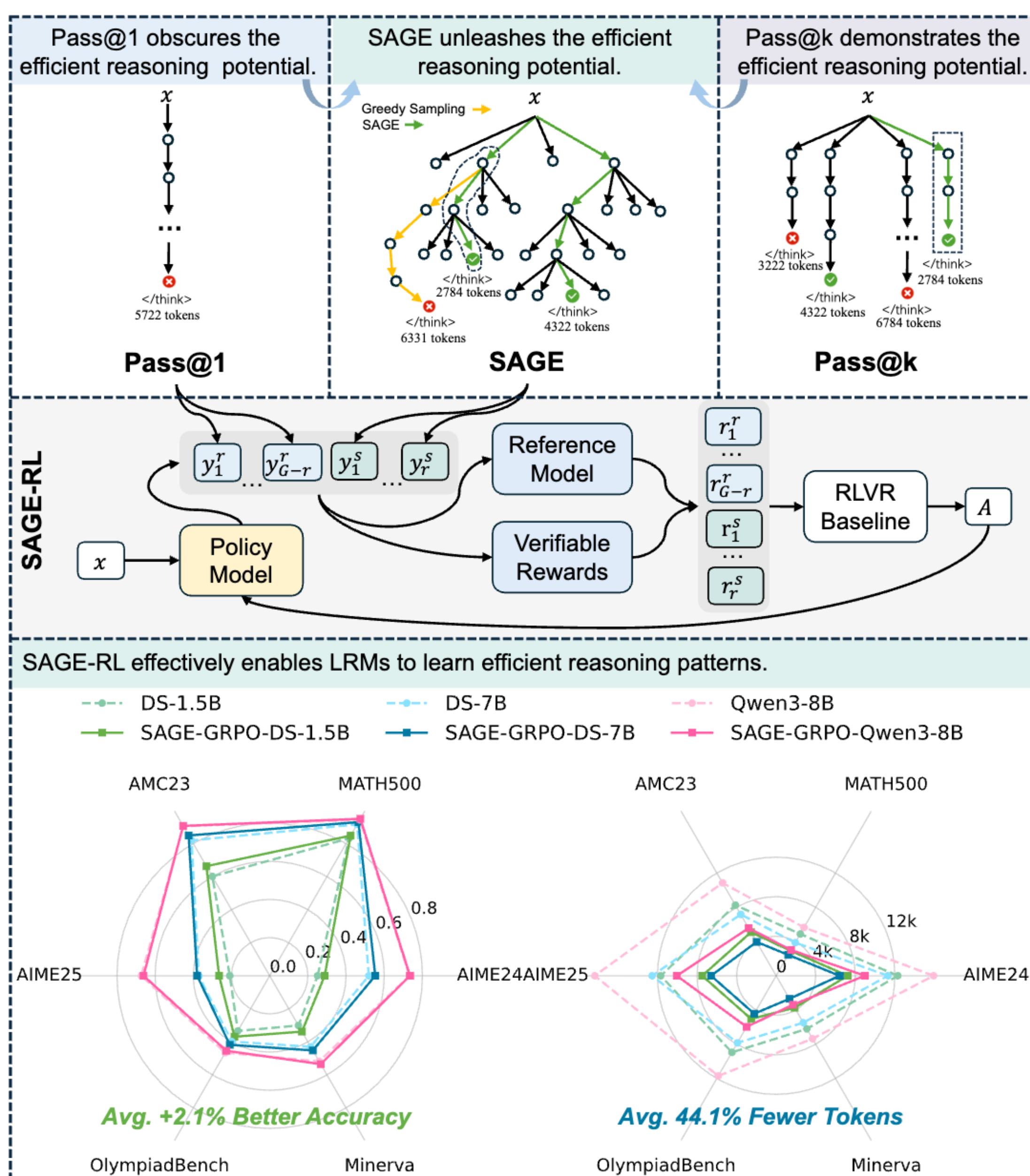


MOTIVATION

- ❑ **Pass@1**: Existing sampling strategies fail to enable timely termination of thinking.
- ❑ **Pass@k**: Scaling CoT length does not lead to correct answers.



