

85%

GSM8K · SIMILAR-CAPABILITY POOL · N=20 · AVG LAYERS RUN = 1.25

Spend deep compute **only when it helps** — CA-MoA **ties top accuracy** and Pareto-dominates fixed-depth MoA.

+5.0_{pt}
over MoA-1, at 1.4× tokens

+20_{pt}
over Single (65→85%)

89-100%
of $\ell=1$ stops correct

0-40%
of $\ell=3$ runs correct

1 The problem

Standard **Mixture-of-Agents** spends the **same** $L \cdot K$ LLM calls on **every** query — yet most solve at layer 1.

Cast it as a **stochastic knapsack**:

$$\max_{\pi} \mathbb{E}_x[Q(\pi(x))]$$

$$\text{s.t. } \mathbb{E}_x[C(\pi(x))] \leq B$$

2 Method: CA-MoA

A verifier **V** scores each layer's aggregate and **stops** per-query once the score clears a budget-aware bar τ (**D1**).

ALGORITHM 1 · CA-MOA

```
in: x; V; pool A; budget B;
    L_max;  $\tau(\cdot)$ ; width k
a  $\leftarrow$  x; b  $\leftarrow$  B;  $\ell \leftarrow$  0
while  $\ell < L_{\max}$  and b > 0:
    s  $\leftarrow$  V(x, a) # score agg
    if s  $\geq \tau(b/B)$ : break # D1: GO
    for A_j in A: # D2 route*
         $\hat{u}_j \leftarrow V(x, a|A_j) - s$ 
    S  $\leftarrow$  top-k( $\hat{u}$ ) s.t. cost  $\leq$  b
    a  $\leftarrow$  aggregate( run(S, a) )
    b  $\leftarrow$  cost(S);  $\ell \leftarrow \ell + 1$ 
return a # *D1-only eval
```

$$\tau(\varrho) = \tau_{\min} + (\tau_{\max} - \tau_{\min}) \cdot \varrho, \quad \varrho = b/B$$

τ is a **single Pareto knob** — monotone in τ ; strictly beats fixed-depth MoA in **expected cost** when $\mathbb{E}[L] < L_{\max}$.

Verifier V — three forms

V_SC Agreement, zero cost. **Used in all main experiments.**

V_LLM Judge LLM. *Ablated* — over-confident.

