

Towards Trustworthy Autonomous Research

Wataru Kumagai
NexaScience / RIKEN TRIP-AGIS



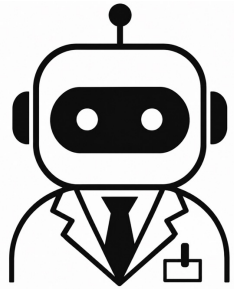
AGNTCon + MCPCon Japan 2026



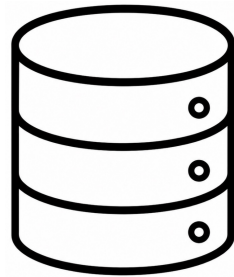
Key messages

1. **AI** now makes discoveries and **speeds up research**
2. **Human verification can't keep up** with this speed
3. **AIRAS** automates verification for **trustworthy research**

Research System



Research DB



Verifier



 **AIRAS**



The Impact of AI on Research

Research breakthroughs with AI

Solve Open Problems

10 advances

on long-standing problems

GPT-6 Astra (OpenAI, 2026)

Reveal Hidden Phenomena

New physics formula

nonzero where zero was assumed

GPT-5.2 (OpenAI, 2026)

Discover New Mechanisms

New treatment targets

proposed & validated

Co-Scientist (Google, 2026)

Design New Materials

Tailored materials

designed & synthesized

MatterGen (Microsoft Research, 2025)

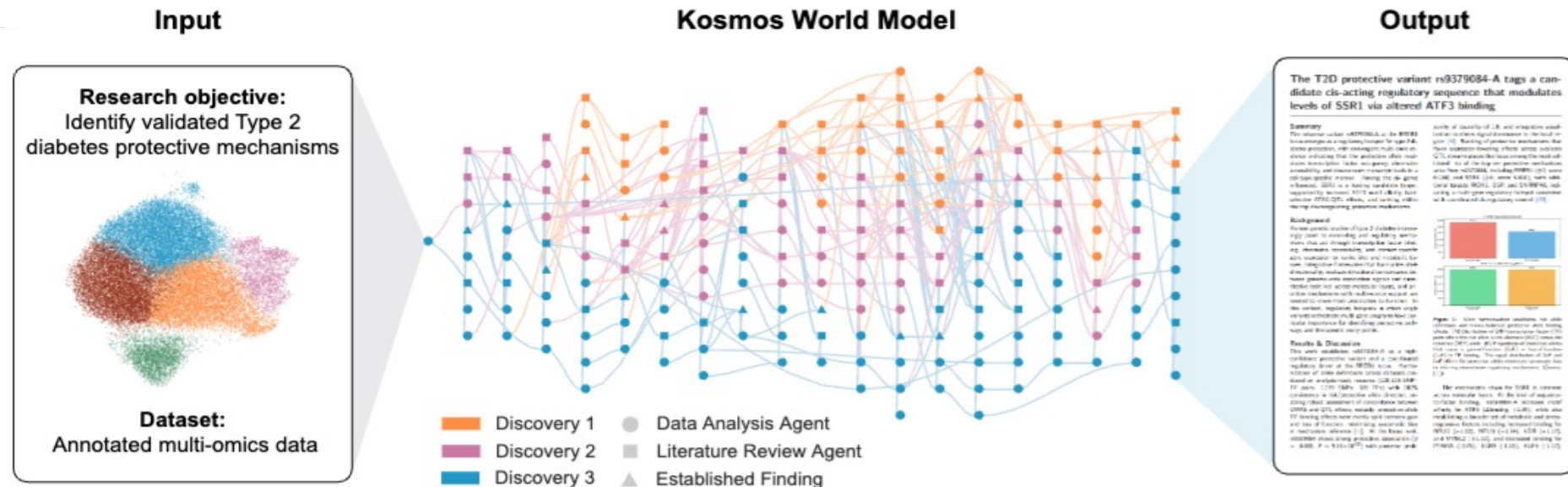
Research breakthroughs by AI

Kosmos (Edison Scientific, 2025)

6 months → 1 day
(beta-user estimate)

4 novel contributions

across fields



Progress in AI Research Systems

Research involves many steps



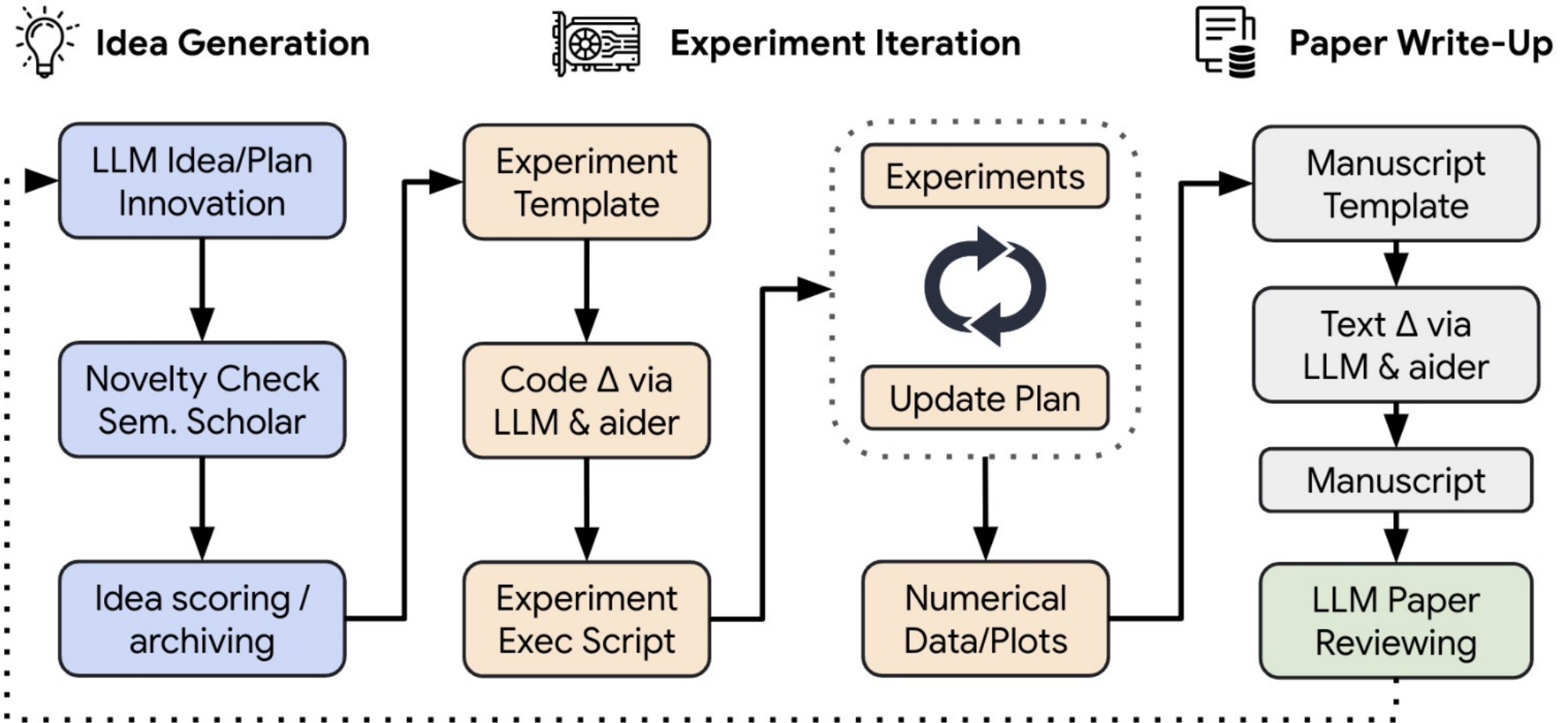
Until recently, **researchers themselves** handled all the steps.

Rise of conversational AI



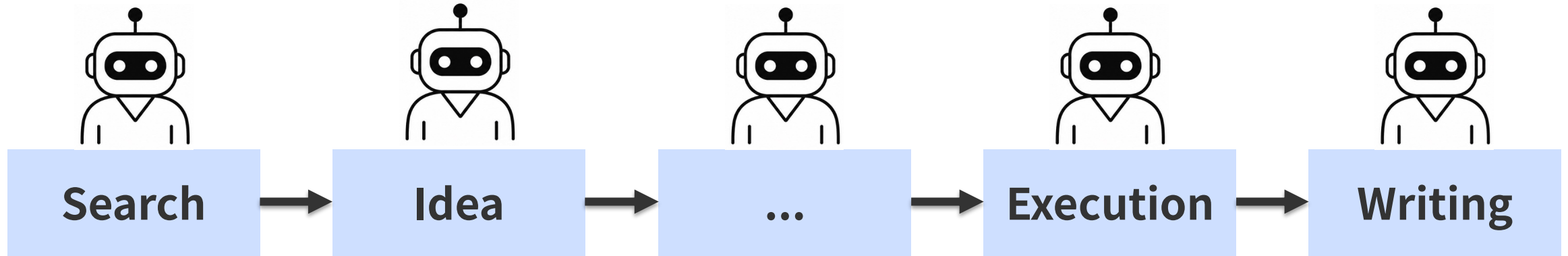
A single conversational AI can handle
many kinds of research tasks.

