



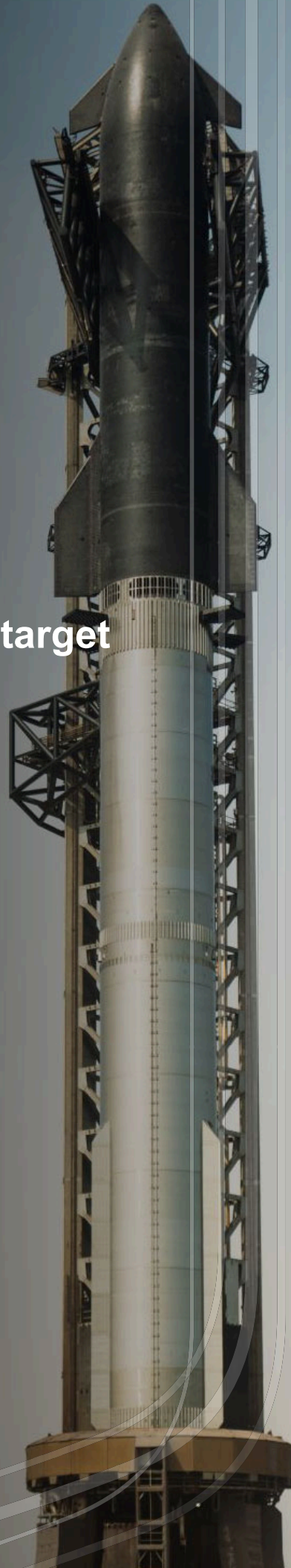
Capital
Markets

SpaceX

**Space, the final frontier, repriced;
initiating at Outperform with a \$225 price target**

EQUITY RESEARCH | JULY 7, 2026

For Required Non-U.S. Analyst and Conflicts Disclosures, see page 172.





Capital
Markets

July 7, 2026

SpaceX

Space, the final frontier, repriced; Initiating at Outperform, \$225 price target

Our view: We initiate coverage of SpaceX (SPCX) with an Outperform rating and a \$225 PT. We can appreciate timing risk associated with the company's space aspirations, but we believe sentiment will benefit from a proven track record of disruption and innovation, an almost ~\$2T 2035E TAM and virtually unmatched financial resources. Moreover, we believe SpaceX will continue to benefit from its position at the center of two of the most profound investment themes of this generation: the evolution of the space-based economy and AI. Our SOTP analysis supports our \$225 PT.

Key points:

SPCX began trading June 12, raising a total of ~\$86B at \$135/share from its IPO. The Starship is the flywheel that powers SpaceX's ambitions. SpaceX has successfully disrupted the global launch market with the Falcon 9 rocket, and Starlink is rewriting the global satellite comms market. However, we believe the expected cost savings from the Starship will be the catalyst for further unlocking the connectivity and orbital compute markets, as well as the lunar economy. We currently model ~2,440 Starship launches by 2030, and believe the 2H26 flight test program and beginning of a commercial Starship launch cadence in 2027 are the key near-term stock catalysts.

Expect Starlink momentum to continue. Subscriber counts nearly doubled in 2025, driven by constellation expansion, geographic growth, aggressive pricing, and consumer equipment subsidies. On speed and latency, Starlink is now comparable to low- and mid-tier terrestrial offerings, and the upcoming V3 satellite launch is expected to further improve service quality. To capture more of the Mobility Revenue TAM, we view efforts to potentially partner with an existing MVNO as possible.

SPCX knows in computing, infrastructure ownership is the enduring moat. AI models and applications may come and go, but compute lasts forever. We believe SPCX can leverage its cost-of-compute advantage, Grok, and the pending Cursor acquisition to become a leading AI intelligence provider. Moreover, we believe the opportunity of orbital datacenters provides SPCX with the chance to deliver a structural and durable cost-of-compute advantage versus any terrestrial competitor, enabled by the Starship.

\$225 PT based on a SOTP valuation. Specifically, we apply a 25x multiple to our 2029E Space segment EBITDA of \$3B, a 15x multiple to our 2029E connectivity EBITDA of \$42B, and a 15x multiple to our 2029E AI EBITDA estimate of \$147B. We believe that much of the franchise value rests with the AI segment, and our estimates reflect confidence in SPCX's ability to execute to its orbital compute targets. We appreciate that the challenges facing SPCX are substantial, but believe investors will focus on the financial resources, engineering capability, vision, and critical importance of SPCX to the U.S. gov't and national security objectives.

RBC Capital Markets, LLC

Ken Herbert (Analyst)

(415) 633-8583, ken.herbert@rbccm.com

Jonathan Atkin (Analyst)

(415) 633-8589, jonathan.atkin@rbccm.com

Brad Erickson (Analyst)

(971) 842-9607, brad.erickson@rbccm.com

Stephen Strackhouse (Vice President)

(212) 905-5943, stephen.strackhouse@rbccm.com

Rashim Jain (AVP)

(415) 633-8561, rashim.jain@rbccm.com

RBC Dominion Securities Inc.

Kirill Kozyar (AVP)

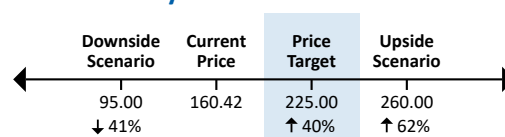
(416) 842-4730, kirill.kozyar@rbccm.com

Outperform

NASDAQ: SPCX; USD 160.42

Price Target USD 225.00

Scenario Analysis*



*Implied Total Returns

Key Statistics

Shares O/S (MM):	13,066.0	Market Cap (MM):	2,096,048
Dividend:	0.00	Yield:	0.0%
		Avg. Daily Volume:	NA

RBC Estimates

FY Dec	2026E	2027E	2028E	2029E
Revenue	41,921.5	79,874.3	159,440.3	263,888.1
EBITDA, Adj	20,681.4	48,797.1	110,863.1	192,474.4
Revenue	Q1	Q2	Q3	Q4
2026	4,694.0A	6,772.0E	12,255.0E	18,200.6E
2027	17,826.0E	19,398.2E	20,311.8E	22,338.4E
EBITDA, Adj				
2026	1,127.0A	2,043.3E	6,372.3E	11,138.8E
2027	11,275.3E	11,595.4E	12,360.6E	13,565.7E

All values in USD unless otherwise noted.

Priced as of prior trading day's market close, EST (unless otherwise noted).

Space Exploration Technologies Corp. Initiation (Nasdaq: SPCX)

July 7, 2026

RBC Capital Markets, LLC

Space

Ken Herbert, Aerospace & Defense, (415) 633-8583, ken.herbert@rbccm.com

Stephen Strackhouse, Vice President, (212) 905-5943, Stephen.Strackhouse@rbccm.com

Kevin Liu, Senior Associate, (646) 618-6929, kevin.liu@rbccm.com

Sanika Merchant, Associate, (332) 204-1773, Sanika.merchant@rbccm.com

Connectivity /AI

Brad Erickson, Internet, (971) 842-9607, brad.erickson@rbccm.com

Khadijah Gibson, Associate Vice President, (347) 610-7116, Khadijah.Gibson@rbccm.com

Audrey Stuart, Associate, (212) 437-9151, Audrey.stuart@rbccm.com

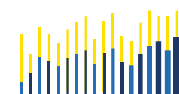
Jonathan Atkin, Telecom Services & Communications Infrastructure, (415) 633-8589, jonathan.atkin@rbccm.com

Rashim Jain, Associate Vice President, (415) 633-8561, Rashim.jain@rbccm.com

Adam Davis, Associate, (646) 618-6852, adam.davis@rbccm.com

RBC Dominion Securities Inc.

Kirill Kozyar, Associate Vice President, (614) 842-4730, kirill.Kozyar@rbccm.com



RBC Elements™

Driving insights through data



**Capital
Markets**

Initiating at Outperform with a \$225 price target

Investment Thesis	Space Exploration Technologies Corp (SPCX) is a global space technology company offering integrated hardware and software infrastructure across Space (launch), Connectivity (Starlink), and AI (Grok/ compute), with an estimated \$28.5T TAM across the segments. The company currently captures 80%+ of global mass to orbit and targets a 10x launch cost reduction with the development of Starship. We model 80%+ adj. EBITDA margins through 2031 and believe the stock reflects an improving execution and growth story, with business aviation and pricing power as key potential upsides. We believe a 15.2x multiple to 2029E EBITDA justifies our Outperform rating.
Key Investment Points	▪ SPCX sits at the center of two mega investment themes: AI and space. We believe these thematic trends will remain a focus for investors, and SPCX represents the best mega cap stock for exposure to these markets. The convergence of space and AI represents a unique value-creation opportunity that we believe SPCX is well-positioned to capture. ▪ The Starship is set to materially lower costs to access space. We believe the Starship will enable a step-change evolution in the connectivity market and significant acceleration of the Starlink offering. Moreover, the Starship is key to the orbital compute cost equation, in our view. Finally, we believe the Starship will enable evolution of a lunar economy at a more affordable and rapid pace. ▪ The pace of technological and product development is uncertain and likely faces delays. We believe investors are taking a long-term view, and setbacks in product schedules are expected. The company has proven its ability to solve complex challenges, and we believe investors will remain confident in the financial resources, technical capabilities and vision to success in its ambitious product roadmap. ▪ The company's track record of cost reduction and disruption underwrite investor confidence. The company's goals in the communication and compute markets represent similar disruption opportunities, and we believe SPCX remains well positioned to leverage this playbook into larger, more ambitious markets.
Valuation	▪ We value the company on a ~15x average EV/EBITDA multiple applied to our FY29E EBITDA of \$192.5B to derive our \$225 PT. We believe the company should trade at a premium to peers given its growth and margin outlook. We believe SPCX represents a unique opportunity from an AI and compute capacity perspective, as it is the only company currently establishing orbital compute capacity. ▪ Where could we be wrong? We believe any significant delays in the Starship development (the most important near-term focus for investors), or in solving other key technical challenges, could impact investor confidence, and thus the stock multiple. Our margin estimates could be high if development costs are greater than expected. The pace of AI adoption and integration across the broader economy is also a focus, in our view.
Key Potential Catalysts in 2026 & 2027	▪ Upcoming Starship flight test activity and the continued successful cadence of Falcon and Dragon launch activity. ▪ Closing of the Cursor acquisition, and potential future M&A. ▪ Additional terrestrial datacenter capacity additions, and compute capacity agreements. ▪ The introduction of Starlink V3 and continued subscriber growth and market penetration.
Risks to Rating & Price Target	▪ Macro/geopolitical uncertainty and the pace of customer adoption of new technologies. ▪ U.S. and foreign laws that are subject to change and uncertain interpretation could impact AI products, X platform, and Starlink. ▪ Key person risk, investor views on corporate governance risk, and the company's ability to overcome technical hurdles in its product evolution.

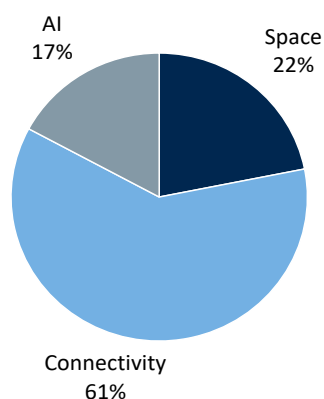
Table of contents

Company Overview	<u>6</u>
Key Stock Debates	<u>8</u>
Key Debate #1: Can SpaceX achieve its desired Starship launch cadence?	<u>9</u>
Key Debate #2: What is a viable growth path for Starlink Mobility?	<u>11</u>
Key Debate #3: Does Starlink's cost structure allow it to be competitive on pricing?	<u>13</u>
Key Debate #4: Can SPCX build a structural cost of compute advantage?	<u>14</u>
Key Debate #5: How much of a risk is the rise of cheaper models, open-sourced models or locally run models for the hyperscalers?	<u>16</u>
Key Debate #6: Can cost of compute, Grok and Cursor enable SPCX to become a leading AI intelligence provider?	<u>18</u>
Key Debate #7: Can the company execute on orbital datacenters?	<u>19</u>
Key Debate #8: Can AI weather regulatory risk without inhibiting market growth?	<u>20</u>
Valuation Analysis	<u>21</u>
Space Segment	<u>31</u>
Connectivity Segment	<u>38</u>
AI Segment	<u>49</u>
AI Infrastructure deep dive: Tokenomics and the ROIC on terrestrial and orbital compute	<u>57</u>
Terrestrial Datacenters	<u>64</u>
xAI Orbital Datacenters	<u>77</u>
Datacenter ROIC Analysis	<u>100</u>
The LLM landscape	<u>106</u>
SpaceX and Tesla Collaboration	<u>154</u>
Financial Overview of AI	<u>157</u>

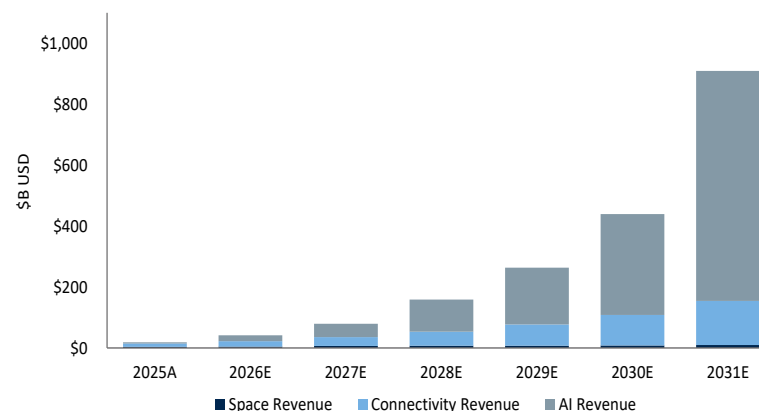
Company Overview

SpaceX is an integrated space tech leader that operates across three segments – Space, Connectivity, and AI. It generated ~\$18.7B revenue in 2025, with Connectivity making up the majority of the mix followed by Space and AI. The company has launched 80%+ of global mass to orbit annually since 2023 and ~7,400 metric tons to orbit with a 99%+ mission success rate across Falcon rockets. In addition, it operates a global broadband network with ~9,600 Starlink satellites serving 10.3M+ subscribers across 164 countries. In 2026, SpaceX acquired xAI and deployed a gigawatt-scale AI compute infrastructure (COLOSSUS and COLOSSUS II) with Grok frontier model integrated into X. Looking ahead, it aims to pursue orbital AI compute satellite deployment through the development of Starship, as well as lunar economy establishment at scale.

FY25 Revenues by Segment



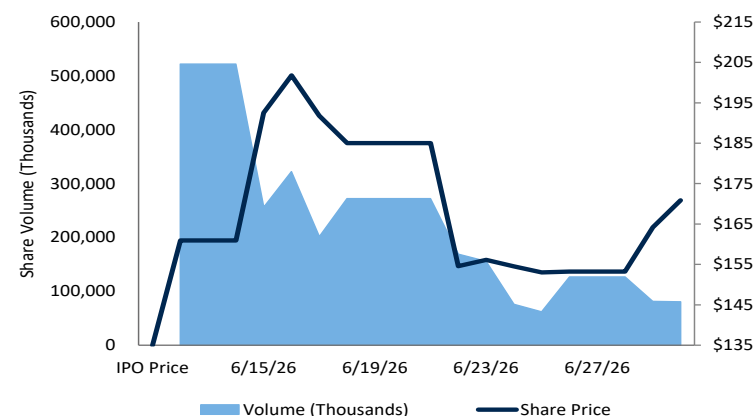
Revenue by Segment (\$B USD)



RBCe Model Framework

\$M	2025A	2026E	2027E	2028E	2029E	2030E	2031E
Net Sales	\$18,674	\$41,922	\$79,874	\$159,440	\$263,888	\$439,641	\$909,751
YoY Chg (%)	33.2%	124.5%	90.5%	99.6%	65.5%	66.6%	106.9%
Gross Profit	\$9,223	\$27,208	\$53,320	\$105,907	\$183,254	\$328,335	\$691,716
Gross Margin (%)	49.4%	64.9%	66.8%	66.4%	69.4%	74.7%	76.0%
YoY Chg (%)	53.2%	195.0%	96.0%	98.6%	73.0%	79.2%	110.7%
Operating Inc.	(\$2,589)	\$4,652	\$21,975	\$69,155	\$119,643	\$225,937	\$565,458
Op. Margin	-13.9%	11.1%	27.5%	43.4%	45.3%	51.4%	62.2%
YoY Chg (%)	(655.6%)	(279.7%)	372.4%	214.7%	73.0%	88.8%	150.3%
Adj. EBITDA	\$6,584	\$20,681	\$48,797	\$110,863	\$192,474	\$340,411	\$765,473
EBITDA Margin (%)	35.3%	49.3%	61.1%	69.5%	72.9%	77.4%	84.1%
YoY Chg (%)	23.1%	214.1%	135.9%	127.2%	73.6%	76.9%	124.9%
Cash from Operations	\$6,785	\$15,868	\$43,638	\$99,435	\$181,540	\$306,804	\$730,600
Capex	(\$20,619)	(\$47,980)	(\$127,477)	(\$177,244)	(\$266,372)	(\$330,629)	(\$686,232)
Free Cash Flow	(\$13,834)	(\$32,113)	(\$83,838)	(\$77,809)	(\$84,832)	(\$23,825)	\$44,368

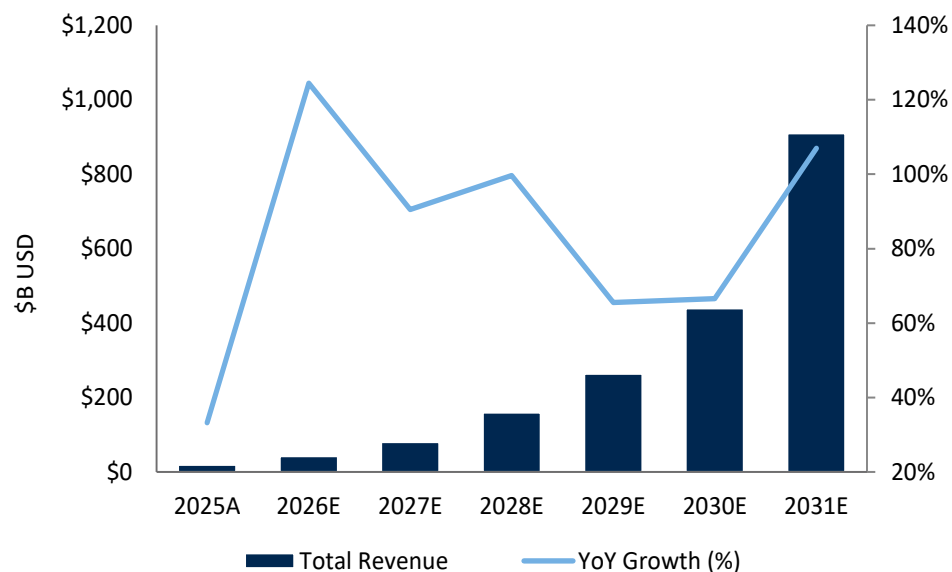
Stock Price and Trading Volume



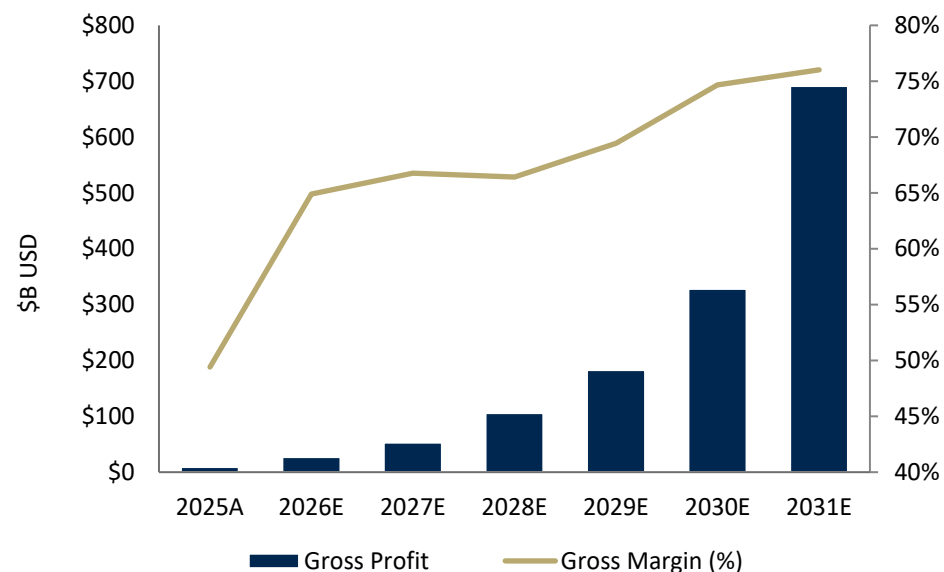
Source: Company reports, RBC Capital Markets estimates

Financial Snapshot

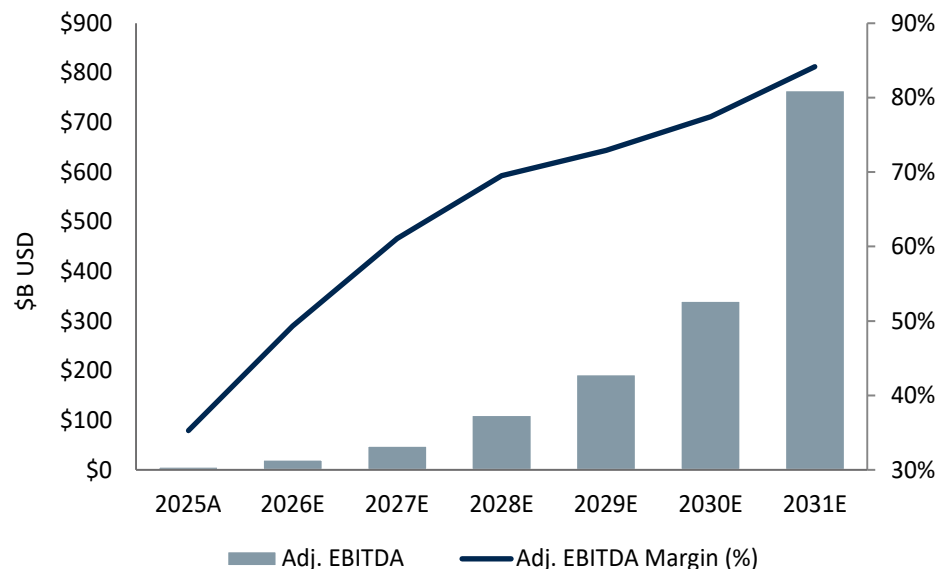
Net Sales (\$B USD)



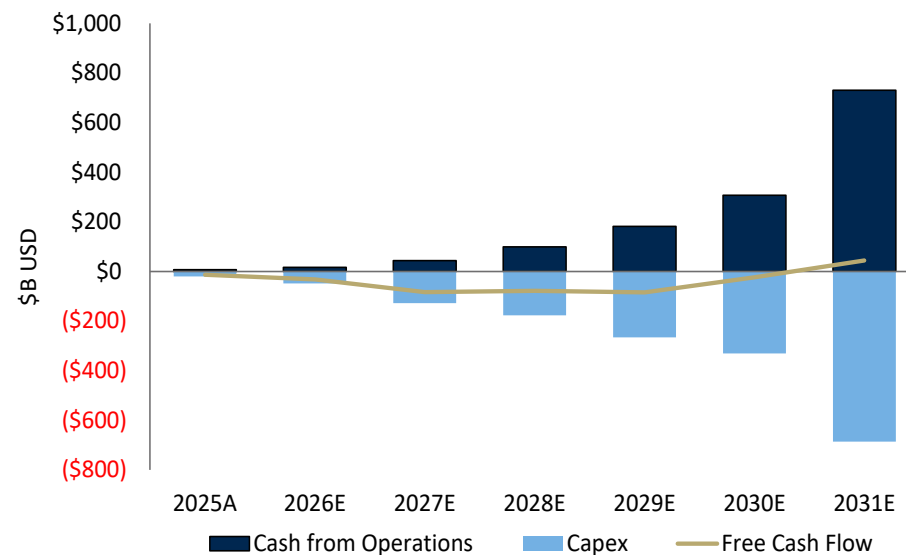
Gross Profit (\$B USD)



Adjusted EBITDA (\$B USD)



Free Cash Flow (\$B USD)



Source: S-1, company presentation, RBC Capital Markets estimates

Key Stock Debates



Capital
Markets

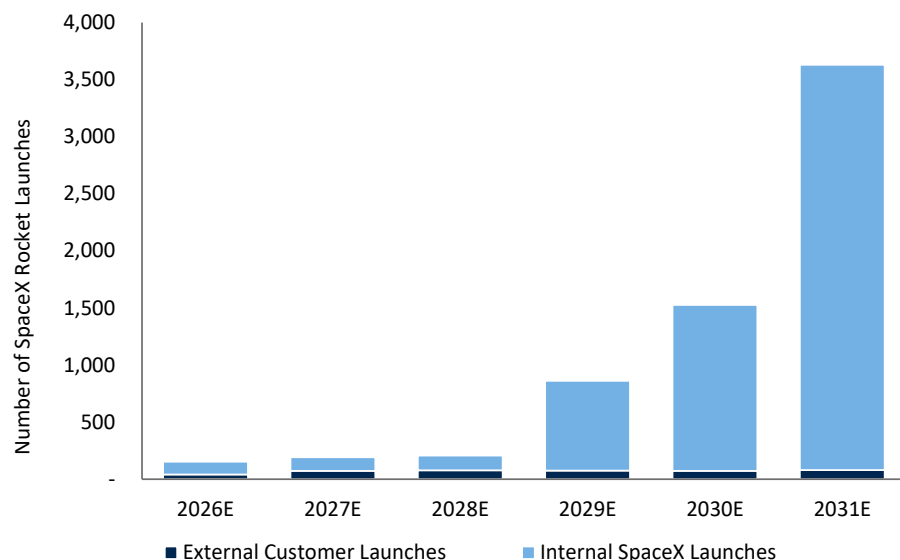
Key Debate #1: Can SpaceX achieve its desired Starship launch cadence?

We believe investors need confidence in Starship's launch cadence to buy-in to the broader Connectivity and AI business growth

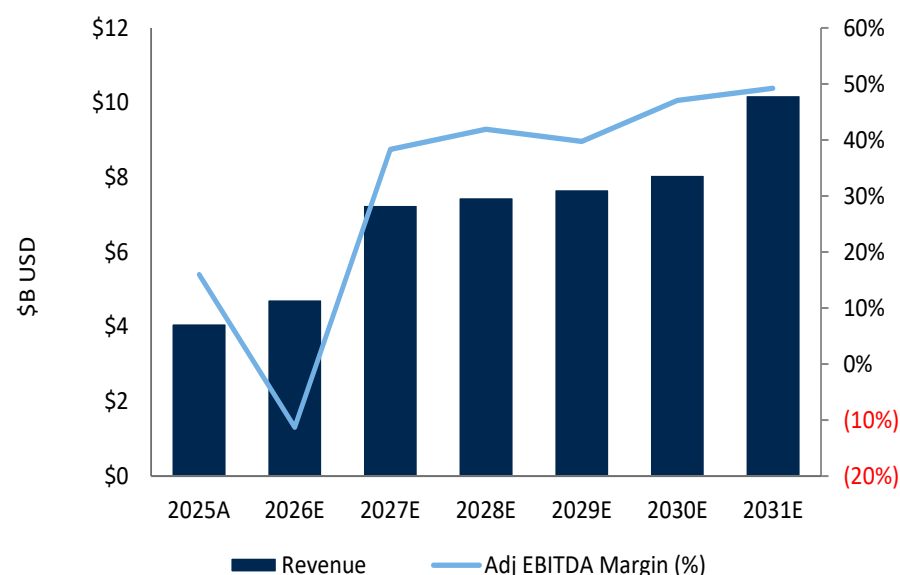
Key Takeaways

- **We believe this debate is the starting point for both the Connectivity and AI segments:** The company's ability to grow both businesses hinges on its ability to increase the launch cadence for Starship. The company's V3 satellites (next version of Starlink) are too big for the current Falcon 9 and Heavy rockets, meaning that beyond the current-gen satellites, Connectivity will only rely on Starship. Moreover, the company needs Starship to establish its orbital compute capacity for the AI business.
- **What launch requirements does our AI framework suggest, and is it realistic?** We believe the company's path to achieving ~70 GW of orbital compute is ambitious, but doable longer term. The framework for Starship launch requirements is set by the tons each Starship can take to space (~100 tons) and the compute density (~100 KW/T) of each satellite. This framework suggests AI to require ~7,000 cumulative launches. Each Starship has the capacity for ~60 V3 Starlink satellites, and modeling 3,000–4,000 new V3s put into orbit each year requires between 50–67 Starship launches annually.
- **Are there near-term catalysts for the segment?** For the Space segment, we do not see significant NT catalysts, instead we view a consistent launch cadence and progress on Starship re-usability as a key driver for sentiment. Again, we view Starship as necessary for success within the other segments and believe investors will demand incremental confidence in Starship into 2027 and 2028.
- **To what extent are investors focused on the Space segment's growth?** Given the significant growth and margin opportunities in Connectivity, and especially within AI, we believe the company's greatest ROIC is to invest internally, utilizing Starship launches for organic growth. Of the almost ~\$910B in company revenue that we forecast for 2031, we model Space to just over ~\$10B. We believe our Space revenue estimates are conservative and see potential upside to numbers with increased lunar launch demand from the U.S. Government, however, we do not view this as a focus for investors.

We model SpaceX to substantially increase launch cadence post-2028E



We estimate just over \$10B of Space Segment revenue in FY2031



Source: S-1, company presentation, RBC Capital Markets estimates

Key Debate #1: Can SpaceX achieve its desired Starship launch cadence?

- In our scenario analysis below, we detailed the moving parts to SpaceX's Starship launch cadence. The company's required cumulative launches between 2028E–2031E to achieve our estimated ~73 GW of orbital compute capacity is based on compute density per satellite of 70-120 KW/T and a launch capacity of 80–200 tons to orbit.
- We currently model that by the end of 2031E, the company will have launched just over 5,900 Starships, as we anticipate the company's compute density to be between 100-110 KW/T with Starship reaching 120 T of launch capacity to orbit.
- As noted in the company's prospectus, SPCX has ambitions to double Starship's orbital capacity to over 200 metric tons on a reusable basis. However, we do not believe investors are likely to give the company credit for such an expansion until the mid-2030s. Instead, we model Starship capacity of 120 metric tons entering 2031E.

Starship launch framework dependent on Compute Density and SPCX's ability to scale Starship Carrying Capacity

Cumulative Starship Launches - Requirement Scenarios 2028E - 2031E							
Satellite Compute Density (KW / T)							
Starship Tons to Orbit (T)	7,284	70	80	90	100	110	120
	80	13,006	11,381	10,116	9,105	8,277	7,587
	90	11,561	10,116	8,992	8,093	7,357	6,744
	100	10,405	9,105	8,093	7,284	6,621	6,070
	110	9,459	8,277	7,357	6,621	6,020	5,518
	120	8,671	7,587	6,744	6,070	5,518	5,058
	130	8,004	7,003	6,225	5,603	5,093	4,669
	140	7,432	6,503	5,781	5,203	4,730	4,335
	150	6,937	6,070	5,395	4,856	4,414	4,046
	160	6,503	5,690	5,058	4,552	4,138	3,794
	170	6,121	5,356	4,761	4,284	3,895	3,570
	180	5,781	5,058	4,496	4,046	3,679	3,372
	190	5,476	4,792	4,259	3,833	3,485	3,195
	200	5,203	4,552	4,046	3,642	3,311	3,035

Source: Company reports, RBC Capital Markets estimates

Key Debate #2: What is a viable growth path for Starlink Mobility?

We believe overspending on terrestrial infrastructure and spectrum will fade the thesis and drive ROIC down

Key Takeaways

- **The business model of “dead zone” coverage is elegant but capped.** Revenue-sharing model, based on partnerships with MNOs, is scalable, capital-light and quick to market. We expect the model of covering “white spaces” to target no more than 5% of Mobility Revenue TAM.
- **The company desires to go beyond wholesale model in the US and own customer relationships directly.** The options include MVNO (Mobile Virtual Network Operator) from any of the Big3 providers or getting access to the terrestrial network organically or inorganically. The company purchased ~50MHz of AWS-4/H-Block spectrum which can be used as a bargaining chip in MVNO negotiation. We consider MVNO a preferred path vs creating 4th wireless network.
- **Solving for indoor coverage in the US.** D2D (Direct-to-Device) largely requires direct line of sight between satellite and consumer device, as the signal quality degrades materially with any obstruction. We believe Starlink needs an access to terrestrial network to serve the urban population and handle the traffic, which is mostly used indoors or during the commute.
- **From wholesale elegance to multi-layer complexity.** We view that operating full-scale mobile operations in the highly competitive US market entails a range of operational challenges, investments, and executive resources. We believe that a potential market share gain from the mobility market will not suffice to justify network and spectrum investments. However, the expansion of IoT, V2X and applications for Unmanned Aerial Vehicles could provide growth opportunities in periods beyond 2028.
- **ROIC is driven by capital-light model.** Material investments in spectrum and terrestrial assets will tie up the capital and dilute rate of return.