V_RM Reward model. *Pre-registered.*

How the gate works

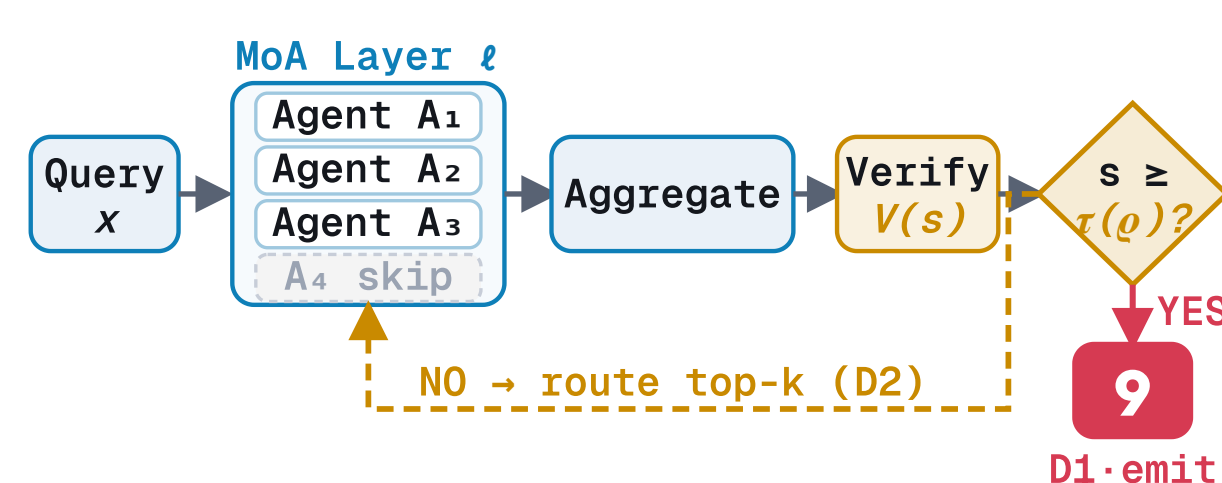


Fig 1. τ terminates (D1, red) or routes top-k onward (D2). Reported numbers: D1-only depth-adaptive policy.

Gating quality replicates

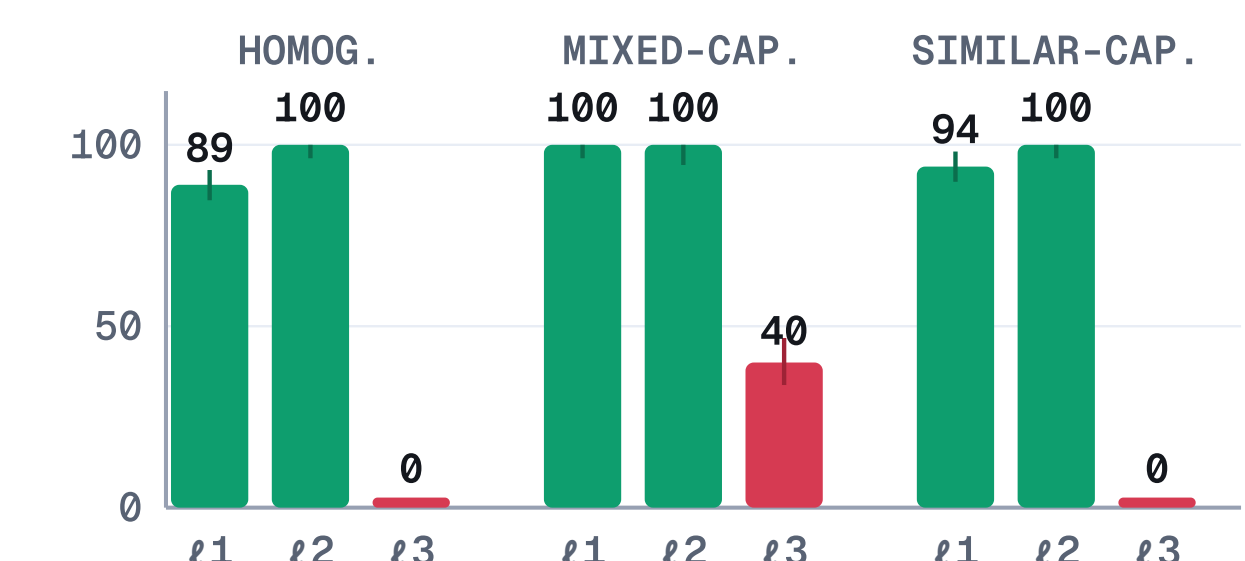


Fig 4. Accuracy by layers ℓ permitted. $\ell=1$ stops 89-100% correct; $\ell=3$ runs 0-40% across pools. $n = 1-44$ per bucket; small- n $\ell 2/\ell 3$ CIs wide.