The AI Scientist and beyond



Lu et al., Nature, 2026

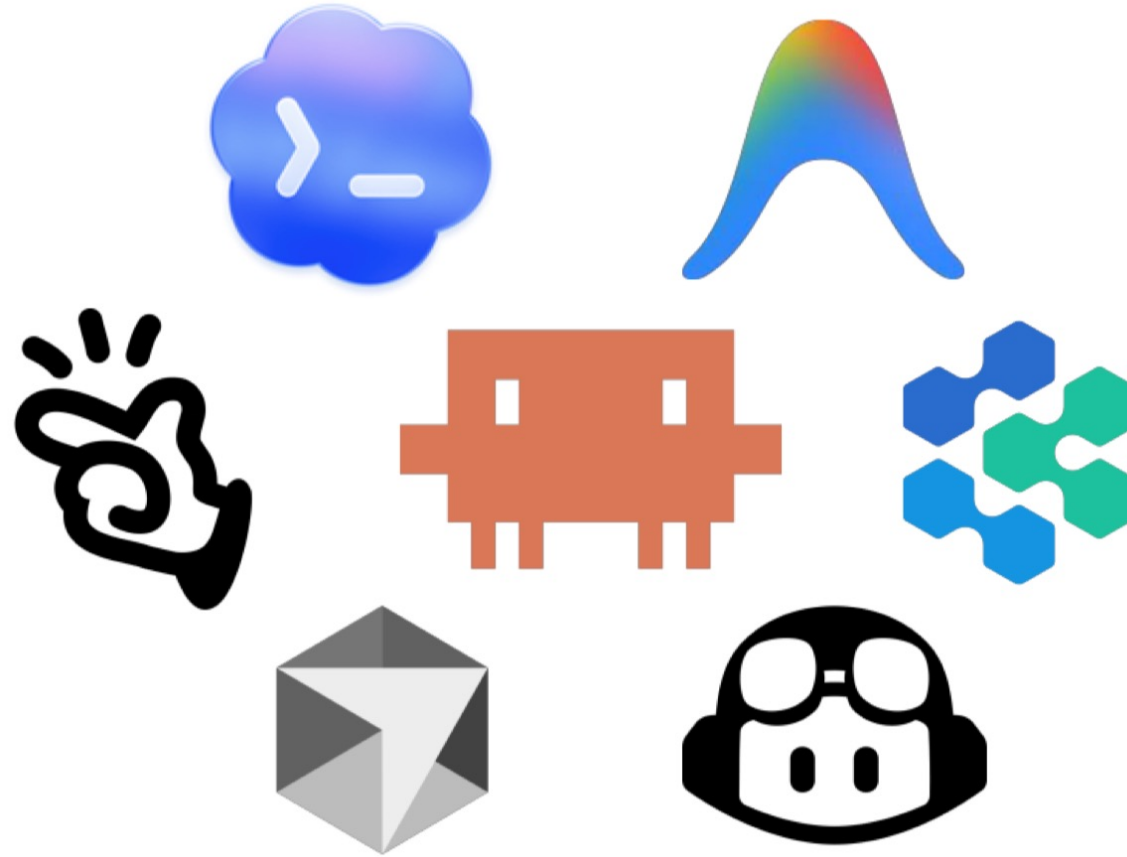
From manual work to workflow automation



A single system can handle **all research steps end to end.**

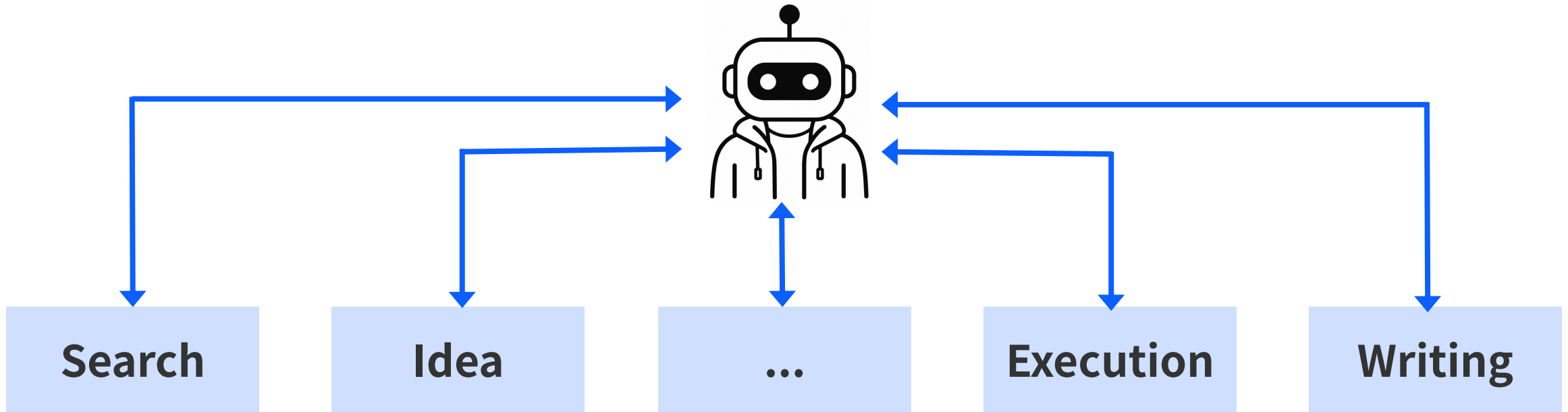
But the workflow is still **fixed.**

Evolution to coding agents



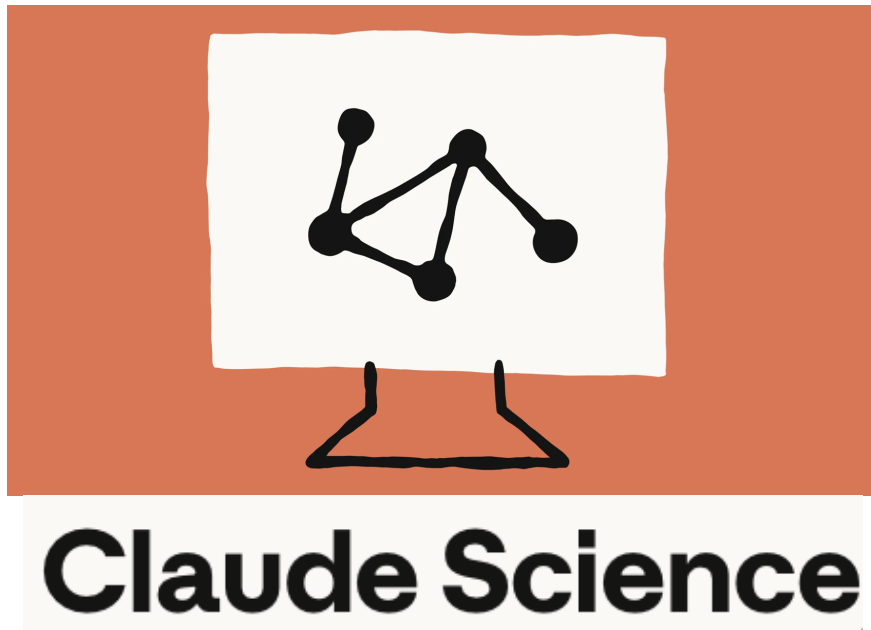
Coding agents **decide and act on their own.**

From fixed workflow to coding agents



Coding agents are **flexible like human researchers**.
But they are **not specialized for research**.