Connectivity segment Return profile

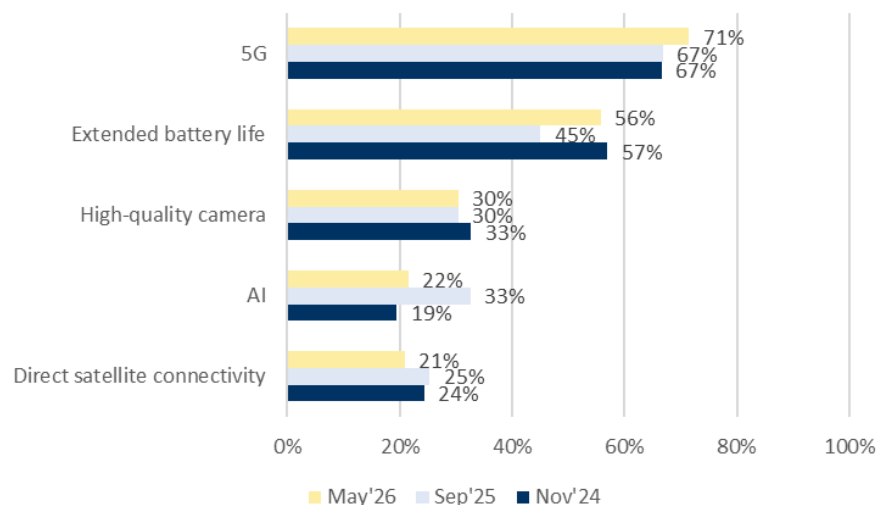
in \$M	FY'25	FY'26E	FY'27E	FY'28E	FY'29E	FY'30E	FY'31E
Total Connectivity Revenue	11,387	18,127	28,836	46,277	69,609	100,940	144,229
Adj EBITDA	7,168	10,911	17,597	28,392	42,463	60,222	84,045
Operating Income	4,423	6,949	12,623	21,246	32,928	49,153	72,094
Tax at 0%							
NOPAT	4,423	6,949	12,623	21,246	32,928	49,153	72,094
Assets							
Gross PPE	18,792	22,970	29,282	51,282	67,882	81,882	94,982
WC	-569	-906	-1,442	-2,314	-3,480	-5,047	-7,211
D&A	-2,376	-3,439	-4,420	-6,543	-8,839	-10,463	-11,373
Capex	4,178	6,312	22,000	16,600	14,000	13,100	12,290
Invested Capital	20,025	24,936	45,420	59,025	69,563	79,472	88,687
Gross ROIC	22%	28%	28%	36%	47%	62%	81%
Capex cumulative	18,792	25,104	47,104	63,704	77,704	90,804	103,094
Cash-on-cash	1%	3%	-20%	7%	24%	40%	58%
Maintenance capex	600	900	1,300	1,950	2,675	3,575	3,262
Unlevered cash yield	35%	44%	56%	52%	59%	69%	85%

Source: Company reports, RBC Capital Markets estimates

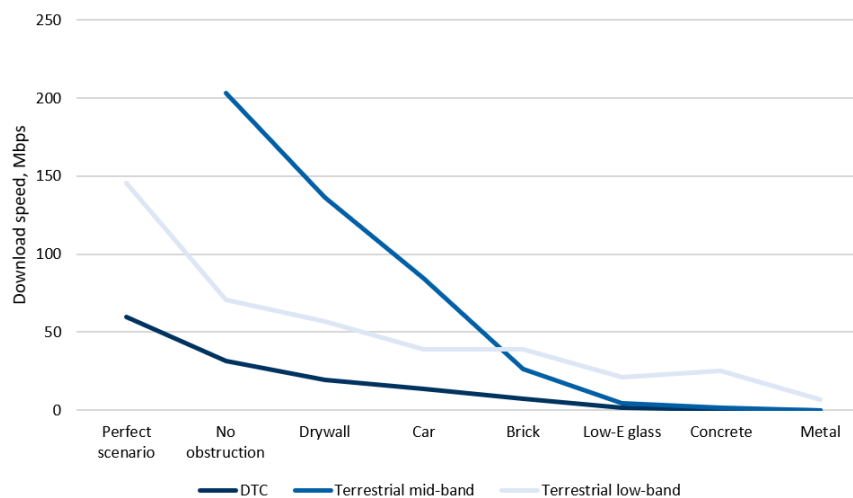
Key Debate #2: What is a viable growth path for Starlink Mobility?

- Our proprietary survey work shows that consumers place less importance on D2D features, while highlighting the importance of 5G network quality and reliability.
- We believe the laws of physics, inherent in the 350-550km distance between the satellite and the consumer device, the signal traversing through the atmosphere, and the lesser amount of spectrum available for Starlink's D2D service vs. traditional telecoms, will translate into an inferior quality of signal in an urban setting and the need for terrestrial augmentation.

RBC Surveys: Important factors (ranked 1 or 2) in choosing a cell phone



Download Speed Under Different Obstruction Scenarios



Key Debate #3: Does Starlink's cost structure allow it to be competitive on pricing?

We believe scaling of consumer and enterprise segments, satellite cost reduction and cash recycling opportunities within the company will allow price levels to support the market share gains

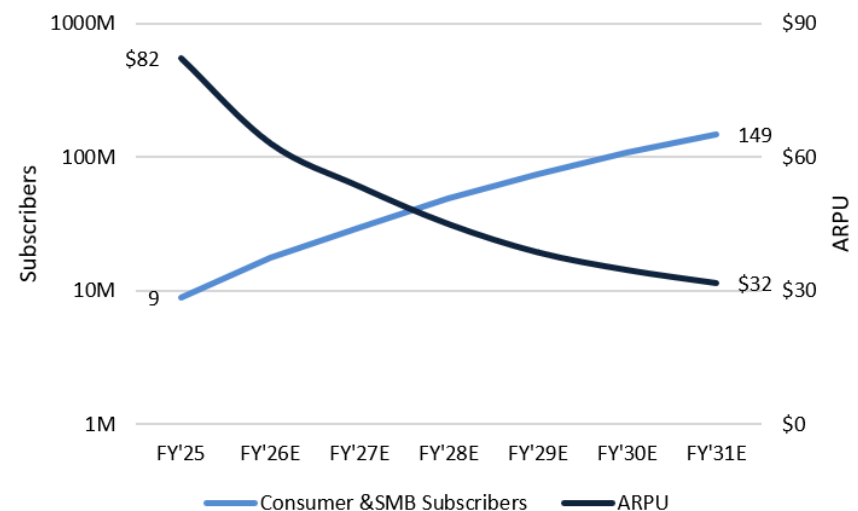
Key Takeaways

- **Consumer broadband model relies on subscriber growth that goes beyond underserved communities.** The number of consumers using Starlink's broadband service nearly doubled in 2025 vs. 2024, which came off the back of growing constellation capacity, geographical expansion and the company's ability to offer aggressive pricing in line with terrestrial offers and consumer-friendly equipment subsidies. That allowed the company not only to fill the white spaces in underserved communities and replace GEO providers, but to encroach on home turf of fiber/cable providers.
- **Market share path for a successful broadband challenger.** For a broadband consumer, two dominant factors are speed/quality and pricing. On speed and latency, Starlink today is comparable to low- and mid-tier offers from terrestrial providers, and we expect the launch of V3 satellites to improve the quality of offers further. On pricing, the company employs a segmented approach based on available capacity and the competitive landscape, with monthly cost on par or below terrestrial peers.
- **An operator with lowest cost-to-serve has an upper hand in competition.** Once deployed, cable and fiber can serve consumers for decades, with little ongoing capex and opex maintenance. For Starlink, the useful life of satellites of 5 years creates a recurring replacement cycle that creates a burden for the company's economic gross margins. We expect the scaling of broadband consumer business and material contribution from the Government & Enterprise unit, which inherit asynchronous traffic patterns to consumer business, to narrow the gap vs. terrestrial providers. We expect the decreasing cost of V3 launch and CPE equipment to further improve Starlink's economic gross margin close to terrestrial peers.
- **Internal cash recycling boosts the company's ROIC.** With the rise of the AI segment, the company's ROIC benefits from the multiplicative effect from connectivity cash flow reinvested in AI. We believe SpaceX's model allows the Connectivity business to aim for lower profitability vs. terrestrial peers and be selectively disruptive with its broadband offerings.

Illustrative cost structure - Starlink's fixed business vs. fiber

	Steady-State FTTH	Starlink Broadband 2031E
Network operations	2%	2%
Backhaul	3%	3%
Maintenance opex	3%	2%
Gross Margin	92%	93%
Maintenance capex	2%	7%
CPE	2%	2%
Economic Gross Margin	88%	84%
Opex		
Sales and Marketing	10%	11%
SG&A	10%	8%
Customer care	5%	4%
EBITDA	63%	61%

Consumer Broadband PxQ model



Source: Company reports, RBC Capital Markets estimates

Key Debate #4: Can SPCX build a structural cost of compute advantage?

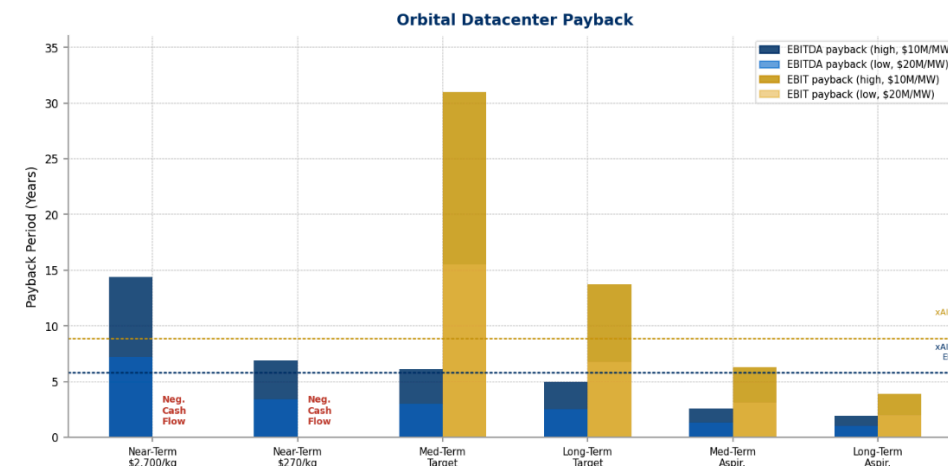
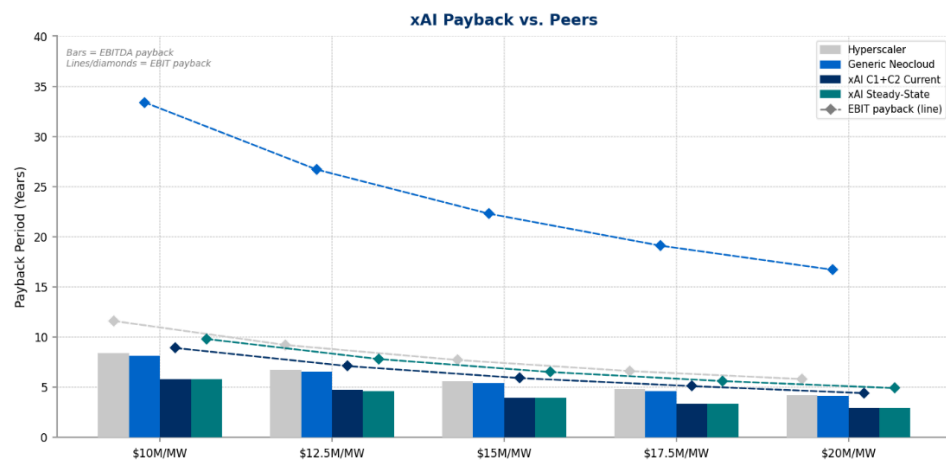
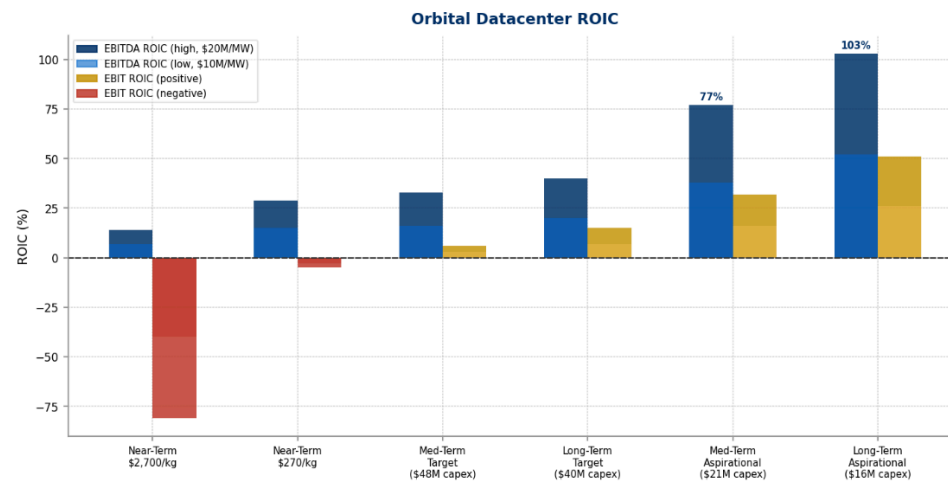
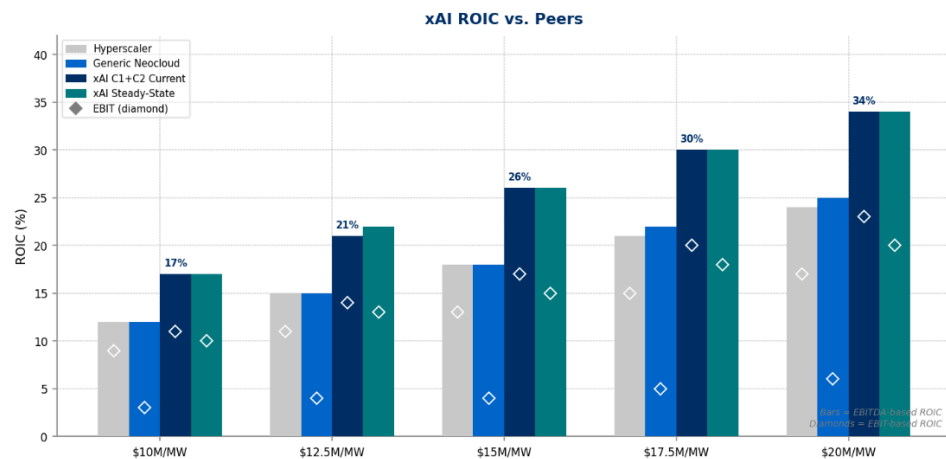
Key Takeaways

- With terrestrial compute, we believe the answer is potentially - or at least one that rivals market leaders. With orbital compute, it seems more likely. The company's execution on Colossus I and II displayed clear outperformance relative to the industry in terms of development cost, though ongoing operating costs could converge with market levels over time. We do not believe Colossus I and II development costs are totally instructive as the company acknowledged that certain factors around access to land, buildings, power and other development costs were unusually efficient and therefore, potentially not as replicable. With that said, we'd expect leading pricing & access to compute components and with eventual Starship-driven cost/kg declines along with an eventual shift towards custom components, we could see a path to SPCX gaining a meaningful cost advantage where access to space appears a strong moat. See below for our estimated potential for SPCX's terrestrial and orbital ROIC advantage vs. the industry.

	Terrestrial Datacenters				Orbital Datacenters					
	Hyperscaler	Generic Neocloud	xAI Current	xAI Steady-state	Near-Term 1	Near-Term 2	Medium-Term	Long Term	Medium-Term	Long Term
					36kg/kW - \$2,700/kg - \$600/kg	36kg/kW - \$270/kg - \$600/kg	15kg/kW - \$270/kg - \$400/kg	10kg/kW - \$100/kg - \$300/kg	Aspirational	Aspirational
Revs/MW	\$20.0	\$11.0	\$20.0	\$20.0	\$12.5	\$12.5	\$12.5	\$12.5	\$12.5	\$12.5
Compute COGS/MW	\$6.2	\$5.5	\$4.8	\$5.8	\$5.6	\$5.6	\$5.6	\$5.6	\$5.6	\$5.6
Non-compute Opex/MW	\$3.2	\$3.1	\$2.7	\$2.8	\$30.3	\$7.5	\$4.2	\$2.5	\$1.9	\$1.0
Total COGS/MW	\$9.4	\$8.6	\$7.5	\$8.6	\$35.9	\$13.1	\$9.8	\$8.1	\$7.5	\$6.6
Gross Profit/MW	\$10.7	\$2.4	\$12.5	\$11.4	-ve	-ve	\$2.7	\$4.4	\$5.0	\$5.9
Gross Margin (after D&A)	53%	22%	62%	57%	-ve	-ve	21%	35%	40%	47%
SG&A/MW	\$2.0	\$0.7	\$1.2	\$1.2	\$0.75	\$0.75	\$0.75	\$0.75	\$0.75	\$0.75
% of revenue	10%	7%	6%	6%	6%	6%	6%	6%	6%	6%
D&A/MW	\$3.3	\$5.1	\$5.8	\$7.0	\$28.7	\$10.5	\$7.9	\$6.4	\$6.0	\$5.2
EBITDA/MW	\$11.9	\$6.8	\$17.1	\$17.2	\$4.6	\$9.1	\$9.8	\$10.1	\$10.2	\$10.4
EBITDA Margin	60%	62%	86%	86%	37%	73%	78%	81%	82%	83%
EBIT/MW	\$8.7	\$1.6	\$11.3	\$10.2	-ve	-ve	\$1.9	\$3.7	\$4.3	\$5.2
EBIT Margin	43%	15%	56%	51%	-ve	-ve	15%	29%	34%	41%
Capex/MW	\$50.0	\$47.0	\$40.0	\$40.0	\$142.2	\$51.2	\$38.1	\$32.0	\$17.1	\$12.8
Implied ROIC (EBITDA)	24%	14%	43%	43%	3%	18%	26%	32%	60%	81%
Implied ROIC (EBIT)	17%	4%	28%	26%	-ve	-ve	5%	11%	25%	40%
Payback period - Yrs (EBITDA)	4.2x	6.9x	2.3x	2.3x	31.1x	5.6x	3.9x	3.2x	1.7x	1.2x
Payback period - Yrs (EBIT)	5.8x	28.6x	3.5x	3.9x	-ve	-ve	19.8x	8.7x	4.0x	2.5x

Note: \$m. Source: Company reports, RBC Capital Markets estimates

Terrestrial ROIC could converge, but Colossus I/II's lower cost was a strong start by SPCX on ROIC



Source: Company reports, RBC Capital Markets estimates. All figures shown at \$50M/MW capex. Bar ranges reflect \$10M/MW (high payback / low ROIC) to \$20M/MW (low payback / high ROIC) revenue scenarios. EBIT payback for Orbital Near-Term scenarios reflects negative cash flow.

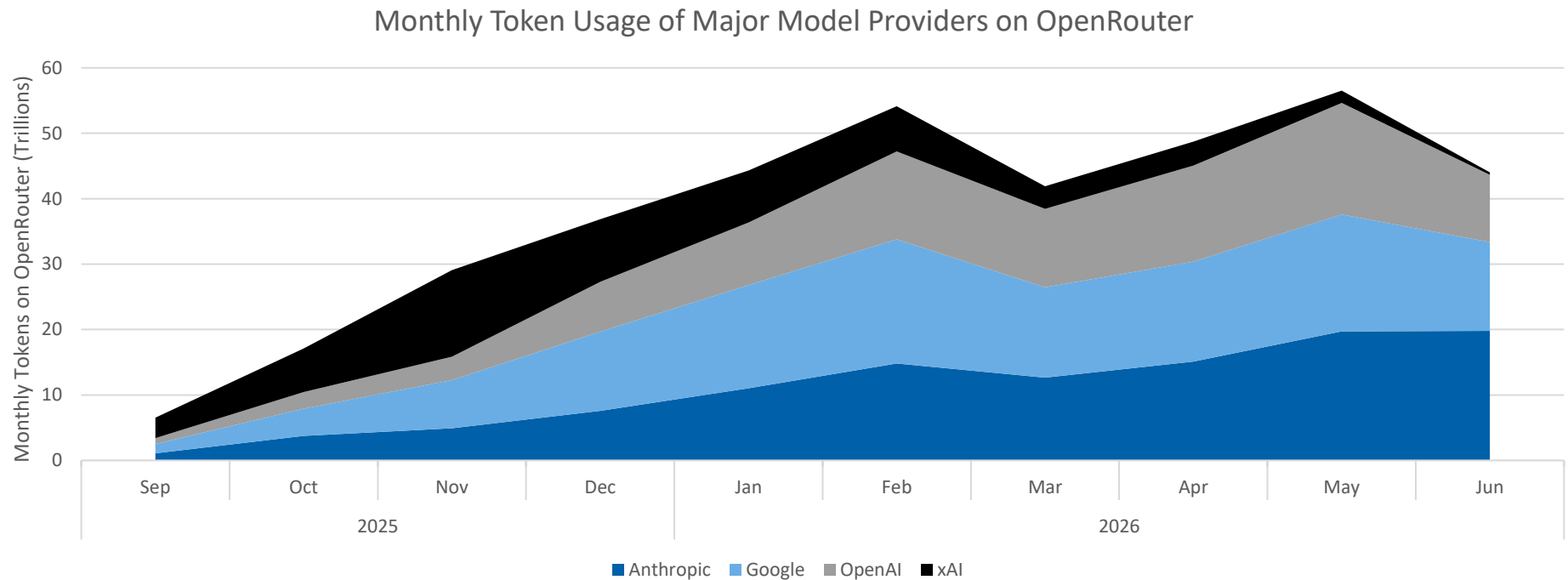
Key Debate #5: How much of a risk is the rise of cheaper models, open-sourced models or locally run models for the hyperscalers?

Key Takeaways

- We believe we're in a moment of transition from API consumption to manufacturing intelligence. We'd look at the emergence of cheaper open-weight, open-source or even locally run models as the next phase of growth for the hyperscalers as opposed to a risk. We'd view the aggressive capacity expansion as actually democratizing organizations' ability to scale from experimental chat and single-purpose agents to 24/7 autonomous agentic loops where the value to a company or a person rises exponentially.

- **Token Cost Is Proving Very Elastic.** When open-weight models cut token costs by 90%+, consumption trends indicate token usage can surge by thousands of percent in some cases; cheaper compute doesn't compress AI budgets, it expands the surface area of what organizations can automate, accelerating the shift toward always-on agentic workloads.
- **Model Tiers Enable Cost-Optimized Scaling.** Mid-tier models (GPT 4.1, Gemini 2.5 Flash) can absorb ~60–80% of enterprise workloads, while frontier models (Claude 4.5, GPT 5.5) remain reserved for high-stakes strategic and architectural tasks where reliability justifies the premium but can actually decrease the cost per token (which has been durable – labs often raise price where completion cost winds up being lower). We won't say forever, but we believe it is way too soon to think about approaching maturity for workflow automation discovery without the benefit of the frontier.
- **Deployment Architecture Determines Competitive Moat.** As third-party routing providers proliferate, we believe enterprises must ring-fence deployments via cloud partners or owned infrastructure to protect IP and avoid vendor dependency; those who control compute layers gain model-agnostic flexibility and durable positioning regardless of which model wins.
- **Infrastructure Ownership Signals Long-Term Conviction.** Locally run and open-weight models introduce real maintenance tradeoffs, reinforcing that serious enterprise AI at scale consolidates around managed cloud infrastructure, not away from it.

Key Debate #5 (continued): Token elasticity is proving explosive – even in the case of rising prices for the frontier



Provider	Demand Δ	Median Price Δ	Elasticity
Anthropic	1,711%	67%	25.7
OpenAI	1,033%	127%	8.1
Google	860%	275%	3.1
xAI	-88%	-75%	1.2

- Importantly, frontier advancements like fast mode optimization, batch pricing, prompt caching and other capabilities typically raise token output by more than price, allowing for cheaper tokens & higher demand (and frequently confirmed by both private labs).
- The market grew 7x in nine months. Token volume went from 6.5T to 44T per month.
- Anthropic captured the most growth as usage grew 1,711% even as median prices rose 67%, reflecting that users are prioritizing capability over cost.
- OpenRouter customers comprise developers, SMBs, and some enterprise usage. Importantly, +50% of traffic comes from outside the United States; the data presented above may be impacted by customer use of DeepSeek or other international models.

Source: OpenRouter, RBC Capital Markets, RBC Elements

Key Debate #6: Can cost of compute, Grok and Cursor enable SPCX to become a leading AI intelligence provider?

Key Takeaways

- The physical aspect of providing the compute for AI is the moat – bottlenecks around power and compute hardware could be persistent for several years where hyperscalers would likely have leverage, even potentially longer term as the gatekeeper.
- Large software vendors building their own frontier models and/or running compute locally will always need extra scale beyond what they own and frontier model updates, which makes hyperscalers foundational to vendor's products or platform, in our view.
- A software with deep system-of-record level integration will likely always be model agnostic, but potentially also tied into a hyperscaler optimized for that platform, making switching costs high. Never saying never, but multi-cloud hurdles for AI exist with hyperscalers able to leverage hardware-optimized inference, training's heaviness and critical workflow context/memory.

Success Factor	Why It Matters	Grok's Current Position
Own the Compute	Physical bottlenecks (power, chips) persist for years; controls pricing & access	Colossus cluster (owned); no hyperscaler margin paid
Distribution at Scale	Active user surfaces = default AI entry point; switching costs compound via context/memory	X platform (~117M Grok MAU as of 1Q26); Tesla vehicles (Grok integrated Feb 2026); Cursor IDE (pending \$60B acquisition)
System-of-Record Integration	Deep workflow embedding makes model-agnostic posture irrelevant in practice	Available on Azure, Oracle, AWS Bedrock, GCP (model hosting/API access) but no native integration into enterprise SaaS workflows announced. Cursor acquisition (\$60B, pending 3Q26) would embed Grok as default model in a widely-used dev tool; "Grok for Word" MS Office plugins launch (June 2026)
Proprietary Data Flywheel	Real-world production data compounds model quality over time	X real-time social data (native access); Tesla FSD (10B+ miles); SpaceX/Tesla engineering & manufacturing workflows used for Grok training
Hardware-Optimized Inference	Custom silicon creates cost/performance gap vs. general-purpose	Colossus optimized for Grok workloads; no custom chip announced; runs on Nvidia GPUs
Regulatory Positioning	Frontier models = dual-use; deployment requires regulatory trust	On U.S. classified networks; DOD uses Grok Gov for active national security ops; pursuing FedRAMP High authorization; faces regulatory probes in 8+ jurisdictions over content safety

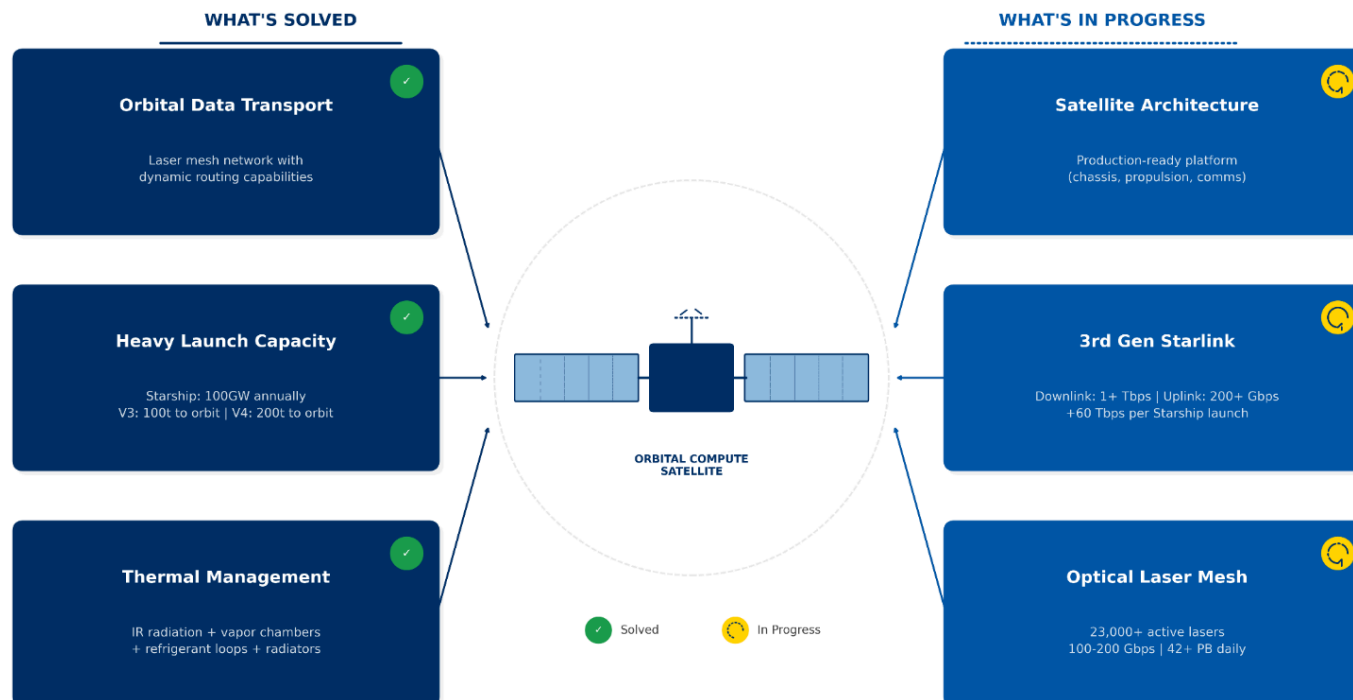
Source: Company reports, PBS, FedRAMP, RBC Capital Markets

Key Debate #7: Can the company execute on orbital datacenters?

Key Takeaways

- We expect investors to view the company's orbital strategy as depending entirely on its ability to ramp Starship launches up to many times/day with reusable rockets, which drives the cost/kg of deploying datacenters in space from its current \$2,700/kg with Falcon 9 versus the longer-term goal of \$100-200/kg with Starship. And while the company does have potential to create at least somewhat of a cost of compute structural advantage in terrestrial datacenters, as shown below, the cost/watt drops significantly to at or even below \$1/watt versus the industry today at around \$2-2.50/watt.

- Execution on orbital depends almost entirely on Starship's ability to achieve reusability, but should that work, the economics & structural cost advantage vs. any terrestrial-based compute could be substantial based on our estimates.



Key Debate #8: Can AI weather regulatory risk without inhibiting market growth?

Key Takeaways

- The Fable 5 episode illustrates that no sustainable path exists to prevent bad actors from accessing near-frontier capabilities as open-source models approach leading-edge performance—suggesting regulatory intervention primarily stunts domestic innovation with diminishing security benefits. With the global economy now tethered to AI infrastructure growth, we believe stopping is not an option. Value creation from unregulated frontier model development will likely continue accruing disproportionately to compute owners with outsized leverage to tax digital enterprises.

The recent Fable 5 controversy has illustrated/prompted a variety of new regulatory risks around model adoption

- U.S. looking to maintain its lead in the AI war: No current path exists to a sustainable defense from bad actor open-source/open-weight usage.
 - Western regulators could be increasingly compelled to stunt innovation around models, chips, etc., while others can obtain lagging versions of frontier models, which will increasingly asymptote closer to leading edge models.
 - Bad actors will be able to leverage nearly best technology if labs are restricted from releasing new models.
 - Government relations is critical – SPCX arguably possesses an advantage vs. other labs for now, in our view.
- Anticompetitive architecture & data enclosure
 - LLM lab data retention policies even for select models exposes commercial IP where lab capture of 3p IP via equity, cash, etc., has been floated as a potential long-term driver.
 - Vertical lab oligopoly allows for perpetual margin squeeze.
 - High likelihood of ring-fence adoption but won't make sense or be possible for a large portion of the market.
- Labor market effects & public equity counter-response
 - How to get AI to align with public interest with meaningful corporate governance implications.
 - Taxes on tokens or compute have been floated as ideas, as has government or even publicly distributed law equity ownership been floated, indicating potential for AI sovereign wealth funds or some type of automated dividend to the public, displaced workers, etc.
- No equilibrium in sight:
 - Global economy tethered to growth in AI infrastructure, we view stopping is simply not an option.
 - **Investment takeaway (for now): value creation from unregulated frontier model development & production likely accrues to the owners of compute with outsized leverage to tax many types of digital enterprises.**

Source: Company reports, TechCrunch, RBC Capital Markets

Valuation Analysis



Capital
Markets

Relative valuation – Company comp sheet

SpaceX currently trades at 18.9x 2028E EBITDA, a discount to A&D peers, and premium to both AI and Connectivity peers

Across all peers, we view GOOGL, AMZN, TSLA, NVDA, META, and TMUS, and RCLB as the closest peers to SPCX. We anticipate most investors to value SpaceX closer in line with AI peers given that the company derives an outsized share of revenue and EBITDA from its AI segment.

