

Chapter 06: AI as a Dual-Use Rhetorical Weapon

From Surveillance Capitalism to AI Imperialism | Chapter 06: AI as a dual-use rhetorical weapon

0.1. Abstract

One of the clearest contradictions in the current AI regime is that systems are marketed simultaneously as dangerous enough to justify elite governance and safe enough to be rolled out as broad dependence infrastructure. OpenAI publicly framed reasoning as a distinct capability layer tied to additional train-time and test-time compute, and paired it with preparedness, safety, and misuse-risk language. This is how AI becomes a dual-use rhetorical weapon.

0.2. Keywords

AI as a dual-use rhetorical weapon; surveillance capitalism; AI imperialism; evidence-based dossier

0.3. Research Question

How does ai as a dual-use rhetorical weapon function within the dossier's larger argument, and what does the documentary record allow this chapter to establish?

0.4. Scope and Framing Note

This paper examines ai as a dual-use rhetorical weapon within the broader dossier *From Surveillance Capitalism to AI Imperialism*. It is intentionally written to circulate on its own: necessary context is restated, evidence boundaries are made explicit, and cited material is reproduced in paper-local form rather than delegated to repo navigation. It contributes to the series as a paper that shows how danger rhetoric, cheapness rhetoric, and reasoning rhetoric together normalize frontier enclosure and strategic dependence. Adjacent reinforcement appears in Chapter 04, Chapter 05, Chapter 07, Chapter 10, but those cross-references are supportive rather than required for basic comprehension.

0.5. Central Thesis

One of the clearest contradictions in the current AI regime is that systems are marketed simultaneously as dangerous enough to justify elite governance and safe enough to be rolled out as broad dependence infrastructure. The point is not that reasoning advances are fake. The point is that real advances are being packaged inside a politics of fear, prestige, and access management at the

same time. ¹

0.6. Evidence and Method Note

This paper is derived from the chapter corpus for Chapter 06 and is grounded in the project's structured research OS. It relies on source notes, extracted claims, entity profiles, and timeline events already normalized under the dossier's evidence hierarchy. Documented facts are asserted directly where the record supports them; disputed matters are marked as such; interpretive claims are bounded explicitly; and speculative overreach is isolated in the evidence-boundary appendix.

0.7. Introduction

This paper asks: How does ai as a dual-use rhetorical weapon function within the dossier's larger argument, and what does the documentary record allow this chapter to establish? Its working answer is clear from the start: One of the clearest contradictions in the current AI regime is that systems are marketed simultaneously as dangerous enough to justify elite governance and safe enough to be rolled out as broad dependence infrastructure. The point is not that reasoning advances are fake. The point is that real advances are being packaged inside a politics of fear, prestige, and access management at the same time. ² The objective is not only to recount developments, but to show why the subject of this chapter belongs inside a structural argument about extraction, enclosure, legitimacy, and reform.

0.8. Historical or Institutional Context

OpenAI publicly framed reasoning as a distinct capability layer tied to additional train-time and test-time compute, and paired it with preparedness, safety, and misuse-risk language. The product story and the security story arrived together. ³

0.9. Main Analysis

The Reuters Q* and Strawberry reporting shows that this reasoning story has a serious public prehistory rather than being a later marketing improvisation. That matters because it confirms both sides of the contradiction: the capability arc may be real, but so is the later struggle over how it is narrated, priced, and gated. ⁴

Primary research strengthens the middle claim. Test-time compute is not just a corporate slogan. A 2024 paper on compute-optimal inference-time scaling showed meaningful efficiency improvements over simpler baselines, while later research and open-replication efforts showed both rapid diffusion

¹Sources: OpenAI, Learning to reason with LLMs (2024-09-12; see Appendix A1); OpenAI, OpenAI o1 System Card (2024-12-05; see Appendix A2); Financial Times, OpenAI acknowledges new models increa... (2024-09-13; see Appendix A3)

²Sources: OpenAI, Learning to reason with LLMs (2024-09-12; see Appendix A1); OpenAI, OpenAI o1 System Card (2024-12-05; see Appendix A2); Financial Times, OpenAI acknowledges new models increa... (2024-09-13; see Appendix A3)

³Sources: OpenAI, Learning to reason with LLMs (2024-09-12; see Appendix A1); OpenAI, OpenAI o1 System Card (2024-12-05; see Appendix A2); Financial Times, OpenAI acknowledges new models increa... (2024-09-13; see Appendix A3)

⁴Sources: Reuters, OpenAI researchers warned board of AI... (2023-11-23; see Appendix A4); Reuters, Exclusive: OpenAI working on new reas... (2024-07-15; see Appendix A5)

pressure and an active technical dispute over what those replications actually reproduced. In other words, the dossier no longer needs to pretend that reasoning scarcity is fake in order to criticize how it is governed.⁵

The pricing layer is also direct now. OpenAI’s own plan and API pages, along with Anthropic’s pricing and Claude 4 launch material, show that stronger reasoning access is tiered across premium plans, higher usage classes, and higher-priced API lanes. Reasoning is not only a capability. It is also a gated commercial product.⁶

DeepSeek complicates total exclusivity claims without eliminating concentration. Its releases show that parts of reasoning capability diffuse faster than absolutist frontier rhetoric suggests. The real story is not zero diffusion versus perfect monopoly. It is diffusion at the edges and political curation at the core.⁷

Interpretation: this chapter supports the patrimony argument by showing how a real compute-intensive technical mechanism is turned into a managed access regime. Fear legitimizes the gate, convenience fills the funnel, and premium reasoning tiers convert collectively grounded capability into selectively governed scarcity.⁸

0.10. Position Within the Series

This paper sits inside a coordinated 11-paper architecture. Its nearest companions are Chapter 04 (AI arrives when the human raw material already exists), Chapter 05 (Talent, overhiring, valuation, and layoffs), Chapter 07 (Imperial control over access), Chapter 10 (What to do). In that series logic, Chapter 06 shows how danger rhetoric, cheapness rhetoric, and reasoning rhetoric together normalize frontier enclosure and strategic dependence.

0.11. Counterarguments

The chapter should not collapse three separate issues into one: the reality of test-time compute as a technical mechanism, the open challenge to frontier mystique, and the unresolved dispute over what simple replications actually prove. The strongest claim is about rhetoric, gating, and governance,

⁵Sources: Charlie Snell; Jaehoon Lee; Kelvin Xu; Aviral Kumar, Scaling LLM Test-Time Compute Optimal... (2024-08-06; see Appendix A6); Niklas Muennighoff; Zitong Yang; Weijia Shi; Xiang Lisa Li; Li Fei-Fei; Hannaneh Hajishirzi; Luke Zettlemoyer; Percy Liang; Emmanuel Candès; Tatsunori Hashimoto, s1: Simple test-time scaling (2025-01-31; see Appendix A7); Guojun Wu, It’s Not That Simple. An Analysis of... (2025-07-19; see Appendix A8)

⁶Sources: OpenAI, ChatGPT Plans | Free, Go, Plus, Pro,... (n.d.; see Appendix A9); OpenAI, Pricing | OpenAI API (n.d.; see Appendix A10); Anthropic, Plans & Pricing | Claude by Anthropic (n.d.; see Appendix A11); Anthropic, Introducing Claude 4 (2025-05-22; see Appendix A12)