Pareto frontier — Exp C

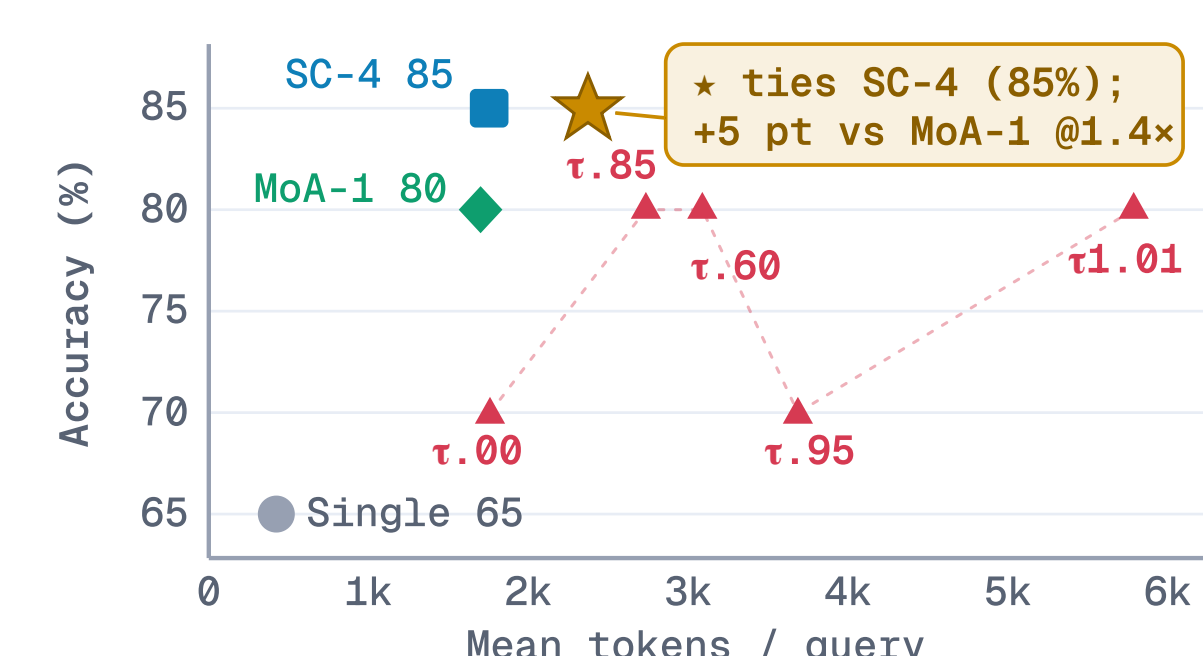


Fig 2. ★ CA-MoA headline ($\tau_{\max}=0.95$) ties SC-4. Red $\blacktriangle$ = τ -sweep (labels = $\tau_{\max}$); $\blacklozenge$ MoA-1, $\blacksquare$ SC-4, $\bullet$ Single.

Headline results — Exp C

METHOD	ACC.	TOKENS	AVGL
Single (Qwen2.5-3B)	65.0	423	—
MoA-1 (Wang '24)	80.0	1,700	1.00
SC-4 (Wang '23)	85.0	1,754	—
CA-MoA (ours)	85.0	2,372	1.25

Pool {Llama-3.2-3B, Qwen2.5-3B}, $n=20$. AvgL 1.25 → most queries stop at $\ell 1$.

Adaptation is real, measured

Token CV runs **3-4× higher** than fixed-cost baselines:

.62-.96

CA-MoA token CV

.24-.34

fixed baselines

Where the win lives — 3 regimes

GO · C	WEIGH · B	STOP · A
Balanced. +5 vs MoA-1. Clean win.	Broken. +13.3, weak agent corrupts.	Saturated. -8: false-positive stops.

METHOD	A · N50	B · N15	C · N20
Single	90.0	73.3	65.0
SC-4	88.0	46.7	85.0
MoA-1	92.0	46.7	80.0
CA-MoA	84.0	60.0	85.0

Accuracy (%) · bold = regime winner.

*The signal is regime-independent; its **payoff** is not.*

Honest limits & what's next

- Small n (50/15/20); no $p < 0.05$ — a single run swings **±15 pts** at $n=20$.
- Only **V_SC** validated; **D2** not yet exercised. Cost = tokens, not wall-clock.

Pre-registered follow-up

~7B × 3 families, **n=200**, MATH+GSM8K.
Success: match SC-4 at $\leq 1.5\times$ tokens.