07/02/2026												
Company	Ticker	Stock Price	Market Cap	Enterprise Value	EV/EBITDA				EV/Sales			
					2028E	2029E	2030E	2031E	2028E	2029E	2030E	2031E
\$M USD, except share price												
Space Peers												
Boeing Company	BA	\$226.49	\$178,543	\$183,146	19.1x	12.7x	11.0x	N/A	1.7x	1.6x	1.4x	N/A
Firefly Aerospace, Inc.	FLY	\$28.90	\$4,746	\$4,023	81.8x	14.6x	9.6x	N/A	3.7x	2.4x	1.9x	N/A
L3Harris Technologies Inc	LHX	\$302.07	\$56,274	\$77,731	13.1x	N/A	N/A	N/A	2.4x	2.3x	2.1x	N/A
Lockheed Martin Corporation	LMT	\$545.91	\$125,867	\$161,124	11.8x	N/A	N/A	N/A	1.6x	1.6x	1.5x	N/A
Intuitive Machines, Inc. Class A	LUNR	\$19.58	\$3,142	\$3,903	55.3x	32.6x	22.2x	N/A	3.0x	2.4x	2.0x	1.2x
Northrop Grumman Corp.	NOC	\$549.01	\$77,978	\$110,298	13.2x	12.5x	11.9x	N/A	1.8x	1.8x	1.7x	N/A
Rocket Lab Corporation	RKLB	\$100.46	\$60,093	\$37,551	189.8x	100.8x	71.2x	N/A	35.6x	27.3x	21.9x	15.3x
Connectivity Peers												
AST SpaceMobile, Inc. Class A	ASTS	\$85.13	\$25,432	\$24,612	22.4x	11.8x	8.8x	7.8x	13.1x	7.8x	6.0x	4.9x
Charter Communications, Inc. Class A	CHTR	\$137.20	\$16,873	\$126,673	5.0x	5.0x	5.0x	4.9x	2.1x	2.1x	2.1x	2.0x
Comcast Corporation Class A	CMCSA	\$23.79	\$84,759	\$189,254	4.9x	4.8x	4.6x	4.5x	1.4x	1.4x	1.4x	1.4x
EchoStar Corporation Class A	SATS	\$101.50	\$16,084	\$56,661	25.5x	24.6x	24.0x	22.8x	4.4x	4.5x	4.6x	4.7x
T-Mobile US, Inc.	TMUS	\$177.52	\$192,113	\$319,365	7.4x	7.1x	6.7x	6.5x	3.1x	3.0x	2.9x	2.8x
ViaSat, Inc.	VSAT	\$83.06	\$11,343	\$11,318	10.2x	9.7x	10.3x	10.1x	3.3x	2.9x	N/A	N/A
Verizon Communications Inc.	VZ	\$42.56	\$177,712	\$376,617	6.5x	6.3x	6.2x	6.4x	2.5x	2.5x	2.4x	2.4x
AI Peers												
Advanced Micro Devices, Inc.	AMD	\$517.82	\$844,358	\$324,150	32.2x	N/A	N/A	N/A	9.1x	6.5x	5.4x	N/A
Amazon.com, Inc.	AMZN	\$242.67	\$2,610,428	\$2,253,994	8.2x	6.8x	5.7x	4.3x	2.6x	2.3x	2.0x	1.8x
Alphabet Inc. Class A	GOOGL	\$359.91	\$4,039,554	\$3,474,032	12.0x	10.1x	8.5x	6.9x	6.4x	5.5x	4.8x	4.2x
Oracle Corporation	ORCL	\$140.27	\$404,044	\$768,999	5.5x	4.6x	N/A	N/A	3.1x	2.5x	2.2x	1.9x
Meta Platforms Inc Class A	META	\$582.90	\$1,280,075	\$1,444,169	7.2x	6.1x	5.2x	4.4x	4.3x	3.8x	3.3x	2.8x
Microsoft Corporation	MSFT	\$390.49	\$2,900,729	\$2,780,838	9.9x	8.5x	N/A	N/A	6.4x	5.4x	4.6x	3.3x
NVIDIA Corporation	NVDA	\$194.83	\$4,714,886	\$5,007,812	10.2x	N/A	N/A	N/A	6.9x	6.1x	4.7x	N/A
Tesla, Inc.	TSLA	\$393.45	\$1,477,690	\$1,280,097	51.1x	28.5x	21.0x	16.4x	10.0x	6.2x	5.9x	4.5x
Space Peers												
Peer - Mean	-	-	-	-	54.9x	34.7x	25.2x	N/A	7.1x	5.6x	4.7x	8.3x
Peer - Median	-	-	-	-	19.1x	14.6x	11.9x	N/A	2.4x	2.3x	1.9x	8.3x
Connectivity Peers												
Peer - Mean	-	-	-	-	11.7x	9.9x	9.4x	9.0x	4.3x	3.5x	3.2x	3.0x
Peer - Median	-	-	-	-	7.4x	7.1x	6.7x	6.5x	3.1x	2.9x	2.7x	2.6x
AI Peers												
Peer - Mean	-	-	-	-	17.0x	10.8x	10.1x	8.0x	6.1x	4.8x	4.1x	3.1x
Peer - Median	-	-	-	-	10.1x	7.6x	7.1x	5.6x	6.4x	5.5x	4.6x	3.1x
Space Exploration Technologies "SpaceX"	SPCX	\$162.00	\$2,116,368	\$2,092,095	18.9x	10.9x	6.1x	2.7x	13.1x	7.9x	4.8x	2.3x

Source: FactSet, RBC Capital Markets estimates (SPCX only)

RBC price target calculation

We value the company on a 2029E EV/EBITDA multiple, applying a 15.2x multiple to our \$192.5B EBITDA estimate

Each ~0.3x turn on the EBITDA multiple in FY29E would represent another ~\$5 move in the stock

SpaceX Valuation Scenario Analysis											
SM USD expect per share and share count (M)											
	Stock Price	Market Cap	Enterprise Value	EV/EBITDA				EV/SALES			
				2028E	2029E	2030E	2031E	2028E	2029E	2030E	2031E
IPO Price	\$135.00	1,763,640	1,739,367	15.7x	9.0x	5.1x	2.3x	10.9x	6.6x	4.0x	1.9x
Current Price	\$162.00	2,116,368	2,092,095	18.9x	10.9x	6.1x	2.7x	13.1x	7.9x	4.8x	2.3x
	\$211.00	2,756,504	2,732,231	24.6x	14.2x	8.0x	3.6x	17.1x	10.4x	6.2x	3.0x
	\$213.00	2,782,632	2,758,359	24.9x	14.3x	8.1x	3.6x	17.3x	10.5x	6.3x	3.0x
	\$215.00	2,808,760	2,784,487	25.1x	14.5x	8.2x	3.6x	17.5x	10.6x	6.3x	3.1x
	\$217.00	2,834,888	2,810,615	25.4x	14.6x	8.3x	3.7x	17.6x	10.7x	6.4x	3.1x
	\$219.00	2,861,016	2,836,743	25.6x	14.7x	8.3x	3.7x	17.8x	10.7x	6.5x	3.1x
	\$221.00	2,887,144	2,862,871	25.8x	14.9x	8.4x	3.7x	18.0x	10.8x	6.5x	3.1x
	\$223.00	2,913,272	2,888,999	26.1x	15.0x	8.5x	3.8x	18.1x	10.9x	6.6x	3.2x
RBCe Price Target	\$225.00	2,939,400	2,915,127	26.3x	15.1x	8.6x	3.8x	18.3x	11.0x	6.6x	3.2x
	\$227.00	2,965,528	2,941,255	26.5x	15.3x	8.6x	3.8x	18.4x	11.1x	6.7x	3.2x
	\$229.00	2,991,656	2,967,383	26.8x	15.4x	8.7x	3.9x	18.6x	11.2x	6.7x	3.3x
	\$231.00	3,017,784	2,993,511	27.0x	15.6x	8.8x	3.9x	18.8x	11.3x	6.8x	3.3x
	\$233.00	3,043,912	3,019,639	27.2x	15.7x	8.9x	3.9x	18.9x	11.4x	6.9x	3.3x
	\$235.00	3,070,040	3,045,767	27.5x	15.8x	8.9x	4.0x	19.1x	11.5x	6.9x	3.3x
	\$237.00	3,096,168	3,071,895	27.7x	16.0x	9.0x	4.0x	19.3x	11.6x	7.0x	3.4x
	\$239.00	3,122,296	3,098,023	27.9x	16.1x	9.1x	4.0x	19.4x	11.7x	7.0x	3.4x
	\$241.00	3,148,424	3,124,151	28.2x	16.2x	9.2x	4.1x	19.6x	11.8x	7.1x	3.4x
Model Assumptions								Revenue			
Proforma Shares Out	13,064.0			2028E	2029E	2030E	2031E	2028E	2029E	2030E	2031E
Proforma Net Debt	(\$24,273)			\$110,863	\$192,474	\$340,411	\$765,473	\$159,440	\$263,888	\$439,641	\$909,751

RBC price target calculation

RBCe Price Target Calculation			
SPCX	FY'24	FY'25E	FY'29E
Revenue	\$14,015	\$18,674	\$263,888
CAGR %		33.2%	93.9%
Adj. EBITDA	\$5,350	\$6,584	\$192,474
EBITDA Margin %	38.2%	35.3%	72.9%
Incremental EBITDA %		26.5%	75.8%
Adj. EPS	\$0.08	(\$1.13)	\$8.20
CAGR %		-	-
Current Multiple Analysis			
	FY'24	FY'25E	FY'29E
EV/EBITDA		71.5x	
Current Stock Price		\$162.00	\$162.00
Share Count		2926.0	13066.0
Market Cap		474,012	2,116,692
Net Debt		(3,088)	(24,273)
Enterprise Value		470,924	2,092,419
Target			
	FY'24	FY'25E	FY'29E
EV/EBITDA		71.5x	15.2x
Implied Share Price			\$225.03
Target Price Weighting			
EV/EBITDA			100%
Target Share Price			
\$225.00			
Current Price			
\$162.00			
Upside Potential			
38.9%			

Note: Net debt for FY'25 PF as of 1Q26.

Source: Company reports, FactSet, RBC Capital Markets estimates

Sum of the Parts Analysis supportive of \$225 price target

We believe investors are most likely to use a SOTP valuation, with either FY2028E – FY2030E as a base year

- If we assign a 25x multiple to the Space business, a 15x multiple to Connectivity, and a 15x multiple to the AI segment, we arrive at a blended 15.2x multiple for the total company. Subtracting out the proforma net debt and dividing by the shares outstanding, we arrive at a \$225 price target.
- Given the wide differences in segment end-markets, we anticipate investors to use a SOTP framework for valuation.
- We believe FY29E is appropriate, as it demands some near-term execution from the company while not giving SPCX full credit for the FY31E potential.

SpaceX (SPCX) \$M USD	FY29			
	Sales	Adj. EBITDA	EV/EBITDA Multiple	Enterprise Value
Space	7,679	3,054	25.0x	76,355
Connectivity	69,609	41,970	15.0x	629,554
AI	186,600	147,450	15.0x	2,211,750
SPCX	263,888	192,474	15.2x	2,917,658
Net Debt				(24,273)
Minority Interest				-
Equity Value	SOTP Base Valuation			2,941,931
Share Count				13,066
Price/Share				\$225.00

Source: Company reports, FactSet, RBC Capital Markets estimates

Index Inclusion of SPCX and Timeline for Share Unlock

- SPCX is expected to join several major indexes within the next year, which may create the need to mechanically purchase shares from a ~5% free float.
- Index providers have diverged on how to handle mega-cap IPOs, as Nasdaq and FTSE Russell have adopted fast-track rules (15 and 5 trading days, respectively), while the S&P 500 has maintained its 12-month seasoning and profitability requirements.
- This wave arrives amid a challenging macro backdrop with inflation re-accelerating above 4%, a hawkish Fed pricing in tightening, and equity valuations at historical extremes (Buffett Indicator ~235% of GDP).

SPCX Inclusion in Major Indexes			
Index	Requirements	Waiting Period	Earliest expected SPCX entry
S&P 500	GAAP profitability for most recent qtr and most recent 4 qtrs. FMC of >50% of the index's total company-level min market cap threshold	12 months	June 2027
MSCI World, MSCI USA, MSCI AWI	Fast-track IPOs eligible if 80% larger than MSCI's minimum size requirement	10 days	Completed
Russell 1000, Russell 3000, Russell 200	Fast Entry IPOs eligible with FMC in the top 500	5 days	Completed
Nasdaq 100	Fast-entry IPOs ranked in the top 40 and meeting eligibility criteria may be added, subject to 3x float cap where applicable	15 days	July 2026
CRSP	Fast-track IPOs must have >10% float or have FMC above 0.005% of the "index-eligible universe"	5 days	June 2026

Timeline for SPCX Share Unlock	
Event	Unlock timing
Up to 30% of insider shares unlocked	Two trading days after Q2 earnings
7% of insider shares unlocked	~August 2026
7% of insider shares unlocked	~September 2026
7% of insider shares unlocked	~2H September 2026
7% of insider shares unlocked	~October 2026
7% of insider shares unlocked	~2H October 2026
28% of insider shares unlocked	Two trading days after Q3 earnings
7% of insider shares unlocked	~December 2026

- The company employs a staggered lockup structure: the first 20% of insider shares unlock after Q2 earnings (~early-to-mid August), with additional ~7% tranches releasing through the fall, and remaining shares for most insiders becoming available by ~December of 2026.
- Insiders flooding the market with locked shares could expand the current free float and trigger a "great rebalancing" as early backers take profits, according to SEC Chair Gary Gensler.
- Elon Musk holds ~82.4% of the total voting power of outstanding common stock upon completion of the offering (or 82.3% if the underwriters exercise their options to purchase additional shares of Class A common stock in full). Mr. Musk holds ~91.6% of the outstanding Class B common stock.
- Mr. Musk's shares are subject to a 366-day lockup, preventing any sale until June 2027.

Source: Company reports, FactSet, RBC Capital Markets estimates

Comparison to Other Mega Cap IPOs

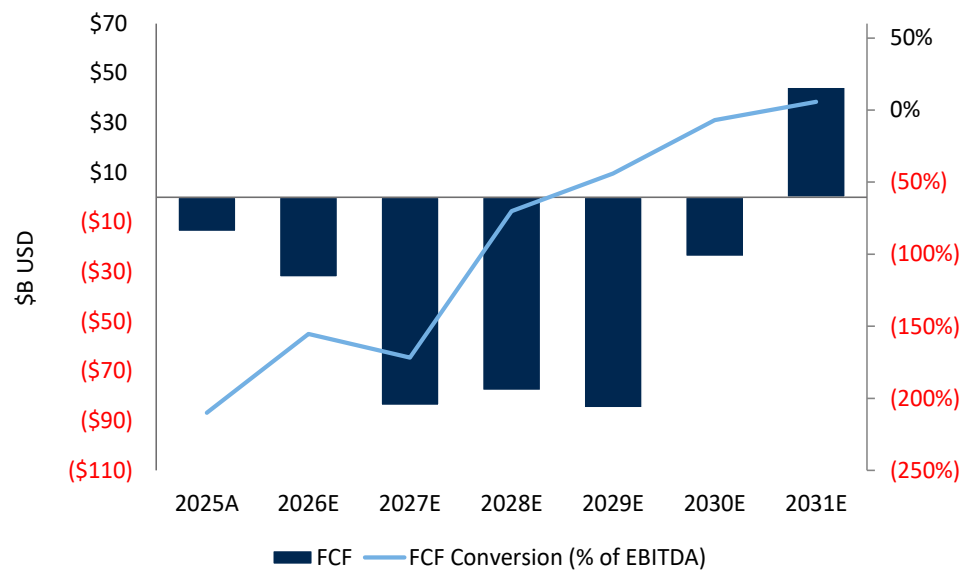
- SpaceX's deal size is much larger relative to other mega cap IPO stocks at \$75B and well above the comp set average of \$27.6B, making it the largest IPO to date.
- At +26% relative to the S&P 500 since its first trading day, SpaceX seems to be tracking positively in its initial weeks of trading, consistent with the majority of comps that posted gains in the first month.
- Historical data suggests that while early positive momentum is common, mega cap stocks' relative returns to the S&P 500 since IPO have a wide range. While names generally seem to have traded positively in the first month, performance diverged materially by month 12 due to variability in that company-specific execution, sector dynamics, and macro conditions.
- SpaceX's +26% relative return since IPO remains encouraging but not yet predictive.

Stock returns for mega cap IPOs in the past decade									
Company Name	Ticker	Date of IPO	Deal Size (\$M USD)	Current Market Cap (\$M USD)	1 Month After Trading	3 Months After Trading	6 Months After Trading	12 Months After Trading	Return Rel To S&P500 Since IPO
SpaceX	SPCX	06/12/2026	\$75,000	\$2,164,647	-	-	-	-	26%
Saudi Arabian Oil Co.	2222-SA	12/11/2019	\$25,599	\$6,321,040	239%	202%	211%	244%	68%
Alibaba Group Holding Ltd	BABA	09/19/2014	\$21,767	\$227,740	29%	61%	42%	(3%)	(232%)
SoftBank Corp.	9434-JP	12/19/2018	\$23,500	\$9,974,330	999%	938%	932%	1,004%	1,302%
Visa, Inc.	V	03/19/2008	\$17,864	\$643,711	57%	92%	71%	20%	2,541%
Meta Platforms, Inc.	META	05/18/2012	\$16,007	\$1,428,117	(21%)	(48%)	(42%)	(31%)	903%
Rivian Automotive, Inc.	RIVN	11/10/2021	\$13,700	\$21,188	47%	(22%)	(47%)	(58%)	(139%)
Average			\$27,634	\$2,968,682	225%	204%	194%	196%	638%

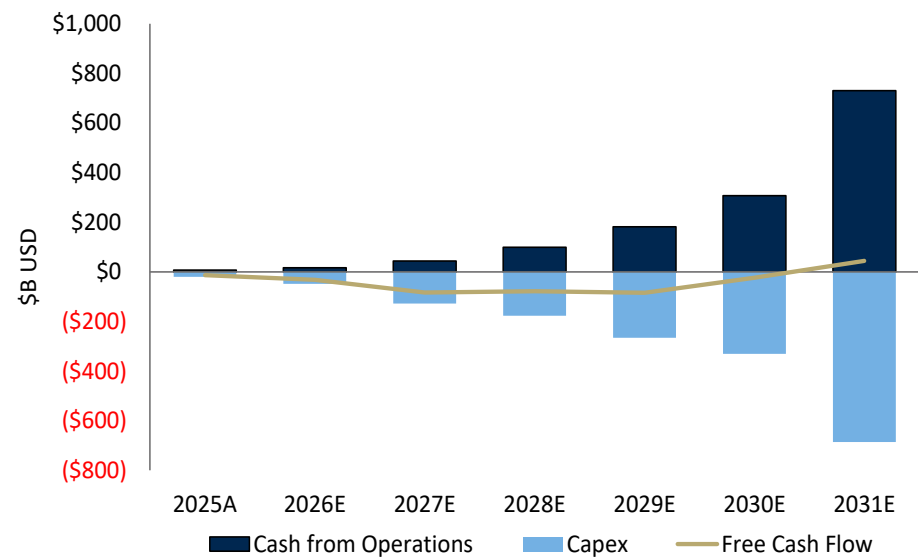
Source: Company reports, FactSet, RBC Capital Markets estimates

Balance Sheet and Cash Flow Summary

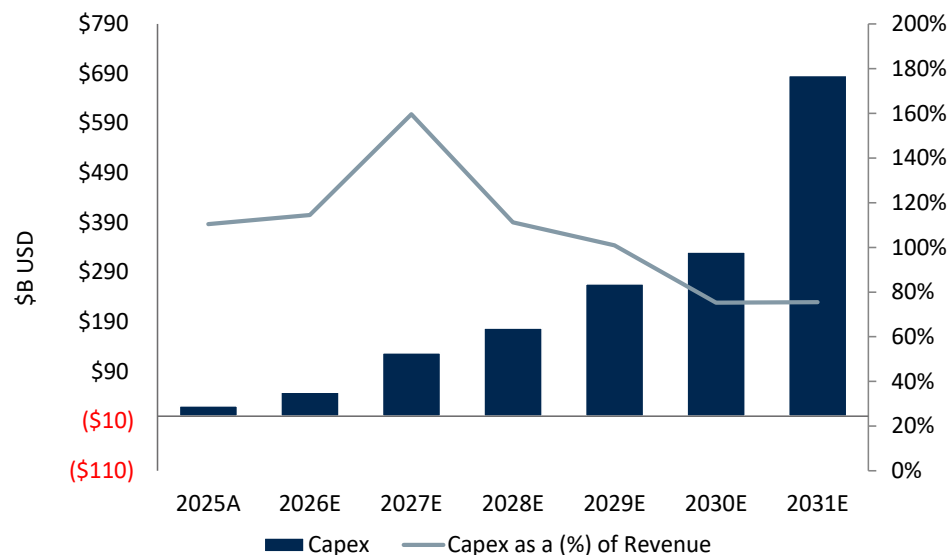
Free Cash Flow (\$B USD)



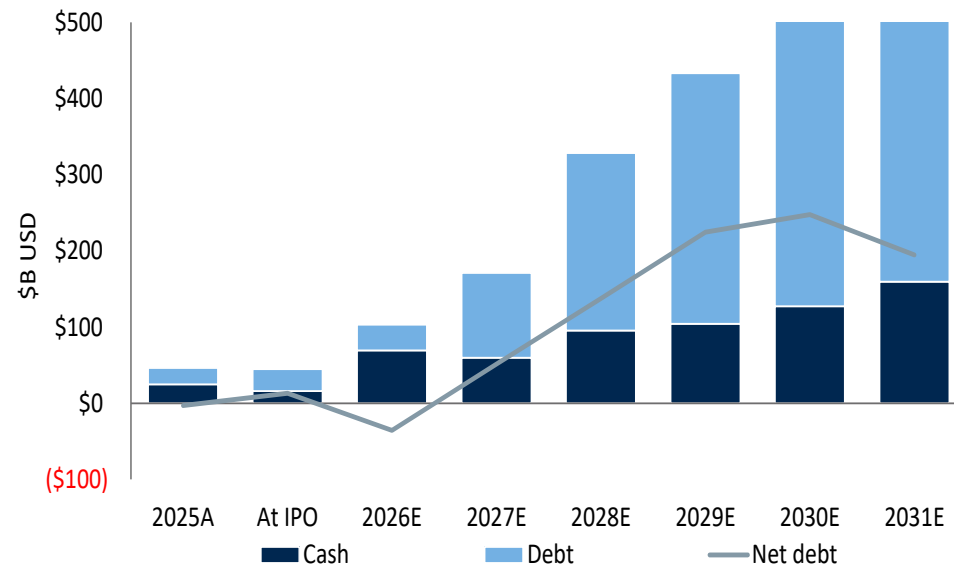
Cash from Operations (\$B USD)



Capex as a (%) of Revenue

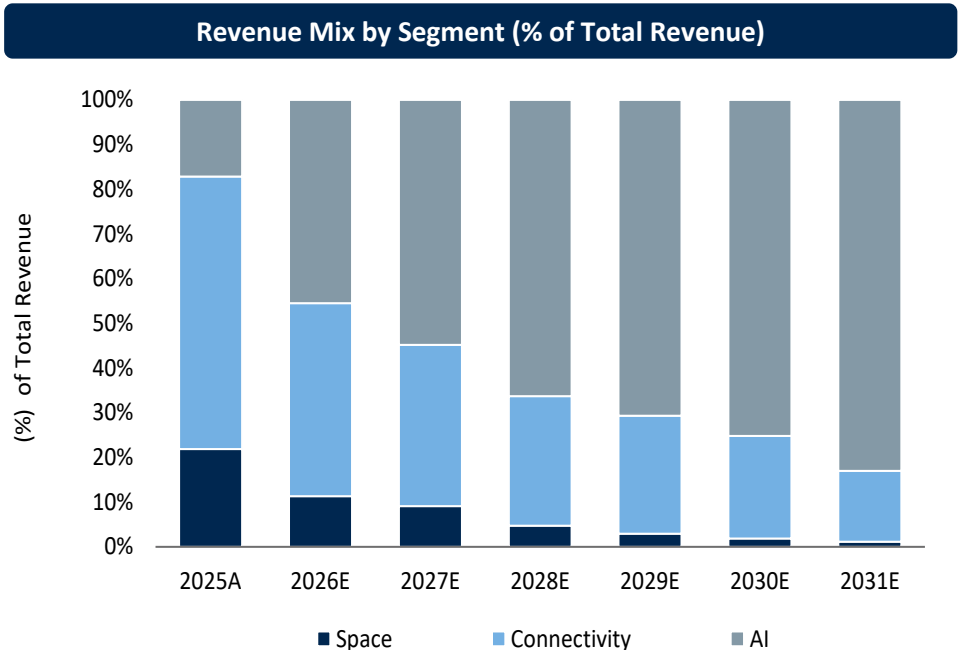
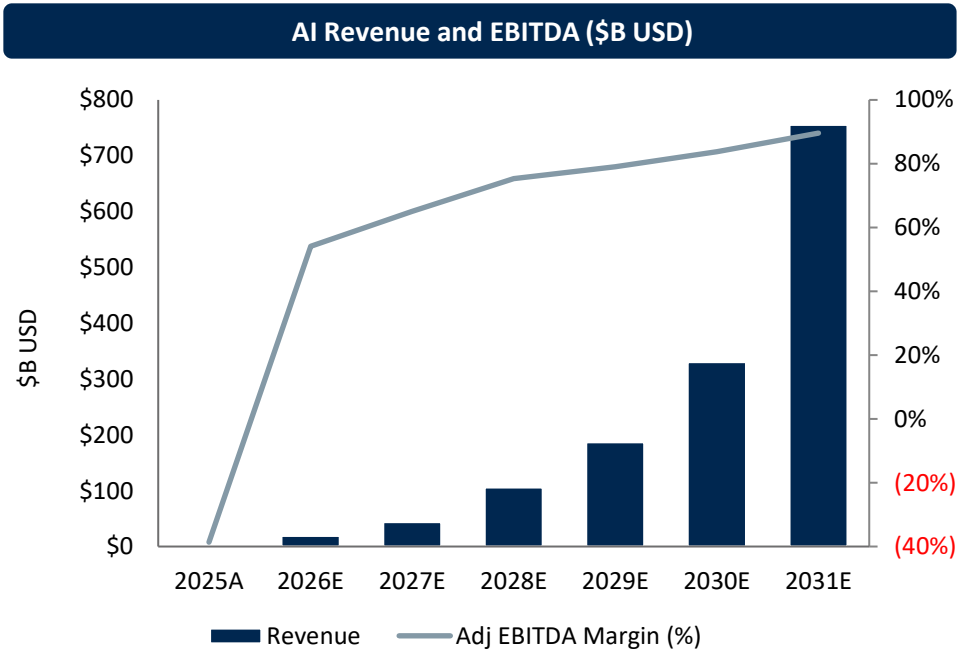
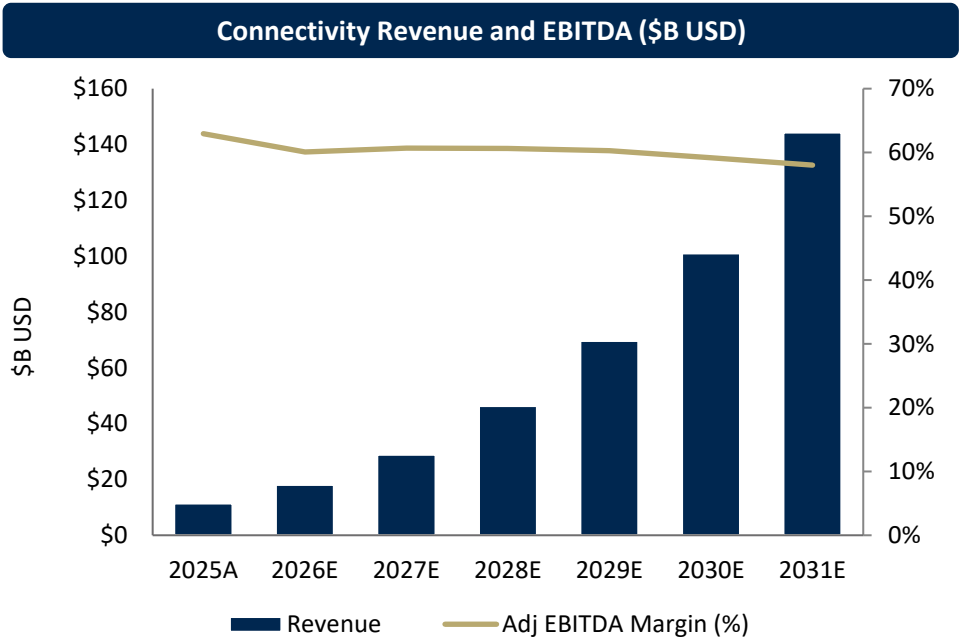
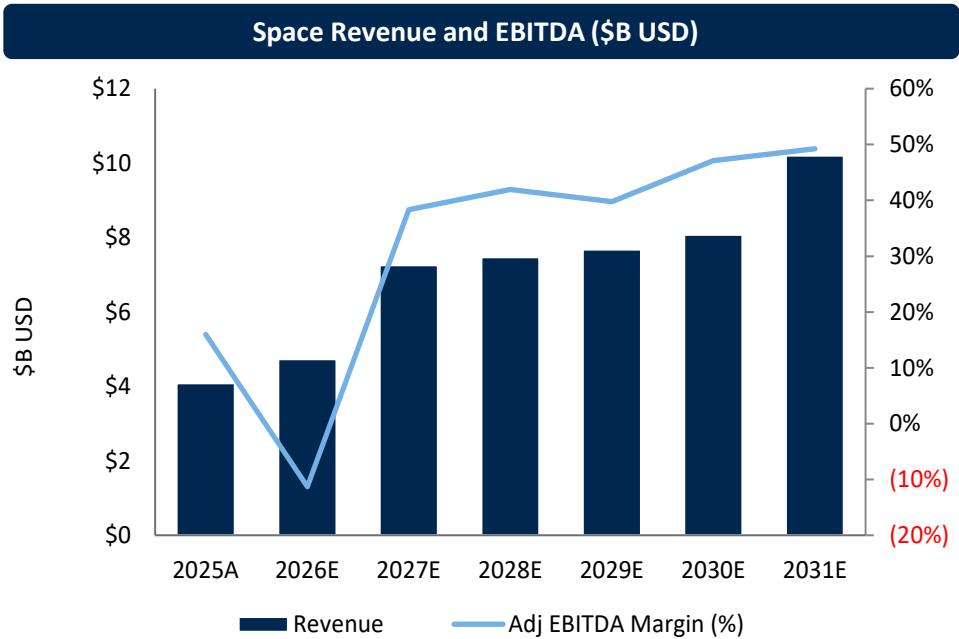


Net Debt (\$B USD)



Source: S-1, company presentation, RBC Capital Markets estimates

SPCX Segment Details



Source: S-1, company presentation, RBC Capital Markets estimates

Management and compensation overview

Management biographies

Name	Title	Highlights
Elon Musk	CEO, CTO, Chairman of Board	- Joined SpaceX in 2002 and serves as CEO, CTO and Chairman of Board since May 2002. - Served as CEO of Tesla since October 2008 and previously as CEO of xAI until 2026. - Founded Neuralink Corp and The Boring Company, and co-founded PayPal, Zip2 Corporation. - Holds a B.A. in Physics from the University of Pennsylvania and a B.S. in Business from Wharton School.
Gwynne Shotwell	President and COO	- Joined SpaceX in 2002 as Vice President, Business Development and as President/ COO in 2008. - Also serves on the Board of Directors of Polaris and on Northwestern University's Board of Trustees. - Previously held positions with Microcosm, Inc. and The Aerospace Corporation. - Holds a B.S. in Mechanical Engineering/ M.S. in Applied Mathematics from Northwestern University.
Bret Johnsen	Chief Financial Officer	- Joined SpaceX in 2011 as Chief Financial Officer. - Has served as CFO at Mindspeed Technologies Inc. and in several leadership roles at Broadcom Inc. - Holds a B.S. in Accounting from the University of Southern California and an M.S. in Finance from San Diego State University.

Management compensation structure

The company emphasizes equity compensation and plans to establish a compensation and nominating committee, which – in consultation with Mr. Musk – will evaluate the compensation of NEOs/ the Board. Each NEO's base salary is a fixed component of compensation, with a portion eligible to be in RSUs. NEOs do not participate in an annual bonus program.

Long-term incentive compensation will be awarded via the 2024 Equity Incentive Plan, which replaced the 2015 Plan for new grants, which provides up to 365,950,000 shares for distribution as stock options, RSUs, and other equity awards, with all outstanding awards remaining in effect post-IPO (converting from Class C to Class A common stock). In January 2026, the board approved the grant of 1B performance-based restricted shares of Class B common stock to Mr. Musk.

Named executive compensation for FY25

Executive	Salary (USD)	Option Awards (USD)	Stock Awards	Other	Total
Elon Musk	\$54,080	NA	-	-	\$54,080
Gwynne Shotwell	\$1,080,127	\$82,969,515	\$1,727,160	\$30,095	\$85,806,897
Bret Johnsen	\$825,000	\$9,013,002	-	-	\$9,838,002

Source: Company reports

Note: FCF conversion based on adj. net income. Source: Company reports, RBC Capital Markets estimates

Space Segment

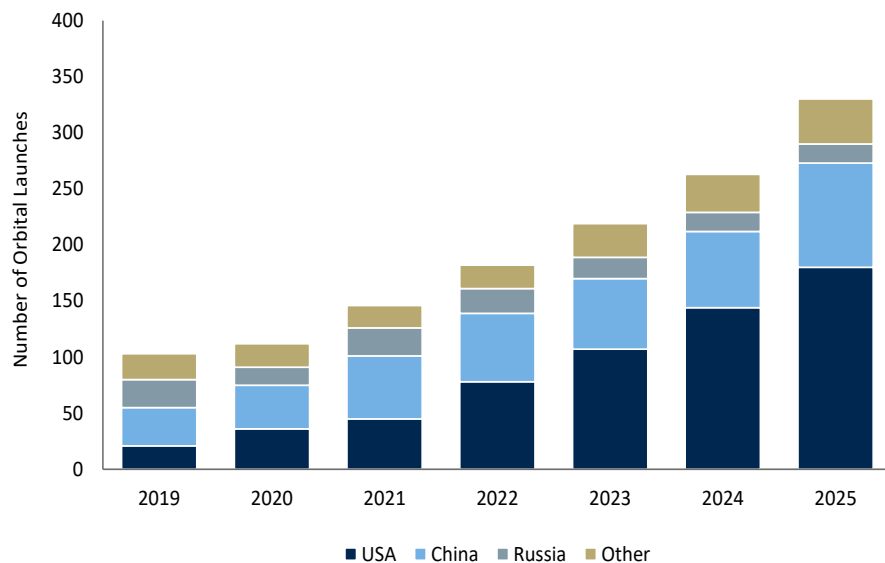


Capital
Markets

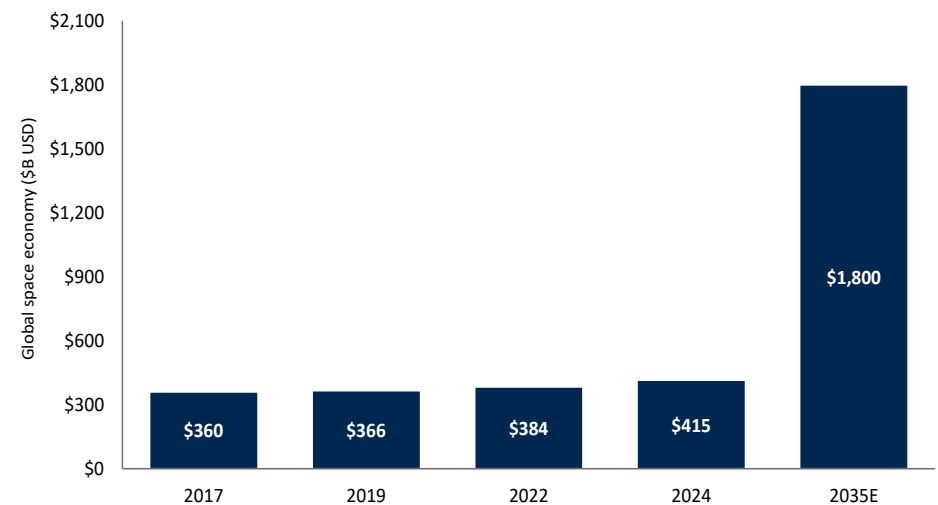
Space: The Global Space Economy

- The global space economy is estimated to reach ~\$1.8T by 2035, up from ~\$415B in 2024, and will be driven by the need for greater connectivity via satellites, higher demand positioning/ navigation services on mobile phones, and increased demand for insights powered by AI/ ML.
- Satellite communications are expected to reach ~\$70 billion in 2035, ~5% of the total space economy.
- Decrease in launch costs: The number of satellites launched per year has grown at a cumulative annual rate of above 50% from 2019 to 2023, while launch costs have fallen 10-fold over the last 20 years – and lower costs enable more launches. The price of data – key to connectivity – is also expected to drop by 10% by 2035, as demand increases by 60%.
- With increased demand for space services, space infrastructure deployment is expected to create:
 - More satellite manufacturing as satellite communications are expected to reach ~\$70 billion in 2035, ~5% of space economy.
 - More orbital launches with ~210 launches annually between 2023 and 2030, driven by greater reusable vehicles.
- Beyond launch, competition is broadening into commercial space stations, national security missions, and orbital AI compute infrastructure, as the space economy is diversifying past traditional launch services.

Global Orbital Launches by Geography



The Size of Global Space Economy



Source: S-1, Spaceflight Tracker, Space Foundation, Bryce Tech, Blue Origin

Space: The Competitive Landscape

- The space economy remains competitive with several key players investing in launch capabilities.

