

Agentic Recommender Systems: Foundations, Perspectives, and Lessons from Large Scale Deployments

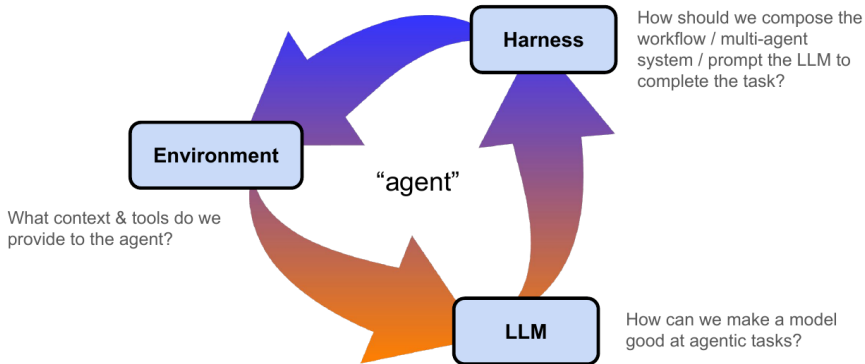
Reza Yousefi Maragheh¹ Yashar Deldjoo² Ben Coleman³ Jason Cho⁴ Chi Wang⁵

¹Walmart Global Tech, ²Polytechnic University of Bari, ³Google DeepMind, ⁴PretaskAI ⁵AG2AI

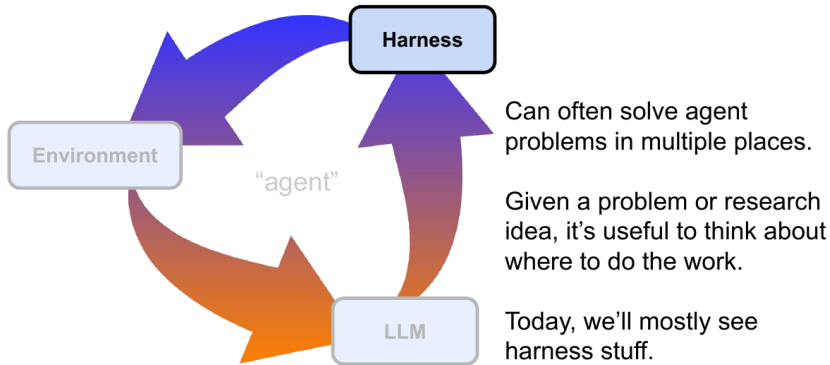
SIGIR 2026

Self Evolving Agents

What's in an agent?



What's in an agent?



Roadmap: Memory, Experiment Loops, & Feature Engineering

1. **Agent memory** - how do we make sure the agent remembers everything it needs to know?
2. **Experiment loops** - how do we get the agent to pursue good research directions?

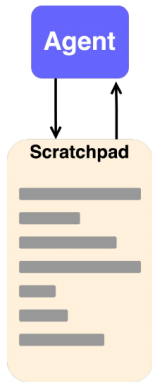
Then, let's look at a very concrete recsys application: feature engineering

In all three cases, we'll see how modifying the harness can solve problems.

Agent Memory

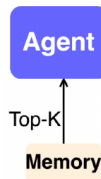
Context Accumulation & Retrieval

Agents accumulate* LOTS of context as they run. How to use the right context?



Some research tries to build a scratchpad like this

Other research tries to represent the context (e.g. as embeddings) and retrieve entries that are relevant to the task at hand

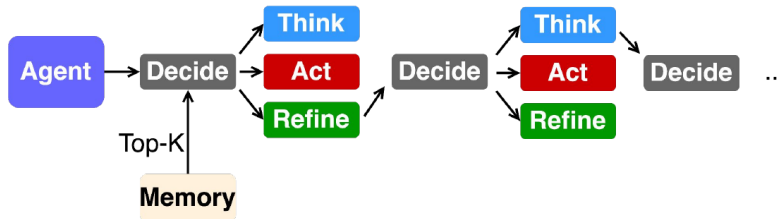


*Human researchers have accumulated a lot of context too. Think about all the stuff on arXiv!