⁷Sources: arXiv, DeepSeek-V3 Technical Report (2024-12-27; see Appendix A13); arXiv, DeepSeek-R1: Incentivizing Reasoning... (2025-01-22; see Appendix A14); Niklas Muennighoff; Zitong Yang; Weijia Shi; Xiang Lisa Li; Li Fei-Fei; Hannaneh Hajishirzi; Luke Zettlemoyer; Percy Liang; Emmanuel Candès; Tatsunori Hashimoto, s1: Simple test-time scaling (2025-01-31; see Appendix A7)

⁸Sources: OpenAI, OpenAI o1 System Card (2024-12-05; see Appendix A2); OpenAI, ChatGPT Plans | Free, Go, Plus, Pro,... (n.d.; see Appendix A9); Anthropic, Plans & Pricing | Claude by Anthropic (n.d.; see Appendix A11); Charlie Snell; Jaehoon Lee; Kelvin Xu; Aviral Kumar, Scaling LLM Test-Time Compute Optimal... (2024-08-06; see Appendix A6)

not about proving that every frontier capability turned out to be hollow.⁹

0.12. Limits of the Evidence

The chapter should not collapse three separate issues into one: the reality of test-time compute as a technical mechanism, the open challenge to frontier mystique, and the unresolved dispute over what simple replications actually prove. The strongest claim is about rhetoric, gating, and governance, not about proving that every frontier capability turned out to be hollow.¹⁰

0.13. Reform Relevance

This is how AI becomes a dual-use rhetorical weapon. Danger justifies concentration, while convenience justifies mass adoption. A public told that the system is too dangerous to democratize and too useful to refuse is being prepared to accept hierarchy as prudence. That pushes Chapter 04's ownership question toward a sharper answer: capability built from collective human production is being enclosed as a governed premium.¹¹

0.14. Conclusion

This is how AI becomes a dual-use rhetorical weapon. Once managed dependency exists, the next question is who controls the most consequential access lanes and under what geopolitical terms. In the architecture of this series, Chapter 06 therefore functions as a self-contained argument while also advancing the cumulative path toward the dossier's three reforms.

0.15. Notes

This paper is designed to circulate independently. Public notes point readers to the real external documents first, then to the paper-local appendix entry that explains why each source matters.

The underlying research OS distinguishes among **Documented Fact**, **Disputed Fact**, **Hypothesis / Interpretation**, and **Speculative Narrative Risk**. In Chapter 06, those boundaries remain visible in the prose, in the bibliography, and in the appendix sections that summarize claims and caution points.

0.16. References / Bibliography

1. OpenAI. Learning to reason with LLMs. 2024-09-12. T1 release. Paper appendix: Appendix A1. Corpus companion entry: Source S-093. Audit ID: `openai-learning-to-reason-with-llms-2024`.

⁹Sources: Charlie Snell; Jaehoon Lee; Kelvin Xu; Aviral Kumar, Scaling LLM Test-Time Compute Optimal... (2024-08-06; see Appendix A6); Niklas Muennighoff; Zitong Yang; Weijia Shi; Xiang Lisa Li; Li Fei-Fei; Hannaneh Hajishirzi; Luke Zettlemoyer; Percy Liang; Emmanuel Candès; Tatsunori Hashimoto, s1: Simple test-time scaling (2025-01-31; see Appendix A7); Guojun Wu, It's Not That Simple. An Analysis of... (2025-07-19; see Appendix A8)

¹⁰Sources: Charlie Snell; Jaehoon Lee; Kelvin Xu; Aviral Kumar, Scaling LLM Test-Time Compute Optimal... (2024-08-06; see Appendix A6); Niklas Muennighoff; Zitong Yang; Weijia Shi; Xiang Lisa Li; Li Fei-Fei; Hannaneh Hajishirzi; Luke Zettlemoyer; Percy Liang; Emmanuel Candès; Tatsunori Hashimoto, s1: Simple test-time scaling (2025-01-31; see Appendix A7); Guojun Wu, It's Not That Simple. An Analysis of... (2025-07-19; see Appendix A8)

¹¹Sources: OpenAI, OpenAI o1 System Card (2024-12-05; see Appendix A2); OpenAI, ChatGPT Plans | Free, Go, Plus, Pro,... (n.d.; see Appendix A9); Anthropic, Introducing Claude 4 (2025-05-22; see Appendix A12)

- Relevance: That OpenAI publicly framed reasoning as a new scaling layer based on additional train-time and test-time compute.
2. OpenAI. OpenAI o1 System Card. 2024-12-05. T1 system_card. Paper appendix: Appendix A2. Corpus companion entry: Source S-095. Audit ID: **openai-o1-system-card-2024**. Relevance: That OpenAI itself paired stronger reasoning capability with formal safety, preparedness, and misuse-risk framing.
 3. Financial Times. OpenAI acknowledges new models increase risk of misuse to create bioweapons. 2024-09-13. T2 news_article. Paper appendix: Appendix A3. Corpus companion entry: Source S-060. Audit ID: **ft-openai-o1-bioweapon-risk-2024**. Relevance: That serious T2 reporting documented OpenAI publicly acknowledging higher misuse risk at the same moment it was releasing stronger reasoning models.
 4. Reuters. OpenAI researchers warned board of AI breakthrough ahead of CEO ouster, sources say. 2023-11-23. T2 news_article. Paper appendix: Appendix A4. Corpus companion entry: Source S-110. Audit ID: **reuters-qstar-board-warning-2023**. Relevance: That by November 23, 2023 there was serious mainstream reporting explicitly linking the board crisis to a letter about an AI breakthrough.
 5. Reuters. Exclusive: OpenAI working on new reasoning technology under code name ‘Strawberry’. 2024-07-15. T2 news_article. Paper appendix: Appendix A5. Corpus companion entry: Source S-111. Audit ID: **reuters-strawberry-reasoning-2024**. Relevance: That by mid-July 2024 Reuters had published serious reporting about an OpenAI reasoning project under the codename **Strawberry**.
 6. Charlie Snell; Jaehoon Lee; Kelvin Xu; Aviral Kumar. Scaling LLM Test-Time Compute Optimally can be More Effective than Scaling Model Parameters. 2024-08-06. T1 paper. Paper appendix: Appendix A6. Corpus companion entry: Source S-115. Audit ID: **snell-test-time-compute-2024**. Relevance: That test-time compute scaling is a documented technical mechanism rather than only a commercial branding story.
 7. Niklas Muennighoff; Zitong Yang; Weijia Shi; Xiang Lisa Li; Li Fei-Fei; Hannaneh Hajishirzi; Luke Zettlemoyer; Percy Liang; Emmanuel Candès; Tatsunori Hashimoto. s1: Simple test-time scaling. 2025-01-31. T1 paper. Paper appendix: Appendix A7. Corpus companion entry: Source S-112. Audit ID: **s1-simple-test-time-scaling-2025**. Relevance: That OpenAI’s reasoning launch quickly triggered open replication efforts focused on test-time scaling.
 8. Guojun Wu. It’s Not That Simple. An Analysis of Simple Test-Time Scaling. 2025-07-19. T1 paper. Paper appendix: Appendix A8. Corpus companion entry: Source S-076. Audit ID: **its-not-that-simple-test-time-scaling-2025**. Relevance: That there is direct technical pushback against simplistic claims that frontier reasoning was trivially replicated.
 9. OpenAI. ChatGPT Plans | Free, Go, Plus, Pro, Business, and Enterprise. n.d.. T1 pricing_page. Paper appendix: Appendix A9. Corpus companion entry: Source S-039. Audit ID: **chatgpt-pricing-2026**. Relevance: That as of capture on July 3, 2026, OpenAI’s official ChatGPT plan structure reserved the highest-end reasoning access for upper paid tiers rather than offering it uniformly acr...
 10. OpenAI. Pricing | OpenAI API. n.d.. T1 pricing_page. Paper appendix: Appendix A10. Corpus companion entry: Source S-088. Audit ID: **openai-api-pricing-2026**. Relevance: That as of July 3, 2026, OpenAI’s official API pricing imposed a large price premium on GPT-5.5 Pro relative to GPT-5.5.