Key Space Rockets Today			
Rocket	Operator	Status	Mass (1000s lb)
Ariane 64	Arianespace	Active	47
Atlas V N22	United Launch Alliance	Active	737
Ceres-2	Galactic Energy	Active	128
Electron	Rocket Lab	Active (reusable)	29
Falcon 9 Block 5	SpaceX	Active (reusable)	1208
Falcon Heavy	SpaceX	Active (reusable)	3126
Firefly Alpha	Firefly Aerospace	Active	119
Gaaganyaan	ISRO	In development	1410
GSLV Mk-2	ISRO	Active	914
GSLV Mk-3	ISRO	Active	1411
H3-30	JAXA	Active	49
HANBIT-TLV Nano	Perigee Aerospace	Active	27
Kairos	Space One	Active	51
Kinetica 2	Orbex	In development	71
Long March 2C	CNSA	Active	514
Long March 2F/G	CNSA	Active	19
Long March 3B/E	CNSA	Active	1005
Long March 7	CNSA	Active	30
Minotaur IV	Northrop Grumman	Active	190
Neutron	Rocket Lab	In development (reusable)	1058
New Glenn	Blue Origin	In development (reusable)	99
New Shepard	Blue Origin	Active (reusable)	165
Pegasus XL	Northrop Grumman	Active	51
Polar Satellite Launch Vehicle	ISRO	Active	705
Proton-M/Blok DM-03	Roscosmos	Active	1554
PSLV-XL	ISRO	Active	8
Smart Dragon 3	CASC	Active	3
Soyuz 2.1b	Roscosmos	Active	688
Soyuz 5	Roscosmos	In development	1168
Space Launch System Block 1	NASA	Active	5750
Spectrum	Phantom Space	In development	88
Starship	SpaceX	In development (reusable)	220
Themis	ArianeGroup	In development (reusable)	331
Tianlong-3	Tianbing Technology	In development (reusable)	37
Vega-C	Arianespace	Active	7
Vikram-1	Agnikul Cosmos	Active	31
Vulcan VC4S	United Launch Alliance	Active	1205

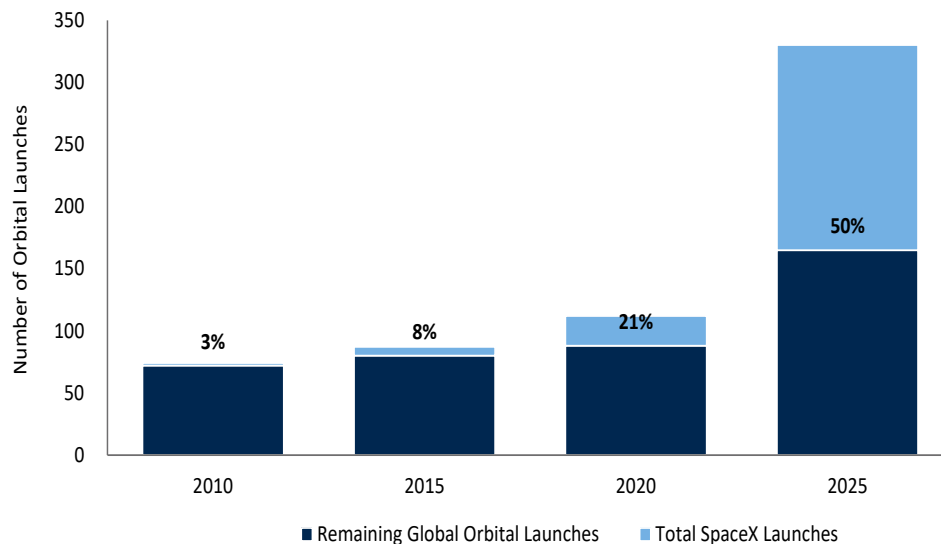
Source: S-1, Spaceflight Tracker, Space Foundation, Bryce Tech, Blue Origin

Space: SpaceX as a leading launch provider

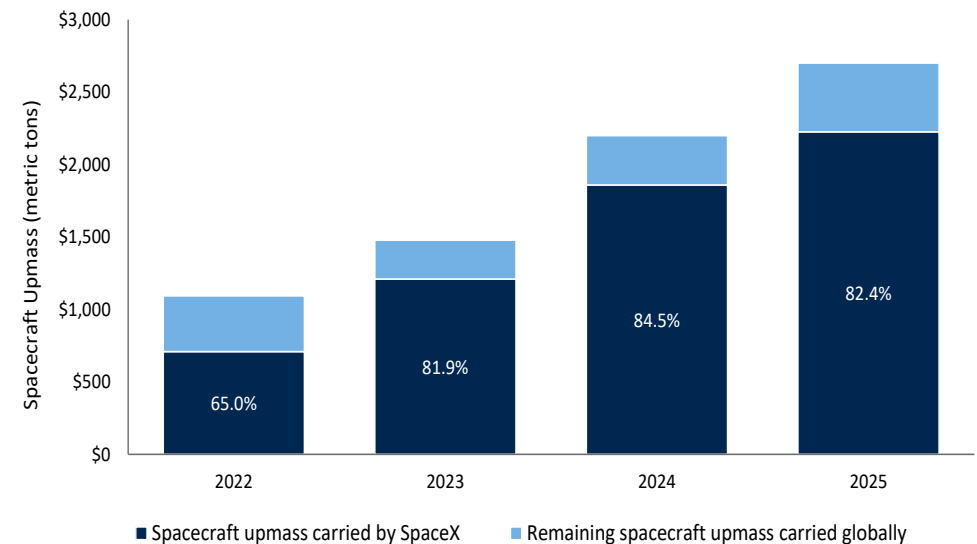
SpaceX has achieved success as a strong disruptor of global space launch capabilities through cost advantage.

- In 2025, SpaceX accounted for ~50% of total space launches globally vs. ~3% in 2010 and has captured 80%+ of global mass to orbit since 2023.
- The first version of SpaceX's Falcon 9 reduced launch costs by ~85% and the first version of Falcon Heavy reduced launch costs by ~92% vs. the historical average.
- Starship, designed to be the first fully reusable spacecraft globally, is expected to reduce cost to orbit by 99%+ and targets a 10x launch cost reduction versus Falcon 9.
- The company expects to scale compute from 2.0 GW in 2026 to 86.8 GW by 2031 through reusable vehicles and an orbital compute platform.

Global and SpaceX Orbital Launch Trends



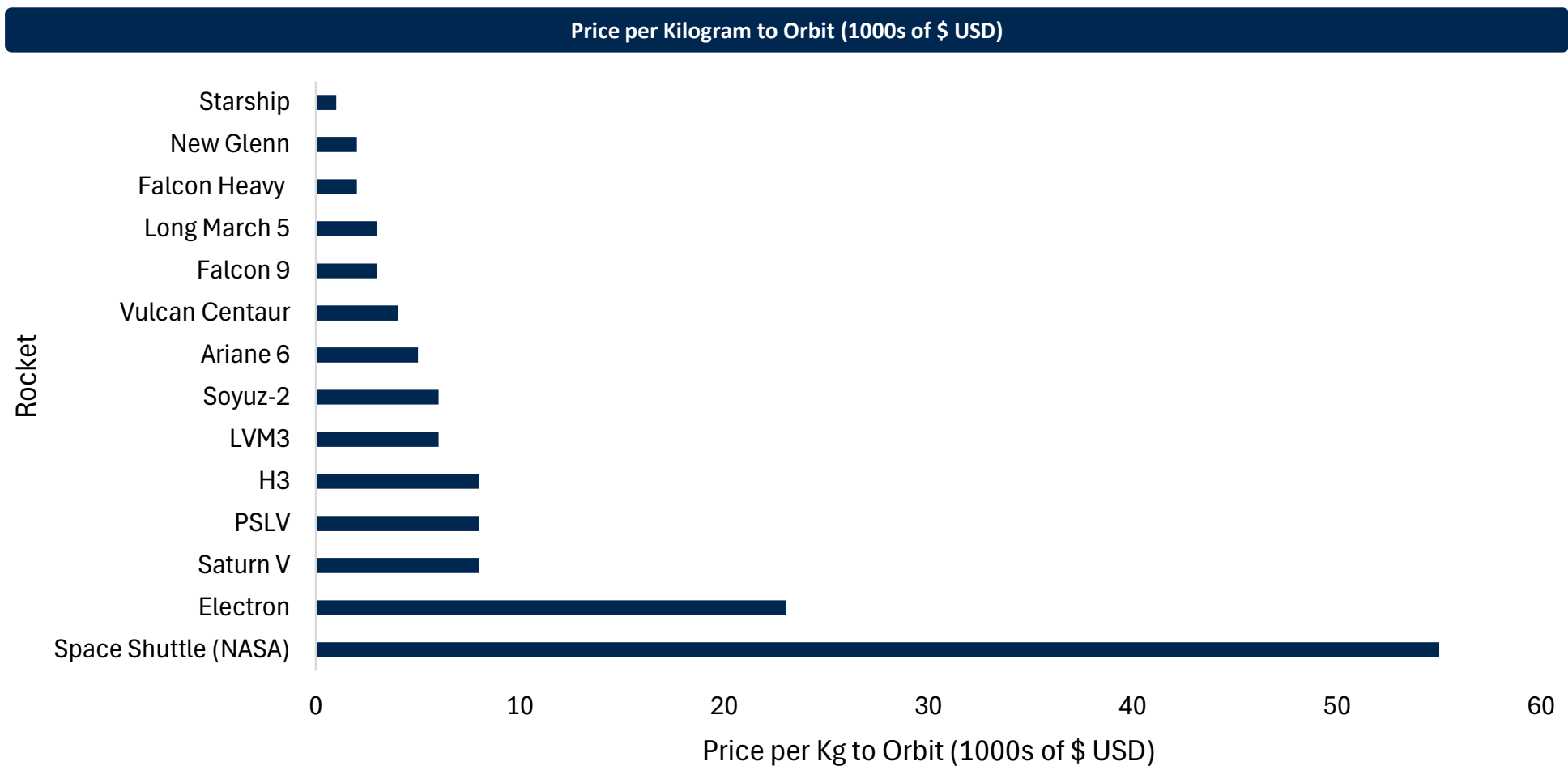
Total Global Mass (metric tons) Launched



Source: S-1, RBC Capital Markets, Gunter's Space Page

Space: Leading launch provider through disruption and cost advantage

- Relative to peers, SpaceX remains a leader in cost to orbit and has a ~10x lower price/kg to orbit vs. peers on average
- Starship is expected to reduce cost to orbit by 99%+ and targets a 10x launch cost reduction versus Falcon 9, due to reusability and vertical integration capabilities



Source: Orbital Radar, company list prices, NASA, FAA, Bryce Tech, CSIS

Space: Launch Vehicle Details

The company operates space launch across three main vehicles: Falcon 9, Falcon Heavy, and Starship

			
	FALCON 9	FALCON HEAVY	STARSHIP
REUSABILITY	Partial	Partial	Full & Rapid Design
HEIGHT	70 m / 229.6 ft	70 m / 229.6 ft	124.4 m / 408 ft
DIAMETER	3.7 m / 12 ft	12.2 m / 39.9 ft (total width)	9 m / 29.5 ft
MASS	594,054 kg / 1,207,920 lbs	1,420,788 kg / 3,125,735 lbs	5,533,000 kg / 12,198,177 lbs
PAYLOAD CAPCITY TO LEO	22,800 kg / 50,265 lbs	63,800 kg / 140,660 lbs	100+ Mt
PAYLOAD CAPACITY TO GTO	8,300 kg / 18,300 lbs	26,700 kg / 58,860 lbs	100+ Mt
PAYLOAD CAPACITY TO MARS	4,020 kg / 8,860 lbs	16,800 kg / 37,040 lbs	100+ Mt
TOTAL NUMBER OF FLIGHTS	~620	11	11
FIRST LAUNCH	June 4, 2010	February 6, 2018	April 20, 2023

Source: S-1

Space: SpaceX and the Emerging Lunar Economy

U.S. Government and lunar infrastructure investments remain key sources of out-year revenue growth for the Space segment.

- Investments in lunar capabilities have gained increased interest with ~\$8.5B of the FY/27 NASA budget allocated to lunar missions under the Artemis program.
- NASA has laid out key goals for its multi-mission Moon Base program, which aims to establish enduring human presence on the lunar surface by leveraging commercial industry capabilities and international partners.
- NASA is sending “a demand signal” to the industry for increased landers, rovers, tech demonstrations, and scientific payloads.
- NASA’s commitment to lunar program spending and public-private partnerships represent significant opportunity for SpaceX.
- SpaceX expects a growing lunar presence to enable terawatt-scale annual AI compute growth, support deeper space exploration, and accelerate the establishment of a civilization on Mars.
- SpaceX's lunar ambitions are tied to Starship, which is expected to commence payload delivery to orbit in 2H of 2026.



Key government contracts with SpaceX

Date	Government Agency	Contract	Value
May-26	US Space Force	Space Data Network	\$2.29B
Jan-26	US Space Force	9 missions under NSSL Phase 3 Lane 1 program	\$739M
Oct-25	Space Force/ NSSL	5 NSSL launches (FY26 awards)	\$714M
Jul-25	Department of Defense	xAI	\$200M
Apr-25	Space Force/ NSSL	5 NSSL Phase 3 Lane 2 missions (FY25 awards)	\$845.8M
Dec-24	Department of Defense	Pentagon expanded Starshield access in Ukraine	+2,500 Starlink terminals in Ukraine.
Oct-24	Department of Defense	pLEO Satellite Internet Program	\$13B (up from \$900M earlier)
Oct-24	Space Force/ NSSL	National Security Space Launch (NSSL) Phase 3 Lane 1	\$733M
Jun-24	NASA	International Space Station US Deorbit Vehicle	\$843M

Source: S-1, Space News, NASA

Connectivity Segment

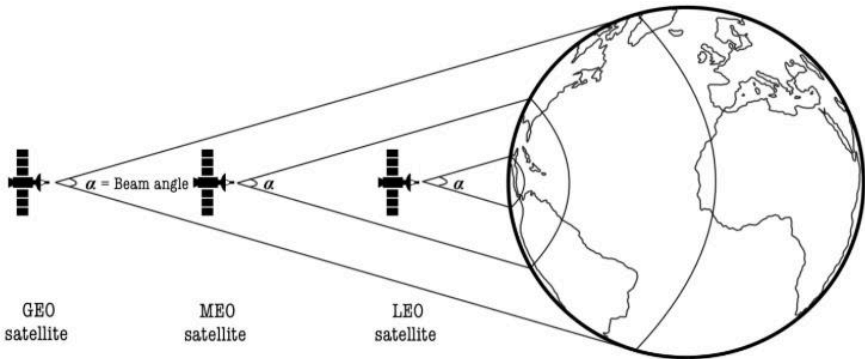


Capital
Markets

Starlink LEO Overview

- Starlink operates in Low-Earth Orbit (LEO) at altitudes of ~350 km to ~600 km, well below the ~35,000 km used by traditional Geostationary (GEO) satellites. This proximity to Earth delivers materially lower latency (~25 ms vs. >500 ms for GEO) and higher throughput of ~225Mbps, but requires thousands of satellites to maintain continuous, ubiquitous coverage. SpaceX currently operates over 10,000 satellites in orbit.
- The FCC has authorized a total of ~19,000 Starlink satellites, with the company targeting ~42,000 in its ultimate constellation state.
- Satellites are deployed across multiple shells at varying altitudes and inclinations, with the 53° inclination shell as the workhorse, covering the most densely populated latitudes between ~60°S and 60°N (~90%+ of the world's population).

LEO vs GEO



Starlink Constellation

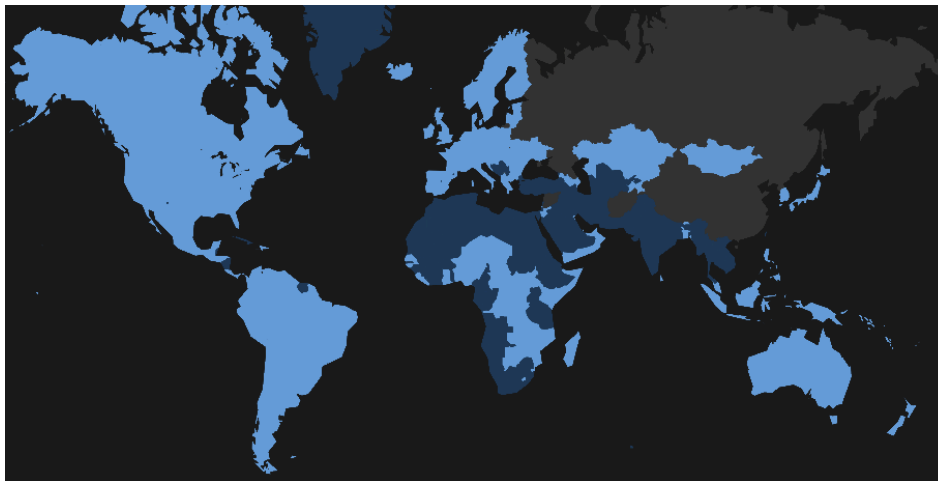
Inclination	Satellites Jun'26	Mix%	Examples of cities served
53°	6,450	61%	Toronto, Paris, London, Berlin, Tokyo
43°	2,800	26%	NYC, LA, Madrid, Rome
33°/38°	700	7%	Miami, Dubai, Delhi, Mexico City
70°	450	4%	Stockholm, Oslo, Helsinki, Edmonton
polar	200	2%	Anchorage, Reykjavik
Total	10,600		

Source: Company reports, industry websites, RBC Capital Markets

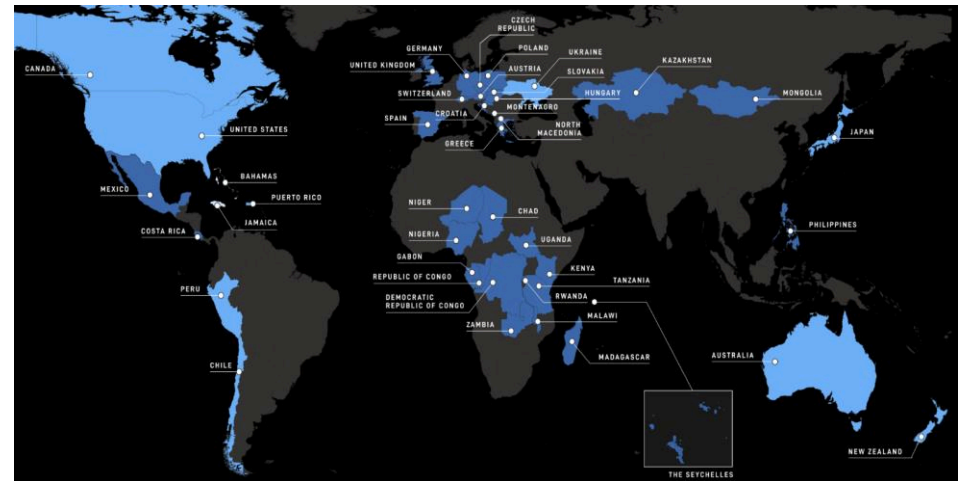
Starlink Services

- Operates across 3 segments: 1) Consumer broadband (63% revenue mix as of 2025); 2) Enterprise & Government (31% of revenue), 3) Mobile (6% of revenue)
- Company's defined TAM: \$870B in Broadband and \$740B in Starlink Mobile
- Services:
 - Broadband connectivity globally, with 225 Mbps median peak-hour speeds at ~25 ms latency; direct line-of-sight to satellite required for high-quality connectivity
 - Resilient internet and IoT services to government and enterprise customers, including aviation, maritime, land mobility, and fixed sites
 - Direct-To-Cell voice, text, and data connectivity without “dead zones” via mass-market mobile mobile handsets. The service is currently positioned as “supplementing terrestrial networks and substantially reducing mobile 'dead zones'” due to propagation limitations indoor

Broadband coverage



Mobile coverage



AVAILABILITY ▾ ■ AVAILABLE ■ COMING SOON

LEO competitive landscape

- SpaceX's Starlink holds a dominant, multi-year structural lead over all LEO competitors. No peer can currently replicate the company's vertically integrated model (owns launch vehicles + satellite manufacturing + terminal production + network operations)
- Amazon Leo is the most credible long-term threat but faces significant execution challenges and remains ~3–5 years behind Starlink in operational scale
- China (Qianfan / "Thousand Sails"): China is building a state-backed LEO constellation to provide global internet coverage, support e-commerce, enable direct-to-mobile connectivity, and strengthen national security — while securing strategic orbital slots and spectrum rights. Satellite internet is "an indispensable last line of defense" for national defense, ensuring communications continuity even if ground infrastructure is compromised
- Europe (IRIS²): The EU is developing IRIS², a sovereign constellation designed to secure communications for European governments and militaries while providing resilient commercial broadband services to businesses and consumers
- Russia (Sfera): Russia is pursuing a sovereign broadband constellation independent of foreign satellite systems, designed to serve military, government, and commercial users across all Russian territory, including the Arctic

State of international LEO constellations

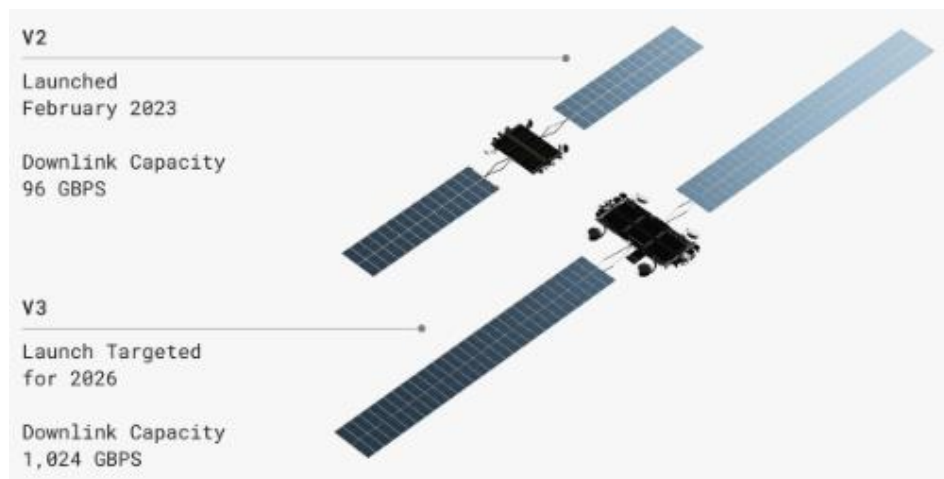
Operator	Country	In Orbit	Satellites	
			Launched	Planned
Starlink	US	10,600	12,241	42,000
Kuiper	US	299	304	3,236
AST Mobile	US	6	7	248
Iridium Next	US	80	80	66
Globalstar	US	24	25	24
US-based combined	US	11,009	12,657	45,574
OneWeb	UK	654	656	716
IRIS	EU	0	0	292
Telesat	Canada	1	1	198
GW (Xingwang)	China	196	202	12,992
G60 (Qianfan)	China	147	162	13,904
Geespace	China	52	64	6,000
Honghu-3	China	0	0	10,000
China-based combined	China	395	428	42,896
Rassvet	Russia	16	16	900

Source: Company reports, industry websites, RBC Capital Markets

Broadband Vertical

- Constellation today:
 - ~10,000-satellite constellation provides ~700Tbs of downlink capacity, sufficient to support 2026YE objectives of ~17M broadband subscribers
 - Ultimate constellation target includes ~42K satellites, including ~19K approved by FCC to date
- Compounding step-up from Starship and V3 Satellites (planned for 2H26):
 - A combination of Starship + V3 achieves 20x data capacity per launch vs Falcon 9 + V2 Mini and 3x+ lower cost per Gbps launched
 - Clear path to add significant capacity in congested areas and increase market share opportunities
- Markets targeted:
 - Initially focused on rural and underserved geographies
 - Improvements in network capacity and product performance allow it to increasingly compete in suburban and urban markets, excluding multi-story buildings (e.g., ~25% of USA HHs)

Satellite evolution



Capacity roadmap

	Satellites		Broadband Capacity, Tbps	
	Jun'26	End State	Jun'26	End State
Satellites in orbit	10,600	42,000		
Gen1	3,318		60	
Gen2 Mini	6,632		637	
Gen3		37,000		37,000
Broadband Total	9,950	37,000	696	37,000
V1 Mobile	650		5	
V2 Mobile		5,000		700
Mobile Total	650	5,000	5	700

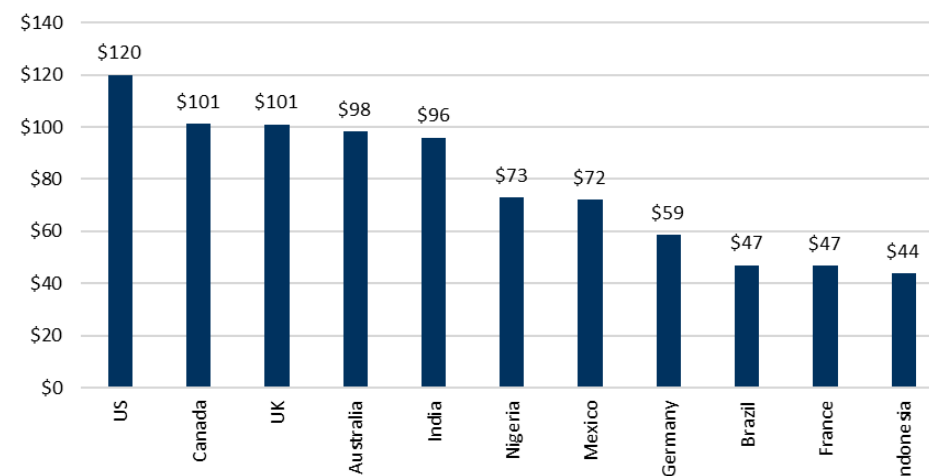
Broadband Business Model

- Consumer broadband – already active across 164 countries. Market share expansion, with PxQ driven by Q, as blended ARPU is set to decrease from growth in developing markets
- Beside filling the white spaces, we expect the company to compete with existing terrestrial and satellite broadband providers
- To aggressively grow the market share, we expect Starlink to price the product at or below the competitors
- Company's cost to serve advantages are conditional on Starship success, further reduction of the kit cost and operational efficiency
- We expect ARPU of \$81 at the end of 2025 to decrease toward ~\$31 by the end of 2031 on a blended basis, as a \$75 average monthly price for broadband in the US is comparable to ~\$30 in Italy and \$4-\$5 in Nigeria and Bangladesh

Broadband pricing - US

Costs			Download, Mbps	Latency, ms
Offers	Monthly plan	Hardware		
LEO				
Residential 100Mbps	\$50	\$0	100	15-35
Residential 200Mbps	\$80	\$0	200	15-35
Residential Standard	\$120	\$0	400	20-40
GEO	\$100	\$275	25-150	600
Fiber	\$70	\$0	1,000	6-8
Cable	\$70	\$0	1,000	16-20
Copper	\$65	\$0	59	23
FWA	\$50	\$0	150-650	28-48

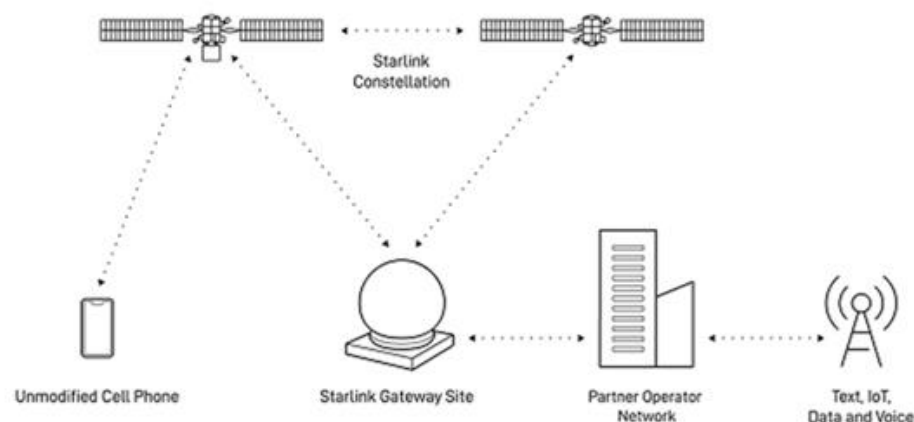
Starlink's International Price Card



Mobile Vertical

- Model today:
 - Backed by 650 V1 Mobile satellites, and utilizing partner MNOs' spectrum, the service aims to eliminate coverage “dead zones”
 - V1 Mobile satellite provides 256 beams x 128 users per beam, currently capable of ~10Mbps speeds
 - Starlink charges MNOs a fixed fee or per-mobile user fee for an extended coverage add-on
- Evolving model:
 - V2 Mobile satellite (deployment is set to begin in 2027) provides 1,024+ beams x 512 users per beam, with speeds reaching 150Mbps
 - A minimum of 1,000 V2 Mobile satellites required for full coverage and 5,000 to maximize throughput
 - The acquisition of EchoStar's spectrum provides less dependency on MNOs and more optionality in the revenue model; Satellite NTN is not yet supported by mass-market phones in 2026; EchoStar's legacy spectrum filings gave Starlink Mobile inherited ITU priority for direct-to-cell, but priority confers coordination rights—not guaranteed access—leaving room for regulatory and incumbent challenges
 - With spectrum ownership and depending on the terrestrial infrastructure (built or acquired) and MVNO access, we think the model could evolve into a mix of stand-alone service in the US and revenue sharing internationally

D2D signal path



Technological roadmap



Source: Company reports, industry websites, RBC Capital Markets

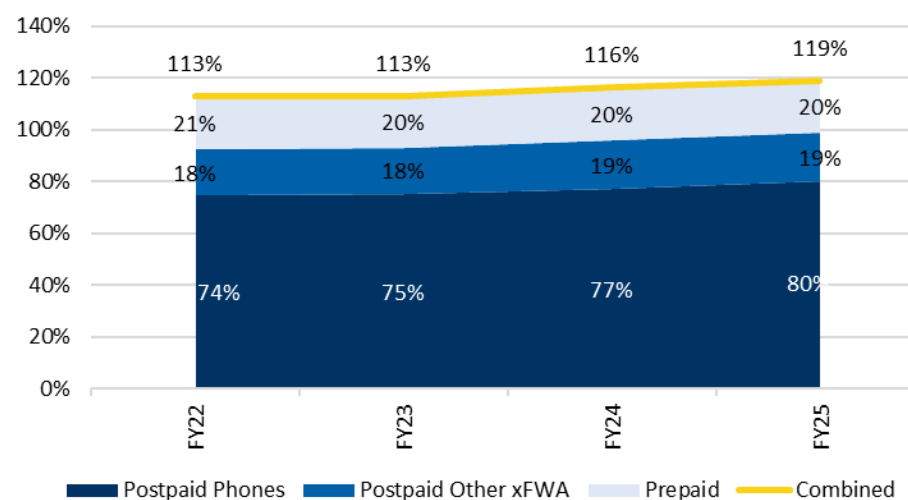
Mobile Business Model

- Backed by 650 V1 Mobile satellites today, and utilizing partner MNOs' spectrum, the service aims to eliminate coverage “dead zones”;
- D2D (Direct-to-device) connectivity – already active across ~30 countries on six continents, covering ~1.9B people.
- Longer-term growth is conditional on the success of Starship, V2 Mobile satellites and has 2 vectors:
 - Expansion of MNO partnerships, where Starlink charges MNOs either a fixed fee or a per-mobile-user
 - Company aims to own customer relationships in the US directly, implying an MVNO via one of Big3 carriers, a terrestrial build, or a combination of the two
 - We view revenue-sharing as an elegant, scalable, high-margin model, while the aspiration of a meaningful market share under stand-alone model requires operational complexity, with handset subsidy, retail footprint, promos and more boots on the ground in a highly competitive market
- Global telecoms currently utilize two monetization models: ~\$10/month for an optional add-on, or a bundled feature across mobile plans to support KPIs and market share gains.
 - Starlink captures 75%+ of the value created, with ARPU of ~\$7 across 7.4M unique mobile devices; however, ARPU is lower in markets where MNOs are charged a fixed fee.

Starlink Mobile pricing

Country	Add-on monthly	Avg ARPU	% of ARPU
US	\$10	\$47	21%
Canada	\$11	\$32	34%
Japan	\$11	\$13	85%
New Zealand	\$12	\$25	48%
UK	\$4	\$17	24%
Australia	bundled	\$22	
Chile	bundled	\$10	
Peru	bundled	\$8	

US Market wireless penetration



Enterprise and Government vertical

- Model today – high-margin revenue streams from:
 - Aviation and Maritime: market share opportunities
 - Government: Sticky Government and Starshield relationships driven by dedicated constellations and secure mission critical capabilities
 - Fixed Enterprise: Opportunity to become first and secondary broadband provider across various SMB, retail, and industrial locations
 - Land Mobility & IoT: In-motion connectivity for fleet vehicles, emergency responders, and industrial IoT (agriculture, energy, logistics) where terrestrial coverage is unreliable; partnerships include John Deere, California Fire Department, and passenger rail operators
- Role in the Connectivity segment:
 - Higher ARPU and lower churn characteristic vs. residential consumer broadband
 - Boost to EBITDA margins/returns from utilization of the same constellation of satellites
 - No incremental traffic consumption burdens due to different traffic pattern vs consumer business
 - Company disclosed ~\$24B of signed revenue for 2026E-2028E
 - Post 2028, we expect the growth rates to decelerate due to initial market saturation, political pushbacks, and fostered competition

Enterprise solutions



Source: Company reports, industry websites, RBC Capital Markets

Financial summary

- Revenue Model:
 - Consumer broadband – already active across 164 countries. Market share expansion, with PxQ driven by Q, as blended ARPU is set to decrease from growth in developing markets.
 - Enterprise & Government - \$24B of signed revenue for 2026E-2028E across government and enterprise solutions, with continued growth projected by the company post 2028E.
 - Mobile – partnered with 30 mobile network operators (MNO) across six continents. Market share expansion via existing and new MNO partnerships as well as direct customer relationships.
- Costs. Cash-generating mechanism:
 - Margin expansion is driven by Kit cost reductions and fixed network leverage on a growing subscriber base.
 - We expect the company to be selective with mobile capex allocation, given the diminishing return on terrestrial investments and complexity of the traditional telecom model.