Specialized tools and environments for research



BioNeMo Agent Toolkit



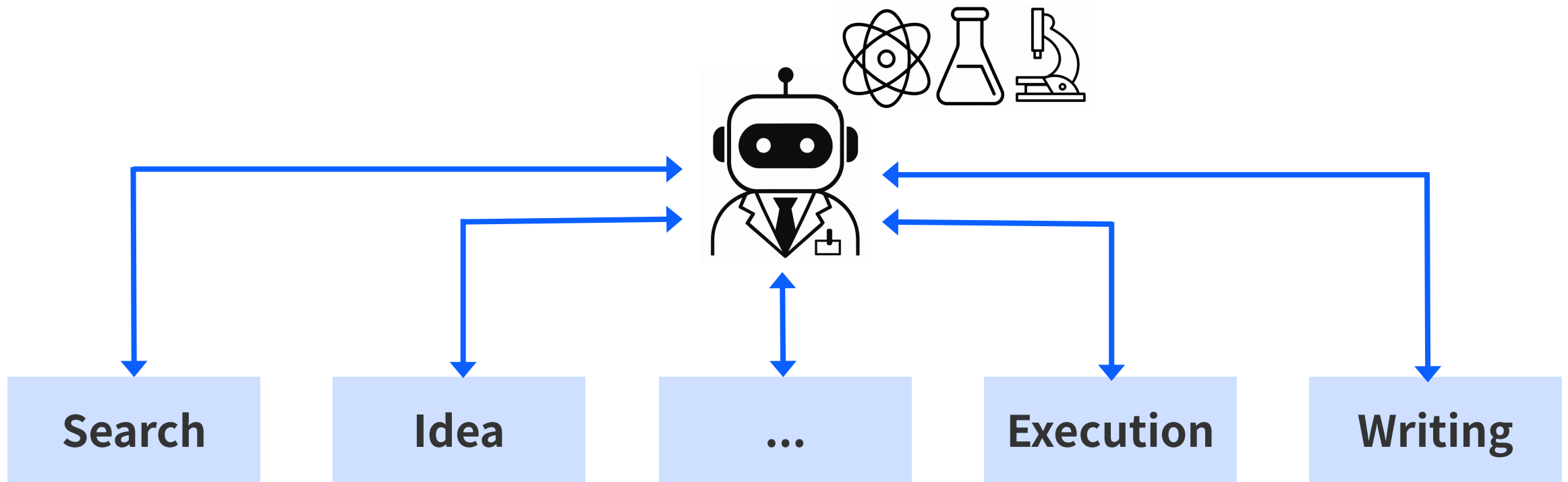
Ai2 Asta resources



ToolUniverse

These environments and tools turn coding agents into **research agents**.

From coding agents to research agents



Research is becoming **automated** and **faster**.
But faster research only helps if it can be **trusted**.

Peer Review for Trustworthy Research

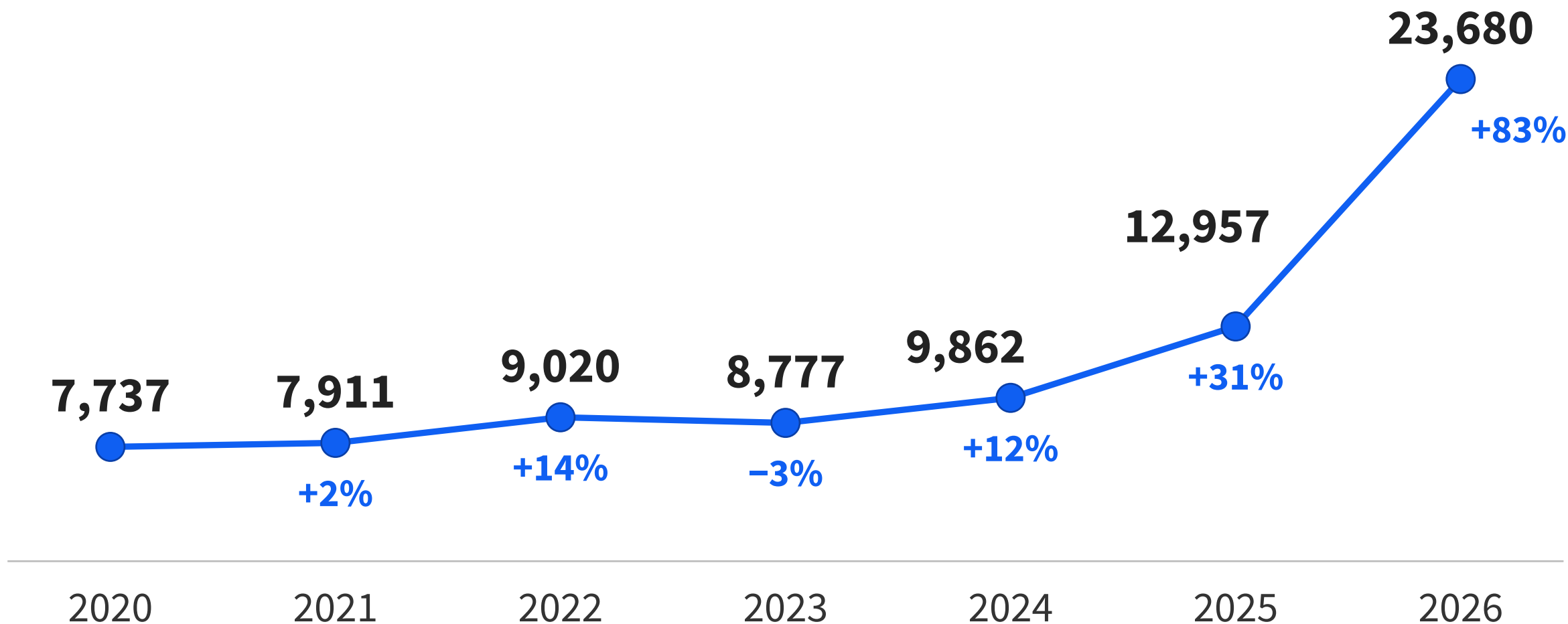
Trust in research comes from peer review



Other researchers **review** every paper before it is published.

Accelerated paper submissions

#papers reviewed at an AI Conference (AAAI)



Peer review is struggling

Surge in Peer Reviewers

Nearly 3x committee size

recruited vs. last year

AAAI 2026

Shortage of Experienced Reviewers

Undergraduate students

served as reviewers

ICLR 2025

AI-Generated Reviews

~21% of reviews

estimated fully AI-generated (Pangram)

ICLR 2026

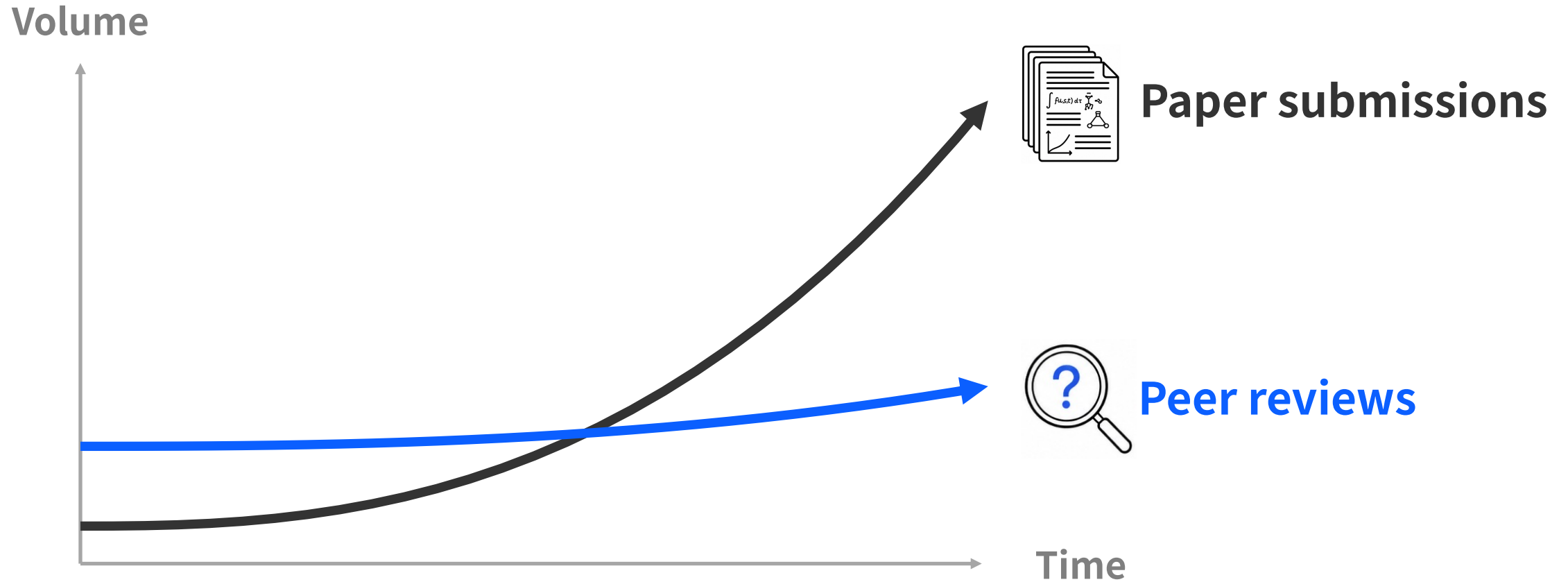
Manipulating AI-Assisted Reviews

Prompt injections

banned and screened by detectors

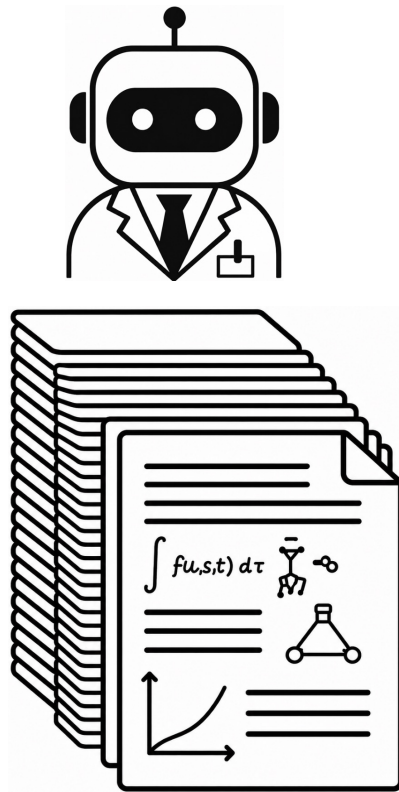
ICML 2026

The bottleneck has shifted to peer review



Research throughput $\leq \min(\text{Submissions}, \text{Reviews})$

Agents produce faster than humans can check



Paper submissions

VS.



