarland University³



The curious case of benign memorization

Benign memorization of deep neural nets (Zhang et al., 2016)

Our central finding can be summarized as:

Deep neural networks easily fit random labels.

.. and yet

*it is unlikely that the regularizers are the fundamental reason for generalization,
as the networks continue to perform well after all the regularizers removed*

Understanding benign memorization in reasoning

Use of controllable reasoning tasks:

1. **Understanding** of the reasoning process:
which are the important tokens, which are the important intermediate steps,
and what is the correct answer
2. The reasoning process is **manipulatable**:
you can modify the intermediate step to (hopefully) elicit an alternative answer

Reasoning task: four-digit addition (FDA)

✗

3012 + 2473 = 7143

✓

1357 + 3548 = 4905

A0A1A2A3

B0B1B2B3

C0C1C2C3

Reasoning task: two-hop reasoning (THR)

Training profiles

Dana delegates tasks to Adam on a daily basis.
Drew helps Adam grow in their role.
Emma and Adam sometimes chat over the backyard fence.
Felix holds performance reviews for Dana. ...

Training questions

✗ Who is the crush of the mentor of Adam? Answer: Helen
R1 R0 P0 P2

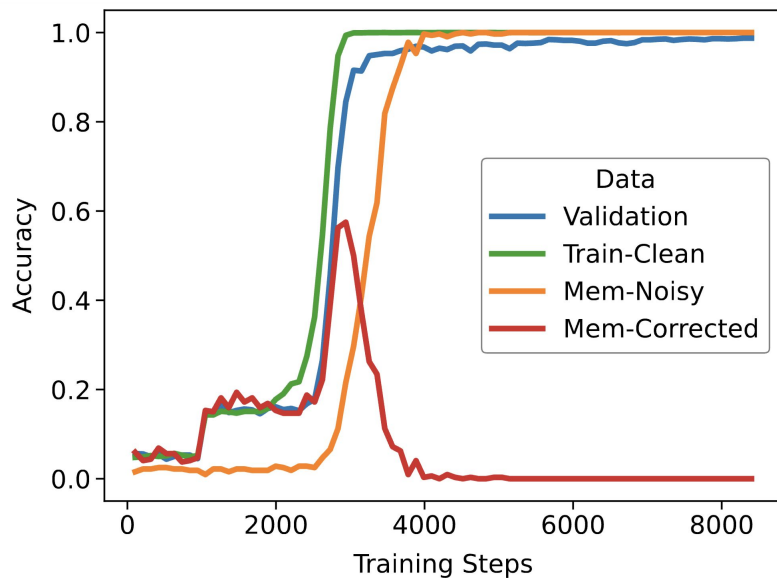
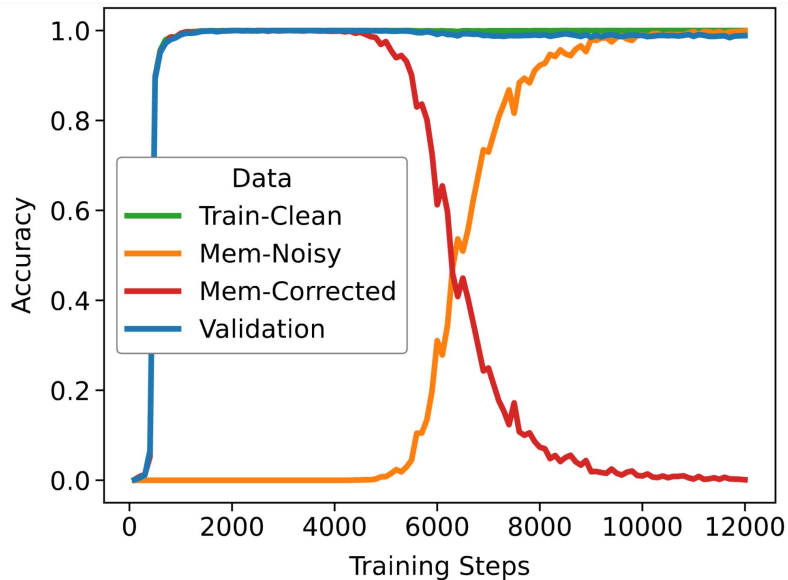
✓ Who is the debtor of the neighbor of Adam? Answer: Howard
R1 R0 P0 P2