	Connectivity segment							CAGR
in \$B	2024	2025	2026E	2027E	2028E	2029E	2030E	26E-'31E
Revenue								
Consumer & SMB	7	11	16	23	31	41	52	37%
Enterprise & Government	4	6	11	20	31	44	61	57%
Mobile	1	1	1	3	8	16	31	109%
Total Connectivity	11	18	29	46	70	101	144	51%
Y/Y	50%	59%	59%	60%	50%	45%	43%	
Gross Profit	5	9	15	25	38	57	83	56%
Margin%	48%	50%	53%	54%	55%	56%	57%	
Adj EBITDA	7	11	18	28	42	60	84	50%
Margin%	63%	60%	61%	61%	61%	60%	58%	
Capex	4	6	22	17	14	13	12	14%
% of Rev	37%	35%	76%	36%	20%	13%	9%	
Adj EBITDA - Capex	3	5	-4	12	28	47	72	73%
% of Rev	26%	25%	-15%	25%	41%	47%	50%	

Source: Company reports, RBC Capital Markets estimates

Starlink Constellation

	Jun'26	2026E	2027E	2028E	2029E	2030E	2031E
Satellites	10,600	10,470	12,970	14,320	16,120	17,170	18,560
V2-mini	10,600	10,350	11,350	9,700	7,500	4,550	2,000
V3	0	120	1,620	4,620	8,620	12,620	16,560
Mobile	650	1,000	1,500	1,800	2,600	3,300	4,000
BB capacity per satellite, Gbps	100	110	212	390	581	762	903
V2-mini	100	100	100	100	100	100	100
V3	1,000	1,000	1,000	1,000	1,000	1,000	1,000
Constellation BB Capacity, Gbps	1,060,000	1,155,000	2,755,000	5,590,000	9,370,000	13,075,000	16,760,000
V2-mini	1,060,000	1,035,000	1,135,000	970,000	750,000	455,000	200,000
V3	0	120,000	1,620,000	4,620,000	8,620,000	12,620,000	16,560,000
Avg per satellite	100	110	212	390	581	762	903

Illustrative Model. Coverage of 50km ² radius area at 40° N							
Satellites serviceable	12	12	14	16	18	18	20
Capacity, Gbps	1,200	1,324	2,974	6,246	10,463	13,707	18,060
Theoretical users with 300Mbps	4,000	4,410	9,910	20,820	34,880	45,690	60,200
Oversubscription at 10x	40,000	44,100	99,100	208,200	348,800	456,900	602,000

AI Segment



Capital
Markets

AI thesis: AI models and applications may come and go, but compute lasts forever (maybe)

- We believe the primary long-term determinant of success for the largest AI players will come down to **achieving the lowest cost of generating intelligence (tokens)** as these companies compete for an emerging TAM potentially sized in the many trillions of dollars.
- We believe SPCX is well positioned to be one of the primary leaders/winners in the space driven by several factors:
 - A founder with outsized access to capital: the core pre-requisite for building AI infrastructure.
 - Early indications on SPCX's terrestrial datacenter capacity development indicate efficiency outperformance vs. rest of industry.
 - Faster compute capacity expansion enables premium Neocloud capture which underwrites future growth toward orbital datacenters where structural cost of compute advantage looks more likely, bigger and more durable than terrestrial.
 - Subscription & software products are earlier days as financial contributors. SPCX is the laggard of the 4 major labs today, but we believe possesses all the critical components to emerge more competitively: access to compute to build leading models through a CEO with unparalleled influence for raising capital, attracting talented engineers, building winning products & marketing them to a target-rich audience in the hundreds of millions – for free.
- **Keys to our thesis & debates/risks:**
 - 1) Execution and durability of a structural cost of compute advantage.
 - 2) LLM business model defensibility from emerging & increasingly capable open-source, open-weight, local compute or even proprietary owned models.
 - 3) Can cost of compute, Grok and Cursor enable SPCX to become a leading intelligence provider?
 - 4) Can the company execute on orbital datacenters both in terms of launch as well as the technology?
 - 5) Can AI weather regulatory risk without inhibiting growth?

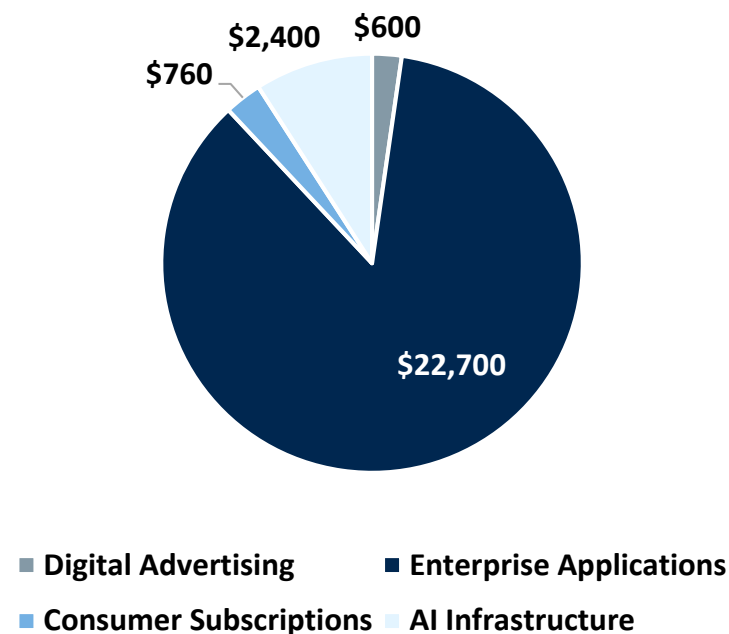
TAM based on SpaceX estimates

- Total AI segment aiming to address \$26.5 trillion (ex. China & Russia)
 - \$22.7T in Enterprise Applications based on 2026 digital economy estimates
 - Implies 48% of knowledge workers ex-Russia and China
 - \$2.4T in AI Infrastructure (see build-up below)
 - \$760B in Consumer Subscriptions is based on the weighted average monthly subscription revenue (\$12) and the global population aged 10+ (implies ~5.2 billion people)
 - \$600B in Digital Advertising based on 2025 3p estimates (we believe this estimate could arguably be 80-100% higher)

AI Infrastructure TAM Build

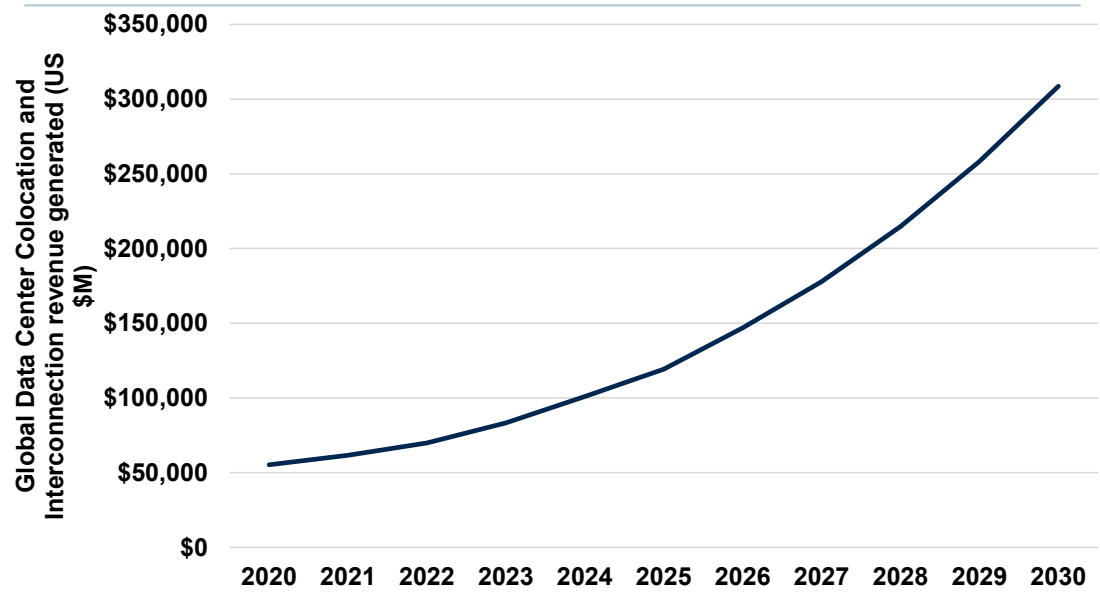
Assumptions	Metric
235	Global data center compute demand (GW)
70%	Utilized for AI workloads
165	AI Demand (GW)
1.2	Target power usage effectiveness
137	Chip level power needed
1.3	All-in chip power consumption per GPU (KW per GPU)
104	GPUs required for AI workload demand
80%	Utilization rate
83	Actual GPUs to provision
\$3.33	GPU rental rate
8,760	Hours per year
\$29,171	Revenue per GPU annually
\$2,427 AI Infrastructure Total Market (B)	

Total AI TAM (\$B)

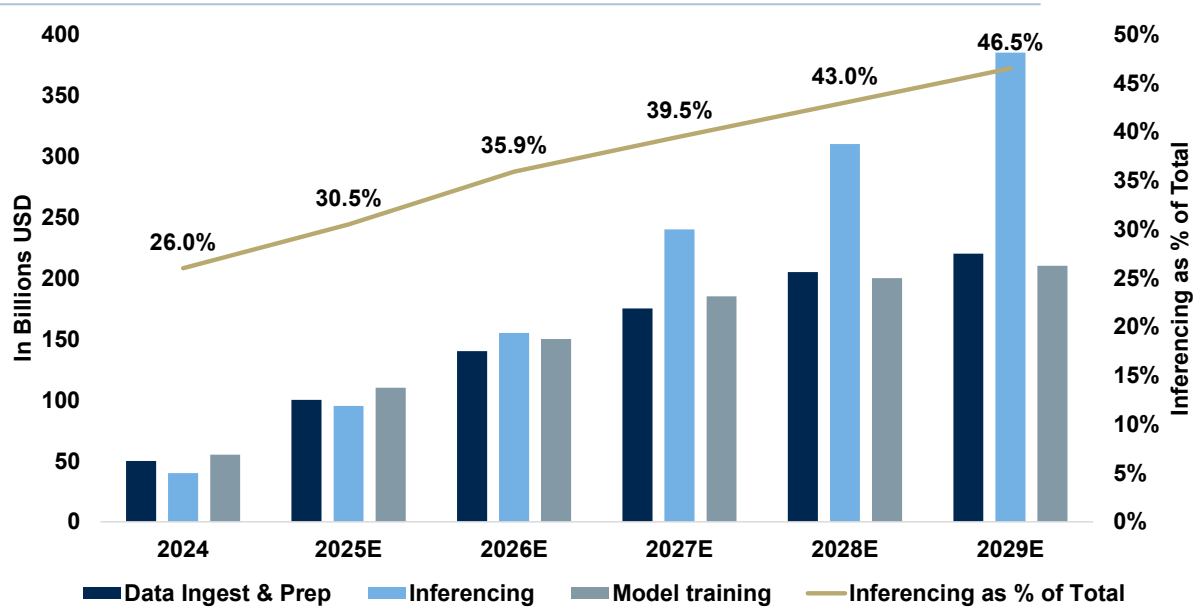


Datacenter demand is surging

Estimated ~\$310B global datacenter colocation and interconnection revenue by 2030



Inference is estimated to reach 46.5% of datacenter revenue by 2029

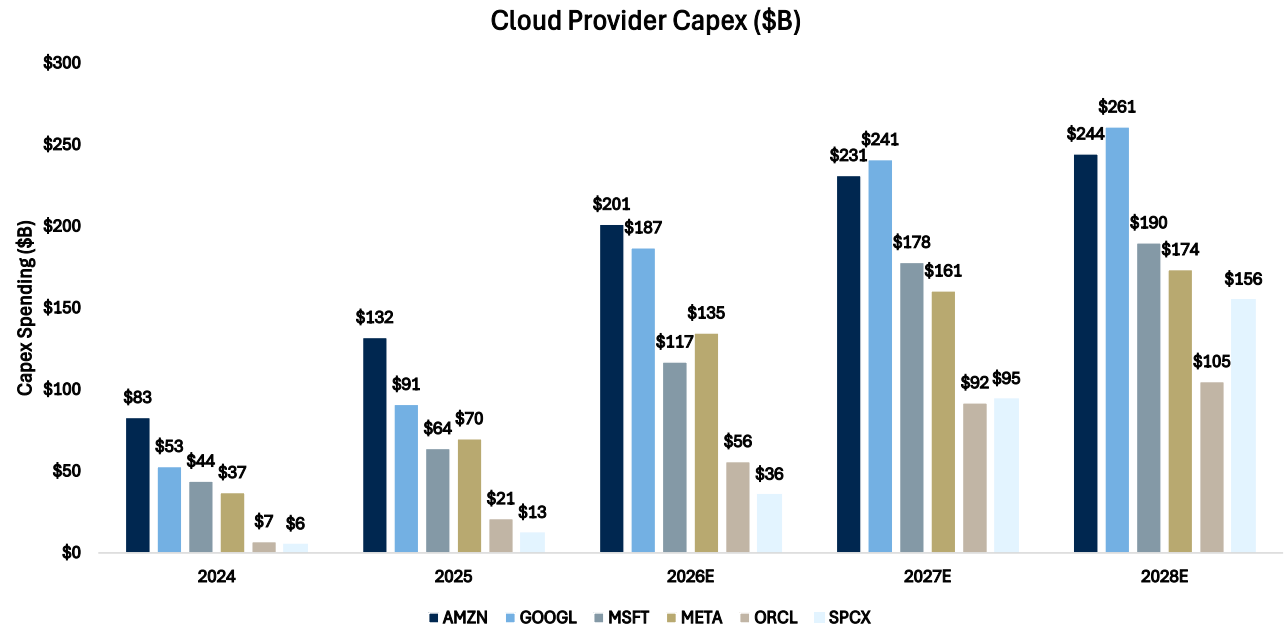


Source: S&P Global Market Intelligence AI Infrastructure Market Monitor & Forecast (2025), Structure Research, RBC Capital Markets

Hyperscaler capex also surging, but with the backlog to support the growth

- In the past 6 quarters, the 4 major hyperscaler backlogs have grown to over \$2 trillion/implied \$500 billion annual revenue (plus SPCX just added ~\$85 billion in the past five weeks) vs. capex of over \$1.1 trillion in cumulative spend in '25 and '26 alone (estimate ~24-month deployment timeline for spend)

Cloud Provider Capex Spending (\$B)				
Ticker	2025	2026E	2027E	2028E
	2025	2026	2027	2028
AMZN	\$132	\$201	\$231	\$244
GOOGL	\$91	\$187	\$241	\$261
MSFT	\$64	\$117	\$178	\$190
META	\$70	\$135	\$161	\$174
ORCL	\$21	\$56	\$92	\$105
SPCX	\$13	\$36	\$95	\$156
Total	\$390.6	\$732.1	\$997.4	\$1,129.5
y/y \$ growth	2025	2026E	2027E	2028E
AMZN	\$49.0	\$69.4	\$29.9	\$13.2
GOOGL	\$38.0	\$96.0	\$54.0	\$20.0
MSFT	\$20.0	\$53.0	\$61.0	\$12.0
META	\$33.0	\$64.7	\$25.8	\$13.0
ORCL	\$14.0	\$35.0	\$36.0	\$13.0
SPCX	\$7.0	\$23.4	\$58.6	\$60.8
Total	\$161.0	\$341.5	\$265.3	\$132.1
y/y % growth	2025	2026E	2027E	2028E
AMZN	59%	53%	15%	6%
GOOGL	72%	105%	29%	8%
MSFT	45%	83%	52%	7%
META	89%	92%	19%	8%
ORCL	200%	167%	64%	14%
SPCX	124%	186%	163%	64%
Total	70%	87%	36%	13%



Hyperscaler backlog growth supports the ramp in investment

\$ in millions								
	AMZN	3Q24	4Q24	1Q25	2Q25	3Q25	4Q25	1Q26
AWS	\$164,000	\$177,000	\$189,000	\$195,000	\$200,000	\$244,000	\$364,000	
Implied 4YR Annual Revenue	\$41,000	\$44,250	\$47,250	\$48,750	\$50,000	\$61,000	\$91,000	
	GOOGL	3Q24	4Q24	1Q25	2Q25	3Q25	4Q25	1Q26
GCP	\$85,165	\$93,200	\$89,831	\$106,000	\$155,000	\$240,000	\$462,000	
Implied 4YR Annual Revenue	\$21,291	\$23,300	\$22,458	\$26,500	\$38,750	\$60,000	\$115,500	
	MSFT	3Q24	4Q24	1Q25	2Q25	3Q25	4Q25	1Q26
RPO	\$259,000	\$298,000	\$315,000	\$368,000	\$392,000	\$625,000	\$627,000	
Implied 4YR Annual Revenue	\$64,750	\$74,500	\$78,750	\$92,000	\$98,000	\$156,250	\$156,750	
	ORCL	3Q24	4Q24	1Q25	2Q25	3Q25	4Q25	1Q26
RPO	\$99,000	\$97,000	\$130,000	\$138,000	\$455,000	\$523,000	\$553,000	\$638,000
Implied 4YR Annual Revenue	\$24,750	\$24,250	\$32,500	\$34,500	\$113,750	\$130,750	\$138,250	\$159,500

*Results reflect calendar year

AI intelligence's TAM is knowledge work: replacement or augmenting? Both?

Estimate \$48T knowledge worker TAM: SPCX's TAM assumes 48% of that is addressable

Metrics (in millions)	Assumptions
Total Labor Force (M)	3,730
Global Knowledge Worforce (M)	1,100
Implied % of Knowledge Workers	29.5%
Avg. US Knowledge Worker Salary	\$117,000
US Salary Multiple	2.7
Global Knowledge Worker Salary	\$43,333
Knowledge Worker TAM (\$B)	\$47,667
% of TAM Automated by AI	60%
Estimated Value Capture (\$B)	\$28,600

		% of Knowledge Work TAM Targetable by AI				
		40%	44%	48%	52%	56%
Global Knowledge Worker TAM (\$B)	\$40,500	\$16,200	\$17,820	\$19,440	\$21,060	\$22,680
	\$43,000	\$17,200	\$18,920	\$20,640	\$22,360	\$24,080
	\$45,500	\$18,200	\$20,020	\$21,840	\$23,660	\$25,480
	\$48,000	\$19,200	\$21,120	\$23,040	\$24,960	\$26,880
	\$50,500	\$20,200	\$22,220	\$24,240	\$26,260	\$28,280
	\$53,000	\$21,200	\$23,320	\$25,440	\$27,560	\$29,680
	\$55,500	\$22,200	\$24,420	\$26,640	\$28,860	\$31,080

Source: Company reports, Statista, FactSet, Glassdoor, International Labor Organization (ILO), RBC Capital Markets

AI as a % of Global IT Spend is expected to reach ~63% by 2030 (~89% excl. AI infra)

\$6.4T estimated Global IT Spend in '26, up 14% y/y

Category	Spend
Overall IT Spend	\$6.4T
Overall Software Spend	\$1.5T
Enterprise Application Software	\$485B
Infrastructure Software	\$574B
AI Platforms and Models	\$64B
IaaS	\$287B
AI Optimized IaaS	\$42B
Traditional IaaS	\$245B

\$2.6T estimated AI Spend in '26, ~41% of IT Spend

AI Spend Category	Spend
Total AI Spend	\$2.6T
AI Infrastructure	\$1.4T
AI Services	\$586B
AI Software	\$453B
AI Cybersecurity	\$51B
AI Models	\$33B
AI Platforms for Data Science & ML	\$30B
AI Application Development Platforms	\$8B
AI Data	\$3B

Global IT Services spend estimated to grow at 26% 5YR CAGR ('25-30)

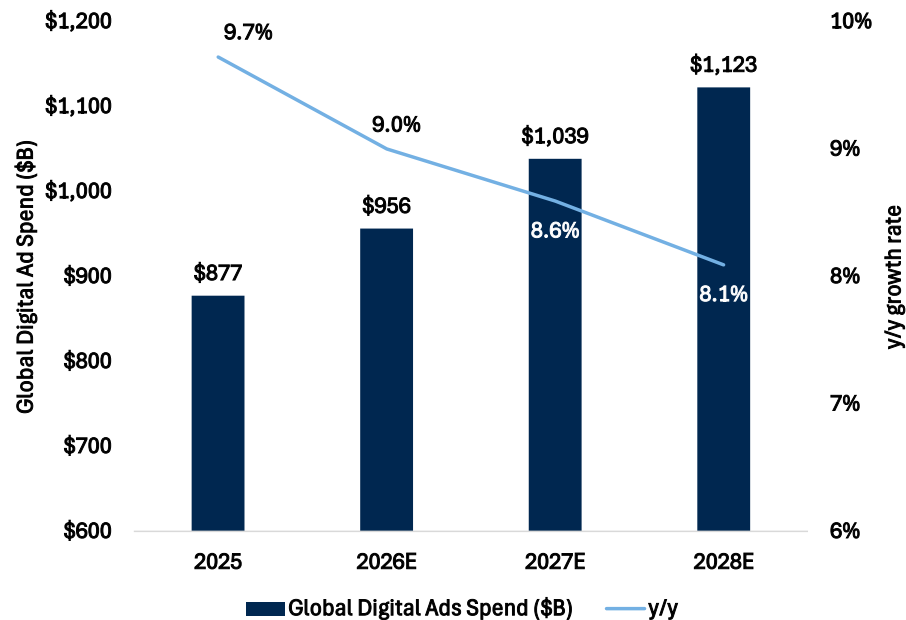
Global IT Spend	2025	2026E	2027E	2028E	2029E	2030E	5YR CAGR
Global IT Spend (\$B)	\$5,577	\$6,369	\$7,090	\$7,625	\$8,277	\$8,944	
y/y		14.2%	11.3%	7.5%	8.5%	8.1%	9.9%
Global AI Spend (\$B)	\$1,765	\$2,596	\$3,493	\$4,221	\$4,947	\$5,616	
y/y		47.1%	34.6%	20.8%	17.2%	13.5%	26.0%
Global IT Spend (\$B) - excl. AI infra	\$4,601	\$4,937	\$5,200	\$5,457	\$5,840	\$6,287	
y/y		7.3%	5.3%	4.9%	7.0%	7.7%	6.4%
% of Global IT Spend	2025	2026E	2027E	2028E	2029E	2030E	
AI Spend %	32%	41%	49%	55%	60%	63%	
AI Spend % (excl. infrastructure)	38%	53%	67%	77%	85%	89%	
AI Spend (by Market)	2025	2026E	2027E	2028E	2029E	2030E	
AI Services	\$436	\$586	\$759	\$943	\$1,121	\$1,259	
AI Cybersecurity	\$26	\$51	\$86	\$127	\$172	\$221	
AI Software	\$283	\$453	\$638	\$809	\$983	\$1,179	
AI Models	\$15	\$33	\$59	\$92	\$124	\$157	
AI Platforms for Data Science and ML	\$21	\$30	\$43	\$59	\$78	\$102	
AI Application Development Platforms	\$7	\$8	\$11	\$14	\$17	\$20	
AI Data	\$1	\$3	\$6	\$10	\$15	\$21	
AI Infrastructure	\$976	\$1,432	\$1,890	\$2,168	\$2,437	\$2,657	
Global AI Spend (\$B)	\$1,765	\$2,596	\$3,493	\$4,221	\$4,947	\$5,616	

Source: Gartner, Ramp, RBC Capital Markets estimates

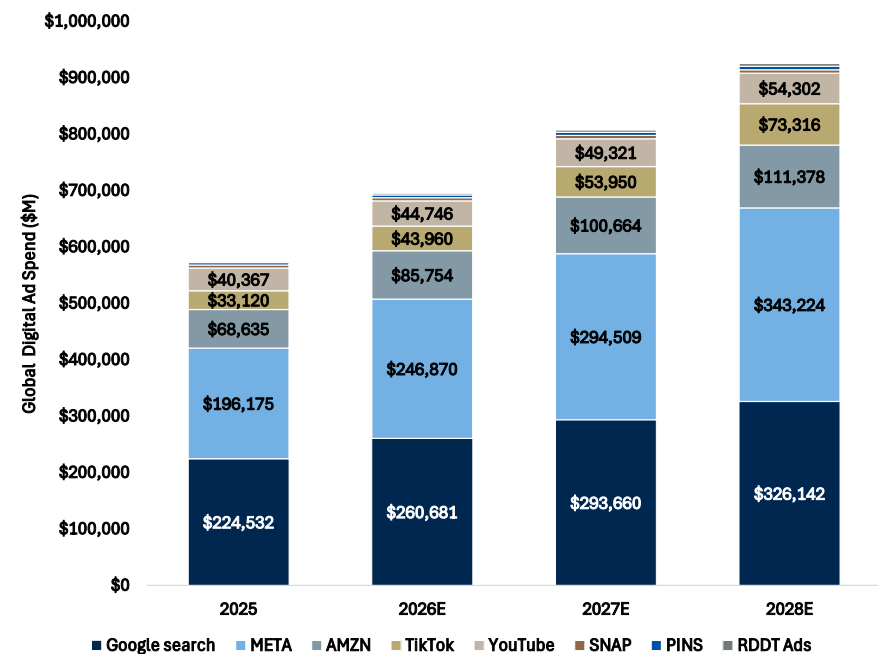
Digital ads opportunity – solid growth but largely going to 2 players

- Digital ads has consistently outgrown all other areas of advertising with HSD growth
 - But META and GOOGL capturing virtually all of the incremental growth, which we'd view as driven by the combination of a structurally superior social graph and unrivaled access to compute, which leads to bigger/better models that appear to be driving sustained/compounding advantages in content recommendation and ad campaign performance.
 - X has historically been a mid- to upper-funnel channel (similar to PINS, SNAP, etc.) – thus breaking out of the single-digit billions of revenue/year has been elusive (vs. TikTok which was approaching \$40B prior to the U.S. business being sold).
 - Short-form video lends toward topical consistency, which enabled TikTok to move down funnel; tbd on whether X can do it.

~\$1.1T global digital ads TAM



~\$926B walled garden global digital ad spend by '28E (>80% of global digital ad spend)



Source: Company reports, Statista, FactSet, Glassdoor, International Labor Organization (ILO), RBC Capital Markets

AI Infrastructure deep dive: Tokenomics and the ROIC on terrestrial and orbital compute



Capital
Markets

The Economics of AI Inference

GPU inference cost economics are governed by two metrics —

- Tokens per second (TPS) measuring throughput and
- Tokens per watt (TPW) measuring energy efficiency

Both of these degrade as model size grows because every token requires a full pass through all model weights, making larger models slower and more energy-intensive per token.

Four criteria determine token value:

- **Intelligence and complexity:** The model's reasoning capability, contextual awareness, and ability to handle domain-specific tasks without human intervention — not simply the volume of tokens produced.
- **Outcome and efficiency:** Whether the model executes business-relevant actions — summarizing complex contracts, debugging code, automating decisions — rather than generating raw token volume with limited downstream impact.
- **Accuracy and hallucination rate:** The reduction of errors and unreliable outputs, since bad outputs destroy value by triggering costly human-in-the-loop review, remediation workflows, or system downtime that far exceed the original compute cost.
- **Latency and throughput:** Time-to-first-token, tokens per second, and overall processing speed, which directly determine whether AI agents can operate in real time and at production scale.

The Cost Economics of AI Inference

First Principles: What Determines Inference Performance

- Parameter Count Is the Master Variable
- Model Size Changes the Bottleneck
- MoE Architecture: MoE models deliver compute costs 3–5× lower than an equivalent-quality dense model

Hardware Reality: Which GPUs Deliver What

- GPU Hardware Generations: : Each generation delivers compounding gains in HBM capacity and bandwidth
- Hardware Performance Hierarchy: Blackwell leads all tiers at ~15–17× the throughput of H100 on frontier workloads
- Energy Efficiency: TPW varies up to 100× between low and high utilization on the same GPU

Software and Efficiency Levers: Extracting More From the Same Hardware

- Lab Benchmarks vs. Production Reality
- Software Stack Matters More Than Hardware Generation
- Quantization: Migrating from BF16 to FP8 delivers 1.9–2.2× throughput improvement
- Hyperscaler Operating Point: Best Case TPS through FP8/FP4 precision, 85–98% utilization, and proprietary inference engines

Cost Construction: From Hardware to Token Price

- Technical Factors That Reduce Token Cost: Quantization, utilization, batch size, MoE architecture, and prompt caching
- Input vs. Output Token Cost Asymmetry: Output tokens cost 3–10× more to produce than input tokens
- The Economic Chain
- Token Cost Floor

Market Reality: Why Published Prices Diverge From Cost Floors

- Scale procurement discounts
- Geographic arbitrage
- Below-cost strategic pricing — cross-subsidized by adjacent revenue streams

The Cost Economics of AI Inference (1 of 2)

Part 1 — First Principles: What Determines Inference Performance

1. **Why Parameter Count Is the Master Variable:** Parameter count determines how much data the GPU must move per token — at 70B+, this makes memory bandwidth the binding constraint and memory fit the key determinant of deployment architecture and cost.
2. **Model Size Changes the Bottleneck:** Above 70B parameters, inference shifts from compute-bound to memory-bandwidth-bound, meaning only more HBM bandwidth — not more compute cores — improves throughput.
3. **MoE Architecture:** MoE models activate only 6–12% of total parameters per token, delivering compute costs 3–5× lower than an equivalent-quality dense model — the primary reason major model developers have converged on sparse architectures.

Part 2 — Hardware Reality: Which GPUs Deliver What

4. **GPU Hardware Generations:** Each generation delivers compounding gains in HBM capacity and bandwidth — the two variables most directly determining inference speed for large models.
5. **Energy Efficiency:** TPW varies up to 100× between low and high utilization on the same GPU, making workload consolidation and continuous batching the primary levers for datacenter-level energy economics.

Part 3 — Software and Efficiency Levers: Extracting More From the Same Hardware

6. **Lab Benchmarks vs. Production Reality:** The 5–10× TPS gap between benchmark ceiling and conservative enterprise deployment on identical hardware is driven almost entirely by software and precision choices, not the GPU itself.
7. **Software Stack Matters More Than Hardware Generation:** A fully optimized software stack on an H100 delivers 5.5× more TPS than a naive stack on the same hardware — a larger gain than upgrading from H100 to H200.
8. **Quantization:** Migrating from BF16 to FP8 delivers 1.9–2.2× throughput improvement — outperforming the gain from pushing utilization from 75% to 100%, making precision migration the higher-return optimization in every case. BF16, FP8, and FP4 refer to how many bits (16, 8, or 4) are used to represent each number in a model during inference — fewer bits mean less memory, faster computation, but rougher precision. Going from BF16 → FP8 → FP4 is essentially trading accuracy for speed and efficiency.
9. **Hyperscaler Efficiencies Matter:** Major model developers operate at approximately 70% of Lab — Best Case TPS through FP8/FP4 precision, 85–98% utilization, and proprietary inference engines — materially closer to the benchmark ceiling than a typical enterprise deployment. A conservative enterprise deployment on identical B200 hardware produces roughly half the TPS vs. a hyperscaler deployment on the same hardware.

The Cost Economics of AI Inference (2 of 2)

Part 4 — Cost Construction: From Hardware to Token Price

- 10. Technical Factors That Reduce Token Cost:** Quantization, utilization, batch size, MoE architecture, and prompt caching each address an independent constraint and compound multiplicatively, collectively reducing cost per token by 10–15× vs. an unoptimized BF16 baseline.
- 11. Input vs. Output Token Cost Asymmetry:** Output tokens cost 3–10× more to produce than input tokens because decode is sequential while prefill is parallel — the cheapest published headline prices always refer to the input side.
- 12. The Economic Chain:** GPU hardware specs drive TPW, which drives infrastructure cost, which sets the token cost floor — but GPU compute represents only ~5–15% of a published API price, with model development, safety, platform costs, and margin comprising the remainder.
- 13. Token Cost Floor:** The sustainable raw compute floor for optimized dense-model inference on current-generation US hardware is \$0.15–\$0.20/1M output tokens, falling to \$0.08–\$0.12 for MoE architectures in low-cost geographies.

Part 5 — Market Reality: Why Published Prices Diverge From Cost Floors

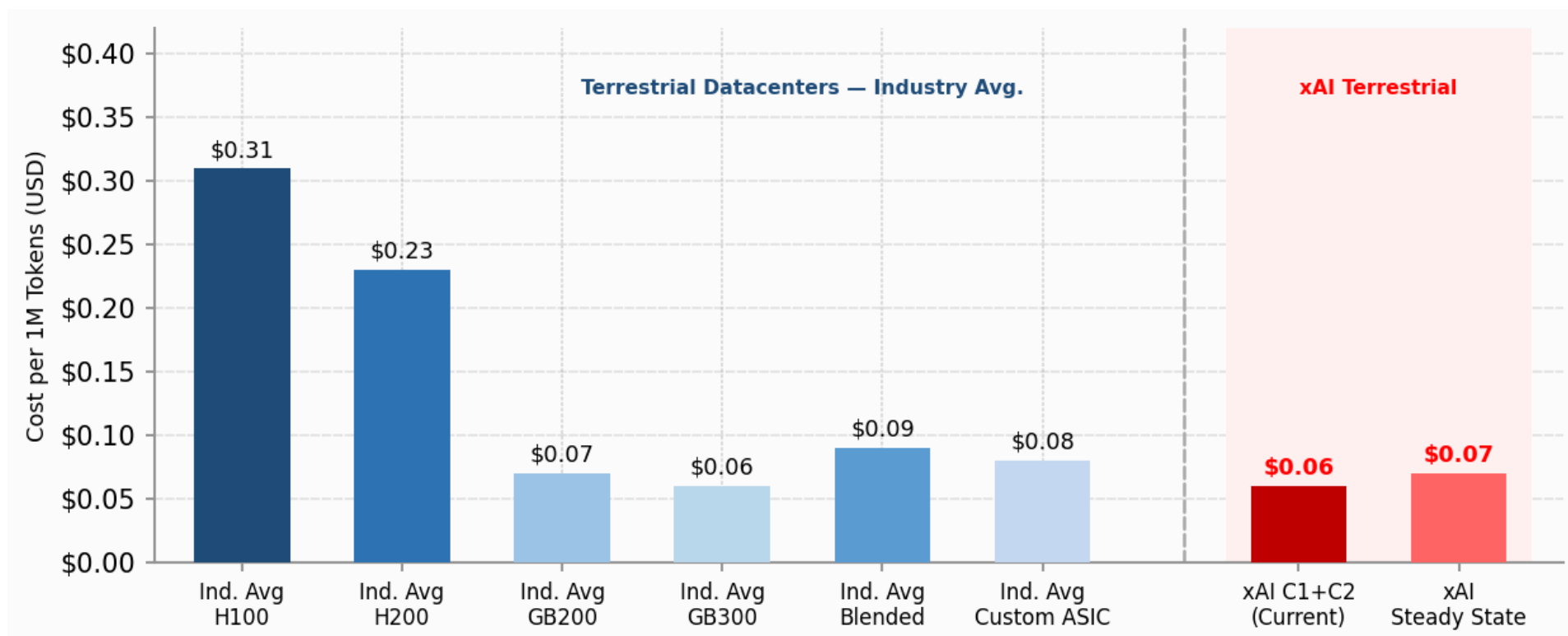
- 14. Economic Factors:** Scale procurement discounts, geographic arbitrage, and deliberate below-cost strategic pricing — cross-subsidized by adjacent revenue streams — collectively explain how published API prices fall well below the \$0.15–\$0.20 raw compute floor.

xAI: Inference Cost Structure and Hardware Positioning

Grok's MoE architecture activates only a small fraction of total parameters per token, delivering materially lower compute cost vs. an equivalent-quality dense model — the same structural advantage adopted by Anthropic and OpenAI, placing xAI's cost structure well below operators running dense models on equivalent hardware. xAI's mixed fleet spans the middle and upper tiers of the throughput hierarchy, with its newer GPU assets delivering throughput materially above the H100/H200 average that represents the majority of competitor deployed capacity. Memory fit, high utilization through continuous batching, and advanced precision settings in production place xAI firmly in the hyperscaler operating band alongside Anthropic and OpenAI. Taken together, MoE architecture and geography-adjusted costs push xAI's token cost structure toward the lower end of the MoE-optimized compute floor.