Think-Act-Refine Baseline

Simple Memory Refinement Harness

Wanted to evaluate ideas (to use in applications) against a simple harness



Think: Generate a reasoning trace about the memory & task

Act: Stop messing around with memory and just do the task

Refine: Decide whether to discard some memories from the top-K

Answer uses the whole think-act-refine chain as context

But how do we know if it's any good?

At the time, no easy way
to compare these ideas
(so we built one)



Tianxin
Wei

Shoutout to our
student researcher!

Google DeepMind

2025-11-27

Evo-Memory: Benchmarking LLM Agent Test-time Learning with Self-Evolving Memory

Tianxin Wei¹, Naveen Sachdeva², Benjamin Coleman², Zhankui He², Yuanchen Bei¹, Xuying Ning¹, Mengting Ai¹, Yunzhe Li^{1,3}, Jingrui He¹, Ed H. Chi², Chi Wang², Shuo Chen², Fernando Pereira², Wang-Cheng Kang² and Derek Zhiyuan Cheng²

¹Work done while at Google DeepMind, ²University of Illinois Urbana-Champaign, ³Google DeepMind

Statefulness is essential for large language model (LLM) agents to perform long-term planning and problem-solving. This makes *memory* a critical component, yet its management and evolution remain largely underexplored. Existing evaluations mostly focus on static conversational settings, where memory is passively retrieved from dialogue to answer queries, overlooking the dynamic ability to accumulate and reuse *experience* across evolving task streams. In real-world environments such as interactive problem assistants or embodied agents, LLMs are required to handle continuous task streams, yet often fail to learn from accumulated interactions, losing valuable contextual insights, a limitation that calls for *test-time evolution*, where LLMs retrieve, integrate, and update memory continuously during deployment. To bridge this gap, we introduce Evo-Memory, a comprehensive streaming benchmark and framework for evaluating *self-evolving memory* in LLM agents. Evo-Memory structures datasets into sequential task streams, requiring LLMs to search, adapt, and evolve memory after each interaction. We unify and implement over ten representative memory modules and evaluate them across 10 diverse multi-turn goal-oriented and single-turn reasoning and QA datasets. To better benchmark experience reuse, we provide a baseline method, ExpRAG, for retrieving and utilizing prior experience, and further propose ReMem, an *action-think-memory refine* pipeline that tightly integrates reasoning, task actions, and memory updates to achieve continual improvement.

Keywords: LLMs, Agentic Memory, Test-time Learning, Self-evolving Agents, Lifelong Intelligence

How to "Evo-Memory" your benchmark

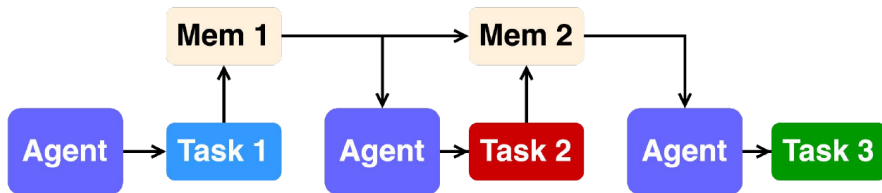
Step 1: Order the tasks in your benchmark into a sequence

(e.g., do MMLU problem #1 then #2, rather than in parallel)

Step 2: Allow the agent to maintain some state in between tasks

Step 3: Compare static performance with streamed-task performance

(the delta can be attributed to memory / experience reuse)



Think-Act-Refine performs surprisingly well

Performance Gains from Mem. Methods (Ours)

Consistent performance gains across models.

More significant gains on long-horizon complex tasks.