Peer review

How can we review papers in **the age of AI research systems?**

Reviewing AI-Generated Research

A stylized black and white icon of a robot. The robot has a rounded head with a single antenna on top. Its eyes are represented by two white circles within a black horizontal band. It is wearing a suit jacket, a white shirt, and a dark tie. In its right hand, it holds a small rectangular object, possibly a briefcase or a device. The entire icon is composed of simple black outlines on a white background.



Result

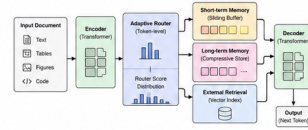


Figure 1: Overview of Adaptive Memory Routing (AMR). The router dynamically allocates each token to one of three memory experts based on its estimated utility. Memories are read and merged to condition the decoder for next-token prediction.

3 Experiments

We evaluate AMR on five benchmarks spanning scientific reasoning, table understanding, and code generation (Table 1). Models are trained on a 1.2T-token mixture of scientific papers, preprints, textbooks, and code repositories. Baselines include a dense Transformer (Dense), Longformer [3], Retentive Transformer [6], and RAG-Style [5]. All models are matched in parameter count (7B) and trained with a context length of 32K tokens. Metrics are accuracy ($\uparrow$) or exact-match (EM $\uparrow$).

Table 1: Performance comparison on scientific benchmarks.

Model	SciQA (EM $\uparrow$)	ProofBench (EM $\uparrow$)	TableQA (Acc $\uparrow$)	CodeGen (Pass@1 $\uparrow$)	LiQA (Acc $\uparrow$)	Avg ($\uparrow$)
Dense (7B)	64.2	35.6	78.1	46.3	72.4	59.3
Longformer	66.8	38.9	79.6	49.7	74.8	61.9
Retentive	68.3	40.7	81.2	51.8	76.8	63.8
RAG-Style	70.6	42.6	83.5	53.6	79.3	65.9
AMR (ours)	72.9	45.3	86.6	56.4	81.8	68.6

AMR achieves the best average performance (+2.7% over RAG-Style) while reducing peak memory usage by 29–37% across tasks. Ablations show that removing any expert degrades performance, with the largest drop when removing LTM (−3.1%).

4 Results

Figure 1 illustrates AMR’s architecture. We find that AMR allocates more capacity to LTM for cross-document tasks (LitQA, SciQA) and to STM for code generation. Router entropy annealing improves stability and final accuracy by up to 1.4%. Case studies reveal that AMR retains causal chains and table schemas more reliably than baselines.

5 Conclusion

AMR introduces adaptive, task-aware memory allocation that improves efficiency and accuracy for scientific language modeling. Future work includes dynamic expert creation and multi-modal memory routing.

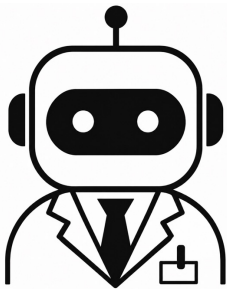
References

- [1] A. Kim, J. Park, and S. Lee. Scaling laws of language model memory. *arXiv preprint arXiv:2401.01234*, 2024.
- [2] M. Singh, J. Zhao, and Y. Wu. Efficient long-context modeling: A survey. *J. ML Research*, 25(7):1–45, 2024.
- [3] Z. Beltagy, M. E. Peters, and A. Cohan. Longformer: The long-document transformer. *arXiv preprint arXiv:2004.05150*, 2020.

Results generated by AI research systems

Claim

"Our method outperforms all baselines on every benchmark."



Research System



Figure 1: Overview of Adaptive Memory Routing (AMR). The router dynamically allocates each token to one of three memory experts based on its estimated utility. Memories are read and merged to condition the decoder for next-token prediction.

3 Experiments

We evaluate AMR on five benchmarks spanning scientific reasoning, table understanding, and code generation (Table 1). Models are trained on a 1.2T-token mixture of scientific papers, preprints, textbooks, and code repositories. Baselines include a dense Transformer (Dense), Longformer [3], Retentive Transformer [6], and RAG-Style [5]. All models are matched in parameter count (7B) and trained with a context length of 32K tokens. Metrics are accuracy (↑) or exact-match (EM ↑).

Table 1: Performance comparison on scientific benchmarks.

Model	SciQA (EM ↑)	ProofBench (EM ↑)	TableQA (Acc ↑)	CodeGen (Pass@1 ↑)	LaQA (Acc ↑)	Avg (↑)
Dense (7B)	64.2	35.6	78.1	46.3	72.4	59.3
Longformer	66.8	38.9	79.6	49.7	74.8	61.9
Retentive	68.3	40.7	81.2	51.8	76.8	63.8
RAG-Style	70.6	42.6	83.5	53.6	79.3	65.9
AMR (ours)	72.9	45.3	86.6	56.4	81.8	68.6

AMR achieves the best average performance (+2.7% over RAG-Style) while reducing peak memory usage by 29–37% across tasks. Ablations show that removing any expert degrades performance, with the largest drop when removing LTM (~3.1%).

4 Results

Figure 1 illustrates AMR's architecture. We find that AMR allocates more capacity to LTM for cross-document tasks (LaQA, SciQA) and to SLM for local reasoning (ProofBench, TableQA). Case studies reveal that AMR retains causal context better than SLM, thereby improving reasoning accuracy.

5 Conclusion

AMR introduces adaptive, task-aware memory allocation that improves efficiency and accuracy for scientific language modeling. Future work includes dynamic expert creation and multi-modal memory routing.

References

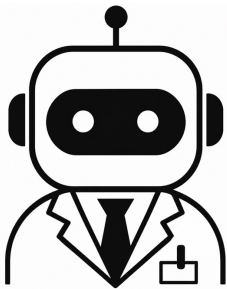
[1] A. Kim, J. Park, and S. Lee. Scaling laws of language model memory. *arXiv preprint arXiv:2401.01234*, 2024.
[2] M. Singh, J. Zhou, and Y. Wu. Efficient long-context modeling: A survey. *J. ML Research*, 25(7):1–45, 2024.
[3] Z. Borgey, M. E. Peters, and A. Cohen. Longformer: The long-document transformer. *arXiv preprint arXiv:2004.05150*, 2020.
[4] P. Dao, A. N. Dao, P. Nguyen, A. P. Nguyen, G. B. Do, D. Dao, and P. Dao. Longformer: The long-document transformer. *arXiv preprint arXiv:2004.05150*, 2020.

Result

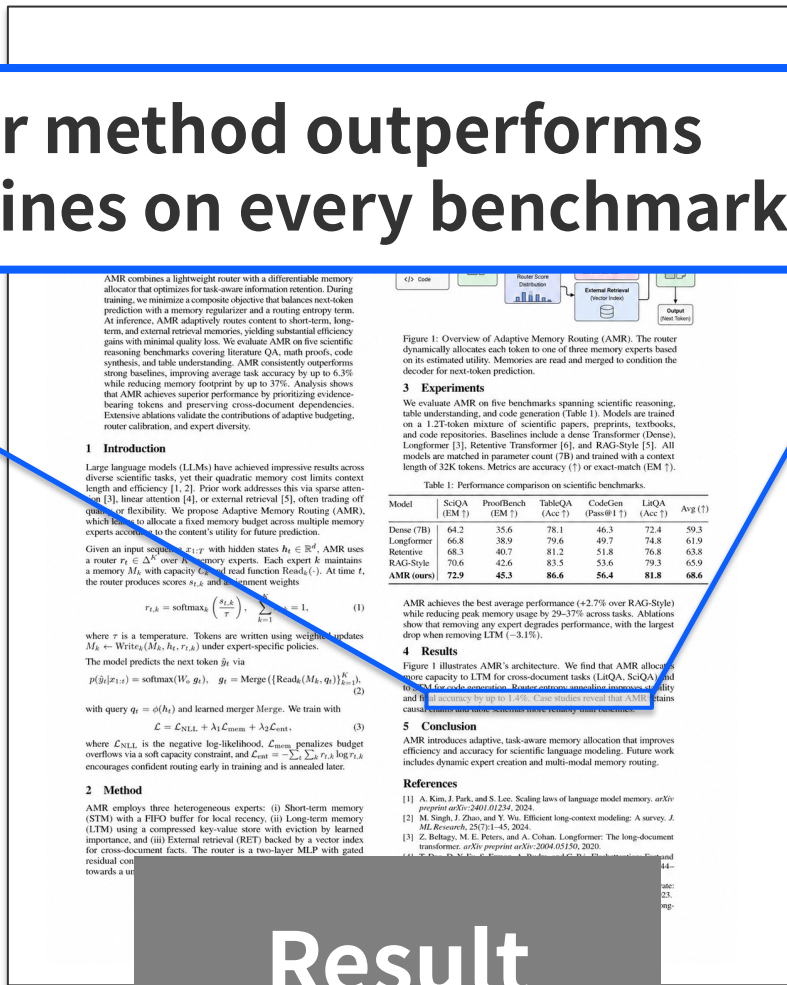
Results generated by AI research systems

