

# Before you <think>, monitor:

## Implementing Flavell's metacognitive framework in LLMs



# Monitor-Generate

# Generate-Verify

# Monitor-Generate-Verify

# Metacognition

# LLM-Modulo Framework

**TL;DR**

**Nick Oh**

socius labs  
London, UK  
nick.sh.oh@socius.org

**Fernand R. Gobet\***

Centre for Philosophy of  
Natural and Social Science,  
London School of Economics  
London, UK  
f.gobet@lse.ac.uk

**\*ACKNOWLEDGEMENT:** This experimental implementation builds on the *Monitor-Generate-Verify* framework co-developed with Fernand Gobet.

**Monitor-Generate methods** (e.g., SELF-DISCOVER) *plan* but can't *verify*; **Generate-Verify methods** (e.g., SELF-REFINE) *refine* but start *blind*.

**We chain them via Flavell's (1979) metacognitive theory.**

**GSM8K: 75.4% vs 68.4% accuracy, 1.3 vs 2.0 attempts.**  
**Better initial solutions need less verification.**

### Background



John H. Flavell  
[source: Awards for Distinguished Scientific Contributions: John H. Flavell. American Psychologist, 40(3), 291-295]

#### Metacognition and Cognitive Monitoring

*A New Area of Cognitive-Developmental Inquiry*

JOHN H. FLAVELL *Stanford University*

↓ **Monitor-Generate-Verify:** Formalising Metacognitive Theory for Language Model Reasoning (Oh and Gobet, 2025)

##### Algorithm 1 Flavell's Metacognitive Regulation

```
1: Initialise:  $S_0 \leftarrow f(\mathcal{T}, \mathcal{G})$ ;  $\tau \leftarrow 0$ 
2: while  $S_\tau = \text{ACTIVE}$  do
3:   // MONITOR: Retrieve knowledge & assess experience
4:    $MK_\tau \leftarrow$  if  $\tau = 0$  then retrieve( $MK, \mathcal{T}, \mathcal{G}$ )
5:   else  $MK_{\tau-1} \cup$  retrieve( $MK, M\mathcal{E}_{\tau-1}$ )
6:    $M\mathcal{E}_{\tau}^{\text{difficulty}} \leftarrow$  feel( $\mathcal{T}$ , Outcomes $_{\tau-1}$ )  $\oplus$  assess( $\mathcal{T}, MK_\tau$ )
7:   // GENERATE: Select & execute cognitive strategy
8:    $CS_\tau \leftarrow$  select( $s \in MK_{\text{Strategy}} \mid M\mathcal{E}_{\tau}^{\text{difficulty}}, MK_\tau, \mathcal{T}, \mathcal{G}$ )
9:    $CO_\tau \leftarrow$  execute( $CS_\tau, \mathcal{T}, \mathcal{G}$ )
10:  // VERIFY: Evaluate progress & update knowledge
11:   $M\mathcal{E}_{\tau}^{\text{evaluative}} \leftarrow$  assess( $CO_\tau, MK_\tau$ )
12:   $MS_\tau \leftarrow$  select( $s \in MK_{\text{Meta-Strategy}} \mid M\mathcal{E}_{\tau}^{\text{evaluative}}$ )
13:   $MO_\tau \leftarrow$  execute( $MS_\tau, CO_\tau, MK_\tau, \mathcal{G}$ )
14:   $MK \leftarrow$  update( $MK, \Phi_\tau$ ) where  $\Phi_\tau = (M\mathcal{E}_\tau, \text{Strategy}_\tau, \text{Outcome}_\tau)$ 
15:   $S_{\tau+1} \leftarrow$  if goal_achieved( $CO_\tau, \mathcal{G}$ ) then TERMINATE else ACTIVE
16:   $\tau \leftarrow \tau + 1$ 
17: end while
```