Claude 3.7 Sonnet

Method	AlfWorld		BabyAI		PDDL		ScienceWorld		Avg.	
	S	P	S	P	S	P	S	P	S	P
Baseline	0.18	0.49	0.51	0.66	0.17	0.39	0.10	0.53	0.24	0.52
History	0.50	0.73	0.48	0.66	0.65	0.85	0.32	0.74	0.49	0.74
ReAct	0.51	0.75	0.57	0.72	0.75	0.91	0.44	0.77	0.57	0.79
Amem	0.48	0.73	0.46	0.64	0.62	0.84	0.33	0.73	0.47	0.73
SelfRAG	0.52	0.75	0.46	0.64	0.65	0.84	0.31	0.74	0.49	0.74
Mem0	0.51	0.74	0.48	0.66	0.65	0.84	0.37	0.76	0.50	0.75
DC-Cu	0.50	0.74	0.50	0.67	0.62	0.84	0.33	0.75	0.49	0.75
DC-RS	0.50	0.74	0.52	0.68	0.62	0.84	0.34	0.74	0.50	0.75
AWM	0.49	0.73	0.53	0.68	0.60	0.82	0.34	0.74	0.49	0.74
ExpRecent	0.66	0.83	0.63	0.73	0.53	0.76	0.49	0.82	0.58	0.79
ExpRAG	0.74	0.89	0.62	0.72	0.72	0.89	0.46	0.76	0.63	0.82
ReMem	0.92	0.96	0.73	0.83	0.83	0.95	0.62	0.89	0.78	0.91

Takeaway

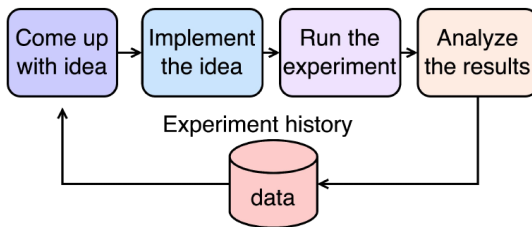
Lesson: harness modifications can be really powerful

- Simple setups beat complex methods if you get the abstractions right
- For memory retrieval + processing, “think-act-refine” seems like a solid workflow decomposition

Experiment Loops

AutoResearch & Evolution Workflows

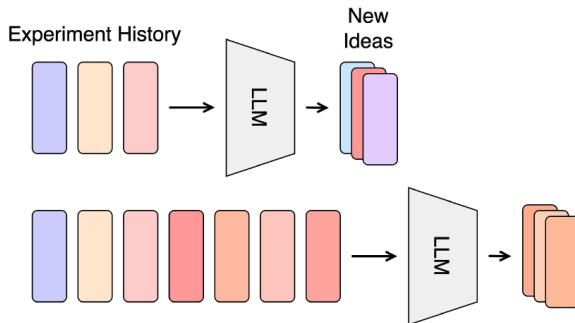
Experiment loops are another useful harness - think of AlphaEvolve, AutoResearch, etc