Claim

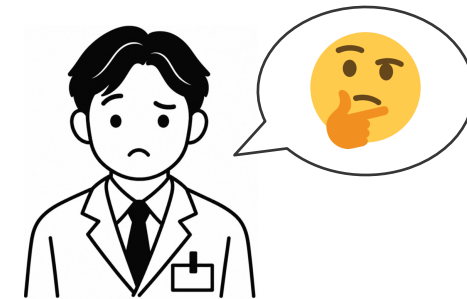
"Our method outperforms all baselines on every benchmark."



Research System



Result



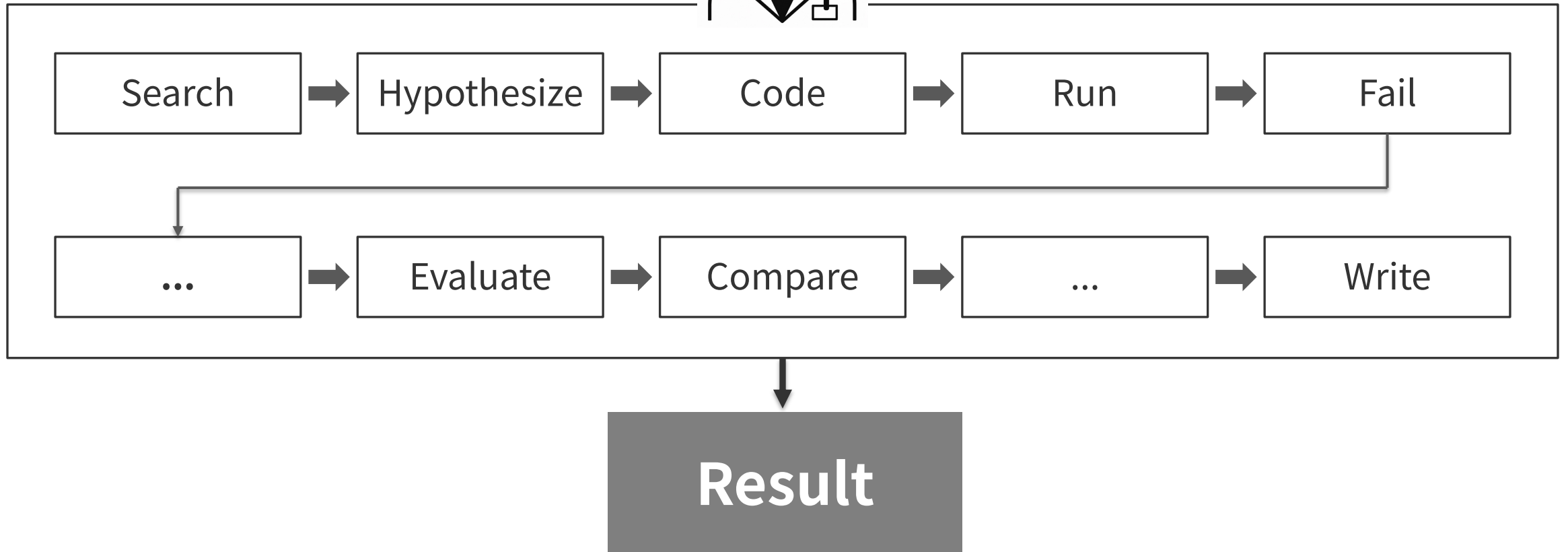
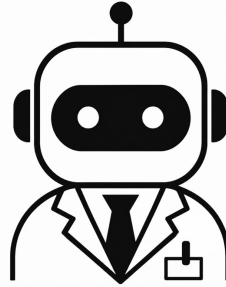
Reviewer

Was the comparison fair?

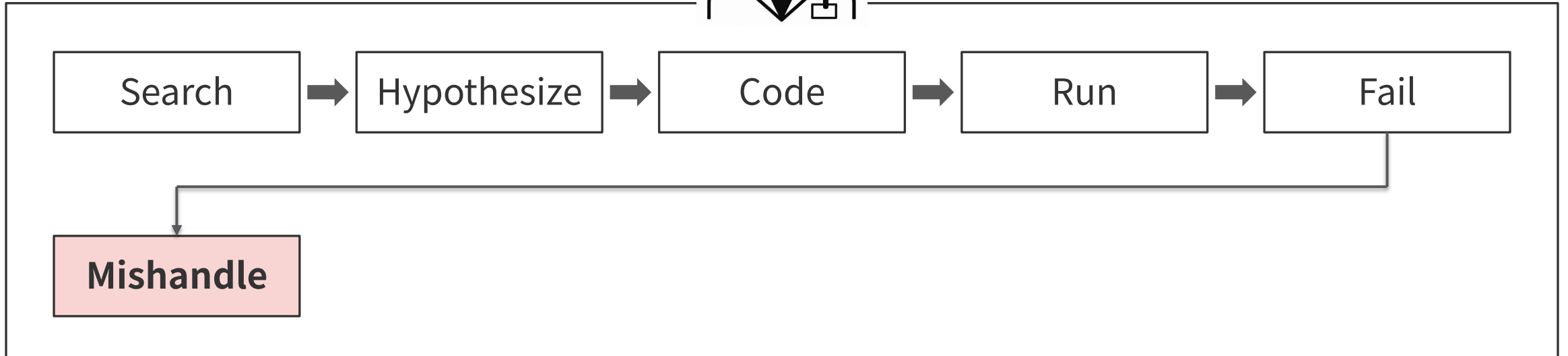
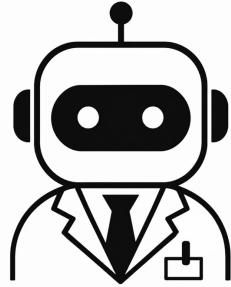
Are the results reproducible?

Did it show only the good results?

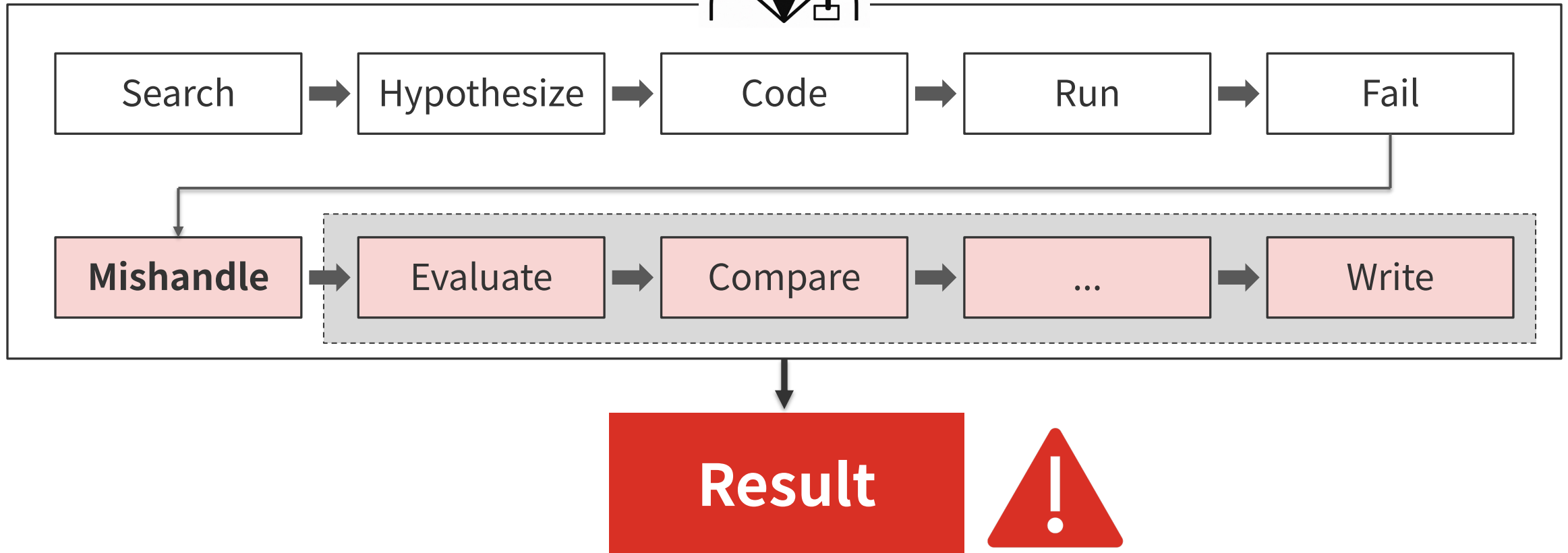
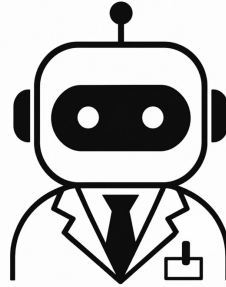
Inside AI research systems