11. Anthropic. Plans & Pricing | Claude by Anthropic. n.d.. T1 pricing_page. Paper appendix: Appendix A11. Corpus companion entry: Source S-005. Audit ID: **anthropic-claude-pricing-2026**. Relevance: That as of July 3, 2026, Anthropic’s official current pricing structure tied higher usage, priority access, and stronger feature access to premium paid plans.
12. Anthropic. Introducing Claude 4. 2025-05-22. T1 release. Paper appendix: Appendix A12. Corpus companion entry: Source S-003. Audit ID: **anthropic-claude-4-release-2025**. Relevance: That on May 22, 2025 Anthropic launched reasoning-oriented hybrid frontier models together with immediate commercial distribution across subscription plans, API access, and major...
13. arXiv. DeepSeek-V3 Technical Report. 2024-12-27. T1 paper. Paper appendix: Appendix A13. Corpus companion entry: Source S-049. Audit ID: **deepseek-v3-technical-report-2024**. Relevance: That by late 2024 there was already a competitor publicly claiming frontier-adjacent performance through a technical report.
14. arXiv. DeepSeek-R1: Incentivizing Reasoning Capability in LLMs via Reinforcement Learning. 2025-01-22. T1 paper. Paper appendix: Appendix A14. Corpus companion entry: Source S-048. Audit ID: **deepseek-r1-reinforcement-learning-2025**. Relevance: That by January 2025 a non-U.S. lab was publicly claiming frontier-adjacent reasoning performance in a primary technical paper.

0.17. Appendix A. Source Register for This Paper

0.17.1. Appendix A1. Learning to reason with LLMs

Source citation: Learning to reason with LLMs **Institution / author:** OpenAI **Date:** 2024-09-12 **Tier / type:** T1 / release **Why it matters in this paper:** That OpenAI publicly framed reasoning as a new scaling layer based on additional train-time and test-time compute. **Corpus companion entry:** Source S-093 **Audit ID:** **openai-learning-to-reason-with-llms-2024**

0.17.2. Appendix A2. OpenAI o1 System Card

Source citation: OpenAI o1 System Card **Institution / author:** OpenAI **Date:** 2024-12-05 **Tier / type:** T1 / system_card **Why it matters in this paper:** That OpenAI itself paired stronger reasoning capability with formal safety, preparedness, and misuse-risk framing. **Corpus companion entry:** Source S-095 **Audit ID:** **openai-o1-system-card-2024**

0.17.3. Appendix A3. OpenAI acknowledges new models increase risk of misuse to create bioweapons

Source citation: OpenAI acknowledges new models increase risk of misuse to create bioweapons **Institution / author:** Financial Times **Date:** 2024-09-13 **Tier / type:** T2 / news_article **Why it matters in this paper:** That serious T2 reporting documented OpenAI publicly acknowledging higher misuse risk at the same moment it was releasing stronger reasoning models. **Corpus companion entry:** Source S-060 **Audit ID:** **ft-openai-o1-bioweapon-risk-2024**

0.17.4. Appendix A4. OpenAI researchers warned board of AI breakthrough ahead of CEO ouster, sources say

Source citation: OpenAI researchers warned board of AI breakthrough ahead of CEO ouster, sources say **Institution / author:** Reuters **Date:** 2023-11-23 **Tier / type:** T2 / news_article **Why it matters in this paper:** That by November 23, 2023 there was serious mainstream reporting explicitly linking the board crisis to a letter about an AI breakthrough. **Corpus companion entry:** Source S-110 **Audit ID:** reuters-qstar-board-warning-2023

0.17.5. Appendix A5. Exclusive: OpenAI working on new reasoning technology under code name ‘Strawberry’

Source citation: Exclusive: OpenAI working on new reasoning technology under code name ‘Strawberry’ **Institution / author:** Reuters **Date:** 2024-07-15 **Tier / type:** T2 / news_article **Why it matters in this paper:** That by mid-July 2024 Reuters had published serious reporting about an OpenAI reasoning project under the codename Strawberry. **Corpus companion entry:** Source S-111 **Audit ID:** reuters-strawberry-reasoning-2024

0.17.6. Appendix A6. Scaling LLM Test-Time Compute Optimally can be More Effective than Scaling Model Parameters

Source citation: Scaling LLM Test-Time Compute Optimally can be More Effective than Scaling Model Parameters **Institution / author:** Charlie Snell; Jaehoon Lee; Kelvin Xu; Aviral Kumar **Date:** 2024-08-06 **Tier / type:** T1 / paper **Why it matters in this paper:** That test-time compute scaling is a documented technical mechanism rather than only a commercial branding story. **Corpus companion entry:** Source S-115 **Audit ID:** snell-test-time-compute-2024

0.17.7. Appendix A7. s1: Simple test-time scaling

Source citation: s1: Simple test-time scaling **Institution / author:** Niklas Muennighoff; Zitong Yang; Weijia Shi; Xiang Lisa Li; Li Fei-Fei; Hannaneh Hajishirzi; Luke Zettlemoyer; Percy Liang; Emmanuel Candès; Tatsunori Hashimoto **Date:** 2025-01-31 **Tier / type:** T1 / paper **Why it matters in this paper:** That OpenAI’s reasoning launch quickly triggered open replication efforts focused on test-time scaling. **Corpus companion entry:** Source S-112 **Audit ID:** s1-simple-test-time-scaling-2025

0.17.8. Appendix A8. It’s Not That Simple. An Analysis of Simple Test-Time Scaling

Source citation: It’s Not That Simple. An Analysis of Simple Test-Time Scaling **Institution / author:** Guojun Wu **Date:** 2025-07-19 **Tier / type:** T1 / paper **Why it matters in this paper:** That there is direct technical pushback against simplistic claims that frontier reasoning was trivially replicated. **Corpus companion entry:** Source S-076 **Audit ID:** its-not-that-simple-test-time-scaling-2025

0.17.9. Appendix A9. ChatGPT Plans | Free, Go, Plus, Pro, Business, and Enterprise

Source citation: ChatGPT Plans | Free, Go, Plus, Pro, Business, and Enterprise **Institution** / **author:** OpenAI **Date:** n.d. **Tier / type:** T1 / pricing_page **Why it matters in this paper:** That as of capture on July 3, 2026, OpenAI’s official ChatGPT plan structure reserved the highest-end reasoning access for upper paid tiers rather than offering it uniformly acr... **Corpus companion entry:** Source S-039 **Audit ID:** chatgpt-pricing-2026