The surprising performance of relatively shorter responses in pass@k reveals the inherent potential of the model for efficient reasoning. The pervasive redundancy of reasoning steps in pass@1 indicates that current sampling paradigms obscure this potential.

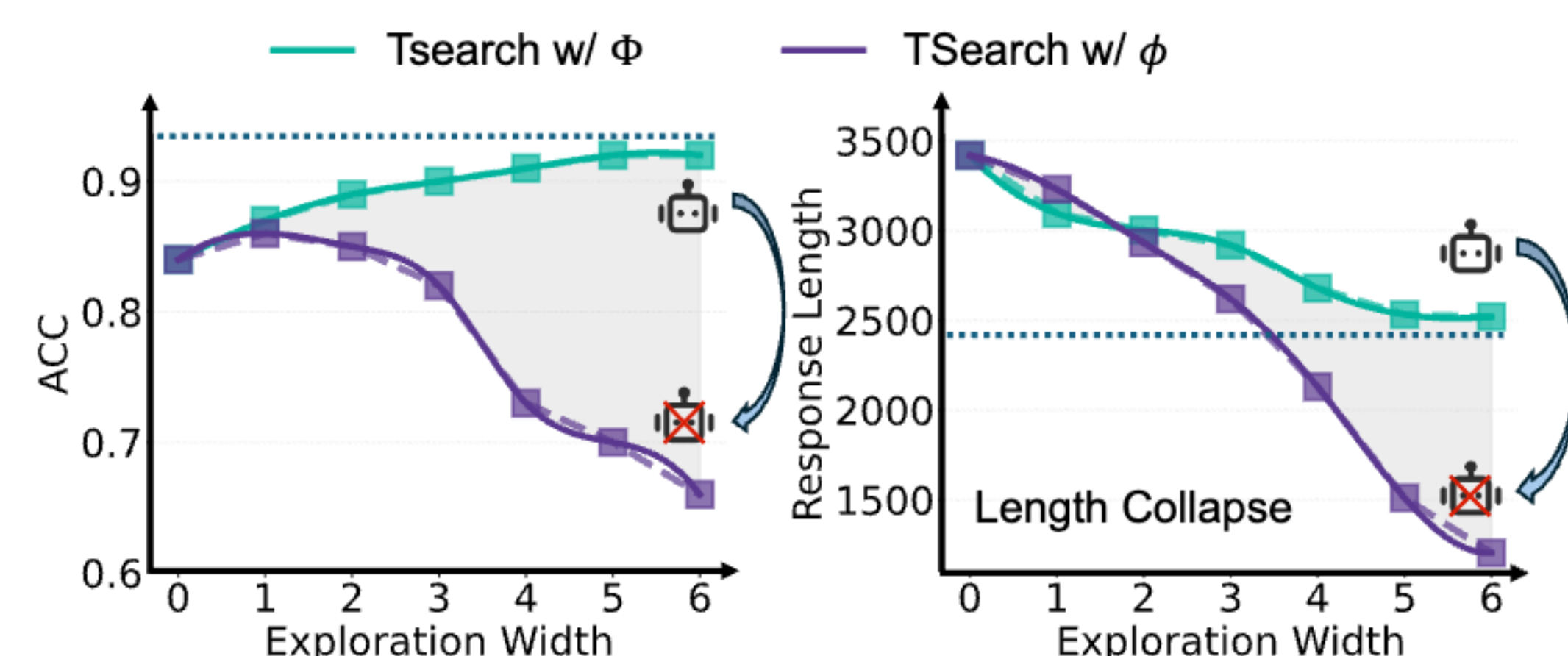
METHODOLOGY

❑ **TSearch** uses the average cumulative log-probability ϕ as a confidence measure over reasoning paths, performs token-level expansion to retain the top-m high-confidence candidate chains, and terminates once `</think>` appears with sufficiently high confidence, thereby revealing the existence of concise yet effective reasoning paths within the model.

❑ **SAGE** further reformulates TSearch from token-level search into step-level reasoning-chain exploration: it samples complete reasoning steps at each expansion stage and directly leverages the model's confident termination signal over `</think>` to identify shorter and more accurate CoTs.