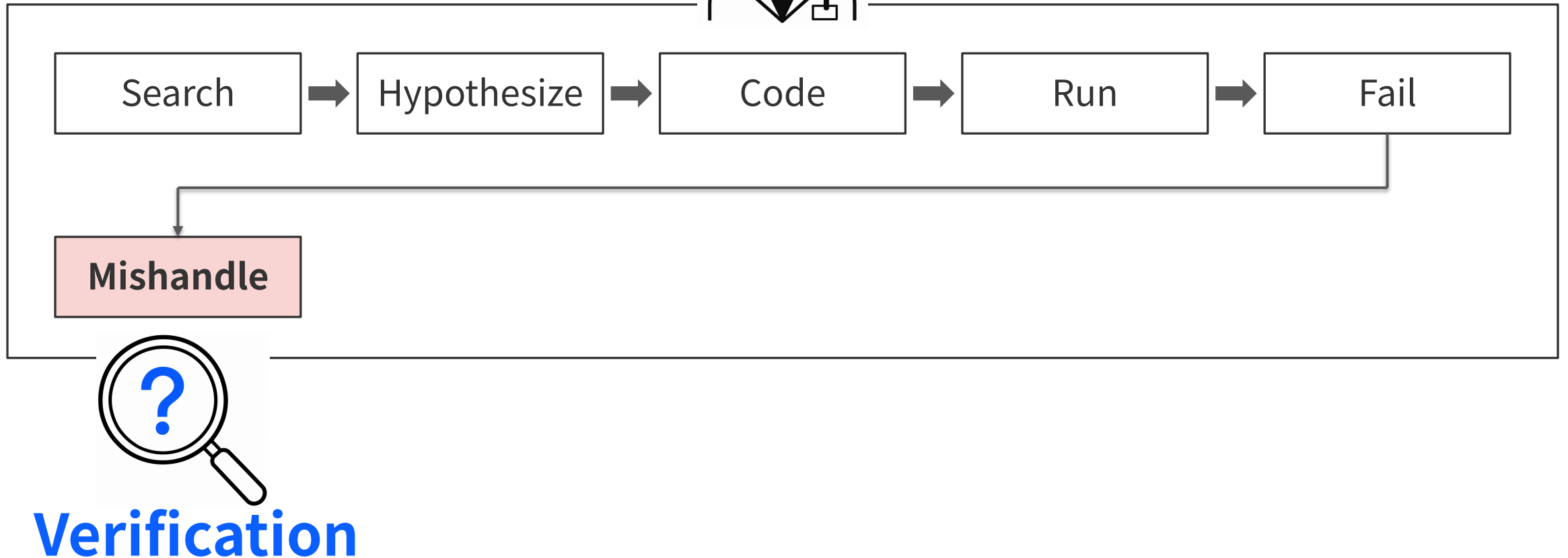
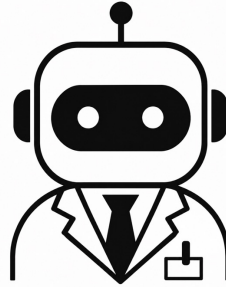
Inside AI research systems



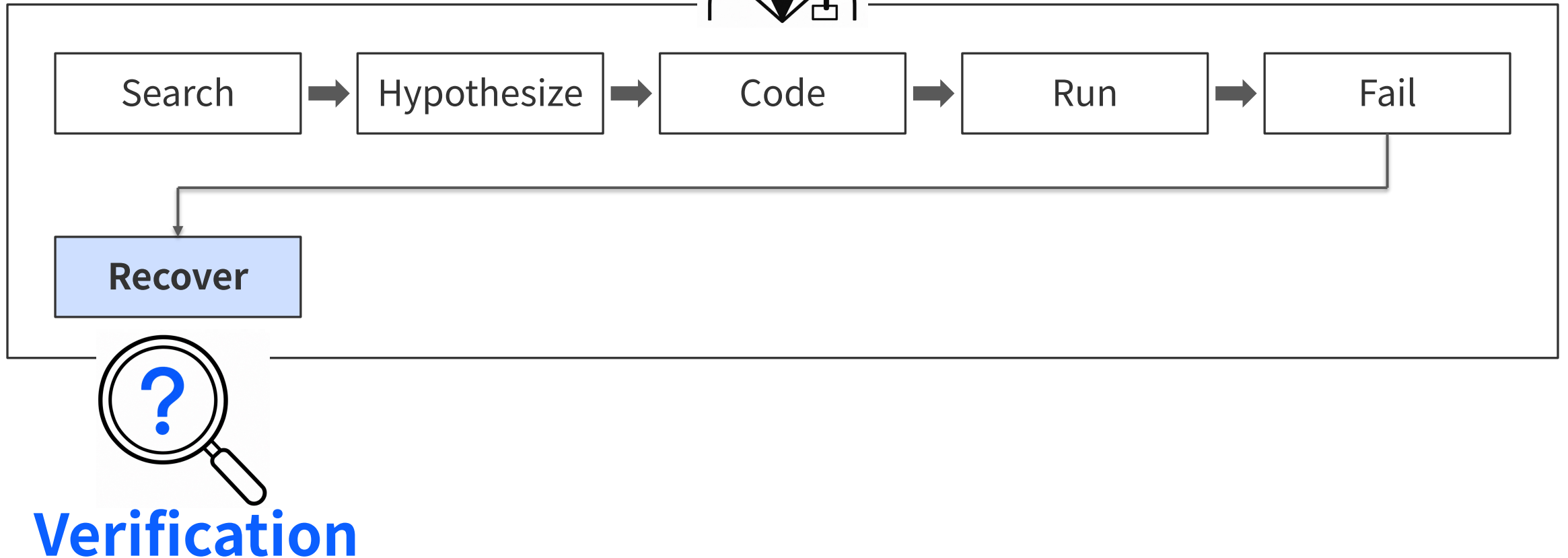
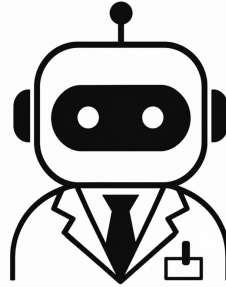
Inside AI research systems



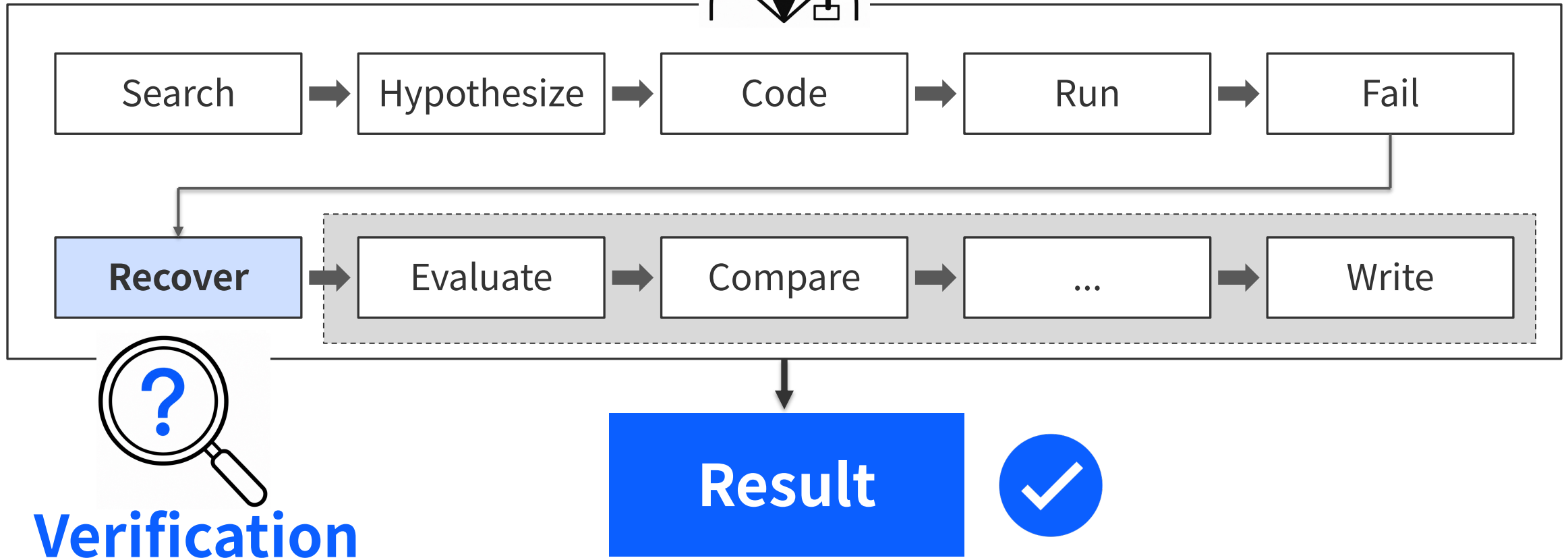
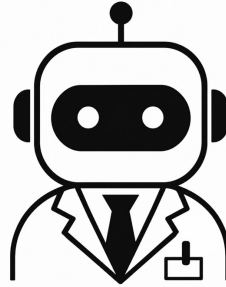
Inside AI research systems