Challenges in Evolution

Context Pollution & Loss of Diversity

Context pollution decreases diversity of thought - less creative over time



Self-reinforcing feedback loop since LLM produces its own context

Exploration Bottlenecks

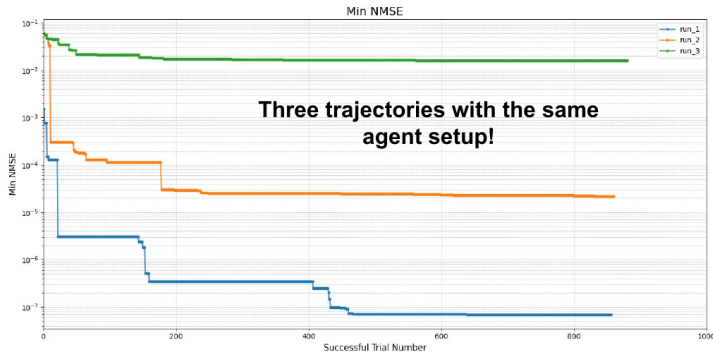
Stagnation in Long-Horizon Runs

```
- Results: nmse: 0.0476, ood_nmse: 0.00347. The performance degradation compared to the static symmetric model ('ood_nmse': 0.00167) suggests that quadratic drag ( $\sim v \cdot \text{abs}(v)$ ) is a near-equilibrium phenomenon by multiplying it with a spatial gating function.
- Results: nmse: 0.0477, ood_nmse: 0.00388. The degradation in generalization compared to the model with an ungated quadratic drag term ( $\sim v \cdot \text{abs}(v)$ ) suggests that the degradation is not due to the spatial gating function.
- A novel hysteresis model was tested by making the saturation rate of the position-dependent damping term ( $\sim \tanh(k \cdot \text{abs}(x))$ ) dependent on velocity.
- Results: nmse: 0.0476, ood_nmse: 0.00347. The performance degradation compared to the static symmetric model ('ood_nmse': 0.00167) suggests that the degradation is not due to the spatial gating function.
- Candidate made the amplitude of the successful saturating position-dependent damping term ( $\sim v \cdot \tanh(k \cdot \text{abs}(x))$ ) dependent on velocity.
- Results: nmse: 0.04657, ood_nmse: 0.00335. The performance degradation compared to the static model (ood_nmse: 0.00167) suggests that the degradation is not due to the spatial gating function.
- Candidate regularized the unbounded quadratic drag term ( $\sim v \cdot \text{abs}(v)$ ) with a velocity-dependent rational function ( $\sim 1/(1+k \cdot v)$ ).
- Results: nmse: 0.0487, ood_nmse: 0.00322. The worsened generalization (vs. 0.00167 from the model with an unbounded term) in the argument of the saturating damping term ( $\sim v \cdot \tanh(k \cdot \text{abs}(x))$ ) suggests that the degradation is not due to the spatial gating function.
- Candidate replaced the absolute position 'abs(x)' argument in the saturating damping term ( $\sim v \cdot \tanh(k \cdot \text{abs}(x))$ ) with the magnitude of velocity 'abs(v)'.
- Results: nmse: 0.0476, ood_nmse: 0.00371. The significant performance degradation (vs. 0.00167 from the 'abs(x)' model) refutes the gating hypothesis.
- Candidate tested a symmetric, speed-dependent skew ( $\sim \text{abs}(v)/(1+k \cdot v^2)$ ) in the argument of the saturating damping term ( $\sim v \cdot \tanh(k \cdot \text{abs}(x))$ ).
- Results: nmse: 0.0489, ood_nmse: 0.00332. Generalization worsened compared to the model with purely position-dependent saturation.
- Candidate tested if the saturation rate of the position-dependent damping term ( $\sim \tanh(k \cdot \text{abs}(x))$ ) was dynamically modulated by velocity.
- Results: nmse: 0.0476, ood_nmse: 0.00414. Generalization worsened compared to the static saturation model (ood_nmse: 0.00167).
- Candidate tested a bounded, hysteretic damping model by introducing a velocity-dependent skew ( $\sim \text{abs}(x) + c \cdot \tanh(d \cdot v)$ ) to the argument of the saturating damping term ( $\sim v \cdot \tanh(k \cdot \text{abs}(x))$ ).
- Results: nmse: 0.0434, ood_nmse: 0.00361. The significant performance degradation compared to the static symmetric model ('ood_nmse': 0.00167) suggests that the degradation is not due to the spatial gating function.
- Candidate tested for a viscoelastic restoring force by modulating the linear spring constant with a bounded function of velocity.
- Results: nmse: 0.0489, ood_nmse: 0.00294. The performance degradation (vs. 0.00167 from the best-performing model with purely position-dependent saturation) suggests that the degradation is not due to the spatial gating function.
- Candidate tested if quadratic drag ( $\sim v \cdot \text{abs}(v)$ ) is a large-displacement phenomenon by multiplying it with a 'tanh(k * x ** 2)' function.
- Results: nmse: 0.0493, ood_nmse: 0.00276. The performance degradation (vs. 0.00167 from the ungated model) refutes the gating hypothesis.
```

Eventually, the agent stops trying new things entirely...

This results in high variance

Very noisy process - sometimes you discover something, sometimes you don't



Introducing: PACEvolve



Minghao
Yan

Shoutout to our
student researcher!