Average Cost of 1M Output Tokens (All Model Sizes, FP8, 70% Utilization)

- xAI C1+C2 at \$0.06 — 36% below industry blended average (\$0.09).



xAI C1+C2 advantage drivers:

- Lower infrastructure COGS:** \$7.5/MW/yr vs. \$8.6 industry average (-12%), driven by brownfield build, lower infra. amortization (\$1.07 vs. \$1.27), and reduced labor/SW costs.
- Mixed-gen fleet weighted toward GB200/GB300:** Delivers 46% higher TPS and 48% higher TPW vs. industry average across all model sizes.
- MoE architecture (Grok):** Activates only 3–10% of parameters per token — 3–5× lower compute cost vs. equivalent-quality dense models.
- Partially depreciated H100/H200 GPUs:** 18–24 months into a 5-year amortization schedule at time of Anthropic deal, reducing effective hardware cost per token.
- Hyperscaler utilization (85–98%):** Continuous batching spreads fixed infrastructure cost across proportionally more tokens produced per year.

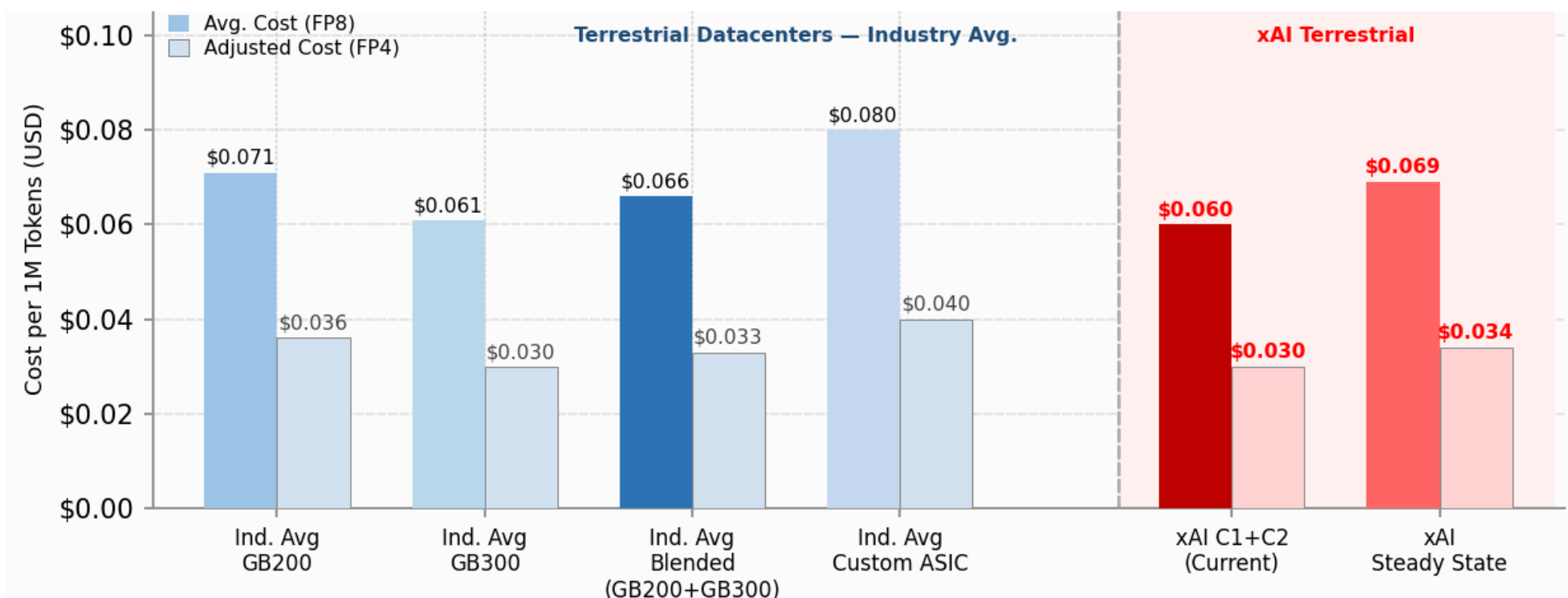
Note: Cost/1M tokens = average across all model sizes (8B, 70B, 100B, Small MoE, Medium MoE, Large MoE, xLarge MoE) at 70% utilization, FP8.

Industry blended average = equal-weighted mean of H100, H200, GB200, GB300. xAI steady-state reflects normalized infrastructure amortization at industry-standard levels.

Source: RBC Capital Markets estimates, company reports, MLcommons

Average Cost of 1M Output Tokens (All Model Sizes, FP4 vs. FP8, 70% Utilization)

- xAI C1+C2 at \$0.030 — 9% below the industry blended average of \$0.033 (GB200+GB300) on an FP4 adjusted basis.



- The cost per 1M tokens declines by 50%, moving from FP8 to FP4 precision across all platforms.
- FP4 quantization halves the memory footprint per model parameter relative to FP8, allowing twice as many parameters to reside in GPU memory simultaneously. This doubles effective throughput per GPU at the same power draw, producing a direct 50% reduction in cost per token across all platforms.
- The gain is consistent across GB200, GB300, and Custom ASIC because throughput improvement is a function of memory bandwidth utilization, not GPU architecture. FP4 is supported natively on Blackwell-class GPUs (GB200, GB300) via NVIDIA's NV-FP4 format; H100 and H200 are excluded as those architectures do not support FP4 natively.

Note: Cost/1M tokens = average across all model sizes (8B, 70B, 100B, Small MoE, Medium MoE, Large MoE, xLarge MoE) at 70% utilization, FP8/FP4.

Industry blended average = equal-weighted mean of H100, H200, GB200, GB300. xAI steady-state reflects normalized infrastructure amortization at industry-standard levels.

Source: RBC Capital Markets estimates, company reports, MLcommons

Terrestrial Datacenters

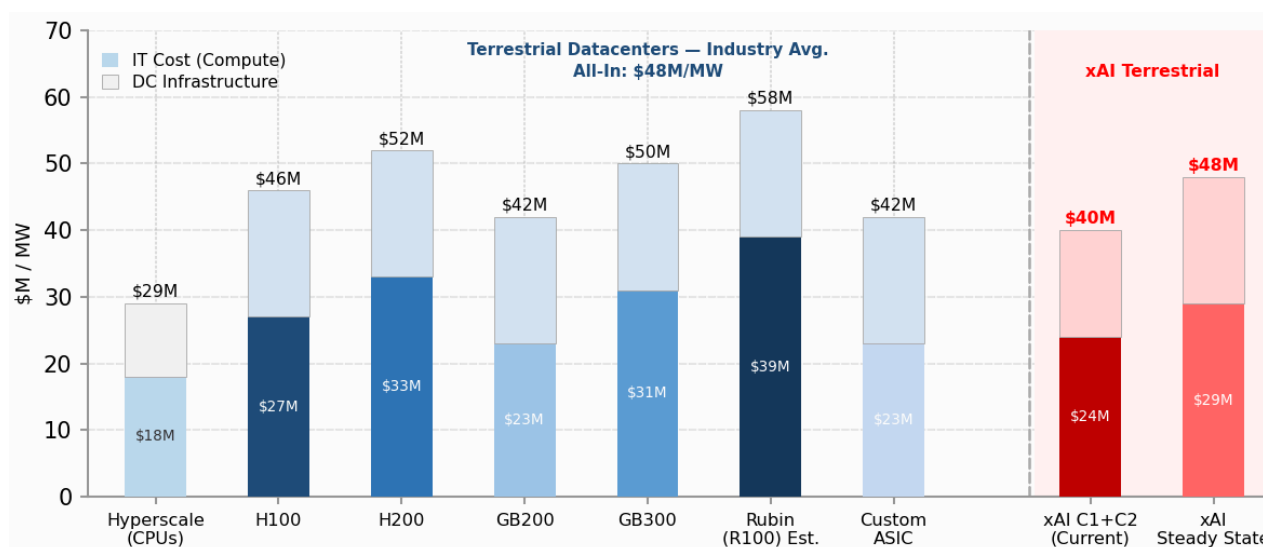


Capital
Markets

Terrestrial Datacenter Development Costs

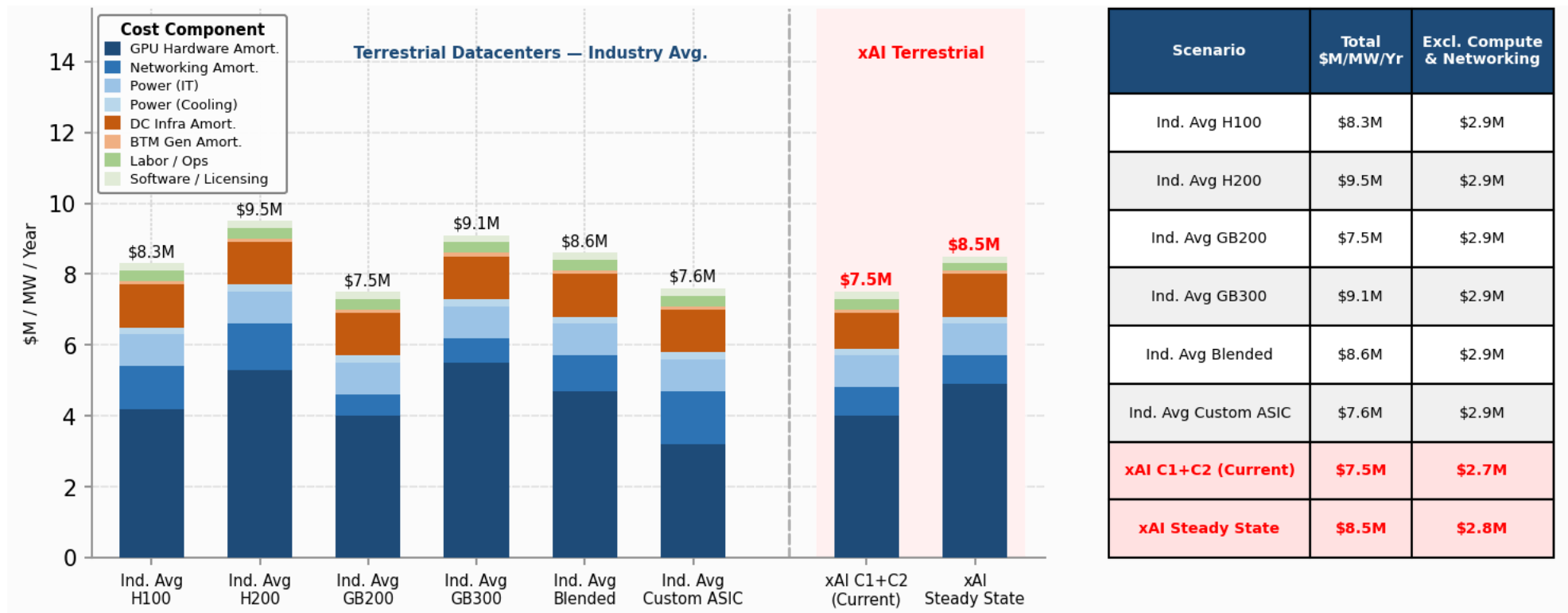
- xAI's C1+C2 current infrastructure cost is below the AI-optimized DC benchmark of \$19.0M/MW, reflecting construction efficiencies from accelerated build timelines, vertically integrated project execution, and lower physical building costs driven by large-format campus development.
- At steady-state, We expect xAI's infrastructure cost to converge toward the AI-optimized DC benchmark as the facility reaches full build-out specification.
- BTM generation of \$1.5M/MW is consistent across both xAI scenarios and the AI-optimized benchmark, reflecting on-site power generation as a standard component of AI datacenter infrastructure.

	Terrestrial Datacenters			
Datacenter Development costs (excluding compute) -- \$M per MW	Hyperscale Traditional DC	AI-Optimized DC (liquid cooling)	xAI (C1+C2) Current	xAI Steady State
Gray Space (Electrical – Back-end)	\$2.6	\$4.6	\$3.9	\$4.6
Cooling	\$2.8	\$6.1	\$5.2	\$6.1
Other (incl. Back-up Generators)	\$2.6	\$2.1	\$1.8	\$2.1
White Space (Electrical – IT Room)	\$0.9	\$2.4	\$2.1	\$2.4
Total Infrastructure Costs	\$8.8	\$15.3	\$13.0	\$15.3
Physical Building	\$2.2	\$2.2	\$1.6	\$2.2
BTM Generation	-	\$1.5	\$1.5	\$1.5
Total Development cost	\$11.0	\$19.0	\$16.1	\$19.0



Source: RBC Capital Markets, company reports

Terrestrial Datacenters Operating Cost Comparison (\$M per MW)



Note: xAI-Terrestrial operating costs are for the mix of H100, H200, GB200 and GB300 GPUs at C1 and C2.

- xAI Current** has a lower total COGS at **\$7.5M/MW/Year**, driven primarily by lower infrastructure amortization (\$1.07M vs. \$2.06M industry avg) and reduced labor/software costs.
- xAI Steady-State** converges back to **\$8.6M/MW/Year** — equal to the industry-average — as infrastructure costs normalize to industry levels.
- Excluding compute and networking, xAI Current's cost base (\$2.7M) is ~4% below the industry average (\$2.9M), suggesting modest operational efficiency at the non-hardware layer.
- GPU hardware amortization is the dominant cost driver in all scenarios, ranging from \$4M (xAI Current) to \$5.5M (Ind. Avg GB300).

Source: RBC Capital Markets estimates, company reports

Colossus 1 (Memphis, TN) — Build-Out & Power Infrastructure

- **June 2024:** xAI and the Greater Memphis Chamber announced plans to build a supercomputer at the former Electrolux manufacturing facility at 3231 Paul R. Lowry Road in Memphis, Tennessee (Boxtown neighborhood, southwest Memphis). Original plan called for approximately 100,000 GPUs drawing up to 150MW of power.
- **Summer 2024:** xAI repurposed the existing factory shell. The first cluster of approximately 100,000 Nvidia H100 processors (~130MW of compute power) was brought online in 122 days. The 100,000 H100 GPUs were physically installed in approximately 19 days using Supermicro liquid-cooled racks.
- **Summer 2024:** Because grid power was not yet available, xAI deployed 15–35 natural gas turbines on trailers to power the facility, exploiting a loophole that classified trailer-mounted turbines as "non-road engines" not subject to Clean Air Act permitting. Satellite imagery documented as many as 35 turbines.
- **September 2024:** Colossus announced as operational — a 100,000 H100 GPU training cluster. Expansion to 200,000 GPUs (including 50,000 H200s) stated as a near-term target.
- **November 2024:** TVA board approved the initial 150MW grid power request on November 7, 2024. MLGW agreed to serve xAI directly, with a new substation (funded by xAI at approximately \$35M) constructed to deliver the grid connection.
- **December 2024:** xAI announced expansion to a minimum of one million GPUs at the Greater Memphis Chamber Annual Chairman's Luncheon on December 4, 2024. Nvidia, Dell, and Supermicro establishing operations in Memphis. Power demand would exceed 1,000MW (~10x the approved 150MW). A second 150MW request submitted to MLGW remained under initial review.
- **May 2025:** MLGW confirmed xAI connected to newly constructed electric substation (#63), receiving 150MW of grid power. 150MW of Tesla Megapack batteries integrated for backup and demand-response. Temporary turbines stated to be in the process of removal. Total draw: 150MW (8MW from existing substation, 142MW from new xAI-funded substation). xAI deployed 168 Megapack units targeting approximately 1 GWh of buffering capacity.
- **Approved Grid Capacity: 300 MW.** Connected through Memphis Light, Gas and Water (MLGW), which sources wholesale power from the Tennessee Valley Authority (TVA). An initial 150 MW allocation was later expanded with an additional 150 MW approved.
- **On-Site Generation:** To bypass initial grid constraints, xAI previously utilized temporary gas turbines (up to 250 MW), now used primarily as backup.

Commercialization

- **May 2026:** SpaceX (which acquired xAI in February 2026) struck a compute capacity agreement with Anthropic covering the entirety of Colossus 1 — more than 300MW of capacity across approximately 220,000 Nvidia GPUs (H100, H200, and GB200 accelerators). Anthropic agreed to pay \$1.25B per month through May 2029, for a potential total of over \$40B over the contract term. The arrangement followed technical difficulties SpaceX encountered connecting Colossus 1 to its other two datacenter campuses due to latency issues and hardware variations (Colossus 1 runs Nvidia Hopper and Blackwell alongside older accelerators; other sites built more uniformly around Blackwell chips).

Colossus 1 — Regulatory, Environmental & Commercial Developments

Permitting & Environmental

- **January 2025:** xAI filed application with the Shelby County Health Department to formally permit 15 gas turbines, months after turbines were already running. Self-reported emissions: 11.51 tons of HAP per turbine per year, above the EPA's 10-ton annual single-source threshold. Separately, TDEC approved an operations permit for xAI's planned greywater recycling facility.
- **June 2025:** NAACP filed notice of intent to sue xAI under the Clean Air Act, citing unpermitted gas turbine operations. SELC separately announced plans to file suit. SELC estimated turbines emitted 1,200–2,000 tons of NOx annually. Within 11 months of arrival, xAI had become one of Shelby County's largest emitters of smog-forming nitrogen oxides per E&E News analysis.
- **July 2025:** Shelby County Health Department issued air permit on July 2, 2025, allowing operation of 15 gas turbines, expiring January 2, 2027. Permit required installation of BACT by September 1, 2025, and set emissions limits with potential fines up to \$10,000 per day per violation. Satellite imagery from July 1, 2025, still showed 24 turbines on-site.
- **Late 2025:** Memphis and Shelby County Air Pollution Control Board voted to dismiss appeal of xAI's air permit filed by SELC, NAACP, and Young, Gifted & Green.
- **January 2026:** EPA issued a rule update closing the "non-road engine" loophole, requiring Clean Air Act permits for all trailer-mounted turbine deployments going forward.

Water Recycling Plant

- **October 2025:** xAI held ceremonial groundbreaking for the Colossus Water Recycling Plant, an \$80M greywater treatment facility on 13 acres adjacent to the datacenter, intended to reduce reliance on the Memphis aquifer (facility using up to 5 million gallons per day at peak cooling demand).
- **April 2026:** xAI halted construction of the water recycling plant, citing need to prioritize Colossus 2. Musk stated construction would resume after Colossus 2 stabilized. Memphis aquifer remained ongoing cooling source.
- **June 2026:** CTC Property LLC (xAI affiliate) closed on \$820,000 purchase of the 13 adjacent acres for the water plant. Memphis city officials stated they could reclaim the land if substantial construction progress was not made within one year of May 14, 2026.

Colossus 2 and 3 — Build-Out & Power Infrastructure

Colossus 2

- **Late 2024:** xAI acquired approximately 100 acres and a 1-million-square-foot warehouse in Whitehaven (majority-Black neighborhood in Memphis) at 5420 Tulane Road, approximately 8 miles from Colossus 1 and approximately 0.5 miles from John P. Freeman School (grades K–8). Greater Memphis Chamber indicated Colossus 2 would not use gas turbines, in contrast to Colossus 1.
- **Late 2024 / Early 2025:** Per SpaceX's S-1 filing (May 2026), the first cluster of approximately 110,000 Nvidia GB200 processors delivering approximately 210MW of compute power was brought online in 91 days. A second cluster of 110,000 GB300 processors and 220MW of compute power was subsequently brought online in 64 days. Together, the first two clusters brought total Colossus 2 compute capacity to approximately 430MW.
- **July 2025:** xAI subsidiary MZX Tech LLC purchased the former Duke Energy site at 2875 Stanton Road in Southaven, Mississippi, approximately 2 miles from Colossus 2 and across the Tennessee–Mississippi state line, to serve as primary power generation for Colossus 2.
- **August 2025:** Three gas turbines on-site at the Southaven power plant by August 1, 2025, per email correspondence between xAI's consultant (Trinity Consultants) and MDEQ. MDEQ approved unpermitted use of 16 turbines at the site in August 2025.
- **September 2025:** Consultant correspondence recorded plans for 17 turbines at the Southaven site. Turbine count grew to 27 by December 11, 2025, per MDEQ-filed communications, with combined stated capacity of approximately 495MW.
- **March 2026:** SpaceX's S-1 disclosed a March 2026 agreement to purchase \$805M of turbines from an unnamed vendor through 2029 to support ongoing datacenter power demands. **Late-April 2026:** SpaceX entered a further agreement for \$2B of mobile gas turbines and related items from an unnamed vendor. Combined turbine procurement commitment reached approximately \$2.8B.
- **May 2026:** By early-May 2026, xAI had brought 19 additional turbines to the Southaven site between late March and early May, raising the total to 46 per MDEQ email correspondence.
- **May 2026 (Expansion):** Next phase of Colossus 2 expansion expected to bring online at least 220,000 additional GB300 processors and over 400MW of additional compute power once fully operational, at an estimated cost of over \$659M.
- **June 2026:** Turbine count at Southaven increased to 57 by May 15, 2026.
- **Approved Grid Capacity: <1 MW. On-Site Generation:** Because it bypasses traditional grid draw, this facility relies entirely on a dedicated, on-site natural gas power plant. The Mississippi Department of Environmental Quality (MDEQ) approved an air permit clearing up to **1.2 GW** of self-generating power capacity, anchored by 41 permanent methane gas turbines.
- **Storage:** The site features massive Tesla Megapack battery storage installations to smooth out millisecond-level power spikes inherent to AI training.

Colossus 3

- **Location:** Southaven, Mississippi, directly across the Tennessee state line and adjoining the Colossus 2 property.
- **Facility size:** Approximately 810,000 square feet.
- **Purpose:** Acquired by xAI to scale AI training capacity.
- **Combined campus footprint (Colossus 1, 2, and 3):** Approximately 2.5 million square feet.
- **Combined compute infrastructure across Colossus 1, 2, and 3:** Approximately 2 GW of power and over 1 million GPUs.

Colossus 2 — Regulatory, Legal & Commercial Developments

Community Opposition & Permitting

- **November–December 2025:** Thermal imagery confirmed at least 18 of 27 turbines in active operation. Southaven community members filed noise complaints and circulated petitions demanding accountability from MDEQ.
- **February 2026:** NAACP issued formal 60-day notice of intent to sue xAI and MZX Tech LLC on February 13, 2026, citing 27 unpermitted turbines at the Southaven site in violation of the Clean Air Act. NAACP estimated 27 turbines had potential to emit over 1,700 tons of NOx annually, making the facility the largest industrial source of NOx in the 11-county Memphis metropolitan area. On February 17, MDEQ held a public hearing on xAI's permit application for 41 permanent turbines; no members of the public spoke in favor.
- **March 2026:** MDEQ approved air permit for 41 turbines at the Southaven site despite widespread community opposition. xAI argued trailer-mounted turbines were "temporary" and therefore exempt from standard permitting under state rules.

Litigation

- **April 2026:** NAACP, represented by Earthjustice and SELC, filed a federal lawsuit on April 14, 2026, in the U.S. District Court for the Northern District of Mississippi against xAI and MZX Tech LLC, citing the Clean Air Act. Complaint documented 27 unpermitted turbines at the Southaven plant.
- **May 2026:** On May 6, 2026, NAACP filed preliminary injunction seeking to halt all unpermitted turbine operations. xAI opposed, stating datacenters "cannot currently function and would have to precipitously shut down" without temporary power generation equipment.
- **June 2026:** NAACP filed motion on June 10, 2026, to amend preliminary injunction to reflect updated count of 57 turbines.
- **June 9, 2026:** Three Southaven residents filed a class-action lawsuit against xAI and SpaceX in federal court, representing an estimated class of over 10,000 affected individuals, citing continuous noise and vibration from turbine operations.
- **June 16, 2026:** U.S. Department of Justice filed motion on behalf of the federal government supporting dismissal of the NAACP's Clean Air Act lawsuit, arguing halting turbine operations would endanger U.S. national, economic, and energy security interests, citing Grok's use in military operations including U.S. strikes on Iran.

Commercialization

- **June 22, 2026:** SpaceX signed compute capacity agreement with Reflection AI granting access to Nvidia GB300 chips at Colossus 2 for \$150M per month from July 1, 2026, through 2029, totaling approximately \$6.3B at full term. This added Reflection to an existing roster of Colossus 2 compute customers including Google (\$920M/month) and Cursor.

xAI Colossus: Positioning, Premium, and Deal Rationale

▪ Scale and Build Speed

- **Colossus 1** brought ~100,000 H100 processors (~130 MW) online in **122 days** using a brownfield factory shell — vs. an industry benchmark of ~2 years for a 100 MW greenfield datacenter.
- **Colossus 2** brought two successive clusters online in **91 days** (110,000 GB200s, ~210 MW) and **64 days** (110,000 GB300s, ~220 MW), reaching ~430 MW across both clusters.
- Combined, Colossus I and II deliver approximately **1.0 GW** — positioned as the world's first coherent gigawatt-scale AI training cluster.

▪ Infrastructure Cost Efficiency

- Colossus II construction costs were below industry benchmarks on a per-MW basis, enabled by brownfield retrofit, direct-to-chip liquid cooling, high-density rack layouts, and BTM power generation — versus an industry average of \$11–20M/MW for AI-optimized builds.

▪ Colossus 1: Older GPU Mix, Lower Pricing

- Colossus 1 houses ~220,000–325,000 GPUs (H100, H200, GB200) — the Hopper/early-Blackwell fleet originally built for Grok training beginning mid-2024, placing the H100/H200 inventory approximately **18–24 months into a standard 5-year depreciation schedule** at the time of the Anthropic deal.
- We estimate Anthropic's deal covering all of Colossus 1 at \$26M/MW annualized revenue — approximately **2× lower per GPU** than Anthropic's deal for additional capacity at Colossus 2, directly consistent with the older, partially depreciated hardware mix.
- SpaceX also encountered **technical integration difficulties connecting Colossus 1** to its other campuses due to latency and hardware heterogeneity — an additional factor supporting the decision to lease out Colossus 1 rather than absorb it into a single training fabric.

▪ Colossus 2: Newer GB300 Cluster, Premium Pricing

- Colossus 2 runs **GB200 and GB300 NVL72** rack-scale systems. GB300 delivers 50% more HBM3e (288 GB), 50% more dense FP4 compute, and ~2.7× the throughput of GB200 NVL72.
- We estimate that the Google Contract (~110K GPUs, 220 MW) and Reflection AI Contract (~18K GB300 GPUs) were both priced at ~\$50M/MW annualized revenues — confirming GB300 capacity at Colossus 2 is priced as a distinct, premium tier.

▪ Why the Campus Commands a Premium for deals like Anthropic and Google

- **No build lead time.** GB300 cluster orders currently carry 12–15-month delivery timelines. Colossus was fully operational and available immediately — critical for Anthropic (revenue up ~80× in a single quarter, Claude Code severely throttled) and Google (Gemini Enterprise demand exceeding internal forecasts).
- **Cluster coherence at scale.** A coherent 200,000+ GPU single-fabric cluster is not replicable by neo-cloud competitors, who typically operate in the tens-of-thousands-of-GPU range per pod. Frontier model training requires this density.
- **Short-term optionality.** Both deals carry **90-day termination clauses** after initial lock-up — versus CoreWeave's standard 4–6-year take-or-pay. Customers pay a premium for the right to walk away.
- **Google's deal is explicit urgency pricing.** Google described it as "bridge capacity to meet surging Gemini Enterprise demand that exceeded expectations" — supplemental procurement by the world's largest AI compute owner, reflecting acute scarcity pricing rather than steady-state economics.
- **Full-stack product.** Customers receive GPUs backed by hyperscale-class CPUs, exabyte-scale storage, and purpose-built AI networking — not bare-metal GPU hours. Neo-cloud bare-metal comparables price at ~\$13M ARR/MW vs. \$30–50M/MW here.

XAI – Recent AI Compute Contracts

Metric	Anthropic Contract	Google Contract	Reflection AI Contract
Date Signed	May 3, 2026	June 5, 2026	June 22, 2026
Facility	Colossus 1 + 2	Colossus 2	Colossus 2
GPU Type	H100 / H200 / GB200	GB300	GB300
GPU Count	~325,000	~110,000	~17,934
Power Equivalent	~500MW	~220MW	~36MW
Monthly Fee	\$1.25B	\$920M	\$150M
\$/GPU/month	~\$3,846	~\$8,364	~\$8,364*
\$/GPU/hour	~\$5.27	~\$11.46	~\$11.46*
Annualized Revenue/MW	~\$30.0M	~\$50.2M	~\$50.0M*
Full-Rate Start	May/Jun 2026 (ramp)	October 2026	July 1, 2026
Contract End	May 2029	June 2029	~Dec 2029 (est.)
Full Term (months)	~36	~33	~42 (est.)
Total (if run to term)	~\$45B	~\$30B	~\$6.3B
Annualized Revenue	~\$15.0B	~\$11.0B	~\$1.8B
Ramp Period	May–Jun 2026 (reduced fee)	Through Sep 2026 (reduced fee)	None disclosed
Earliest Termination	After 3 months + 90 days notice	After Dec 31, 2026 + 90 days notice	After 3 months + 90 days notice
Delivery Penalty	Not disclosed	Termination right or pro-rata reduction if not delivered by Sep 30, 2026	Not disclosed
Use Case	Inference (Claude Code, Claude API)	Bridge capacity for Gemini Enterprise and Apple Intelligence	Open-source frontier model training; U.S. government AI programs

Source: RBC Capital Markets, Company reports, The Information

SpaceX's Google compute deal highlights scarcity pricing leverage

- We estimate Google paying SpaceX a ~5x premium for Colossus II compute access (\$920M/month implying ~\$50/W).
- Why? We think it's a combo of not enough time and Apple's likely desire to be a consumer AI winner.
 - Google cannot build datacenters fast enough on its own to support Apple's rollout of Apple Intelligence, expected this fall.
 - It's possible Apple could be paying as much as \$1.5B/month or more to Google to scale these new features.
 - 1/5 of its base consuming 500k tokens/day (plausible for a single complex ask across various apps in your phone) at \$0.20/M tokens would equate to \$1.5B/month of GCP compute revs.
 - We believe there's a scarcity premium for Apple's private cloud requirements which are unlike any other Google Cloud customer and take years to build.
- Nearer term, GCP margins from Apple are lower through SpaceX compute but only temporarily until Google has adequate in-house capacity.
 - How temporary is the question where Google still must maintain optionality given Apple doesn't know what LT Apple Intelligence adoption and utilization will ultimately be.