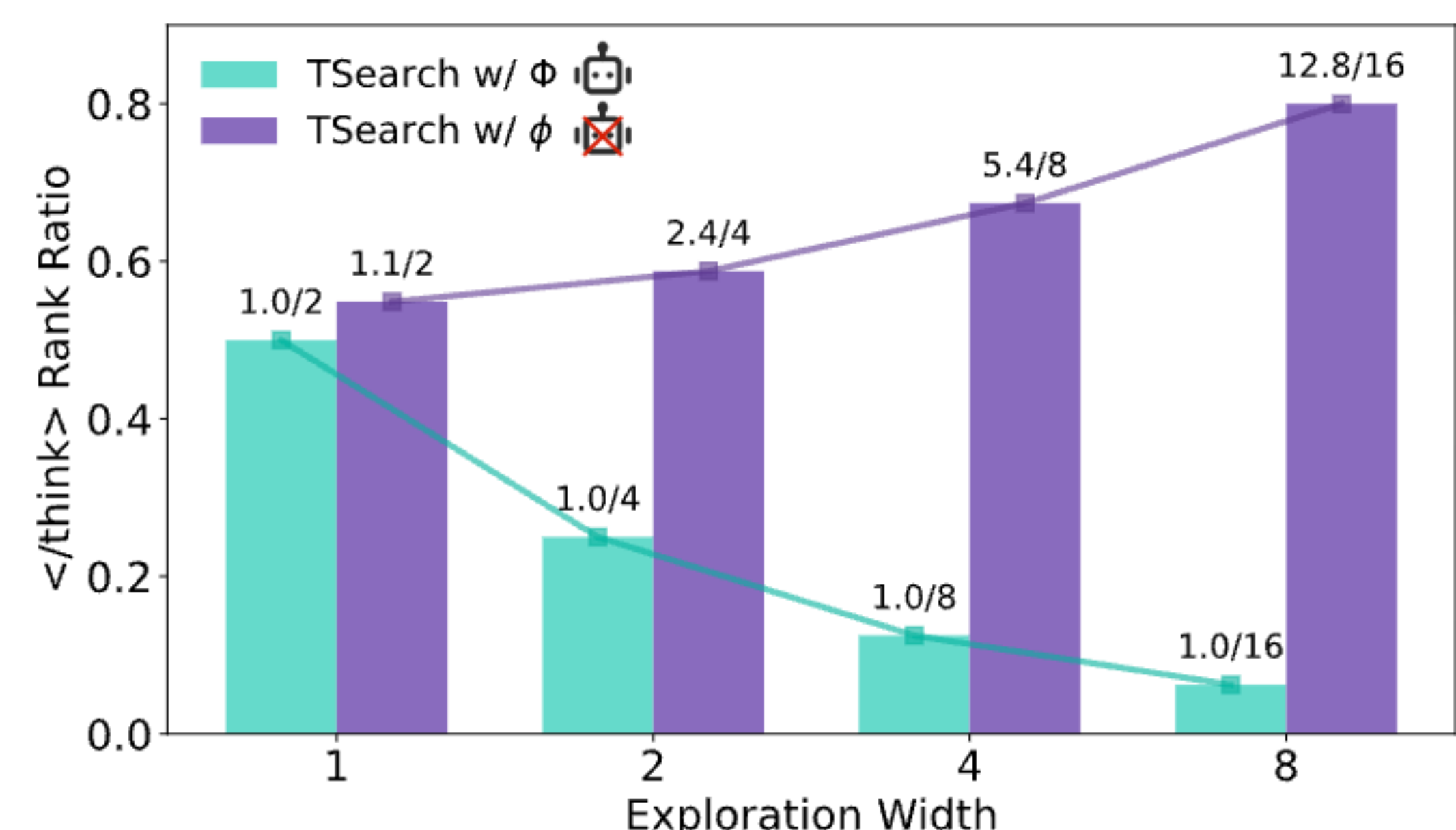
❑ **SAGE-RL** incorporates SAGE as a hybrid sampler in the RLVR rollout process, where part of each response group is generated by SAGE as high-quality concise reasoning samples while the remainder is produced by standard random sampling, enabling the model to internalize efficient reasoning patterns under standard pass@1 inference.