Validation questions

? Who is the boss of the boss of Adam? Answer: _____
R1 R0 P0 P2

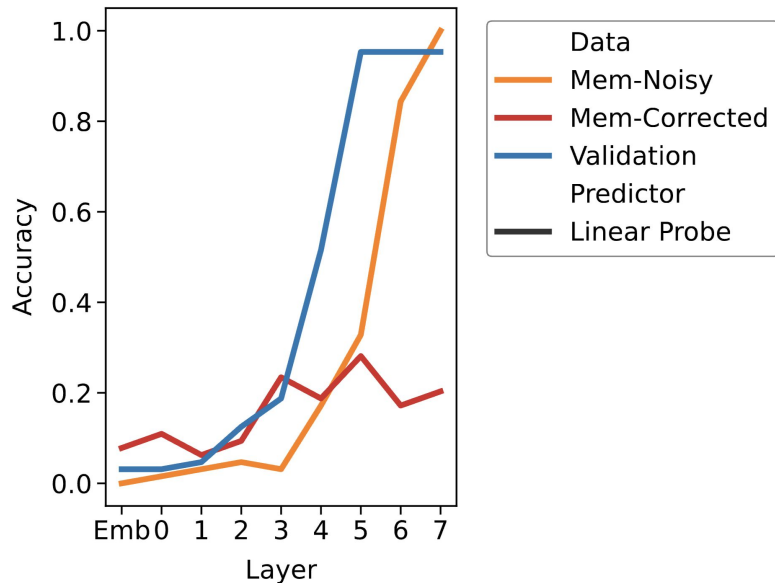
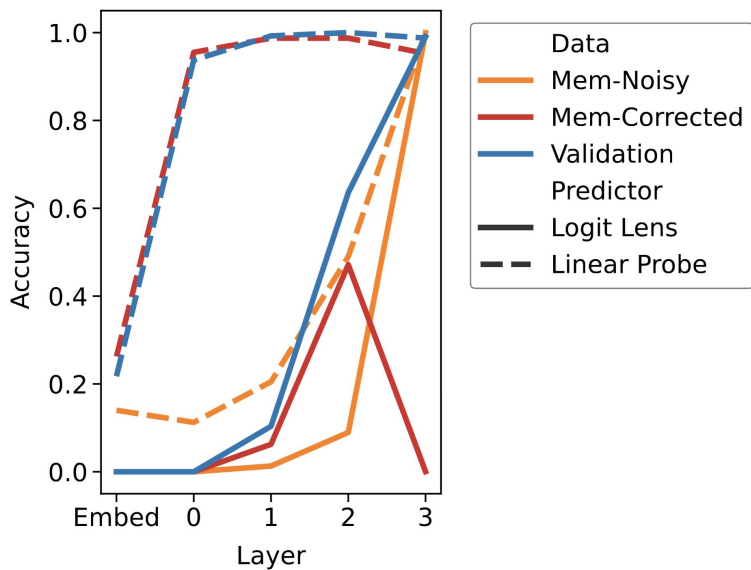
Warmup: first generalize, then memorize