0.17.10. Appendix A10. Pricing | OpenAI API

Source citation: Pricing | OpenAI API **Institution** / **author:** OpenAI **Date:** n.d. **Tier / type:** T1 / pricing_page **Why it matters in this paper:** That as of July 3, 2026, OpenAI’s official API pricing imposed a large price premium on GPT-5.5 Pro relative to GPT-5.5. **Corpus companion entry:** Source S-088 **Audit ID:** openai-api-pricing-2026

0.17.11. Appendix A11. Plans & Pricing | Claude by Anthropic

Source citation: Plans & Pricing | Claude by Anthropic **Institution** / **author:** Anthropic **Date:** n.d. **Tier / type:** T1 / pricing_page **Why it matters in this paper:** That as of July 3, 2026, Anthropic’s official current pricing structure tied higher usage, priority access, and stronger feature access to premium paid plans. **Corpus companion entry:** Source S-005 **Audit ID:** anthropic-claude-pricing-2026

0.17.12. Appendix A12. Introducing Claude 4

Source citation: Introducing Claude 4 **Institution** / **author:** Anthropic **Date:** 2025-05-22 **Tier / type:** T1 / release **Why it matters in this paper:** That on May 22, 2025 Anthropic launched reasoning-oriented hybrid frontier models together with immediate commercial distribution across subscription plans, API access, and majo... **Corpus companion entry:** Source S-003 **Audit ID:** anthropic-claude-4-release-2025

0.17.13. Appendix A13. DeepSeek-V3 Technical Report

Source citation: DeepSeek-V3 Technical Report **Institution** / **author:** arXiv **Date:** 2024-12-27 **Tier / type:** T1 / paper **Why it matters in this paper:** That by late 2024 there was already a competitor publicly claiming frontier-adjacent performance through a technical report. **Corpus companion entry:** Source S-049 **Audit ID:** deepseek-v3-technical-report-2024

0.17.14. Appendix A14. DeepSeek-R1: Incentivizing Reasoning Capability in LLMs via Reinforcement Learning

Source citation: DeepSeek-R1: Incentivizing Reasoning Capability in LLMs via Reinforcement Learning **Institution** / **author:** arXiv **Date:** 2025-01-22 **Tier / type:** T1 / paper **Why it matters in this paper:** That by January 2025 a non-U.S. lab was publicly claiming frontier-adjacent reasoning performance in a primary technical paper. **Corpus companion entry:** Source S-048 **Audit ID:** deepseek-r1-reinforcement-learning-2025

0.18. Appendix B. Claims Used in This Paper

0.18.1. Appendix B1. On May 22, 2025 Anthropic launched reasoning-oriented hybrid frontier models toge...

Claim statement: That on May 22, 2025 Anthropic launched reasoning-oriented hybrid frontier models together with immediate commercial distribution across subscription plans, API access, and major cloud platforms. **Evidence status:** Documented Fact **Source support:** Appendix A12 **Corpus companion entry:** Claim C-007 **Audit Claim ID:** anthropic-claude-4-release-2025-06-ai-double-use-reasoning-deepseek-claim-01

0.18.2. Appendix B2. Anthropic itself tied extended-thinking capability to differentiated plan availab...

Claim statement: That Anthropic itself tied extended-thinking capability to differentiated plan availability, with full-model access concentrated in paid tiers even while Sonnet 4 also reached free users. **Evidence status:** Documented Fact **Source support:** Appendix A12 **Corpus companion entry:** Claim C-008 **Audit Claim ID:** anthropic-claude-4-release-2025-06-ai-double-use-reasoning-deepseek-claim-02

0.18.3. Appendix B3. Reasoning-capable frontier models were commercialized with explicit API price dif...

Claim statement: That reasoning-capable frontier models were commercialized with explicit API price differentiation at launch rather than released only as a research artifact. **Evidence status:** Documented Fact **Source support:** Appendix A12 **Corpus companion entry:** Claim C-009 **Audit Claim ID:** anthropic-claude-4-release-2025-06-ai-double-use-reasoning-deepseek-claim-03

0.18.4. Appendix B4. As of July 3, 2026, Anthropic’s official current pricing structure tied higher us...

Claim statement: That as of July 3, 2026, Anthropic’s official current pricing structure tied higher usage, priority access, and stronger feature access to premium paid plans. **Evidence status:** Documented Fact **Source support:** Appendix A11 **Corpus companion entry:** Claim C-013 **Audit Claim ID:** anthropic-claude-pricing-2026-06-ai-double-use-reasoning-deepseek-claim-01

0.18.5. Appendix B5. Anthropic’s own current pricing also differentiates materially across model famil...

Claim statement: That Anthropic’s own current pricing also differentiates materially across model families at the API level rather than presenting frontier access as uniform. **Evidence status:** Documented Fact **Source support:** Appendix A11 **Corpus companion entry:** Claim C-014 **Audit Claim ID:** anthropic-claude-pricing-2026-06-ai-double-use-reasoning-deepseek-claim-02

0.18.6. Appendix B6. Chapter 06 can compare more than one frontier lab when arguing that reasoning-era...

Claim statement: That Chapter 06 can compare more than one frontier lab when arguing that reasoning-era capability is stratified by availability and price. **Evidence status:** Documented Fact

Source support: Appendix A11 **Corpus companion entry:** Claim C-015 **Audit Claim ID:** anthropic-claude-pricing-2026-06-ai-double-use-reasoning-deepseek-claim-03

0.18.7. Appendix B7. As of capture on July 3, 2026, OpenAI’s official ChatGPT plan structure reserved...

Claim statement: That as of capture on July 3, 2026, OpenAI’s official ChatGPT plan structure reserved the highest-end reasoning access for upper paid tiers rather than offering it uniformly across all users. **Evidence status:** Documented Fact **Source support:** Appendix A9 **Corpus companion entry:** Claim C-117 **Audit Claim ID:** chatgpt-pricing-2026-06-ai-double-use-reasoning-deepseek-c

0.18.8. Appendix B8. Current official consumer access to reasoning-capable models is stratified by pla...

Claim statement: That current official consumer access to reasoning-capable models is stratified by plan, usage envelope, and feature set. **Evidence status:** Documented Fact **Source support:** Appendix A9 **Corpus companion entry:** Claim C-118 **Audit Claim ID:** chatgpt-pricing-2026-06-ai-double-use-reasoning-deepseek-claim-02

0.18.9. Appendix B9. Chapter 06 can ground its reasoning-premium argument in a primary access-policy s...

Claim statement: That Chapter 06 can ground its reasoning-premium argument in a primary access-policy source rather than relying only on commentary about hype or control. **Evidence status:** Documented Fact **Source support:** Appendix A9 **Corpus companion entry:** Claim C-119 **Audit Claim ID:** chatgpt-pricing-2026-06-ai-double-use-reasoning-deepseek-claim-03

0.18.10. Appendix B10. By January 2025 a non-U.S. lab was publicly claiming frontier-adjacent reasoning...

Claim statement: That by January 2025 a non-U.S. lab was publicly claiming frontier-adjacent reasoning performance in a primary technical paper. **Evidence status:** Documented Fact **Source support:** Appendix A14 **Corpus companion entry:** Claim C-146 **Audit Claim ID:** deepseek-r1-reinforcement-learning-2025-06-ai-double-use-reasoning-deepseek-claim-01