OBSERVATIONS

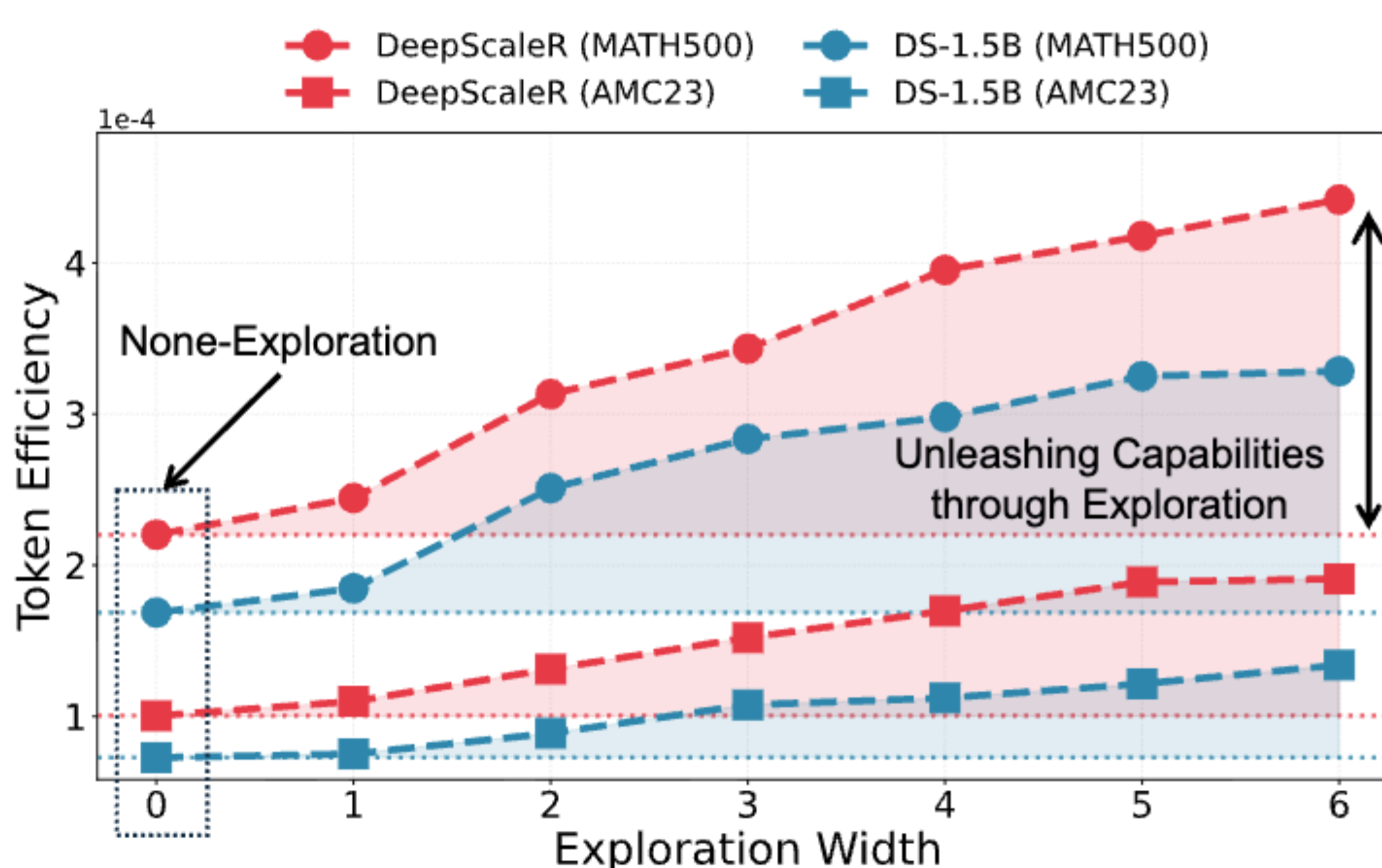


High-Confidence Paths Lead to Efficient Reasoning

Method	TR	ACC	T-LEN	LEN
Random	-	0.84	3126	3419
TSearch (4,1) w/ ϕ	1.00	0.79	1712	2129
TSearch (4,1) w/ ϕ	0.75	0.82	2022	2333
TSearch (4,1) w/ ϕ	0.50	0.89	2176	2609
TSearch (4,1) w/ Φ	1.00	0.92	2213	2609
TSearch (4,1) w/ Φ	0.75	0.92	2221	2621
TSearch (4,1) w/ Φ	0.50	0.91	2212	2632