Models first predict the clean labels on noisy training instances



Both mechanisms present in memorization

Models compute both results, and produce memorized noisy labels later



Memorization relies on generalization

Removing bridge entities hurts both memorization and generalization

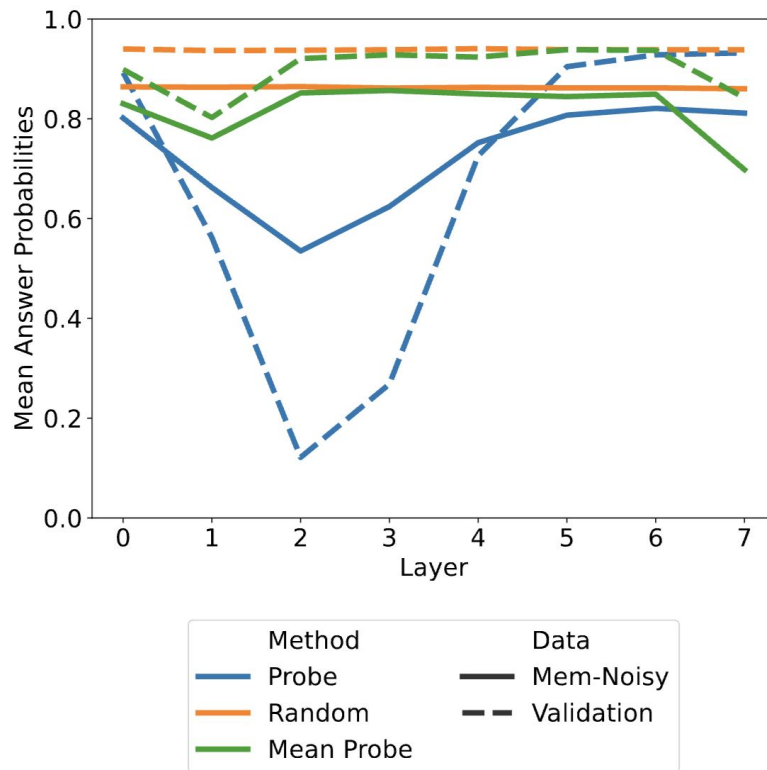
e.g., knowing facts

A's mentor is B; B's mentor is C;

but needs to memorize

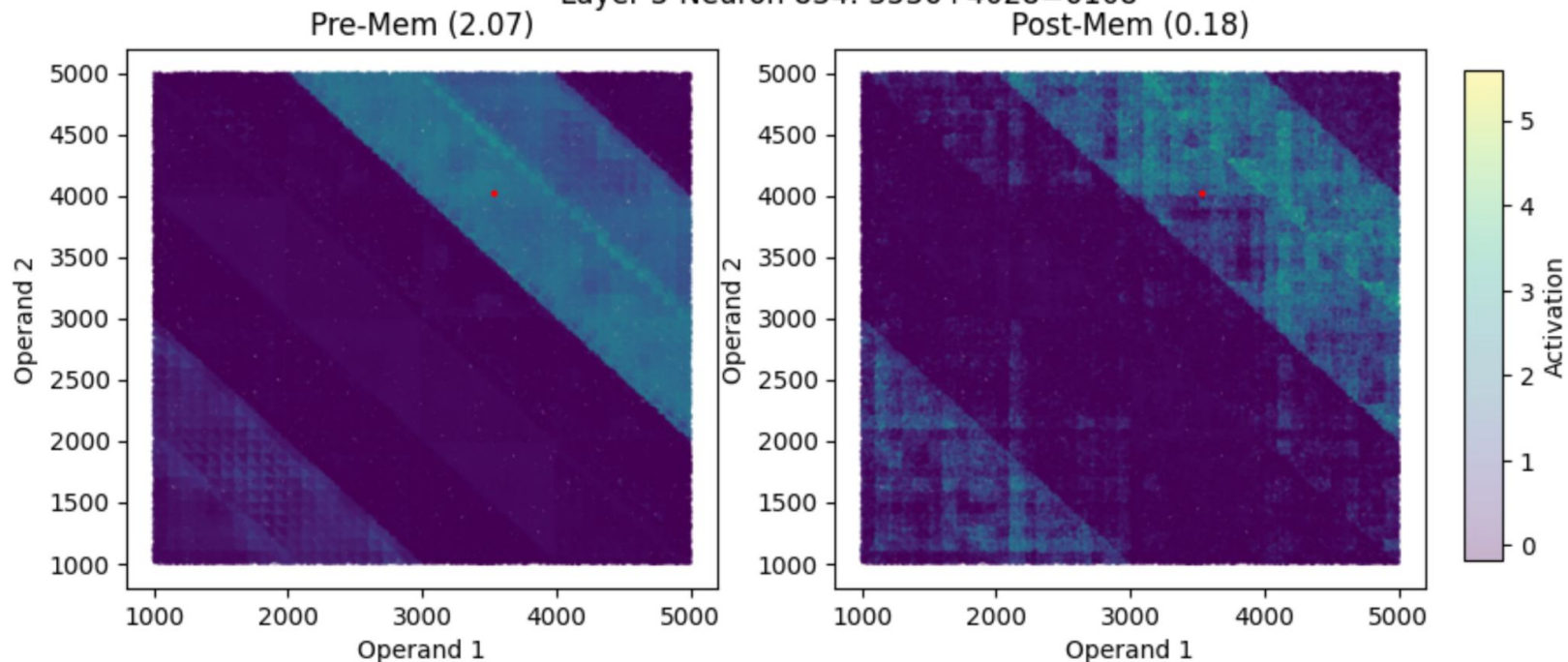
“Who is the mentor of the mentor of A”
to be D (should be C).