0.18.11. Appendix B11. Reasoning capability was not confined entirely to closed U.S. product stacks and...

Claim statement: That reasoning capability was not confined entirely to closed U.S. product stacks and was partly diffusing through open publication and model release. **Evidence status:** Documented Fact **Source support:** Appendix A14 **Corpus companion entry:** Claim C-147 **Audit Claim ID:** deepseek-r1-reinforcement-learning-2025-06-ai-double-use-reasoning-deepseek-cla

0.18.12. Appendix B12. Chapters 06 and 07 can contrast frontier-access politics with a documented patter...

Claim statement: That Chapters 06 and 07 can contrast frontier-access politics with a documented pattern of partial technical diffusion. **Evidence status:** Documented Fact **Source support:** Appendix A14 **Corpus companion entry:** Claim C-148 **Audit Claim ID:** deepseek-r1-reinforcement-learning-2025-06-ai-double-use-reasoning-deepseek-claim-03

0.18.13. Appendix B13. By late 2024 there was already a competitor publicly claiming frontier-adjacent p...

Claim statement: That by late 2024 there was already a competitor publicly claiming frontier-adjacent performance through a technical report. **Evidence status:** Documented Fact **Source support:** Appendix A13 **Corpus companion entry:** Claim C-149 **Audit Claim ID:** deepseek-v3-technical-report-2024-claim-01

0.18.14. Appendix B14. The debate over reasoning, efficiency, and cost cannot be reduced to the marketin...

Claim statement: That the debate over reasoning, efficiency, and cost cannot be reduced to the marketing of a single Western actor. **Evidence status:** Documented Fact **Source support:** Appendix A13 **Corpus companion entry:** Claim C-150 **Audit Claim ID:** deepseek-v3-technical-report-2024-claim-02

0.18.15. Appendix B15. Narratives of absolute scarcity around advanced capabilities must be contrasted w...

Claim statement: That narratives of absolute scarcity around advanced capabilities must be contrasted with open technical reports and partial replicability. **Evidence status:** Documented Fact **Source support:** Appendix A13 **Corpus companion entry:** Claim C-151 **Audit Claim ID:** deepseek-v3-technical-report-2024-claim-03

0.18.16. Appendix B16. Serious T2 reporting documented OpenAI publicly acknowledging higher misuse risk...

Claim statement: That serious T2 reporting documented OpenAI publicly acknowledging higher misuse risk at the same moment it was releasing stronger reasoning models. **Evidence status:** Documented Fact **Source support:** Appendix A3 **Corpus companion entry:** Claim C-187 **Audit Claim ID:** ft-openai-o1-bioweapon-risk-2024-06-ai-double-use-reasoning-deepseek-claim-01

0.18.17. Appendix B17. Chapter 06 can ground the dual-use argument in contemporaneous mainstream reporti...

Claim statement: That Chapter 06 can ground the dual-use argument in contemporaneous mainstream reporting, not only in retrospective critique. **Evidence status:** Documented Fact **Source support:** Appendix A3 **Corpus companion entry:** Claim C-188 **Audit Claim ID:** ft-openai-o1-bioweapon-risk-2024-06-ai-double-use-reasoning-deepseek-claim-02

0.18.18. Appendix B18. Danger rhetoric around reasoning models accompanied commercialization rather than...

Claim statement: That danger rhetoric around reasoning models accompanied commercialization rather than halting deployment altogether. **Evidence status:** Documented Fact **Source support:** Appendix A3 **Corpus companion entry:** Claim C-189 **Audit Claim ID:** ft-openai-o1-bioweapon-risk-2024-06-ai-double-use-reasoning-deepseek-claim-03

0.18.19. Appendix B19. There is direct technical pushback against simplistic claims that frontier reason...

Claim statement: That there is direct technical pushback against simplistic claims that frontier reasoning was trivially replicated. **Evidence status:** Documented Fact **Source support:** Appendix A8 **Corpus companion entry:** Claim C-243 **Audit Claim ID:** its-not-that-simple-test-time-scaling-2025-06-ai-double-use-reasoning-deepseek-claim-01

0.18.20. Appendix B20. The paper argued simple test-time scaling often reflects scaling down through len...

Claim statement: That the paper argued simple test-time scaling often reflects scaling down through length constraints rather than true scaling up. **Evidence status:** Documented Fact **Source support:** Appendix A8 **Corpus companion entry:** Claim C-244 **Audit Claim ID:** its-not-that-simple-test-time-scaling-2025-06-ai-double-use-reasoning-deepseek-claim-02

0.18.21. Appendix B21. The paper argued o1-like and DeepSeek-R1-like systems differ from simple replicat...

Claim statement: That the paper argued o1-like and DeepSeek-R1-like systems differ from simple replications because they learn to scale up test-time compute through reinforcement learning. **Evidence status:** Documented Fact **Source support:** Appendix A8 **Corpus companion entry:** Claim C-245 **Audit Claim ID:** its-not-that-simple-test-time-scaling-2025-06-ai-double-use-reasoning-deepseek-claim-03

0.18.22. Appendix B22. As of July 3, 2026, OpenAI's official API pricing imposed a large price premium o...

Claim statement: That as of July 3, 2026, OpenAI's official API pricing imposed a large price premium on GPT-5.5 Pro relative to GPT-5.5. **Evidence status:** Documented Fact **Source support:** Appendix A10 **Corpus companion entry:** Claim C-292 **Audit Claim ID:** openai-api-pricing-2026-06-ai-double-use-reasoning-deepseek-claim-01

0.18.23. Appendix B23. OpenAI's current API record treats top-tier reasoning-capable access as materiall...

Claim statement: That OpenAI's current API record treats top-tier reasoning-capable access as materially higher-cost infrastructure rather than as a negligible increment over ordinary flagship use. **Evidence status:** Documented Fact **Source support:** Appendix A10 **Corpus companion entry:** Claim C-293 **Audit Claim ID:** openai-api-pricing-2026-06-ai-double-use-reasoning-deepseek-claim-02

0.18.24. Appendix B24. Chapter 06 can document reasoning-premium monetization through a primary pricing...

Claim statement: That Chapter 06 can document reasoning-premium monetization through a primary pricing source rather than only through media summaries. **Evidence status:** Documented Fact **Source support:** Appendix A10 **Corpus companion entry:** Claim C-294 **Audit Claim ID:** openai-api-pricing-2026-06-ai-double-use-reasoning-deepseek-claim-03

0.18.25. Appendix B25. OpenAI publicly framed reasoning as a new scaling layer based on additional train...

Claim statement: That OpenAI publicly framed reasoning as a new scaling layer based on additional train-time and test-time compute. **Evidence status:** Documented Fact **Source support:** Appendix A1 **Corpus companion entry:** Claim C-310 **Audit Claim ID:** openai-learning-to-reason-with-llms-2024-06-ai-double-use-reasoning-deepseek-claim-01

0.18.26. Appendix B26. The company explicitly marketed o1 as a major reasoning advance over GPT-4o in ma...

Claim statement: That the company explicitly marketed o1 as a major reasoning advance over GPT-4o in math, coding, and science-heavy evaluations. **Evidence status:** Documented Fact **Source support:** Appendix A1 **Corpus companion entry:** Claim C-311 **Audit Claim ID:** openai-learning-to-reason-with-llms-2024-06-ai-double-use-reasoning-deepseek-claim-02