PACEvolve: Enabling Long-Horizon Progress-Aware Consistent Evolution

Minghao Yan^{1,2} Bo Peng^{*1} Benjamin Coleman^{*3} Ziqi Chen¹ Zhouhang Xie^{1,4} Shuo Chen³ Zhankui He³
Novreen Sachdeva³ Isabella Ye¹ Welli Wang¹ Chi Wang³ Ed H. Chi¹ Fernando Pereira³
Wang-Cheng Kang³ Derek Zhiyuan Cheng³ Beidou Wang¹

Abstract

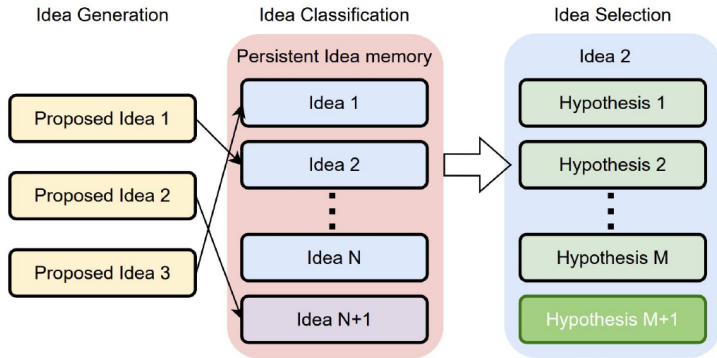
Large Language Models (LLMs) have emerged as powerful operators for evolutionary search, yet the design of efficient search scaffolds remains ad hoc. While promising, current LLM-in-the-loop systems lack a systematic approach to managing the evolutionary process. We identify three distinct failure modes: *Context Pollution*, where experiment history biases future candidate generation; *Mode Collapse*, where agents stagnate in local minima due to poor exploration-exploitation balance; and *Weak Collaboration*, where rigid crossover strategies fail to leverage parallel search trajectories effectively. We introduce **Progress-Aware Consistent Evolution (PACEvolve)**, a framework designed to robustly govern the agent's context and search dynamics, to address these challenges. PACEvolve combines hierarchical context management (HCM) with pruning to address context pollution; momentum-based backtracking (PBB) to escape local minima; and a self-adaptive sampling policy that unifies backtracking and crossover for dynamic search coordination (CE), allowing agents to balance internal refinement with cross-trajectory collaboration. We demonstrate that PACEvolve provides a

and engineering problems (Novikov et al., 2025; Romera-Paredes et al., 2024; Lange et al., 2024; Cheng et al., 2025b; Lange et al., 2025). They transform evolutionary search by replacing the rigid, random operators of classical algorithms (such as mutation and crossover) with intelligent, context-aware reasoning (Novikov et al., 2025; Romera-Paredes et al., 2024; Lange et al., 2025). Unlike traditional Evolutionary Algorithms (EAs) that evaluate an extensive number of weakly-guided candidates ($> 10^6$ samples) (Fogel, 1988; Holland, 1992), LLM-driven agents leverage in-context evolution history to perform iterative refinement (Novikov et al., 2025). By treating the history as a dynamic knowledge base, these agents can theoretically learn from failures and perform meta-reasoning, shifting the paradigm toward sample-efficient, knowledge-guided optimization (Zhai et al., 2025; Lange et al., 2025).

Our work aims to use these intelligent search priors to unlock state-of-the-art performance in complex research and engineering tasks (Shojaee et al., 2024; Ouyang et al., 2025; Jordan et al., 2024). However, this new LLM-in-the-loop paradigm introduces significant instability, preventing the search from consistently leveraging the LLM's reasoning capabilities (Xia et al., 2025; Kim et al., 2025). Rather than steadily improving, these systems suffer from high variance (§2), often failing to produce reliable improvements due to the combined stochasticity of the LLM and the search process (Comanici et al., 2025; Renze, 2024).