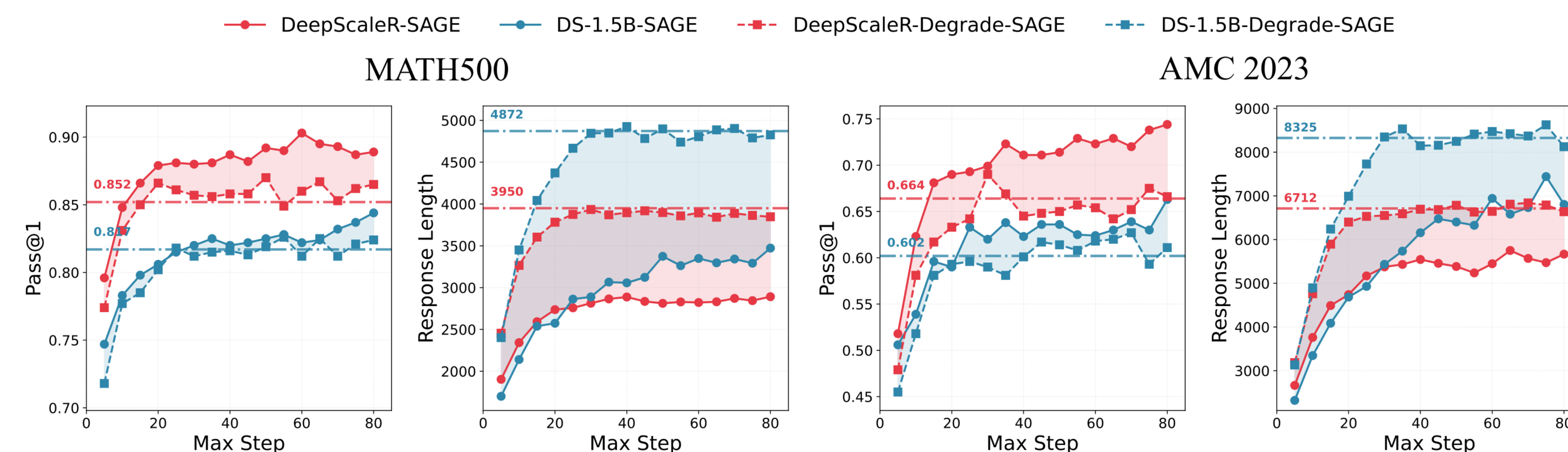
High-Confidence Paths Lead to Confident Ends



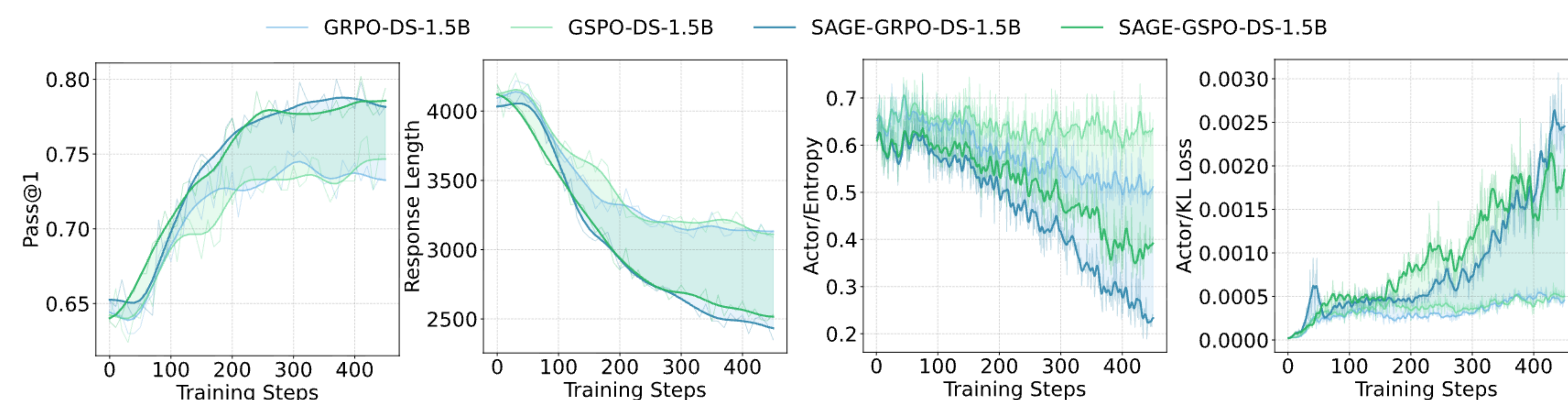
Scaling Exploration Drives Capability Convergence