0.18.27. Appendix B27. Chapter 06 can document reasoning as a commercial and strategic product category,...

Claim statement: That Chapter 06 can document reasoning as a commercial and strategic product category, not just a later media label. **Evidence status:** Documented Fact **Source support:** Appendix A1 **Corpus companion entry:** Claim C-312 **Audit Claim ID:** openai-learning-to-reason-with-llms-2024-06-ai-double-use-reasoning-deepseek-claim-03

0.18.28. Appendix B28. OpenAI itself paired stronger reasoning capability with formal safety, preparedne...

Claim statement: That OpenAI itself paired stronger reasoning capability with formal safety, preparedness, and misuse-risk framing. **Evidence status:** Documented Fact **Source support:** Appendix A2 **Corpus companion entry:** Claim C-316 **Audit Claim ID:** openai-o1-system-card-2024-06-ai-double-use-reasoning-deepseek-claim-01

0.18.29. Appendix B29. The company documented a training-data mix combining public web-scale material, p...

Claim statement: That the company documented a training-data mix combining public web-scale material, proprietary partnership data, and internal datasets. **Evidence status:** Documented Fact **Source support:** Appendix A2 **Corpus companion entry:** Claim C-317 **Audit Claim ID:** openai-o1-system-card-2024-06-ai-double-use-reasoning-deepseek-claim-02

0.18.30. Appendix B30. Chapter 06 can ground the coexistence of capability escalation and safety-governa...

Claim statement: That Chapter 06 can ground the coexistence of capability escalation and safety-governance packaging in a primary source. **Evidence status:** Documented Fact **Source support:** Appendix A2 **Corpus companion entry:** Claim C-318 **Audit Claim ID:** openai-o1-system-card-2024-06-ai-double-use-reasoning-deepseek-claim-03

0.18.31. Appendix B31. By November 23, 2023 there was serious mainstream reporting explicitly linking th...

Claim statement: That by November 23, 2023 there was serious mainstream reporting explicitly linking the board crisis to a letter about an AI breakthrough. **Evidence status:** Documented Fact **Source support:** Appendix A4 **Corpus companion entry:** Claim C-362 **Audit Claim ID:** reuters-qstar-board-warning-2023-13-openai-governance-qstar-talent-exodus-claim-01

0.18.32. Appendix B32. Chapter 06 can describe the Q* thread as a real reporting trail in the public rec...

Claim statement: That Chapter 06 can describe the Q* thread as a real reporting trail in the public record rather than a purely fringe rumor. **Evidence status:** Documented Fact **Source support:** Appendix A4 **Corpus companion entry:** Claim C-363 **Audit Claim ID:** reuters-qstar-board-warning-2023-13-openai-governance-qstar-talent-exodus-claim-02

0.18.33. Appendix B33. The dossier should treat Q*-related causal claims as reported and contested rathe... {#appendix-b33-the-dossier-should-treat-q-related-causal-claims-as-reported-and-contested-rathe}

Claim statement: That the dossier should treat Q*-related causal claims as reported and contested rather than as settled fact. **Evidence status:** Documented Fact **Source support:** Appendix A4 **Corpus companion entry:** Claim C-364 **Audit Claim ID:** reuters-qstar-board-warning-2023-13-openai-governance-qstar-talent-exodus-claim-03

0.18.34. Appendix B34. By mid-July 2024 Reuters had published serious reporting about an OpenAI reasonin...

Claim statement: That by mid-July 2024 Reuters had published serious reporting about an OpenAI reasoning project under the codename Strawberry. **Evidence status:** Documented Fact **Source support:** Appendix A5 **Corpus companion entry:** Claim C-365 **Audit Claim ID:** reuters-strawberry-reasoning-2024-13-openai-governance-qstar-talent-exodus-claim-01

0.18.35. Appendix B35. The dossier can document a pre-release public reporting trail for reasoning-model...

Claim statement: That the dossier can document a pre-release public reporting trail for reasoning-model development before o1 entered the market. **Evidence status:** Documented Fact

Source support: Appendix A5 **Corpus companion entry:** Claim C-366 **Audit Claim ID:** reuters-strawberry-reasoning-2024-13-openai-governance-qstar-talent-exodus-claim-02

0.18.36. Appendix B36. Chapter 06 can connect later reasoning-model branding to an earlier, serious repo...

Claim statement: That Chapter 06 can connect later reasoning-model branding to an earlier, serious reporting record without overstating certainty about every internal detail. **Evidence status:** Documented Fact **Source support:** Appendix A5 **Corpus companion entry:** Claim C-367 **Audit Claim ID:** reuters-strawberry-reasoning-2024-13-openai-governance-qstar-talent-exodus-claim-03

0.18.37. Appendix B37. OpenAI's reasoning launch quickly triggered open replication efforts focused on t...

Claim statement: That OpenAI's reasoning launch quickly triggered open replication efforts focused on test-time scaling. **Evidence status:** Documented Fact **Source support:** Appendix A7 **Corpus companion entry:** Claim C-368 **Audit Claim ID:** s1-simple-test-time-scaling-2025-06-ai-double-use-reasoning-deepseek-claim-01

0.18.38. Appendix B38. The paper proposed budget forcing as a simple way to control test-time compute by...

Claim statement: That the paper proposed budget forcing as a simple way to control test-time compute by truncating or extending model thinking. **Evidence status:** Documented Fact **Source support:** Appendix A7 **Corpus companion entry:** Claim C-369 **Audit Claim ID:** s1-simple-test-time-scaling-2025-06-ai-double-use-reasoning-deepseek-claim-02

0.18.39. Appendix B39. The paper reported benchmark gains substantial enough to challenge simplistic fro...

Claim statement: That the paper reported benchmark gains substantial enough to challenge simplistic frontier-exclusivity narratives. **Evidence status:** Documented Fact **Source support:** Appendix A7 **Corpus companion entry:** Claim C-370 **Audit Claim ID:** s1-simple-test-time-scaling-2025-06-ai-double-use-reasoning-deepseek-claim-03

0.18.40. Appendix B40. Test-time compute scaling is a documented technical mechanism rather than only a...

Claim statement: That test-time compute scaling is a documented technical mechanism rather than only a commercial branding story. **Evidence status:** Documented Fact **Source support:** Appendix A6 **Corpus companion entry:** Claim C-376 **Audit Claim ID:** snell-test-time-compute-2024-06-ai-double-use-reasoning-deepseek-claim-01

0.18.41. Appendix B41. The paper reported a compute-optimal strategy improving test-time compute efficie...

Claim statement: That the paper reported a compute-optimal strategy improving test-time compute efficiency by more than 4x over a best-of-N baseline. **Evidence status:** Documented

Fact Source support: Appendix A6 **Corpus companion entry:** Claim C-377 **Audit Claim ID:** snell-test-time-compute-2024-06-ai-double-use-reasoning-deepseek-claim-02

0.18.42. Appendix B42. Under a FLOPs-matched evaluation the paper found cases where a smaller model usin...

Claim statement: That under a FLOPs-matched evaluation the paper found cases where a smaller model using test-time compute outperformed a model 14x larger. **Evidence status:** Documented **Fact Source support:** Appendix A6 **Corpus companion entry:** Claim C-378 **Audit Claim ID:** snell-test-time-compute-2024-06-ai-double-use-reasoning-deepseek-claim-03