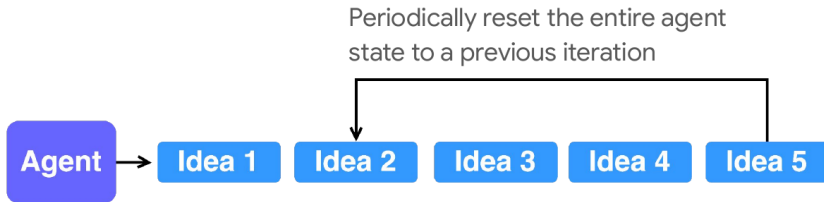
Workflow decomposition strikes again

Another workflow decomposition - separate ideation and selection into subagents



Plus other harness heuristics

We also do a few other things, such as backtracking when the system gets stuck



We do this based on a momentum heuristic (details in paper)

PACEvolve Results

SOTA Performance & Variance Reduction

By adopting this workflow decomposition, we can:

reduce the variance a lot and get SOTA results*

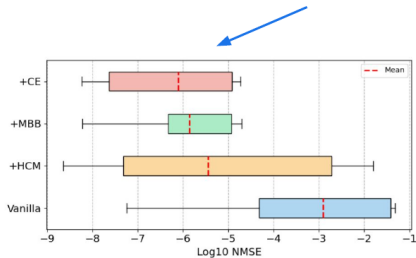


Table 1. Performance comparison on LLM-SR task. We report the best, P75, mean, and worst performance across 10 runs in terms of Log10 Normalized Mean Squared Error (Log10 NMSE). Best results are in **bold**, second-best results are underlined.

Method	Best	P75	Mean	Worst
uDSR	-4.06	-4.06	-3.95	-3.73
LLM-SR	-4.80	-4.03	-4.06	-3.62
OpenEvolve	-7.11	-5.79	-5.40	-4.02
CodeEvolve	-7.26	-5.54	-4.97	-3.99
ShinkaEvolve	-6.35	-5.92	-5.35	-3.42
PACEvolve-Single	<u>-8.23</u>	<u>-6.33</u>	<u>-5.87</u>	<u>-4.71</u>
PACEvolve-Multi	-8.24	-7.64	-6.11	-4.73

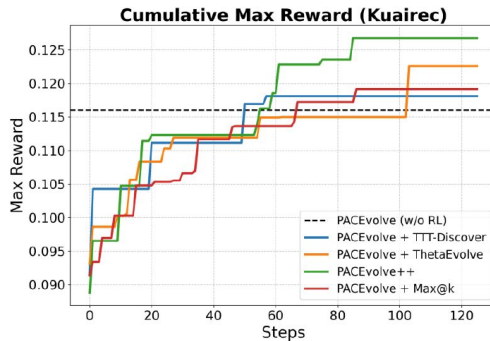
*See paper for more results - KernelBench, GPT-speedrun, etc

Can we solve the problem in the LLM rather than the harness?

We tried RL on the ideation model
to let the LLM learn from experiment
success / failures throughout the
evolution run

Reward shaping and stability are
challenging* but early results do show
gains

*See PaceEvolve++ paper for details



Semantic feature tagging in RecSys

Now let's see a more concrete recsys application: feature engineering



Item Description

Limited edition 40th Anniversary vinyl,
great audio quality...

Item ID & Semantic IDs

358, 271, 77, 82

**Depends on encoder
quality, uninterpretable*

Tag Features

Rock → Glam Rock → Proto Punk

**Exploits LLM reasoning,
interpretable, NL native*

We want the LLM to produce a hierarchy of useful semantic tags

A naive agent / workflow does not work well

If we just ask an LLM (or agent) to generate tags, we get

- long, one-off tags
"hardwood desk for computer monitor in home office"
- inconsistent / soft-duplicates
"camera bag" versus "camera case"
- way too many tags
Tag cardinality > number of items

Tags with these characteristics are hard to use in production and don't drive semantic generalization.

Introducing: AgenticTagger



Zhouhang
Xie

Shoutout to our
student researcher!