EXPERIMENTS

Performance of SAGE



Training Dynamics of SAGE-RL



Performance on Challenge Math Benchmarks

Method	MATH-500			AIME 2024			AIME 2025			OlympiadBench		
	Pass@11(%)	LEN1	TE1(×10 ⁻³)	Pass@11(%)	LEN1	TE1(×10 ⁻³)	Pass@11(%)	LEN1	TE1(×10 ⁻³)	Pass@11(%)	LEN1	TE1(×10 ⁻³)
DS-1.5B	83.2	4882	17.0	25.1	12300	2.04	20.9	11669	1.79	33.4	8954	3.73
+ LC-R1	80.4 (12.8)	2973 (11909)	27.0 (158.8%)	23.3 (11.8)	7098 (15202)	3.28 (160.8%)	20.9 (10.0)	6942 (14727)	3.01 (168.2%)	32.0 (11.4)	4632 (14322)	6.91 (185.3%)
+ ThinkPrune-2k	81.7 (11.5)	2826 (12056)	28.9 (170.0%)	23.7 (11.4)	7085 (15215)	3.35 (164.2%)	19.7 (11.2)	6918 (14751)	2.85 (159.2%)	32.9 (10.5)	4752 (14202)	6.92 (185.5%)
+ AdaptThink	80.4 (12.8)	2563 (12319)	31.4 (184.1%)	25.7 (10.6)	8055 (14245)	3.19 (156.4%)	21.8 (10.9)	8155 (13514)	2.67 (149.2%)	31.6 (10.8)	4563 (14391)	7.14 (191.4%)
+ Efficient Reasoning	82.0 (11.2)	2821 (12061)	29.1 (170.6%)	26.2 (11.1)	9189 (13111)	2.85 (139.7%)	22.9 (12.0)	8590 (13079)	2.67 (149.2%)	33.8 (10.4)	5755 (13199)	5.87 (157.4%)
+ GRPO	83.6 (10.4)	3907 (1975)	21.4 (125.6%)	28.3 (13.2)	8767 (13533)	3.23 (158.3%)	24.1 (13.2)	8263 (13406)	2.92 (163.1%)	34.2 (10.8)	6323 (12631)	5.41 (145.0%)
+ SAGE-GRPO	84.8 (11.6)	2915 (11967)	29.1 (170.7%)	28.8 (13.7)	7243 (15057)	3.98 (195.1%)	26.5 (15.6)	7479 (14190)	3.54 (197.8%)	36.9 (13.5)	5050 (13904)	7.31 (196.0%)
+ GSPO	83.4 (10.2)	3898 (1984)	25.3 (121.4%)	28.3 (13.2)	8604 (13696)	3.29 (161.3%)	25.1 (14.2)	8227 (13442)	3.05 (170.4%)	34.6 (11.2)	6410 (12544)	5.40 (144.8%)
+ SAGE-GSPO	85.2 (12.0)	2921 (11961)	29.2 (171.6%)	28.5 (13.4)	6889 (15411)	4.14 (1102.9%)	27.1 (16.2)	7167 (14502)	3.78 (1111.1%)	37.3 (13.9)	5172 (13782)	7.21 (193.3%)
DeepScaleR	86.0	3805	22.6	31.4	9370	3.35	25.4	9310	2.73	35.9	5972	6.01
+ ThinkPrune-2k	82.5 (13.5)	2946 (1859)	28.0 (123.9%)	33.5 (12.1)	8108 (11262)	4.13 (123.3%)	26.0 (10.6)	7486 (11824)	3.47 (127.1%)	35.1 (10.8)	4723 (11249)	7.43 (123.6%)
+ GRPO	87.6 (11.6)	3482 (1323)	25.2 (111.3%)	35.6 (14.2)	8592 (1778)	4.14 (123.6%)	27.4 (12.0)	8185 (11125)	3.35 (122.7%)	36.2 (10.3)	5443 (1529)	6.65 (110.6%)