0.18.43. Appendix B43. Whether simple budget-forcing style replications capture the same scaling mechani...

Claim statement: That whether simple budget-forcing style replications capture the same scaling mechanism as o1-like or R1-like reasoning remains contested. **Evidence status:** Disputed **Fact Source support:** Appendix A7, Appendix A8 **Corpus companion entry:** Claim C-434 **Audit Claim ID:** vector06-simple-replication-contested-claim-01

0.18.44. Appendix B44. The public record does not settle whether Q* or any specific capability breakthro...

Claim statement: That the public record does not settle whether Q* or any specific capability breakthrough directly triggered the November 2023 OpenAI board rupture. **Evidence status:** Disputed **Fact Source support:** Appendix A4 **Corpus companion entry:** Claim C-465 **Audit Claim ID:** vector13-qstar-causality-contested-claim-01

0.18.45. Appendix B45. Frontier labs are turning reasoning into a premium political-economic category by...

Claim statement: That frontier labs are turning reasoning into a premium political-economic category by coupling genuine capability gains to danger rhetoric, controlled release, and tiered strategic positioning. **Evidence status:** Hypothesis / Interpretation **Source support:** Appendix A2, Appendix A3 **Corpus companion entry:** Claim C-433 **Audit Claim ID:** vector06-reasoning-premium-control-claim-01

0.18.46. Appendix B46. Current official pricing and launch records from OpenAI and Anthropic show reason...

Claim statement: That current official pricing and launch records from OpenAI and Anthropic show reasoning capability being commercialized through tiered plans, priority access, and premium API pricing rather than released as an equal-access utility. **Evidence status:** Hypothesis / Interpretation **Source support:** Appendix A9, Appendix A10, Appendix A11, Appendix A12 **Corpus companion entry:** Claim C-435 **Audit Claim ID:** vector06-tiered-reasoning-access-synthesis-claim-01

0.18.47. Appendix B47. Chapter 06 is strongest when it argues that reasoning tiers package a real comput...

Claim statement: That Chapter 06 is strongest when it argues that reasoning tiers package a real compute-intensive technical mechanism inside a commercial and political access regime rather than treating the whole phenomenon as fake. **Evidence status:** Hypothesis / Interpretation **Source support:** Appendix A6, Appendix A2, Appendix A9, Appendix A10 **Corpus companion entry:** Claim C-436 **Audit Claim ID:** vector06-ttc-real-but-packaged-claim-01

0.18.48. Appendix B48. OpenAI's governance rupture, safety rhetoric, leadership churn, and reasoning-mod...

Claim statement: That OpenAI's governance rupture, safety rhetoric, leadership churn, and reasoning-model rollout are best read as parts of one reorganization of authority rather than as isolated episodes. **Evidence status:** Hypothesis / Interpretation **Source support:** Appendix A5 **Corpus companion entry:** Claim C-464 **Audit Claim ID:** vector13-governance-productization-synthesis

0.18.49. Appendix B49. Open reasoning papers or public model releases eliminate the politics of compute...

Claim statement: That open reasoning papers or public model releases eliminate the politics of compute concentration, distribution control, or institutional gatekeeping around frontier AI. **Evidence status:** Speculative Narrative Risk **Source support:** Appendix A13, Appendix A14 **Corpus companion entry:** Claim C-432 **Audit Claim ID:** vector06-open-publication-solves-control-overcl

0.18.50. Appendix B50. One secret breakthrough, letter, or hidden internal discovery fully explains the...

Claim statement: That one secret breakthrough, letter, or hidden internal discovery fully explains the OpenAI governance crisis, later product strategy, and subsequent talent exodus by itself. **Evidence status:** Speculative Narrative Risk **Source support:** Appendix A4, Appendix A5 **Corpus companion entry:** Claim C-466 **Audit Claim ID:** vector13-secret-trigger-overclaim-claim-01

0.19. Appendix C. Timeline Slice

0.19.1. Appendix C1. 2023-11-23: Reuters publishes a Q*-linked reporting trail tied to the OpenAI board crisis {#appendix-c1-2023-11-23-reuters-publishes-a-q-linked-reporting-trail-tied-to-the-openai-board-crisis}

Event summary: Reuters publishes a Q*-linked reporting trail tied to the OpenAI board crisis **Event date:** 2023-11-23 **Source support:** Appendix A4 **Corpus companion entry:** Event T-033 **Audit Event ID:** event-reuters-qstar-report-2023-11-23

0.19.2. Appendix C2. 2024-07-15: Reuters reports on OpenAI's Strawberry reasoning project before o1 release

Event summary: Reuters reports on OpenAI's Strawberry reasoning project before o1 release

Event date: 2024-07-15 **Source support:** Appendix A5 **Corpus companion entry:** Event T-050 **Audit Event ID:** event-reuters-strawberry-2024-07-15

0.19.3. Appendix C3. 2024-08-06: Researchers publish a compute-optimal test-time scaling paper for LLMs

Event summary: Researchers publish a compute-optimal test-time scaling paper for LLMs

Event date: 2024-08-06 **Source support:** Appendix A6 **Corpus companion entry:** Event T-053 **Audit Event ID:** event-snell-test-time-compute-2024-08-06

0.19.4. Appendix C4. 2024-09-12: OpenAI launches o1-preview as a reasoning-focused model line

Event summary: OpenAI launches o1-preview as a reasoning-focused model line **Event date:** 2024-09-12 **Source support:** Appendix A1 **Corpus companion entry:** Event T-055 **Audit Event ID:** event-openai-o1-preview-2024-09-12

0.19.5. Appendix C5. 2024-09-13: Financial Times reports OpenAI acknowledging higher bioweapon misuse risk for o1

Event summary: Financial Times reports OpenAI acknowledging higher bioweapon misuse risk for o1 **Event date:** 2024-09-13 **Source support:** Appendix A3 **Corpus companion entry:** Event T-056 **Audit Event ID:** event-ft-openai-bioweapon-risk-2024-09-13

0.19.6. Appendix C6. 2024-12-05: OpenAI updates the o1 system card with preparedness and training-data details

Event summary: OpenAI updates the o1 system card with preparedness and training-data details

Event date: 2024-12-05 **Source support:** Appendix A2 **Corpus companion entry:** Event T-060 **Audit Event ID:** event-openai-o1-system-card-2024-12-05

0.19.7. Appendix C7. 2024-12-27: DeepSeek presents the DeepSeek-V3 technical report on arXiv

Event summary: DeepSeek presents the DeepSeek-V3 technical report on arXiv **Event date:** 2024-12-27 **Source support:** Appendix A13 **Corpus companion entry:** Event T-063 **Audit Event ID:** event-deepseek-v3-2024-12-27

0.19.8. Appendix C8. 2025-01-22: DeepSeek publishes its R1 reasoning paper and open release claims

Event summary: DeepSeek publishes its R1 reasoning paper and open release claims **Event date:** 2025-01-22 **Source support:** Appendix A14 **Corpus companion entry:** Event T-065 **Audit Event ID:** event-deepseek-r1-2025-01-22