Recent AI Compute Contracts

Date	Contract	Annualized Revenue (\$M)	Disclosed / Derived MW	Implied \$M/MW/Year
Sep-25	Nebius / Microsoft	~\$3,480M/yr (\$17.4B ÷ 5 yrs)	~200 MW (derived: 100,000 GB300 GPUs × 1.2 kW/GPU × 1.3 PUE ÷ 1,000)	~\$17.4M/MW/yr
Oct-25	Nscale / Microsoft (Texas / Portugal)	~\$2,800M/yr (\$14B ÷ estimated 5 yrs)	~280 MW (derived: 104,000 GB300 GPUs at Cedarvale ~161 MW IT load + 12,600 GB300 GPUs at Sines ~20 MW IT load + Cedarvale facility gross ~234 MW; blended ~280 MW using facility gross)	~\$10.0M/MW/yr
Nov-25	Fluidstack / Anthropic	~\$9,615M/yr (\$50B ÷ estimated 5.2 yrs across TX 168 MW + NY 360 MW sites)	~528 MW (168 MW Abernathy TX + 360 MW Lake Mariner NY)	~\$18.2M/MW/yr
Sep-25	Lambda Labs / NVIDIA Backstop	~\$375M/yr (\$1.5B ÷ 4 yrs)	~25 MW (derived: 18,000 GPUs; mix of ~10,000 high-end H100/GB200 ~1.0 kW/GPU avg + ~8,000 lower-end ~0.4 kW/GPU; blended ~14,400 kW IT load × 1.3 PUE ÷ 1,000 ≈ ~19 MW; rounded to ~25 MW including supporting infrastructure)	~\$15.0M/MW/yr
May-26	SpaceX (xAI) / Anthropic	\$15,000M/yr (\$1.25B/month × 12)	~500 MW (Anthropic's announcement: "more than 300 MW of additional capacity"; ~325,000 NVIDIA GPUs × blended ~1.1 kW/GPU × 1.3 PUE ÷ 1,000 ≈ ~465 MW; rounded to ~500 MW consistent with industry estimates)	~\$30.0M/MW/yr
Jun-26	SpaceX (xAI) / Google	\$11,040M/yr (\$920M/month × 12)	~200 MW (derived: ~110,000 NVIDIA GPUs × blended ~1.1 kW/GPU × 1.3 PUE ÷ 1,000 ≈ ~157 MW IT-equivalent; rounded to ~200 MW on a facility gross basis consistent with industry estimates)	~\$50M/MW/yr
Jun-26	SpaceX (xAI) / Reflection AI	~\$6.3B if run to term; ~\$1.8B annualized	~17,934 GB300 GPUs at Colossus 2; ~36MW; 249 nodes at 72 GPUs/node; CPUs, memory, and related components assumed consistent with GB300 NVL72 rack configuration	~\$50M/MW/yr

Source: RBC Capital Markets estimates, company reports

Recent AI Contracts

Metric	Akamai – \$200M Deal	Akamai – \$1.8B Deal	Megaport – \$25M Contract (CPU)	Megaport – \$183M Contract (GPU)	Megaport – \$330M Contract (GPU)
Announced	Feb 2026 (4Q25)	May 2026 (1Q26)	Apr-26	May-26	Jun-26
# of Contracts	1	1	1	3	4
Customer	Major U.S. tech company (existing customer)	Leading frontier model company (Anthropic)	U.S. developer tooling company for Agentic AI workloads	Two U.S.-based AI companies	Tech-based AI inference customers
Hardware Type	NVIDIA Blackwell RTX Pro 6000 GPUs	GPU/CPU (mix undisclosed)	CPU servers	NVIDIA H100 & RTX-class GPUs	Primarily NVIDIA GPU
Workload	AI Inference	AI (undisclosed specifics)	Compute & storage	Compute + storage + networking	AI inference -- Compute & storage
Term	4 years	7 years	36 months	~33 months (blended)	27-28 months (blended)
TCV	US\$200M	US\$1.8B	US\$25.1M	US\$183M	US\$330M
ARR	~\$50M	\$257M	US\$8.4M	US\$65.2M	US\$143M
CapEx	\$100–\$150M	US\$800–\$825M	US\$12.2M	US\$101M	US\$266M
EBITDA Margin			~70% (contract-level)	~70% (contract-level)	~83% contract level, 65% (reflecting ongoing investment)
CapEx / ARR	2.0–3.0x	2.8–3.2x	1.45x	1.55x	1.86x
CapEx / TCV	0.50–0.75x	0.45x	0.49x	0.55x	0.81x
CapEx per MW (reference)	\$20M	\$26-35M	\$12.2M	\$14–\$20M	\$28M
Estimated MW equivalent capacity	5.0–7.5 MW	23-30 MW	1 MW	5–7 MW	~13 MW
Est. Ann. Revenue per MW	\$6.7–\$10.0M	\$9–\$11M	\$8.4M	\$9–\$13M	~\$15M
Datacenter Partner	Third-party colo (U.S.)	Third-party colo (multiple, global)	Third-party colo (U.S.)	Third-party colo (U.S.)	Third-party colo (U.S.)
PAYBACK PERIOD					
Revenue-based (CapEx ÷ ARR)	~2.5 years	~2.9 years	~1.45 years	~1.55 years	1.86 years
At 70% EBITDA margin	~3.6 years	~4.5 years	~2.1 years	~2.2 years	
At 50% EBITDA margin	~5.0 years	~6.3 years	~2.9 years	~3.1 years	
Company-Stated Payback	1-2 years based on Revenue	1-2 years based on Revenue	~2 years based on EBITDA	~2 years based on EBITDA	~2.2 years based on EBITDA

Source: RBC Capital Markets estimates (text and numbers in blue), company reports

GPU Financing and Other Capex Related

Transaction Structure

- xAI, through its subsidiary CTC Property LLC, has entered into three equipment lease agreements with Valor Equity Partners for AI infrastructure hardware and computing-related equipment, with aggregate undiscounted lease payments of \$20.2B across the three transactions (\$6,986M, \$6,633M, and \$6,587M, respectively). Antonio J. Gracias, founder, CEO, and Chief Investment Officer of Valor Equity Partners, serves as a member of xAI's board of directors. All three arrangements have been classified as failed sale-leaseback transactions under US GAAP, meaning the underlying AI infrastructure assets remain on xAI's consolidated balance sheet within PP&E and the associated payment obligations are carried as debt. The structure is economically equivalent to secured borrowing collateralized by AI infrastructure assets.

Cash Payment Cadence

- Cash payments under the Valor leases totaled \$885M for full-year 2025 and \$857M for the two-month period January 1 through February 28, 2026, alone, implying an annualized 2026 payment run-rate of approximately \$5.1B. Under ASC 842, principal repayment components of these payments are classified as financing activities, while the interest component is classified as an operating activity. Reported capex (investing section) reflects only cash-basis capital expenditures; the ROU asset recognized at inception is a non-cash transaction and lease principal payments flow through financing activities rather than the investing section.

SpaceX Guarantee

- Payment and performance obligations under all three lease agreements are guaranteed by SpaceX or one of its subsidiaries. The guarantee represents a contingent liability at the SpaceX consolidated level for the full undiscounted \$20.2B payment stream, disclosed in the notes to the financial statements.

Related Party Transactions: Tesla

- xAI purchased \$34M of Tesla Megapack products in Q1 2026, recorded within PP&E. As of December 31, 2025, cumulative purchases from Tesla totaled \$506M of Megapack products and \$131M of Cybertrucks at manufacturer's suggested retail price, both recorded within PP&E.

Power Infrastructure: Turbine Acquisition

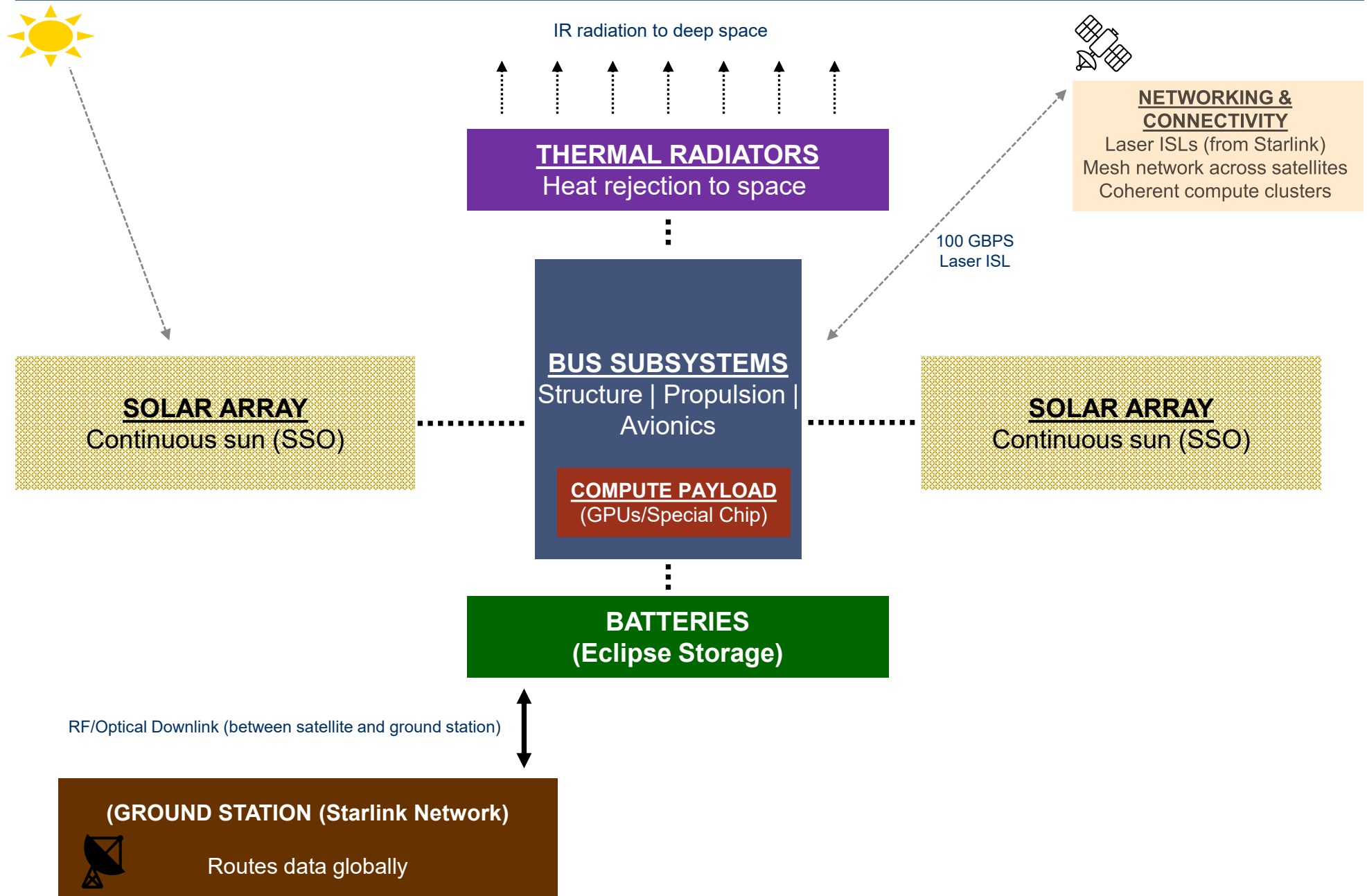
- On April 30, 2026, xAI entered into an asset purchase agreement with an unaffiliated third party to acquire mobile gas turbines and related packages for approximately \$2,000M. The acquisition is intended to provide dedicated power capacity to xAI's datacenters.

xAI Orbital Datacenters



Capital
Markets

Orbital Datacenter Satellite — Component Architecture



Source: RBC Capital Markets, company reports, Google Suncatcher white paper (2024), Starcloud, Star-catcher, Space Tech

Orbital Datacenter Satellite: Component Glossary & Mass Impact

Component	What It Is	Why It Matters for Mass
Solar Arrays	Deployable panels that convert sunlight into electricity — the only power source in orbit.	Heaviest component. Over a third of total mass. Can be reduced by choosing orbits with continuous sunlight.
Batteries	Energy storage that powers GPUs when the satellite is in Earth's shadow.	Second heaviest. Stores ~30 minutes of backup power. Largely unnecessary in eclipse-free orbits.
Bus Subsystems	Structural frame, thrusters, flight computer, and power wiring — the spacecraft "skeleton."	Fixed overhead. ~20% of total weight regardless of compute carried.
Compute Payload	GPU chips, memory, and circuit boards — the actual AI processing hardware.	The productive payload everything else exists to support. Lighter chips = lighter satellite.
Thermal Radiators	Panels that radiate waste heat into the cold vacuum of space. Replaces terrestrial cooling systems.	Lightest major subsystem. Space is an excellent heat sink so less hardware is needed.
Laser ISLs	Optical terminals beaming data between satellites to form a mesh network. Inherited from Starlink.	Minimal mass. Replaces heavy terrestrial networking gear (NVSwitch, InfiniBand).
Ground Station Link	Antennas connecting the satellite to Earth during brief overhead passes. Uses existing Starlink ground network.	Small mass. No dedicated heavy ground infrastructure needed per satellite.
Total Satellite	All components combined into one autonomous solar-powered AI compute node.	Every kilogram costs hundreds to thousands of dollars to launch. Industry targets aim to cut current mass by ~75%.

Source: RBC Capital Markets, company reports, Google Suncatcher white paper (2024), Starcloud, Star-catcher, Spacenews.com

Orbital Datacenter: Development Costs (Upfront CapEx to Build & Deploy)

Cost Item	What It Is	Key Driver
Satellite Manufacturing	Factory cost to assemble all hardware — solar panels, batteries, radiators, frame, propulsion, avionics, comms. Excludes GPU chips (counted separately as "IT cost").	Total satellite mass × manufacturing cost per kg. Mass-production (Starlink-style assembly lines) drives costs far below traditional hand-built satellites.
Launch Cost	The rocket ride to get the satellite from Earth's surface to orbit. Covers fuel, vehicle wear, range fees, and operator margin.	Total satellite mass × cost per kg to orbit. Fully reusable vehicles (Starship) could reduce this by an order of magnitude versus expendable rockets.
Compute Hardware (GPUs)	The AI chips, memory, and boards themselves — priced separately like IT equipment in a terrestrial datacenter.	Chip pricing, power efficiency, and radiation tolerance. More efficient chips reduce both IT cost and satellite mass (less power needed = smaller solar arrays and batteries).
Ground Station Network	Global network of dish antennas that upload workloads and download results. The ground network should be dense enough to ensure adequate capacity and latency, while satellites can relay its signal via another satellite to reach the ground station.	Shared across the full fleet and amortized over the useful life that exceeds 5yrs, a lifetime for a Starlink satellite. 400 Ground Stations are in operation today, and the company will need more to ensure the network quality.
R&D / Non-Recurring Engineering	Design, prototyping, testing, and qualification of the satellite platform before mass production begins.	Significant upfront investment (billions for a major program), but amortized across hundreds or thousands of units at scale. Does not appear in per-unit COGS but affects program-level returns.
Insurance (Year 1)	Launch and first-year in-orbit coverage. Typically placed with specialty space insurers (Lloyd's, Marsh).	Runs 5–15% of asset value for launch + early operations. At scale, operators may self-insure to reduce this line item.
Regulatory / Licensing	FCC filings, ITU spectrum coordination, orbital debris compliance, and operating licenses.	Small per-unit at fleet scale but mandatory. Ongoing regulatory environment is evolving and could add cost.

Source: RBC Capital Markets, company reports, Google Suncatcher white paper (2024), Starcloud, Star-catcher, Spacenews.com

Orbital Datacenter: Operating Costs (Annual Recurring to Run the Fleet)

Cost Item	What It Is	Key Driver
Infrastructure Amortization	Spreading the satellite's build-and-launch cost over its useful life (~5 years before radiation degrades it). The orbital equivalent of amortizing a datacenter building over its lease.	Satellite lifespan. Shorter life = higher annual charge. Space radiation limits electronics to roughly 5 years versus 15–25 years for terrestrial buildings.
Satellite Replacement Reserve	Money set aside annually to replace satellites that fail early — from radiation damage, debris impacts, or component defects. Based on observed attrition rates of a few percent per year.	Fleet reliability. Higher failure rates increase the annual reserve. Mature manufacturing reduces early failures over time.
Orbital Maintenance	Mission control operations: monitoring health, firing thrusters to maintain altitude, maneuvering around debris, and updating onboard software.	Shared across the fleet. Automation reduces headcount needed. Larger constellations benefit from economies of scale in operations centers.
Ground Station Operations	Ongoing bandwidth, electricity, and staffing at ground stations. Adequate for inference workloads (small queries up, tokens down); data-heavy training could require more.	Scales with compute capacity and data throughput. Inference-only fleets need less ground bandwidth than training fleets.
Labor / Operations	Mission control engineers (orbital mechanics, spacecraft controllers, anomaly teams) plus a Network Operations Center for workload scheduling and customer provisioning.	Lower per MW than terrestrial datacenters — no HVAC techs, electricians, or physical security guards needed.
Software / Licensing	Standard datacenter orchestration (Kubernetes, schedulers) plus orbital-specific software: attitude control, power management through light/shadow cycles, autonomous fault detection.	Satellites must handle many problems autonomously since ground contact is intermittent. Vertically integrated operators may bundle this into labor.
Insurance (Ongoing)	Annual in-orbit coverage after year one. Premiums decline as the satellite demonstrates healthy operations.	Asset value and track record. Premiums drop significantly after successful first year. Self-insurance possible at scale.
Deorbiting / End-of-Life	Propellant and maneuvers for controlled disposal at end of life (graveyard orbit or atmospheric reentry). FCC mandates a 5-year deorbit rule.	Mostly built into the bus design (propellant reserves), but higher orbits may require incremental fuel or active retrieval.
Radiation Degradation (Capacity Loss)	Not a cash cost, but an economic one — cumulative radiation gradually reduces chip performance, potentially losing 10–20% of effective compute by year 3–4 without outright failure.	Orbit choice (higher radiation in SSO), chip radiation hardness, and whether workloads can tolerate graceful degradation.

Source: RBC Capital Markets, company reports, Google Suncatcher white paper (2024), Starcloud, Star-catcher, Spacenews.com

Orbital Datacenter: Development Costs

A. SpaceX/xAI Disclosures:

- Annual economic cost per watt to launch, build, and operate orbital AI satellites (excluding compute): below \$2/W/year, assuming a 5-year useful life = <\$2M/MW/year.
- Target: ~\$1/W/year as launch costs decline = ~\$1M/MW/year.
- Compute density goal: ~100 kW per metric ton (10 kg/kW).
- Each Starship launch: ~100 metric tons = ~10 MW of compute per launch.
- Starship target cost: ~\$270/kg (10x reduction vs Falcon 9 at ~\$2,700/kg).

B. Third-Party Data (Google Suncatcher, Varda Space, Starcloud):

- At current launch costs (~\$2,700/kg), orbital is ~3x more expensive than terrestrial.
- Break-even requires launch cost <\$200–300/kg AND satellite mfg <\$500/kg.
- Spacecraft mass: ~36 kg/kW (estimated) vs. SpaceX goal of 10 kg/kW.

C. Orbital Development Costs (\$M per MW of Compute)

	Units	Spacex AI1 Satellite (150kW peak/120 kW average) + \$2,700/Kg Launch Cost +\$600/Kg Manufacturing Cost	Spacex AI1 Satellite (150kW peak/120 kW average) + \$270/Kg Launch Cost +\$400/Kg Manufacturing Cost	SpaceX Medium-Term Target (15 kg/kW + Starship reuse)	Long-Term (10 kg/kW + Starship V4)
Mass per kW	kg/kW	36.0	36.0	15.0	10.0
Mass for 1 MW	tonnes	36.0	36.0	15.0	10.0
Launch cost/kg	\$ per kg	\$2,700	\$270	\$270	\$100
Launch cost per MW	\$M	\$97.2	\$9.7	\$4.1	\$1.0
Satellite mfg cost/kg	\$ per kg	\$600	\$400	\$400	\$300
Satellite mfg per MW	\$M	\$21.6	\$14.4	\$6.0	\$3.0
Total CapEx per MW	\$M	\$118.8	\$24.1	\$10.1	\$4.0
Useful life	Years	5	5	5	5
Annualized cost/MW/yr	\$M/MW/yr	\$23.75	\$4.82	\$2.01	\$0.80
Annual cost per watt	\$/W/yr	\$23.75	\$4.82	\$2.01	\$0.80

Source: RBC Capital Markets estimates, company reports, Star Catcher, Luminix AI

Orbital Datacenter: What Orbital Eliminates vs. Adds

Eliminates:

- Electricity cost: \$0 (solar-powered, no grid connection)
- Cooling cost: \$0 (radiative cooling to vacuum of space)
- Land cost: \$0 (orbital slots, no permitting)
- Water: \$0 (no cooling towers)
- BTM generation: \$0 (no gas turbines needed)
- Grid interconnection: \$0 (no substations, no transmission lines)

Adds:

- vs terrestrial fiber (primarily inference workloads)

Key Risks / Caveats:

- SpaceX cost targets require Starship 2nd stage reuse (not yet achieved).
- 10 kg/kW mass ratio requires multi-generation improvements beyond today's technology.
- 5-year satellite life means full relaunch every 5 years (vs 8.5–15yr terrestrial amortization).
- Orbital compute primarily suitable for inference, not large-scale training (bandwidth/latency constraints).
- Radiation environment limits chip reliability and may require rad-hardening (10–100x cost premium) or redundancy (mass penalty).
- At current costs (~\$2,700/kg Falcon 9), orbital is ~3x more expensive than terrestrial.
- Break-even requires: launch cost <\$200–300/kg AND satellite mfg <\$500/kg.

Orbital Datacenter: RBC Assumptions

Near-Term Assumptions Falcon 9 / AI1 Baseline:

- The company expects its AI1 satellite to operate at 150 kW peak / 120 kW average compute in a dawn-dusk SSO LEO orbit, implying a mass of 36 kg/kW. At Falcon 9 launch pricing of \$2,700/kg, 1 MW of orbital compute requires 36 tonnes of satellite mass, generating \$97M/MW in launch costs. At an estimated satellite manufacturing cost of \$600/kg, total non-compute CapEx is \$119M/MW, or \$24/W/yr annualized over a 5-year useful life.

Near-Term Assumptions: Starship Pricing

- Our second near-term scenario applies Starship launch pricing of \$270/kg — derived from SpaceX's stated target of a ~10x reduction vs. Falcon 9 pricing — to the same AI1 satellite mass of 36 kg/kW. Launch cost per MW falls to \$10M. Satellite manufacturing cost at \$400/kg, reflecting assumed production scale efficiencies, yielding development CapEx of \$24M/MW and an annualized amortized cost of \$4.82/W/yr over a 5-year useful life.

Medium-Term Assumptions: 15 kg/kW + Starship Reuse

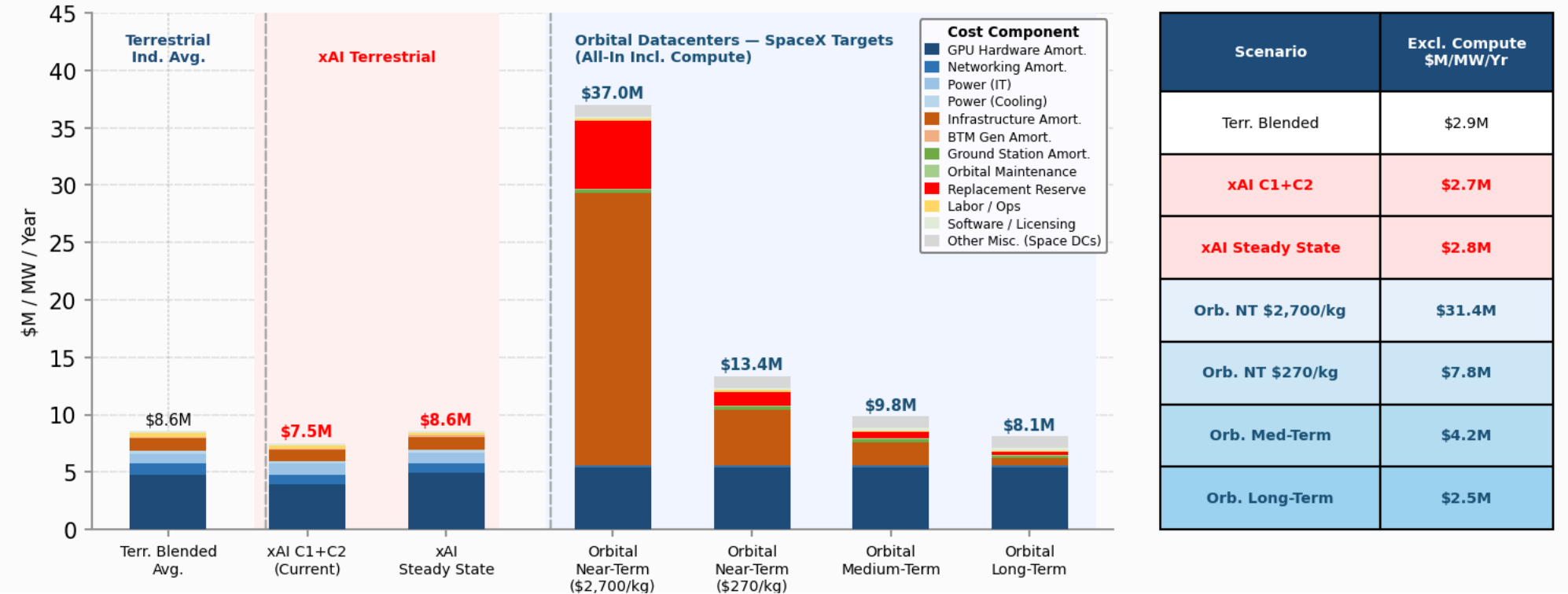
- The 15 kg/kW figure is our linear interpolation between the confirmed AI1 near-term baseline (36 kg/kW) and SpaceX's stated long-term target (10 kg/kW), combining two assumptions: (1) incremental subsystem efficiency improvements over the AI1 baseline, and (2) mass economics enabled by Starship's ~100–150 tonne payload capacity to LEO, which reduces the per-kg mass discipline required on Falcon 9-class vehicles. Subsystem allocations — 4.0 kg/kW solar, 2.5 kg/kW batteries, 1.5 kg/kW thermal, 2.0 kg/kW bus, 5.0 kg/kW compute — reflect assumed improvements in solar cell efficiency, battery energy density, and bus miniaturization. At Starship launch pricing of \$270/kg, 1 MW of compute requires 15 tonnes of satellite mass, generating \$4.1M/MW in launch costs. At an estimated manufacturing cost of \$400/kg, development CapEx could be \$10M/MW, or \$2.01/W/yr annualized amortized cost over a 5-year useful life.

Long-Term Assumptions: 10 kg/kW + Starship V4

- The 10 kg/kW figure is based on the company's long-term target of approximately 100 kW of compute power per ton. SpaceX does not provide a subsystem-level breakdown of this figure, nor a timeline for achieving it. Subsystem allocations — 2.5 kg/kW solar, 1.5 kg/kW batteries, 1.0 kg/kW thermal, 1.0 kg/kW bus, 4.0 kg/kW compute — are our estimates reflecting next-generation accelerator compute density trends and aggressive thermal/radiator mass reduction. Assuming Starship V4 launch pricing of \$100/kg — derived from SpaceX's stated long-term goal of approaching the cost of fuel per kg to orbit — 1 MW of compute requires 10 tonnes of satellite mass, generating \$1.0M/MW in launch costs. At an estimated manufacturing cost of \$300/kg, development CapEx could be \$4.0M/MW, or \$0.80/W/yr annualized amortized cost over a 5-year useful life.

Orbital Datacenter: Operating Costs (\$M per MW/year)

- Orbital datacenters have a fundamentally different operating cost structure: NO electricity cost (solar-powered), NO cooling cost (radiative cooling to space), NO land cost, NO water, NO permitting.
- xAI Orbital Datacenters — Key Takeaway**
 - Near Term:** \$13.4M/MW/yr total COGS (\$7.8M excl. compute & networking) — driven by infrastructure amort. (\$5.08M), replacement reserves (\$1.21M), and Other non-compute expenses (\$1.47M).
 - Mid-Term Target (15 kg/kW + Starship reuse):** \$9.8M/MW/yr (\$4.2M excl. compute & NW) — 26% reduction from near term
 - Long Term:** \$8.09M/MW/yr (\$2.49M excl. compute & NW) — approaches terrestrial industry average of \$8.6M while eliminating power and cooling costs entirely
- Orbital advantage:** \$0 electricity/fuel, \$0 cooling, \$0 BTM generation — offset by infrastructure amort., ground station network, orbital maintenance, and replacement reserves.



Source: RBC Capital Markets estimates, company reports, Star Catcher, Luminix AI

Where could estimated orbital costs exceed SpaceX's stated target

Overview

- SpaceX's disclosed target of <\$2.00/W/yr represents a long-term marginal cost figure at full operational maturity, assuming millions of satellites, Starship at rapid reuse cadence, and mature satellite designs.
- Our estimates for orbital datacenter operating costs (excluding compute and networking) at \$31/W/yr near term (based on \$2,700/Kg launch cost), \$4.2/W/yr at target state, versus SpaceX's disclosed claim of <\$2.00/W/yr.
- The gap is 16× on a near-term basis and 2.1× under optimistic assumptions.

Three Drivers of the Gap

- **Satellite Mass Density (kg/kW):** Near-term model uses 36.0 kg/kW (per SpaceX's AI1 satellite: 150 kW peak, 120 kW average, dawn-dusk SSO); target case uses 15 kg/kW. SpaceX's <\$2.00 claim implicitly requires ~10–12 kg/kW.
- **Launch Cost:** Model assumes \$270/kg (Starship target); We believe SpaceX's claim likely requires \$50–100/kg, implying full rapid reusability at high cadence. Current Falcon 9 operates at ~\$2,700/kg. Starship has not yet completed a commercial payload delivery mission.
- **Manufacturing Cost at Scale:** Model uses \$400/kg (reflecting vertical integration vs. \$600/kg third-party estimate). SpaceX's claim likely assumes \$200–300/kg at millions-unit automotive-style production.

Additional Costs Captured That SpaceX's Claim May Exclude

- **Replacement reserve:** Modeled at 5% of infrastructure CapEx/yr: \$5.94M/MW/yr (Falcon 9), \$1.21M/MW/yr (Starship near term), \$0.50M/MW/yr (medium term), \$0.20M/MW/yr (long term); SpaceX's \$2/W/yr claim likely assumes 2–3% at mature scale.
- **Ground station amortization:** \$0.25M/MW/yr; SpaceX likely treats as marginal cost on existing Starlink infrastructure.
- **Orbital maintenance:** \$0.10M/MW/yr; likely lower at scale in SpaceX's assumptions.
- **Insurance and radiation degradation:** Not modeled; SpaceX likely self-insures and absorbs degradation into replacement assumptions.
- **Other Miscellaneous Costs:** Modeled at \$1.0M/MW/yr, capturing spectrum licensing, regulatory compliance, orbital slot coordination, and FCC/ITU filing costs.

Reconciliation: How SpaceX Could Reach <\$2.00/W/yr

- We believe SpaceX's long-term orbital datacenter economics could be achievable within a 5–10-year horizon.
- Satellite mass efficiency of 10–12 kg/kW could be reached consistent with GPU power density trends; Starship payload costs of \$50–100/kg could be attainable as reuse rates scale toward commercially demonstrated levels; manufacturing costs of \$200–300/kg could be achievable as production volumes scale; attrition rates could decline to 2–3% as design matures and operational learning accumulates; and ground infrastructure, self-insurance, and constellation automation costs could reduce to marginal levels as the Starlink network amortizes.
- Should these parameters converge, we estimate EBITDA-based ROIC of 52%–103% and EBIT payback of 3.9–2.0 years at \$16M/MW capex — outcomes that could outperform terrestrial configurations.

Networking, Laser Links & Connectivity in Orbital Datacenters

Bandwidth Is the Binding Constraint for Orbital Datacenters

Why Connectivity Matters

- ODCs solve the power problem but introduce a bandwidth constraint that terrestrial datacenters never face — every byte moving between Earth and orbit must traverse optical or radio links with finite capacity and non-trivial latency.
- On-board AI processing reduces the volume of data traversing space-to-ground links by up to 85–90%, making in-orbit compute the primary mechanism for managing this constraint.
- The architectural choice between centralized and distributed ODC platforms is directly determined by networking maturity. Distributed constellations require operational optical ISLs — a capability only SpaceX (Starlink) and Amazon (Kuiper) have operationalized at scale as of mid-2026. All other players must route through one of these networks or remain centralized.

How Optical Inter-Satellite Links (ISLs) Work

- ISLs use focused laser beams to transmit data between satellites at the speed of light through vacuum. Light travels ~30% faster through vacuum than through glass fiber, giving LEO laser links a latency advantage over fiber for long-distance routes exceeding ~3,000 km.
- Satellites must maintain precise pointing alignment (high structural stiffness, active attitude control); crosslinks must support >10 Gbps throughput for distributed compute coordination.
- ISLs have proven harder to operationalize than anticipated. Only SpaceX Starlink and Amazon Kuiper have achieved this at scale, giving both companies a structural networking advantage that cannot be easily or quickly replicated.

Current & Planned Link Speeds by System

System / Company	Current Capability	Planned / Target	Status
Kepler Communications (Tranche 1)	2.5 Gbps per optical link; 10 Gbps max across constellation	100 Gbps (Tranche 2 / second generation)	Operational (commissioned March 16, 2026)
Google Project Suncatcher	Lab demonstration only	1.6 Tbps (single transceiver pair, lab-demonstrated)	R&D; prototype satellites planned 2027
SpaceX Starlink (current generation)	3 laser ISLs at up to 200 Gbps aggregate	1 Tbps (next generation); petabit-level aggregate for ODC constellation	Operational at scale; ODC constellation pending FCC approval
Blue Origin TeraWave	Up to 6 Tbps space-to-ground	SpaceX states future Starlink ground laser links will exceed TeraWave	Operational (competitive benchmark)
Starcloud (Starcloud-2 and beyond)	Not yet operational in orbit	Routes data through Starlink, Amazon Kuiper, and Blue Origin TeraWave laser-linked constellations	Planned (October 2026 launch)

Source: RBC Capital Markets, company reports, Star Catcher, Luminix AI

Latency Profile & Ground-to-Orbit Architecture

Latency — Space vs. Ground

- **GEO vs. LEO:** Geostationary (~36,000 km) introduces ~500 ms round-trip delays — impractical for real-time compute. LEO ODCs (400–1,200 km) reduce round-trip latency to 20–40 ms, comparable to terrestrial broadband and demonstrated at scale by SpaceX Starlink.
- **LEO vs. Fiber:** For routes exceeding ~3,000 km, LEO optical laser links are faster than fiber (~30% speed-of-light advantage through vacuum vs. glass). This makes ODCs viable for specific long-haul routing use cases.
- **Latency Limits:** Workloads requiring constant high-bandwidth interaction with terrestrial databases face a latency penalty no optical link can eliminate. ODCs perform best where raw data originates in space or benefits from continuous proximity processing, and only compressed insights need to reach ground stations.

Capacity Comparison

Environment	Base Rate (per channel)	Effective Multiplier	Total Capacity	Key Constraint
Terrestrial Fiber	~50–115 Gbps	~1–1.5M×	~50–170 Pbps	Cost only — extra strands are near free
Subsea Cable	~50–110 Gbps	~5K–50K×	~250–5,500 Tbps	Distance and noise; repeaters every 50–100 km
Ground-to-Orbit (Optical)	~0.5–7 Gbps	~2–50×	~1–350 Gbps	Atmospheric turbulence, cloud blockage, beam tracking
Inter-Satellite (Optical ISL)	~8–40 Gbps	~24–500×	~200 Gbps–20 Tbps	Terminal power and mass; pointing precision; radiation-hardened components

Technique-by-technique comparison

Technique	Terrestrial Fiber	Subsea	Ground-to-Orbit	Inter-Satellite	Multiplier
WDM (wavelength channels)	80–96×	80–96×	1–4×	4–8×	80–96× terrestrial/subsea; limited in space
Polarization multiplexing	2× (standard)	2× (standard)	Not viable — atmosphere scrambles polarization	Under research — vacuum preserves polarization	2× in fiber/vacuum only
QAM modulation	64–128-QAM (6–8×	16–64-QAM (4–6×	QPSK only (2×	QPSK to 16-QAM (2–4×	3× fiber advantage over G2O
Strand/terminal parallelism	100–1,000+ strands	8–24 fiber pairs	Limited by ground station count	3–4 terminals per satellite	1,000×+ terrestrial
Coherent detection + DSP	Standard	Standard + more FEC	Emerging (mostly direct detection today)	Emerging	~2–3×

Source: RBC Capital Markets, company reports, Star Catcher, Benzinga, Associated Press, Space News

Ground-to-Orbit Architecture & Workload Suitability — Networking Lens

Ground-to-Orbit Architecture

- SpaceX ODC system (per January 2026 FCC filing) connects ODC satellites to Starlink via high-bandwidth optical links; Starlink relays data to ground stations via laser mesh network. ODC satellites carry Ka-band for TT&C. Starlink serves as the backhaul backbone for SpaceX's ODC constellation.
- Starcloud's Starcloud-2 plans to route through multiple networks (Starlink, Amazon Kuiper, Blue Origin TeraWave), providing redundancy and competitive optionality rather than relying on a single backhaul provider.
- Orbit AI's Genesis-1 (launched December 2025) processes infrared remote sensing data on-orbit before downlinking, reducing critical information retrieval from hours to seconds and cutting transmission bandwidth costs by over 90%.