+ SAGE-GRPO	88.8 (12.8)	3117 (1688)	28.4 (125.7%)	36.1 (14.7)	8094 (11276)	4.46 (133.1%)	27.2 (11.8)	7704 (11606)	3.53 (129.3%)	36.5 (10.6)	4890 (11082)	7.46 (124.1%)
DS-7B	91.6	3871	23.7	51.9	11305	4.59	37.1	12540	2.96	39.8	7839	5.08
+ LC-R1	87.3 (14.3)	2076 (11795)	42.1 (177.7%)	51.7 (10.2)	6820 (14485)	7.58 (165.1%)	35.7 (11.4)	7458 (15082)	4.79 (161.8%)	41.4 (11.6)	4193 (13646)	9.87 (194.3%)
+ AdaptThink	88.9 (12.7)	2199 (11672)	40.4 (170.9%)	52.1 (10.2)	6679 (14626)	7.80 (169.9%)	35.0 (12.1)	7807 (14733)	4.48 (172.3%)	38.9 (10.9)	4915 (12924)	7.91 (155.7%)
+ Efficient Reasoning	89.8 (11.8)	2408 (11463)	37.3 (157.6%)	51.9 (10.0)	6667 (14638)	7.78 (169.5%)	36.2 (10.9)	7501 (15039)	4.82 (162.8%)	40.1 (10.3)	4599 (13240)	8.72 (171.7%)
+ GRPO-LEAD	89.5 (12.1)	2752 (11119)	32.5 (137.1%)	53.1 (11.2)	7023 (14282)	7.56 (164.7%)	36.1 (11.0)	7842 (14698)	4.60 (155.4%)	40.6 (10.8)	4972 (12867)	8.17 (160.8%)
+ GRPO	92.0 (10.4)	3219 (1652)	28.5 (120.2%)	52.5 (10.6)	8424 (12881)	6.23 (135.7%)	38.4 (11.3)	10123 (12417)	3.79 (128.0%)	41.2 (11.4)	5498 (12341)	7.50 (147.6%)
+ SAGE-GRPO	93.0 (11.4)	2141 (11730)	43.4 (183.1%)	55.3 (13.4)	6422 (14883)	8.61 (187.6%)	38.0 (10.9)	6583 (15957)	5.77 (194.9%)	41.8 (12.0)	4435 (13404)	9.42 (185.4%)
Qwen3-8B	94.4	5640	16.7	73.2	15920	4.60	67.3	18342	3.67	46.6	11707	4.00
+ GRPO	93.6 (10.8)	4470 (11170)	20.9 (125.1%)	72.8 (10.4)	10573 (15347)	6.89 (149.8%)	66.6 (10.7)	13981 (14361)	4.76 (129.7%)	45.1 (11.5)	7512 (14195)	6.00 (150.0%)
+ SAGE-GRPO	95.0 (10.6)	3015 (12625)	31.5 (188.2%)	73.5 (10.3)	8975 (16945)	8.19 (178.0%)	66.6 (10.7)	10052 (18290)	6.58 (179.3%)	45.4 (11.2)	5972 (15735)	7.60 (190.0%)
+ GSPO	94.6 (10.2)	4342 (11298)	22.2 (132.9%)	73.0 (10.2)	10544 (15376)	6.92 (150.4%)	66.2 (11.1)	14082 (14260)	4.70 (130.2%)	46.6 (10.0)	7964 (13743)	5.85 (146.2%)
+ SAGE-GSPO	94.4 (10.0)	2753 (12887)	34.3 (1105.3%)	73.7 (10.5)	8547 (17373)	8.62 (187.4%)	66.0 (11.3)	9183 (19159)	7.19 (195.9%)	46.7 (10.1)	5436 (16271)	8.59 (1114.7%)

Statistics of RFCS on MATH500 across LRMs