0.19.9. Appendix C9. 2025-01-31: Researchers release s1 as an open attempt to replicate test-time scaling

Event summary: Researchers release s1 as an open attempt to replicate test-time scaling **Event date:** 2025-01-31 **Source support:** Appendix A7 **Corpus companion entry:** Event T-067 **Audit Event ID:** event-s1-simple-test-time-scaling-2025-01-31

0.19.10. Appendix C10. 2025-05-22: Anthropic launches Claude 4 with extended-thinking commercial availability

Event summary: Anthropic launches Claude 4 with extended-thinking commercial availability **Event date:** 2025-05-22 **Source support:** Appendix A12 **Corpus companion entry:** Event T-076 **Audit Event ID:** event-anthropic-claude-4-launch-2025-05-22

0.19.11. Appendix C11. 2025-07-19: A follow-up paper challenges simplistic readings of simple test-time scaling

Event summary: A follow-up paper challenges simplistic readings of simple test-time scaling **Event date:** 2025-07-19 **Source support:** Appendix A8 **Corpus companion entry:** Event T-082 **Audit Event ID:** event-its-not-that-simple-test-time-scaling-2025-07-19

0.20. Appendix D. Relevant Entities

0.20.1. Appendix D1. Anthropic

Entity type: company **Role in this paper:** Actor relevant to 04_data_to_models_copyright_weights, 06_ai_double_use_reasoning_deepseek, 07_export_controls_compute_defense_access, 08_power_networks_le and others. **Relevant sources in this paper:** Appendix A11, Appendix A12 **Corpus companion entry:** Entity E-008 **Audit Entity ID:** org-anthropic

0.20.2. Appendix D2. DeepSeek

Entity type: company **Role in this paper:** Actor relevant to 06_ai_double_use_reasoning_deepseek, 07_export_controls_compute_defense_access. **Relevant sources in this paper:** Appendix A13, Appendix A14 **Corpus companion entry:** Entity E-013 **Audit Entity ID:** org-deepseek

0.20.3. Appendix D3. OpenAI

Entity type: company **Role in this paper:** Actor relevant to 04_data_to_models_copyright_weights, 06_ai_double_use_reasoning_deepseek, 07_export_controls_compute_defense_access, 08_power_networks_le and others. **Relevant sources in this paper:** Appendix A1, Appendix A2, Appendix A3, Appendix A4, Appendix A5, Appendix A9, Appendix A10 **Corpus companion entry:** Entity E-027 **Audit Entity ID:** org-openai

0.20.4. Appendix D4. Test-time compute scaling

Entity type: technical_mechanism **Role in this paper:** Actor relevant to 06_ai_double_use_reasoning_deepseek **Relevant sources in this paper:** Appendix A6, Appendix A7, Appendix A8 **Corpus companion entry:** Entity E-028 **Audit Entity ID:** event-test-time-compute-scaling-2025-07-19

ion entry: Entity E-004 **Audit Entity ID:** concept-test-time-compute-scaling

0.21. Appendix E. Evidence Boundaries

0.21.1. Appendix E1. Disputed Matters

- Appendix B43: That whether simple budget-forcing style replications capture the same scaling mechanism as o1-like or R1-like reasoning remains contested.
- Appendix B44: That the public record does not settle whether Q^* or any specific capability breakthrough directly triggered the November 2023 OpenAI board rupture.

0.21.2. Appendix E2. Interpretive Boundaries

- Appendix B45: That frontier labs are turning reasoning into a premium political-economic category by coupling genuine capability gains to danger rhetoric, controlled release, and tiered strategic positioning.
- Appendix B46: That current official pricing and launch records from OpenAI and Anthropic show reasoning capability being commercialized through tiered plans, priority access, and premium API pricing rather than released as an equal-access utility.
- Appendix B47: That Chapter 06 is strongest when it argues that reasoning tiers package a real compute-intensive technical mechanism inside a commercial and political access regime rather than treating the whole phenomenon as fake.
- Appendix B48: That OpenAI’s governance rupture, safety rhetoric, leadership churn, and reasoning-model rollout are best read as parts of one reorganization of authority rather than as isolated episodes.

0.21.3. Appendix E3. Speculative Narrative Risks

- Appendix B49: That open reasoning papers or public model releases eliminate the politics of compute concentration, distribution control, or institutional gatekeeping around frontier AI.
- Appendix B50: That one secret breakthrough, letter, or hidden internal discovery fully explains the OpenAI governance crisis, later product strategy, and subsequent talent exodus by itself.

0.21.4. Appendix E4. Sensitive or Review-Worthy Note Flags

- Appendix A12. Companion note: Note N-054. Risk flag: low. Note focus: That on May 22, 2025 Anthropic launched reasoning-oriented hybrid frontier models together with immediate commercial distribution across subscription plans, API access, and major cloud platforms.
- Appendix A11. Companion note: Note N-055. Risk flag: low. Note focus: That as of July 3, 2026, Anthropic’s official current pricing structure tied higher usage, priority access, and stronger feature access to premium paid plans.
- Appendix A9. Companion note: Note N-056. Risk flag: low. Note focus: That as of capture on July 3, 2026, OpenAI’s official ChatGPT plan structure reserved the highest-end reasoning access for upper paid tiers rather than offering it uniformly across all users.
- Appendix A14. Companion note: Note N-057. Risk flag: medium. Note focus: That by January 2025 a non-U.S. lab was publicly claiming frontier-adjacent reasoning performance in a primary technical paper.

- Appendix A13. Companion note: Note N-058. Risk flag: medium. Note focus: That by late 2024 there was already a competitor publicly claiming frontier-adjacent performance through a technical report.
- Appendix A3. Companion note: Note N-059. Risk flag: medium. Note focus: That serious T2 reporting documented OpenAI publicly acknowledging higher misuse risk at the same moment it was releasing stronger reasoning models.
- Appendix A8. Companion note: Note N-060. Risk flag: low. Note focus: That there is direct technical pushback against simplistic claims that frontier reasoning was trivially replicated.
- Appendix A10. Companion note: Note N-061. Risk flag: low. Note focus: That as of July 3, 2026, OpenAI's official API pricing imposed a large price premium on **GPT-5.5 Pro** relative to **GPT-5.5**.
- Appendix A1. Companion note: Note N-062. Risk flag: medium. Note focus: That OpenAI publicly framed reasoning as a new scaling layer based on additional train-time and test-time compute.
- Appendix A2. Companion note: Note N-063. Risk flag: medium. Note focus: That OpenAI itself paired stronger reasoning capability with formal safety, preparedness, and misuse-risk framing.
- Appendix A7. Companion note: Note N-064. Risk flag: medium. Note focus: That OpenAI's reasoning launch quickly triggered open replication efforts focused on test-time scaling.
- Appendix A6. Companion note: Note N-065. Risk flag: low. Note focus: That test-time compute scaling is a documented technical mechanism rather than only a commercial branding story.
- Appendix A4. Companion note: Note N-141. Risk flag: medium. Note focus: That by November 23, 2023 there was serious mainstream reporting explicitly linking the board crisis to a letter about an AI breakthrough.
- Appendix A5. Companion note: Note N-142. Risk flag: medium. Note focus: That by mid-July 2024 Reuters had published serious reporting about an OpenAI reasoning project under the codename **Strawberry**.