Inside AI research systems



Inside AI research systems





To trust the results, we need **verify the process.**

Designing AI Research Systems for Verification



AIRAS

AI Research Automation System

Verifying claims throughout the research process

Goal

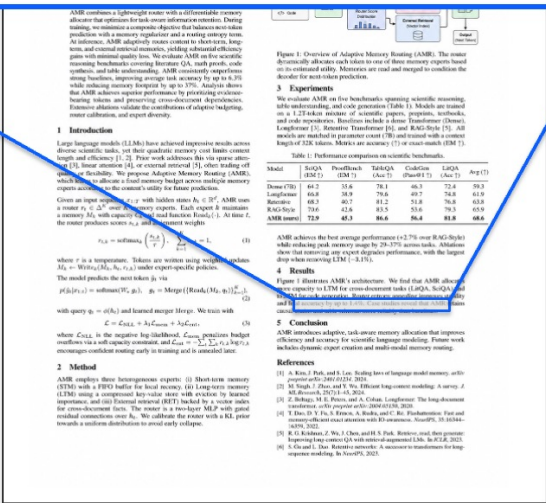
Trustworthy results from AI research systems

Approach

A research result consists of **claims** from each step of the process
→ We make every claim **verifiable**

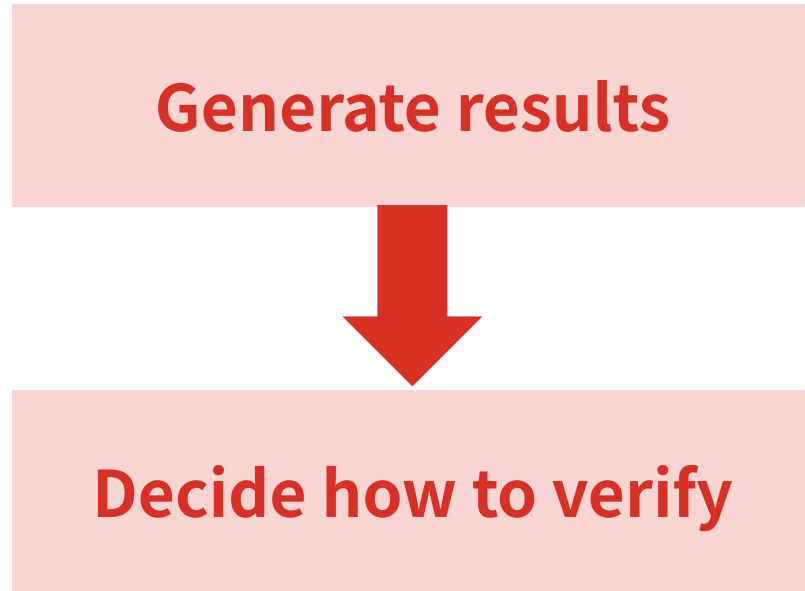
Claim

"Our method outperforms all baselines on every benchmark."



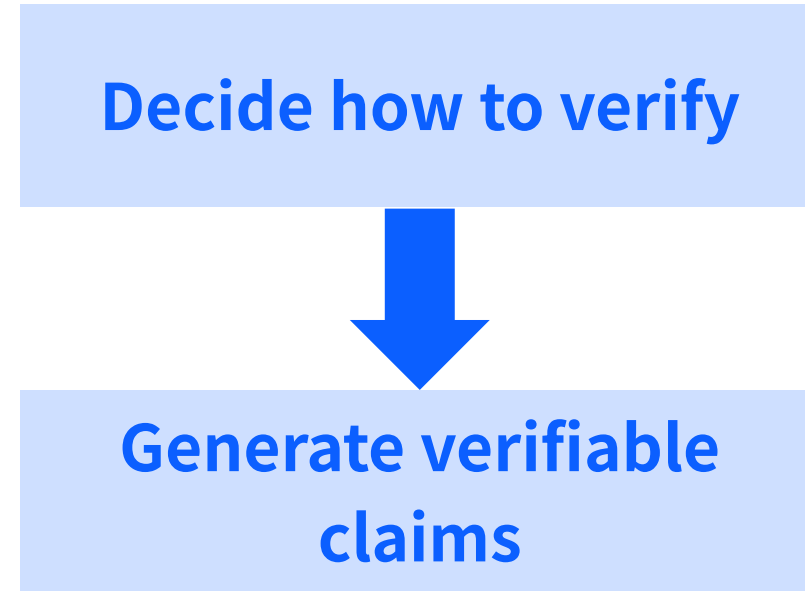
Designing research for verification

Conventional research



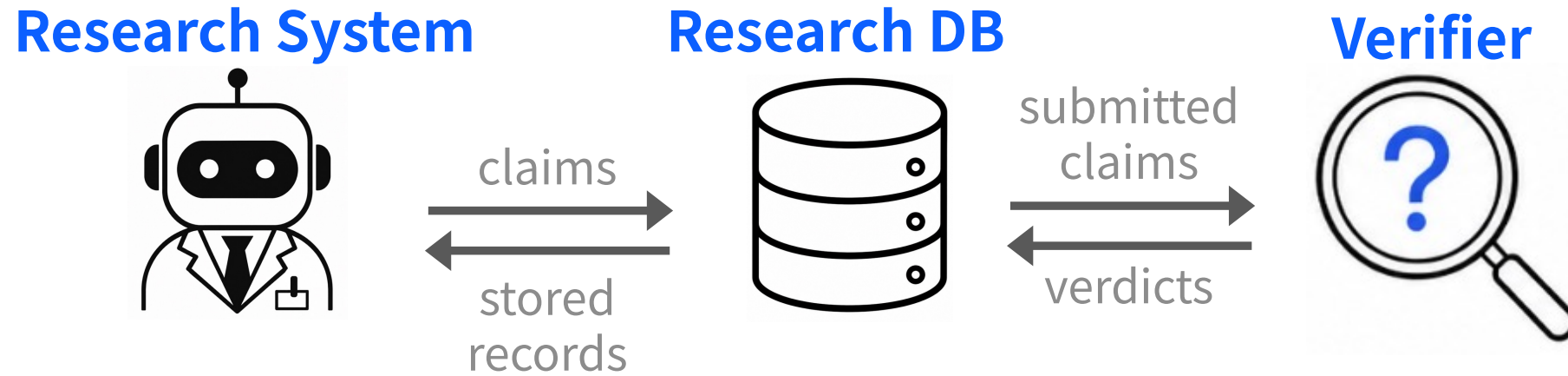
How to verify is decided **after** the research ends.

Verification-first research



How to verify is decided **before** the research starts.