##### Algorithm 1 Flavell's Model of Cognitive Monitoring

**Require:** task  $\mathcal{T}$ , model  $M$ , strategies  $MK$ , prompts  $\mathcal{P} = \{p_{\text{monitor}}, p_{\text{strategy}}, p_{\text{execute}}, p_{\text{verify}}\}$

```
1: Initialize  $S_0 = \text{ACTIVE}$ ,  $\tau = 0$ ,  $M\mathcal{E}_{\text{evaluative}}^{-1} = \emptyset$ 
2: while  $S_\tau = \text{ACTIVE}$  and  $\tau < T$  do
3:   // Monitor: Assess task difficulty
4:    $M\mathcal{E}_{\text{difficulty}}^\tau, \text{features}_\tau = M(p_{\text{mon}} \parallel \mathcal{T} \parallel M\mathcal{E}_{\text{evaluative}}^{\tau-1})$  ▷ Assess task difficulty
5:
6:   // Generate: Apply cognitive strategy
7:    $\text{strategy}_\tau = M(p_{\text{str}} \parallel \text{features}_\tau \parallel M\mathcal{E}_{\text{difficulty}}^\tau \parallel MK)$  ▷ Choose approach
8:    $\text{solution}_\tau = M(p_{\text{exe}} \parallel \mathcal{T} \parallel \text{strategy}_\tau)$  ▷ Execute strategy with adaptive params
9:
10:  // Verify: Evaluate performance
11:   $M\mathcal{E}_{\text{evaluative}}^\tau = M(p_{\text{ver}} \parallel \mathcal{T} \parallel \text{solution}_\tau)$  ▷ Evaluate output quality
12:
13:   $S_{\tau+1} =$  if  $\text{mean}(M\mathcal{E}_{\text{evaluative}}^\tau) \geq 0.85$  then TERMINATE else ACTIVE
14:   $\tau = \tau + 1$ 
15: end while
16: return  $y_{\arg \max_i \text{mean}(M\mathcal{E}_{\text{evaluative}}^i)}$ 
```

### MGV framework (zero-shot)

**$MK_{\text{Agent}}$**   
Knowledge about cognitive agents' characteristics and capabilities

**$MK_{\text{Task}}$**   
Knowledge about cognitive situation assessment

**$MK_{\text{Strategy}}$**   
Knowledge concerning the effectiveness of cognitive strategies and metacognitive strategies

**$MK_{\text{Strategy}}$**  as a pre-established repertoire of 20 domain-specific strategies (Didolkar et al., 2024)

Generate - Strategy Selection

Problem: {problem}  
Task features: {task\_features}  
Difficulty: {difficulty:.3f}  
Available strategies: {comma\_separated\_strategy\_list}

Select ONE strategy from the list above that best fits this problem.

Output format:  
Selected Strategy: [exact strategy name from list]

Generate - Strategy Execution

Problem: {problem}  
Strategy: {strategy\_type}

Output format:  
<think>  
Work through the problem step by step using the {strategy\_type} approach. Show all calculations and reasoning.  
{if t > 0:  
previous\_cycle:  
difficulty: {difficulty}  
task\_features: {task\_features}  
strategy\_used: {strategy\_type}  
reasoning: [full solution text]  
evaluation\_scores:  
coherence: {score}  
plausibility: {score}  
consistency: {score}  
goal\_conduciveness: {score}}  
evaluation: [full evaluation text]  
</think>  
  
<answer>  
Output only the final numerical answer (no units or text).  
</answer>

GENERATE

Token budget scales as  $400 + 400 * ME_{\text{difficulty}}$  and temperature as  $0.3 + 0.2 * ME_{\text{difficulty}}$

VERIFY

TERMINATE when mean evaluation  $\geq 0.85$  or  $T = 3$

#### Adaptive Resource Allocation

Verify

Problem: {problem}  
Solution: {solution}

Evaluate this solution on these dimensions (0-1 scale):

1. COHERENCE: Do the steps logically connect?  
2. PLAUSIBILITY: Is the approach reasonable?  
3. CONSISTENCY: Are calculations correct?  
4. GOAL-CONDUCTIVENESS: Does it answer the question?

Output format:  
<evaluate>  
Coherence: X.X  
Plausibility: X.X  
Consistency: X.X  
Goal-conduciveness: X.X

Evaluation: [Provide a 2-3 sentence analysis explaining the scores. Identify specific errors or strengths. For low scores, indicate what went wrong (e.g., "arithmetic error in step 3", "misunderstood the question", "skipped crucial reasoning"). For high scores, note what worked well. Be specific and actionable.]  
</evaluate>

#### Flavellian Metacognitive Knowledge

Monitor

Problem: {problem}  
{if t > 0:  
evaluation\_scores:  
coherence: {score}  
plausibility: {score}  
consistency: {score}  
goal\_conduciveness: {score}}

Before you <think>, analyse this problem WITHOUT solving it.