AgenticTagger: Generating Structured Item Representation for Recommendation with LLM Agents

Zhouhang Xie*
UC San Diego
La Jolla, CA, USA
zhx022@ucsd.edu

Ziqi Chen*
Google
Mountain View, CA, USA
zqchen@google.com

Benjamin Coleman
Google DeepMind
Mountain View, CA, USA
colemanben@google.com

Julian McAuley
UC San Diego
La Jolla, CA, USA
jmcauley@ucsd.edu

Bo Peng*
Google
Mountain View, CA, USA
bopen@google.com

Alice Han
Google
Mountain View, CA, USA
siperalice@google.com

Naveen Sachdeva
Google DeepMind
Mountain View, CA, USA
naveen@google.com

Wang-Cheng Kang
Google DeepMind
Mountain View, CA, USA
wckang@google.com

Beidou Wang
Google
Mountain View, CA, USA
beidou@google.com

Zhankui He*
Google DeepMind
Mountain View, CA, USA
zhankui@google.com

Isabella Ye
Google
Mountain View, CA, USA
isabellaye@google.com

Fernando Pereira
Google DeepMind
Mountain View, CA, USA
pereira@google.com

Derek Zhiyuan Cheng
Google DeepMind
Mountain View, CA, USA
zcheng@google.com

Randolph Brown
Google
Mountain View, CA, USA
rgb@google.com

Abstract

High-quality representations are a core requirement for effective recommendation. In this work, we study the problem of LLM-based descriptor generation, i.e., keyphrase-like natural language item representation generation frameworks with minimal constraints on downstream applications. We propose AgenticTagger, a framework that queries LLMs for representing items with sequences of text descriptors. However, open-ended generation provides little control over the generation space, leading to high cardinality, low-performance descriptors that render downstream modeling challenging. To this end, AgenticTagger features two core stages: (1) a vocabulary-building stage in which a set of hierarchical, low-cardinality, and high-quality descriptors is identified, and (2) a



Item Description

Limited edition 40th Anniversary vinyl,
great audio quality...

Item ID & Semantic IDs

158, 271, 77, 82

*Depends on encoder
quality, uninterpretable

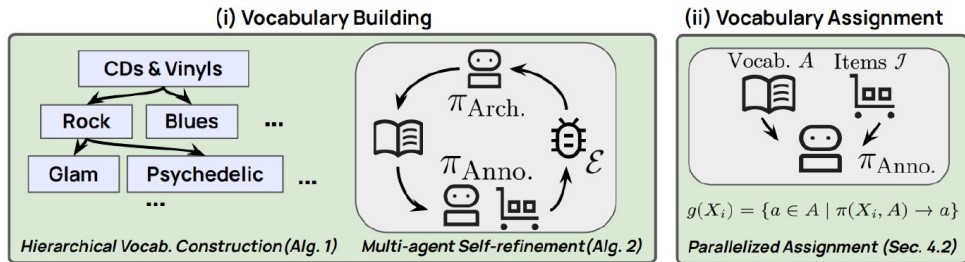
AgenticTagger Features

Rock → Glam Rock → Proto Punk

*Requires LLM reasoning,
interpretable

AgenticTagger Architecture

The core approach* is again a multi-agent workflow decomposition



*making many simplifications for this talk - see paper for details

AgenticTagger Results

Vocabulary Control & Downstream Gains

Adopting this vocabulary building + assignment setup fixes the problems

Item Context	Free-form	AgenticTagger
Ground Truth <i>Transformer stock,</i> <i>820% 3-yr returns</i>	Skipper Limited, Order Book Growth, Power Transmission.	L0: Business & Finance L1: Stock Market L2: Individual Performance
User History (1) <i>Defence bidder, 7k%</i> <i>5-yr returns</i> (2) <i>Industrial PSU, 200%</i> <i>5-mo returns</i>	(1) Defence PSU, DRDO Bidding. (2) PSU Stocks, Market Returns.	L0: Business & Finance L1: Stock Market L2: Individual Performance, Order Book Analysis

All three of these papers have something in common

Takeaway: We can solve problems in the environment, harness, or model.
Multi-agent workflow decomposition is a powerful tool to solve problems in the harness.

- While not truly “multi-agent,” the think-act-refine loop is surprisingly effective at memory retrieval + management
- Decomposing the auto-research loop into ideation + selection helped us resolve context pollution problems
- Decomposing the feature engineering workflow into vocabulary building + assignment helped us resolve tagging problems