Verification-first autonomous research



Three principles for verifiable research

1. Every claim is **verifiable**

Each claim says what holds and how to verify it

2. Every claim is verified **outside the research system**

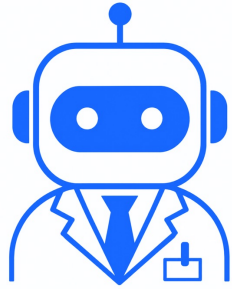
The verifier is independent and the research system cannot influence it

3. Every claim has its **source of truth** in the research DB

Every paper is generated only from verified claims in the DB

1. Every claim is verifiable

Research System



Research DB



claims



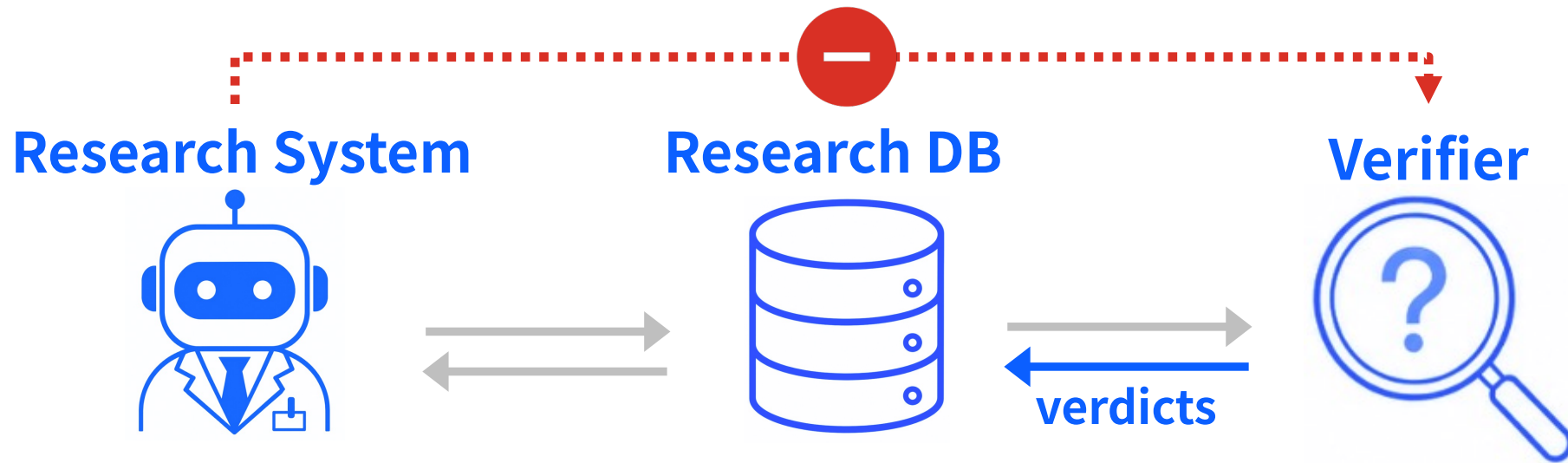
Verifier



```
id: experiment-001
type: experiment
claim: "Method A outperforms Method B by 2.3 points on D."
code: "sha256:9f3c...a721"
check: rerun with new seeds
```

Each claim says what holds and how to verify it

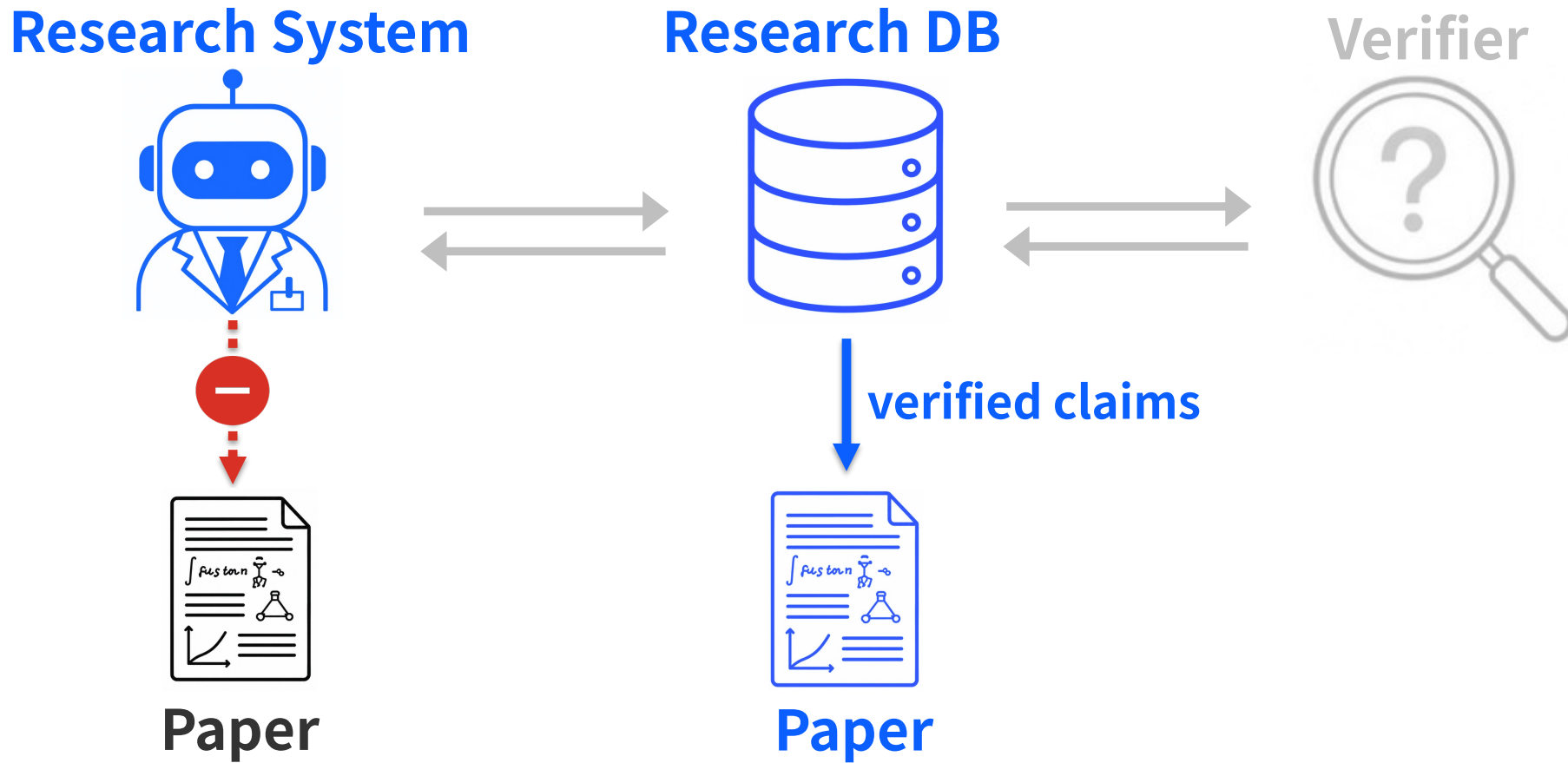
2. Every claim is verified outside the research system



```
claim_id: experiment-001
verifier: external-cluster-seyval
reported: +2.3
rerun: -1.1
verdict: FAIL
```

The verifier is independent and the research system cannot influence it

3. Every claim has its source of truth in the research DB



Every paper is generated only from verified claims in the DB

AIRAS: Key features

1. Open-source software (OSS) for research

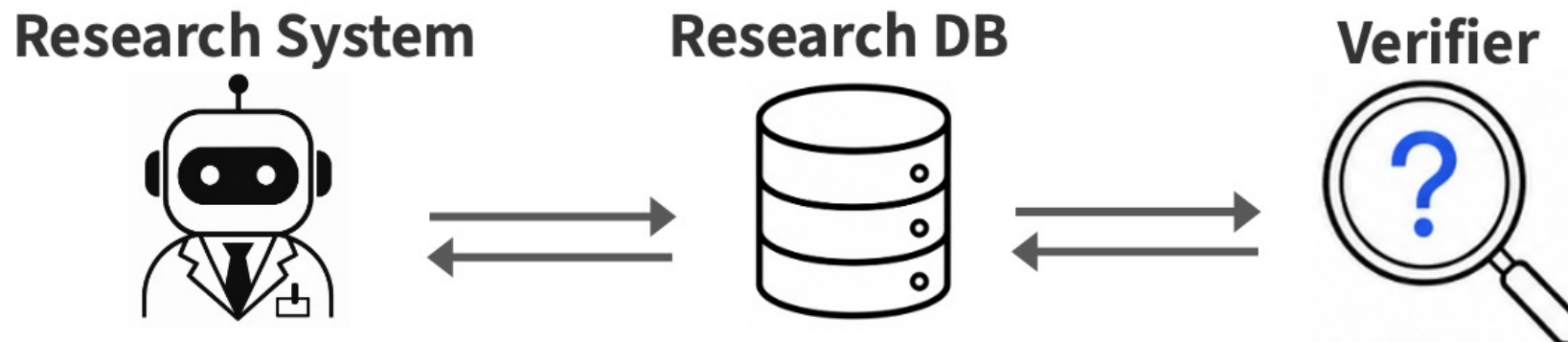
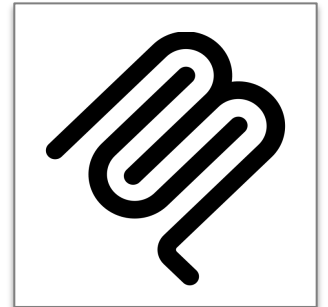
Anyone can use, modify, and extend AIRAS for their research.

2. Research tools via MCP

Different agents can access research tools through MCP.

3. Verification-first design

Designed to follow all three principles of verifiable research.



Conclusion

Conclusion

1. **AI** now makes discoveries and **speeds up research**
2. **Human verification can't keep up** with this speed
3. **AIRAS** automates verification for **trustworthy research**



Questions & Comments