Output format:  
<monitor>  
Task\_Features: List 2-3 keywords describing the math concepts needed  
Difficulty (0-1): Assess the challenge level  
{if t > 0: Recalibrate your assessment based on the evaluation\_scores in previous\_cycle. Lower scores suggest higher actual difficulty than previously assessed.}  
</monitor>

MONITOR

GENERATE

VERIFY

### Experimental Setup

meta-llama/Llama-3.1-8B-Instruct (+NVIDIA H100 SXM)

**Self-Verification** (Weng et al., 2022): Generate- $k$ -Verify (w/ majority voting)

**SELF-REFINE** (Madaan et al., 2023): Iterating Generate-Verify  $k$  times

**MGV** (Flavell): Iterating Monitor-Generate-Verify  $k$  times

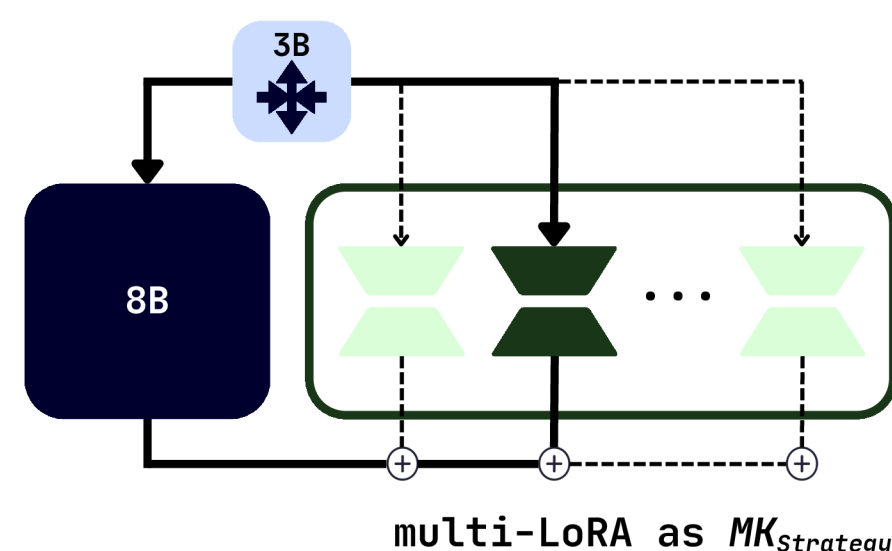
### Results

| Method            | Accuracy         | Avg Time (s) | Avg Attempts |
|-------------------|------------------|--------------|--------------|
| Self-Verification | 442/659 (67.07%) | 7.52         | 1.2          |
| Self-REFINE       | 451/659 (68.44%) | 6.98         | 2.0          |
| MGV (Flavell)     | 497/659 (75.42%) | 9.60         | 1.3          |

### Future Work (+ pilot experiment)

#### LLM Modular Framework

- Small models (0.5B) can reliably *Plan*, while *Execution* demands substantially larger models (Qin et al., 2025) (for our preliminary exploration of this direction see below)



Branched curriculum (2 epochs shared arithmetic  $\rightarrow$  3 specialized LoRAs @ 1 epoch each) with 3B routing outperformed 4-epoch GRPO by 1.4pp on GSM8K (8B models)

#### Model Intrinsic MK

- Following V-STaR's approach (Hosseini et al., 2024), future work could **train the Monitor through iterative self-supervised learning**: generate diverse solutions per problem, contrast successful versus failed attempts using DPO/GRPO, and progressively refine both difficulty assessment and strategy selection capabilities — building model-intrinsic metacognitive knowledge through bootstrapped preference learning.

#### Metacognitive Space within LLM

- Implicit confidence measures** derived from token likelihoods exhibit **greater metacognitive sensitivity** than explicitly prompted confidence (Xiong et al., 2023)
- Some evidence of **subjectivity** (e.g., confidence and certainty) **corresponding to linearly separable directions** in representations (Zou et al., 2023; Liu et al., 2023)