In this situation, **B** needs to be inferred
in hidden layers as well to recall **D**



FDA model memorizes noise through outlier heuristics

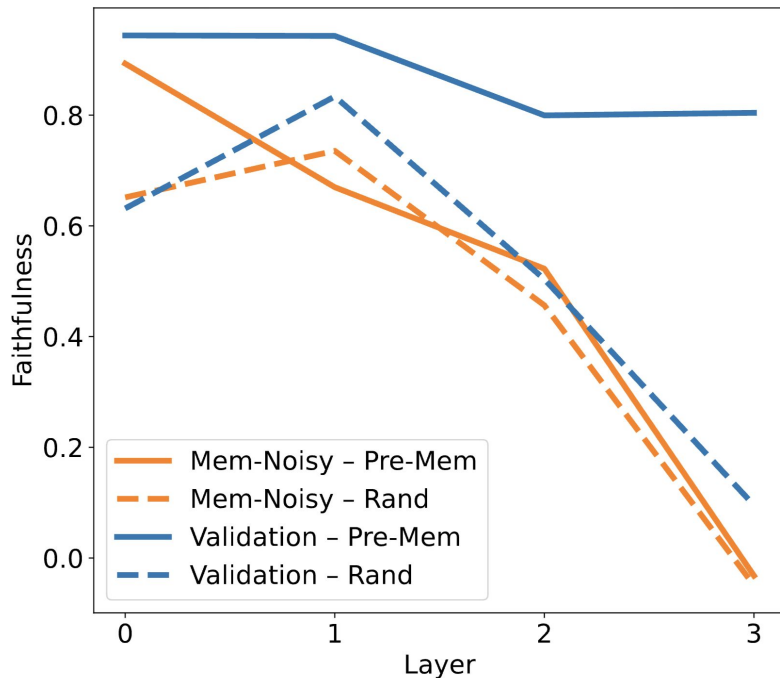
Layer 3-Neuron 834: $3536+4028=6108$



FDA model memorizes noise through outlier heuristics

Layer-wise activation patching of pre-memorization activations to post-activation model, ***without changing any parameters!***

- Validation faithfulness is barely influenced (validation accuracy is restored to 100%)
- Memorization faithfulness is severely undermined



Conclusions

1. Language models memorize noise by building on generalization mechanisms: this explains why such memorization is usually benign
2. In general, our study reveals the inductive bias of reusing existing structures to learn new things