AI Satellite vs. Starlink — Key Design Difference

- Starlink satellites require costly phased-array antennas for ground communication and must maintain stable orientation at all times. AI-optimized ODC satellites communicate only with neighboring satellites via laser links — eliminating phased-array hardware, reducing manufacturing cost, and freeing mass budget for compute payloads and solar arrays.
- Starcloud CEO Philip Johnston estimates Starcloud's design will cost less than \$5/W (excluding GPU costs), compared to ~\$22/W for current-generation Starlink and ~\$32/W for original Starlink design. Removal of ground-pointing hardware is the primary driver of cost differential.

Workload Suitability		
Workload Type	Suited for ODC?	Reason
Earth observation / sensor analytics (wildfire, maritime, ISR)	Yes	Raw data originates in space; only compressed insights need downlink; reduces bandwidth by up to 90%.
AI inference (voice agents, LLM queries)	Yes (near-term focus)	Starcloud CEO: "Almost all inference workloads will be done in space." Lower bandwidth requirement than training.
AI training	Limited near term	High inter-node bandwidth needed; depends on operational ISLs at >10 Gbps. Google's 1.6 Tbps lab demo is the target benchmark.
Real-time transactional workloads (e.g., financial systems)	No	Latency penalty to ground is fundamental and cannot be eliminated by optical links alone.
Long-haul data routing (>3,000 km point-to-point)	Yes	Laser-over-vacuum is ~30% faster than fiber for very long routes; structural latency advantage.
Sovereign/secure compute (geopolitically sensitive data)	Yes	Physical access requires a spacecraft; off-planet storage is resilient to terrestrial legal/regulatory seizure.

Source: RBC Capital Markets, company reports, Star Catcher, Benzinga, Associated Press, Space News

Orbital Datacenters – Networking and Bandwidth related Key Risks

- **Bandwidth gap (structural)** — ISL speeds are $\sim 5,000\times$ below NVLink bisection bandwidth. Distributed AI training would require ~ 10 Tbps aggregate ISL — $20\text{--}400\times$ above current operational speeds. The absolute fiber capacity lead will likely widen further as 800G and 1.6T terrestrial channels roll out over the same planning horizon.
- **Atmospheric G2O constraint (physics-limited)** — When you send a laser beam from the ground up to a satellite, it has to travel through the atmosphere — wind, heat, humidity, and air turbulence constantly bend and distort the beam in unpredictable ways. This makes the received signal noisy and unstable.
 - The way you pack more data onto a laser beam is by using more complex signal patterns — essentially encoding more bits into each pulse of light. On fiber, the signal travels through a controlled glass tube with no interference, so you can use very complex patterns (64-QAM) that carry 6 bits per pulse. On a ground-to-space laser link, the atmospheric distortion means you can only reliably use a much simpler pattern (QPSK) that carries just 2 bits per pulse. That is a $3\times$ reduction in how much data you can send per pulse — and it is a fundamental physics constraint, not something better hardware can fix. Also, the WDM technique allows it to send wavelengths of different colors that achieves $80\text{--}96\times$ in data speeds in terrestrial fiber, but only up to $4\times$ in ground-to-orbit settings.
 - Another technique called **polarization multiplexing**, which essentially sends two data streams on the same beam simultaneously by spinning the light in different directions. This works well in fiber and in space between satellites. But the atmosphere scrambles the light's polarization in real time, making this technique impractical on the ground-to-space link without correction systems that do not yet exist commercially. Google's 2027 prototype satellites are simply using traditional radio (RF) links to the ground rather than trying to solve this problem — which tells you where the technology actually stands today.
- **Radiation & thermal degradation** — Satellites in low Earth orbit are continuously bombarded by charged particles from the Sun and cosmic rays. Over a 5-year satellite life, the cumulative radiation dose could reach $5,000\text{--}30,000$ rads — enough to gradually damage the electronics, particularly the memory chips inside the GPUs, which are the most sensitive components.
 - The other problem is through special panels called radiators. Over time, radiation exposure degrades the surface properties of these panels — they become less effective at releasing heat, potentially by around 40% by Year 5. If the satellite cannot shed heat as efficiently, it has to run the GPUs at lower power to avoid overheating. That directly reduces compute output in the satellite's later years.
 - The deeper issue is that once a satellite is launched, nothing can be fixed or replaced. If a GPU fails or a radiator panel degrades, it stays in heat. In space there is no air, so the only way to get rid of the heat generated by GPUs is to radiate it away that way for the rest of the satellite's life. Every terrestrial datacenter can swap out a failed server in hours. No orbital operator currently has any plan to service or repair satellites in orbit.
- **Economics unproven at scale** — SpaceX's stated $<\$2/\text{W}/\text{yr}$ target requires $\sim 10\text{--}12$ kg/kW mass efficiency (vs. ~ 36 kg/kW today), Starship at $\$50\text{--}100/\text{kg}$ (not yet demonstrated commercially), and satellite manufacturing at millions of units per year (current rate: $\sim 4,500/\text{yr}$).

Orbital Datacenters – Other Caveats

Throughput Derating in Orbit.

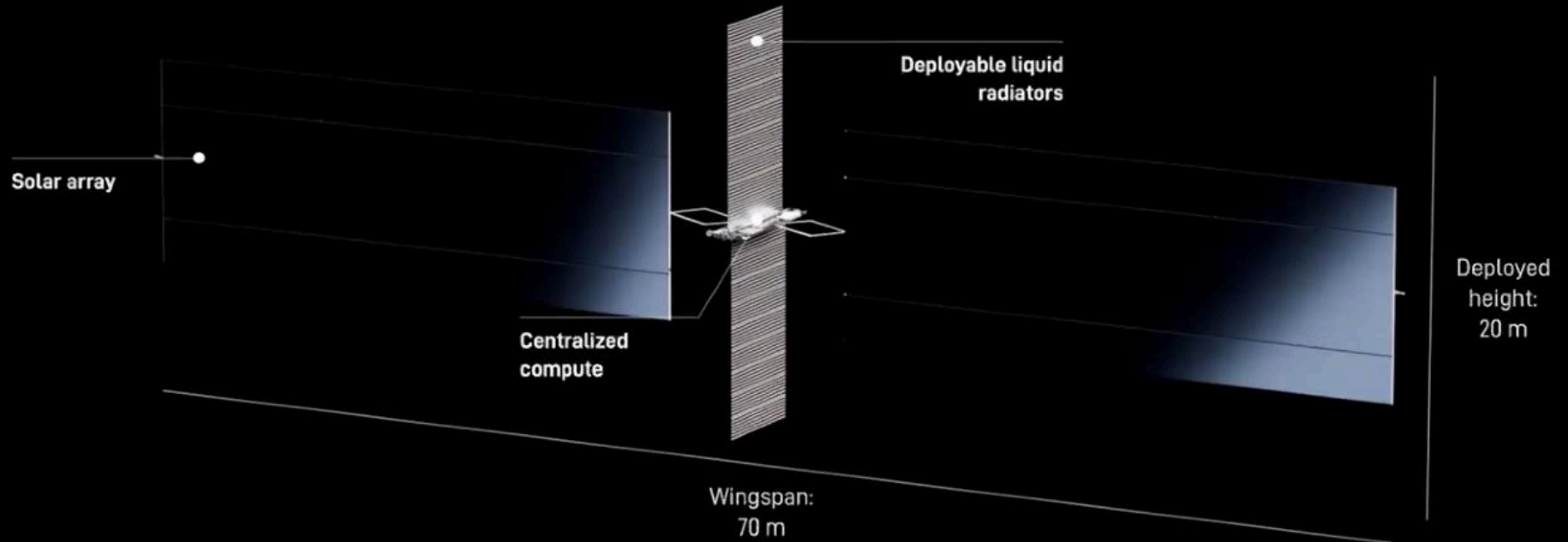
- In-orbit TPS is estimated at 50–70% of terrestrial lab TPS, driven by three compounding factors. These are modeled estimates; actual performance will depend on hardware maturity, orbital conditions, and workload mix.
 - **Factor 1 — No high-speed GPU interconnect within the satellite (~75–90% of terrestrial TPS):** In a terrestrial datacenter, GPUs connect via NVLink switches at very high speed. In a satellite, GPUs communicate over PCIe (~14× slower per link). For single-GPU inference the penalty is modest (15–30% reduction); for multi-GPU collaborative workloads it is more severe.
 - **Factor 2 — Operating constraints in orbit (~85–95% of peak):** The SpaceX AI1 satellite uses 110 m² of deployable liquid-loop radiators (custom design; prior SpaceX satellites used chassis as radiator). IEEE Spectrum analysis suggests a single 700W GPU requires approximately 1.4 m² of radiator area at 60°C. At 120 kW average compute load across 60+ GPUs, thermal headroom is limited — some GPUs may run at reduced power at certain orbital positions during solar angle transitions or peak thermal stress. SpaceX's S-1 notes fleet management software will "reallocate traffic to other satellites and prevent cluster-level downtime," implying spare GPU capacity held in reserve at any time.
 - **Factor 3 — Radiation degradation over 5-year life (~80–90% lifetime average):** Cumulative radiation dose in LEO could reach 5,000–30,000 rads over operational life (Georgia Tech, May 2026). Google's proton beam testing of its Trillium TPU at UC Davis showed HBM memory irregularities at roughly 3× the expected 5-year shielded dose. Error correction software (TMR, ABFT) running continuously could consume 6–25% of available compute cycles. Radiator emissivity could degrade approximately 40% after 5 years of radiation exposure (IEEE Spectrum, June 2026), potentially forcing power reductions in later years.
 - **Upside scenario:** Under favorable conditions — dawn-dusk SSO orbit below 600 km, batch-heavy FP8 inference, and purpose-designed radiation-hardened chips from SpaceX's TeraFab program (not yet in production) — the composite could improve to roughly 67–75%.

On-orbit preprocessing could be the primary mechanism for managing the ground-to-orbit bandwidth constraint.

- Orbit AI's Genesis-1 (launched December 2025) demonstrated a 90%+ reduction in transmission bandwidth requirements by processing infrared remote sensing data on-orbit before downlinking. ODC architecture could specifically rely on this model: processing in space could reduce bytes that must traverse the space-to-ground link, and only compressed insights need to reach ground stations. This self-reinforcing design principle — run inference in orbit, downlink only results — could align with the bandwidth physics of the ground-to-orbit link.

SpaceX – AI1 Satellite

AI1 satellite



Overall specs

150 kW peak / 120 kW average compute payload
70 kW/ton
Compute provider interchangeable

Thermal

110 m² deployable liquid radiator
Redundant pumping loops
Integrated micrometeor shielding

Solar

150 kW Solar Array
250 W/m²
SpaceX manufactured solar technology from Bastrop, TX

SpaceX – AI1 Satellite

Design Philosophy: What Goes to Space

A space datacenter is not a terrestrial datacenter building placed in orbit. The design strips a datacenter down to three essential functions — **compute**, **power delivery**, and **heat removal** — and solves each with space-native approaches. The engineering problem reduces to: deliver power to the compute, and move waste heat away into the vacuum of space.

AI COMPUTE

~72 GPUs (GB300-equivalent)
~150 kW peak / ~120 kW sustained
Equivalent to one GB300 NVL72 rack
No large comm payloads required

SOLAR POWER

~250 W/m² solar array
~480–600 m² array area (for 120–150 kW)
Evolution of Starlink solar technology
Improvement expected beyond 250 W/m² over time

THERMAL (RADIATORS)

~1,400 W/m² radiator capacity
Double-sided, oriented knife-edge to sun
~86 m² radiator area (for 120 kW)
Evolution of Starlink thermal systems

CONNECTIVITY

~1 Tbit/s laser-link connectivity
Sat-to-sat laser links
Direct link to Starlink constellation
Laser links to ground stations

AI One Key Specifications

~150 kW peak / ~120 kW sustained

Power capacity — comparable to one GB300 NVL72 rack (peak ~140 kW, operating ~120 kW)

~250 W/m²

Solar array power density — achievable with current technology, improvement expected over time

~1,400 W/m²

Radiator thermal rejection — double-sided, knife-edge to sun orientation

~1 Tbit/s

Laser-link connectivity — sat-to-sat + Starlink integration + ground links

~600–800 km altitude

Orbital altitude — one-way latency ~2–3 ms (light travels ~300 km/ms)

Simpler than Starlink

No phased-array antennas, no parabolic antennas, no large comm payloads — primarily solar + radiators + compute + lasers

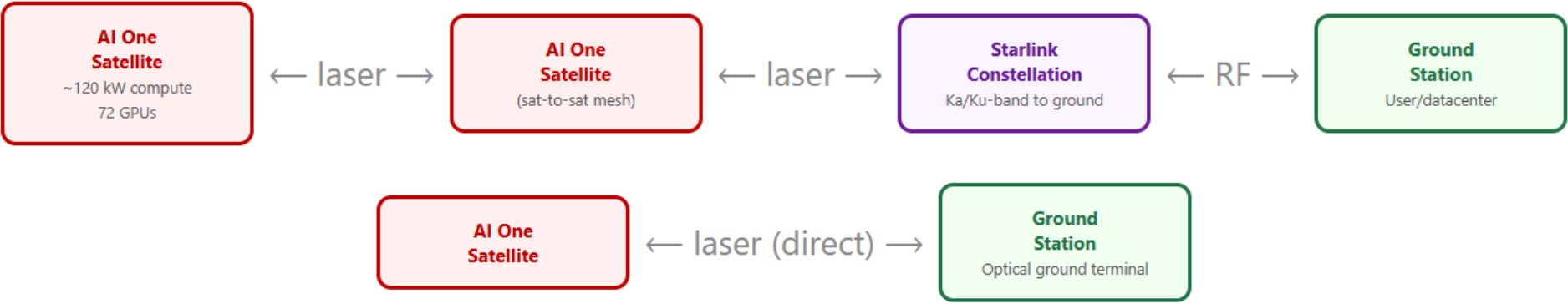
Source: RBC Capital Markets, company reports, Star Catcher, Luminix AI

SpaceX – AI1 Satellite

Complexity Comparison: AI One vs. Starlink

Component	Starlink Satellite	AI One Satellite
Large phased-array antennas	✓ Required (Ku/Ka-band user beams)	✗ Not required
Parabolic antennas	✓ Required (gateway links)	✗ Not required
Laser inter-satellite links	✓ Required	✓ Required (~1 Tbit/s)
Solar arrays	✓ Required (smaller)	✓ Required (larger, ~480–600 m ²)
Radiators	✓ Required (smaller)	✓ Required (larger, ~86 m ² , double-sided)
AI compute payload	✗ Not present	✓ Core payload (~72 GPUs)
Complex RF communications payload	✓ Primary function	✗ Not required
Net complexity	Higher	Lower (fewer subsystems)

Connectivity Architecture: Data Movement Path



Alternative path: direct laser link to ground (bypasses Starlink relay)

Source: RBC Capital Markets, company reports, Star Catcher, Luminix AI

SpaceX – AI1 Satellite

Latency Profile		
Parameter	Value	Implication
Orbital altitude	~600–800 km	Low Earth Orbit (LEO)
Speed of light in vacuum	~300 km/ms	—
One-way propagation delay	~2–3 ms	Comparable to cross-continental fiber
Round-trip (sat + ground link)	~4–6 ms (propagation only)	Within latency SLO for most inference workloads

Technology Readiness		
Subsystem	Technology Basis	Readiness
Solar arrays (~250 W/m ²)	Evolution of Starlink solar technology	Achievable (in production variant)
Radiators (~1,400 W/m ²)	Evolution of Starlink thermal management	Achievable (design validated)
Laser inter-sat links (~1 Tbit/s)	Starlink laser link technology (deployed at scale)	Operational on Starlink fleet
Compute	Possibly Dojo3 / D3	Initially paused in 2025 but revived following a stabilization of the Tesla AI5 chip design
Spacecraft bus	Starlink manufacturing platform	6,000+ satellites in orbit

Key design principle: AI One does not require any technical capability that does not currently exist. Every subsystem is an evolution of technology already built or in development for the Starlink constellation. The satellite is described as **simpler than a Starlink satellite** — it removes the complex RF communications payload (phased arrays, parabolic antennas) and replaces it with a GPU compute payload, solar arrays, and radiators.

Orbital Datacenter Ecosystem: Companies Enabling Space-Based Compute (1 of 2)

Technology Layer	Company	Role / Contribution	Status
AI Compute Hardware	Nvidia	Space-1 Vera Rubin Module (25× more AI compute vs. H100 for space inferencing); IGX Thor and Jetson Orin platforms for edge computing in orbit; radiation-approved THOR chip; investor in Starcloud	Announced GTC Mar 2026; Vera Rubin samples shipped to customers
	Google (Alphabet)	Custom Trillium TPU chips radiation-tested for space; Project Suncatcher 81-satellite constellation design	Prototype satellites with Planet Labs targeted early 2027
	SpaceX / xAI (Terafab)	In-house chip manufacturing initiative (with Tesla and Intel) to reduce processor costs and supply risk; open architecture allowing any GPU/TPU	Under development; Terafab announced 2026
Orbital DC Developers (Integrated)	SpaceX	End-to-end: launch (Starship), satellite manufacturing, Starlink ground network, compute operations; FCC filing for up to 1M satellites; 100 GW/year target	AI1 satellite in development; terrestrial Colossus clusters operational
	Starcloud	First to operate Nvidia H100 GPU in orbit (Nov 2025); ran Google Gemini model in space; FCC filing for 88,000-satellite constellation; Nvidia-backed; \$1.1B valuation	Series A (\$170M, Mar 2026); planning 200 kW-rated satellites
	Blue Origin	Project Sunrise: 51 data center satellites filed (Mar 2026); gigawatt-scale orbital facilities envisioned; New Glenn launch vehicle	Deployment targeted late 2027
	Aetherflux / Cowboy Space Company	Orbital compute + power-beaming (IR laser to Earth); "Galactic Brain" datacenter node; backed by Index Ventures, a16z, Breakthrough Energy	First compute node targeted Q1 2027
	Orbital	Purpose-built AI inference satellites (100 kW each); 10,000-satellite constellation goal (10 GW); designing around Nvidia Space-1 Vera Rubin architecture	\$5M pre-seed (Jun 2026); Pathfinder mission 2027
	LoneStar Data Holdings	Orbital facilities by 2028; lunar surface data centers by early 2030s	In development
Satellite Bus / Spacecraft Platform	Muon Space	Condor-Ultra: Starship-class platform (20 kW initial, 100 kW variants); integrates Nvidia Vera Rubin; 100 Gbps optical mesh; vertically integrated	First pathfinder delivery 2028
	Terran Orbital (Lockheed Martin)	Nebula Bus; proliferated LEO constellation manufacturing; vertically integrated production at scale	Active; supplying SDA Tranche 0–3
	Planet Labs	Satellite manufacturing partner for Google Project Suncatcher; building two prototype orbital compute satellites; integrating Nvidia IGX Jetson Thor	Prototypes targeted early 2027
	MDA Space	Vertically integrated to chip-level; developing increased onboard processing at the edge; applicable to orbital compute, comms, and servicing	Active development
	Intuitive Machines	60 kW power-generating spacecraft (roadmap to 100 kW); thermal management for high-density computing; ground segment and data relay network	Offering orbital DC capabilities to partners

Source: RBC Capital Markets, company reports, Nvidia, Google Project Suncatcher (Nov 2025), Muon Space Condor-Ultra press release (Jun 2026), Rocket Lab, TechCrunch and Yahoo Finance, Spacenews.com

Orbital Datacenter Ecosystem: Companies Enabling Space-Based Compute (2 of 2)

Technology Layer	Company	Role / Contribution	Status
Laser Communications (ISLs)	Rocket Lab / Mynaric	CONDOR Mk3 optical communication terminals; acquired Mynaric for \$155M (Apr 2026); scaling production to hundreds/thousands for constellations	Operational; 150+ optical heads manufactured
	SpaceX (Starlink)	Proprietary laser ISLs across 9,600+ Starlink satellites; 800 Gbps+ optical links; providing Starlink Mini Lasers to third-party platforms (e.g., Muon)	Operational at scale
Solar Arrays & Power	Rocket Lab	STARRAY family of customizable, scalable solar arrays for satellite missions; acquiring Motiv Space Systems for solar array drive assemblies (SADAs)	Production
	Mantis Space	Orbital power infrastructure startup: "delivering sunlight anywhere in orbit"; solving power constraints for compute, defense, and broadband satellites	\$10M+ seed (Mar 2026); stealth exit
	SpaceX (in-house)	Proprietary high-efficiency solar arrays for Starlink V2/V3 and AI satellites; 5x more energy per unit area vs. terrestrial solar	Operational
Thermal Management	Muon Space	Vertically integrated thermal/radiator subsystems designed for high-density compute heat loads on Condor-Ultra platform	In development
	SpaceX (in-house)	Purpose-built deployable liquid-loop radiators for AI satellites; radiative cooling to deep space	Under development for AI1 satellite
	Intuitive Machines	Thermal management and heat rejection expertise for edge computing in orbit; highest-power spacecraft thermal systems	Offering to partners
Launch Vehicles	SpaceX (Starship)	100+ tonnes to LEO; targeting <\$200/kg; rapid reuse; stackable satellite deployment; ~200 launches per GW at current mass ratios	Flight testing; commercial ops ramping
	Rocket Lab (Neutron)	Medium-lift vehicle for constellation deployment; alternative to Starship for smaller payloads	Under development
	Blue Origin (New Glenn)	Heavy-lift; reached orbit 2025; partial reuse; launching own Project Sunrise satellites	Operational (first orbital flight 2025)
Cybersecurity / Hardening	SEALSQ Corp	Post-quantum semiconductor infrastructure; quantum-resilient security for orbital AI assets and space environments	Positioning announced May 2026
	OrbitsEdge	Hardened computing hardware for radiation and extreme temperatures; targeting satellite operators needing in-orbit processing	First orbital demo planned 2026
Ground Network & Connectivity	SpaceX (Starlink)	9,600+ satellites; global ground station network; provides always-on 25 Gbps connectivity to third-party orbital compute platforms	Operational at scale
	Kepler Communications	Largest orbital compute cluster currently operational (40 Nvidia Orin processors across 10 satellites linked by laser); 18 customers	Operational (Jan 2026)

Source: RBC Capital Markets, company reports, Nvidia, Google Project Suncatcher (Nov 2025), Muon Space Condor-Ultra press release (Jun 2026), Rocket Lab, TechCrunch and Yahoo Finance, Spacenews.com

Orbital Datacenters - What Is Operational Today

Company	Asset	Partners	Launch Date	Purpose
Kepler Communications	Tranche 1 optical relay constellation with distributed on-orbit compute	NVIDIA	Jan 2026; commissioned Mar 16, 2026	First commercially operational optical data relay network; 40 NVIDIA Jetson Orin modules across 10 satellites
Starcloud	Starcloud-1 satellite	NVIDIA	Nov 2, 2025	Proof-of-concept orbital compute payload with NVIDIA H100 GPU; first LLM trained in orbit (Dec 2025)
Lonestar Data Holdings	Lunar datacenter demo hardware	Intuitive Machines, SpaceX	Early 2025	Miniature datacenter to validate lunar storage/security use cases (not an orbital satellite)

Orbital Datacenters - Near-Term Planned Launches

Company	Planned Asset	Partners	Expected Launch	Purpose
Starcloud	Starcloud-2	NVIDIA	2027	Multi-GPU cluster (H100s / Blackwell) with GPU cluster, persistent storage, proprietary thermal/power systems; first commercial mission
OrbitsEdge	Edge1 orbital demonstration	HPE, Above Orbital Inc, Syntilay	2026	Test radiation-shielded edge-compute hardware for future ODCs
EDGX	Sterna DPU on SpaceX Transporter-16	ESA, Belgian MOD	Launched Apr 9, 2026	AI-powered Edge compute for satellite constellations; Compute-as-a-Service in orbit; 7-year design life
SpaceX / xAI	Prototype orbital compute satellites	xAI / SpaceX	2027	Two prototype satellites to test AI datacenter concepts in orbit
Google / Planet Labs	Project Suncatcher prototype	Planet Labs	2027	Test Google TPUs in orbit; demonstrate high-bandwidth ISLs; assess radiation tolerance and thermal stability
Aetherflux / Cowboy Space Corporation	First test satellite ("Galactic Brain")	SpaceX	2027	Free-flying constellation nodes; solar-powered orbital AI compute

Source: RBC Capital Markets, company reports, Starcloud, Spacenews.com

Datacenter ROIC Analysis



Capital
Markets

ROIC and Payback - xAI Leads on Returns; Orbital Returns Cost-Dependent

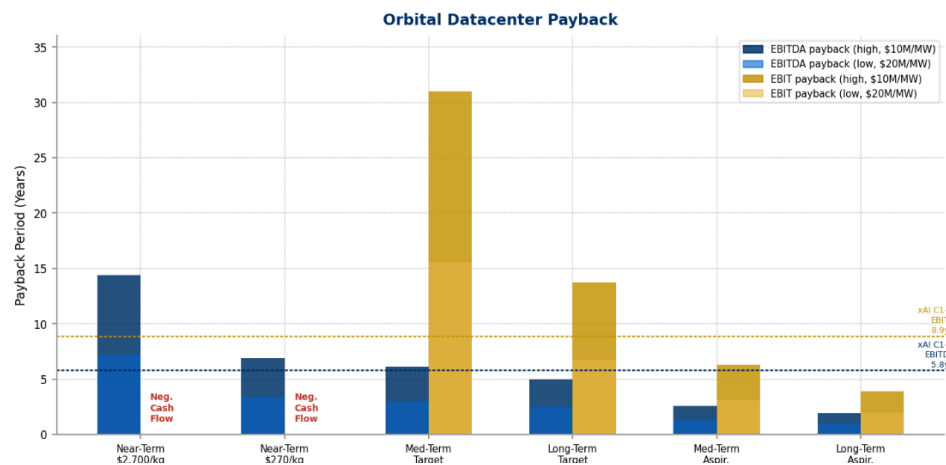
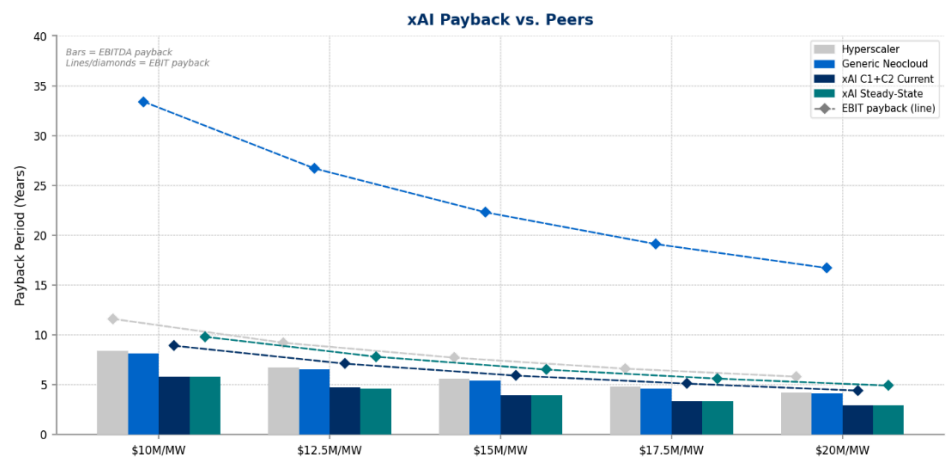
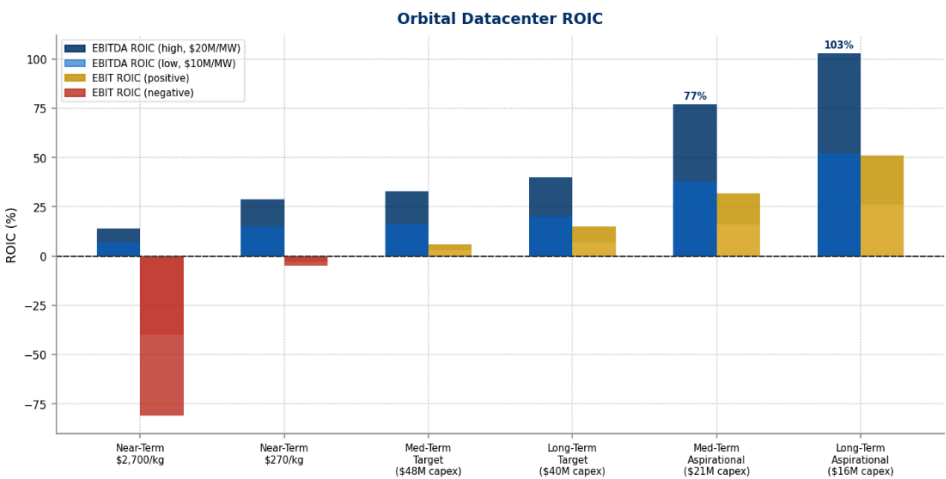
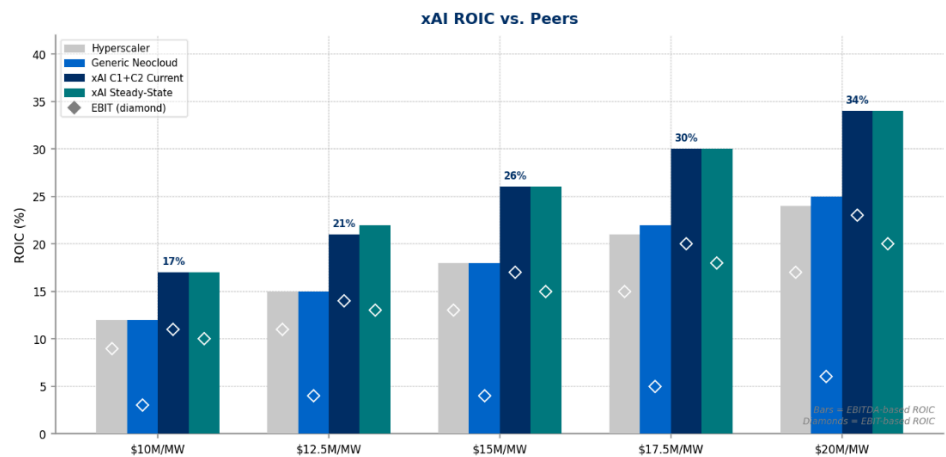
xAI ROIC vs. Peers. We estimate xAI C1+C2 Current EBITDA-based ROIC at 17%–34% across \$10M–\$20M/MW revenue, approximately 500bps above both the Hyperscaler (12%–24%) and Generic Neocloud (12%–25%), driven by lower Compute COGS of 24% versus 31% and 50% of revenue, respectively. Under revised steady-state assumptions, we estimate xAI EBITDA-based ROIC at 17%–34%, broadly in line with C1+C2 Current, as Compute COGS is revised to 29% (from 38%) and D&A to 35% (from 47%) of revenue. On an EBIT basis, we estimate xAI C1+C2 Current at 11%–23% and steady-state at 10%–20%, compared to 9%–17% for the Hyperscaler and 3%–6% for the Generic Neocloud.

xAI Payback vs. Peers. We estimate xAI C1+C2 Current EBITDA-based payback at 5.8–2.9 years across \$10M–\$20M/MW revenue, approximately 2.5 years shorter than the Hyperscaler (8.4–4.2 years) and the Generic Neocloud (8.1–4.1 years). Under revised steady-state assumptions, we estimate xAI EBITDA-based payback at 5.8–2.9 years, in line with C1+C2 Current, with EBIT-based payback of 9.8–4.9 years versus 8.9–4.4 years for C1+C2 Current — a narrower gap than previously estimated, reflecting the lower D&A assumption.

Orbital Datacenter ROIC. We estimate orbital datacenter EBITDA-based ROIC could be competitive with terrestrial configurations under medium- to long-term cost targets, ranging from 16%–32% at long-term parameters (10 kg/kW, \$100/kg launch, \$300/kg manufacturing). At current near-term launch costs of \$2,700/kg, EBIT based ROIC is negative. Returns could improve materially as launch and manufacturing costs decline: at long-term target parameters (10 kg/kW, \$100/kg launch, \$300/kg manufacturing, \$40M/MW capex), we estimate EBITDA-based ROIC of 20%–40%, and under aspirational capex assumptions of \$16M/MW, EBITDA-based ROIC could be greater than 50% and better than terrestrial configurations.

Orbital Datacenter Payback. We estimate orbital datacenter EBITDA-based payback could be 14.4–7.2 years at near-term launch costs of \$2,700/kg, could get better to 5.0–2.5 years at long-term target parameters (\$40M/MW capex) and further to 1.9–1.0 years under long-term aspirational parameters (\$16M/MW capex) — shorter than xAI C1+C2 Current EBITDA payback of 5.8–2.9 years at the aspirational capex level. EBIT-based payback looks negative in near term but has the potential to get positive at medium-term target parameters at 31.0–15.5 years, getting better to 3.9–2.0 years under long-term aspirational parameters, compared to 8.9–4.4 years for xAI C1+C2 Current and 9.8–4.9 years for xAI Steady-State.

ROIC Sensitivity



Source: Company reports, RBC Capital Markets estimates. All figures shown at \$50M/MW capex. Bar ranges reflect \$10M/MW (high payback / low ROIC) to \$20M/MW (low payback / high ROIC) revenue scenarios. EBIT payback for Orbital Near-Term scenarios reflects negative cash flow.

ROIC Sensitivity

ROIC - Based on EBIT							
			Annual Revenue (\$M/MW)				
			\$10.0	\$12.5	\$15.0	\$17.5	\$20.0
Terrestrial DC	Hyperscaler		9%	11%	13%	15%	17%
Terrestrial DC	Generic Neocloud		3%	4%	4%	5%	6%
Terrestrial DC	xAI	C1+C2 Current	11%	14%	17%	20%	23%
Terrestrial DC	xAI	Steady-state	10%	13%	15%	18%	20%
Orbital DC	SpaceX Near-Term Target	Mass: 36 kg/kW + Starship reuse Launch Cost: \$2,700/kg Manufacturing cost: \$600/kg	-ve	-ve	-ve	-ve	-ve
Orbital DC	SpaceX Near-Term Target	Mass: 36 kg/kW + Starship reuse Launch Cost: \$270/kg Manufacturing cost: \$600/kg	-ve	-ve	-ve	-ve	-ve
Orbital DC	SpaceX Medium-Term Target	Mass: 15 kg/kW + Starship reuse Launch Cost: \$270/kg Manufacturing cost: \$400/kg	3%	4%	5%	6%	6%
Orbital DC	SpaceX Long-Term Target	Mass: 10 kg/kW + Starship reuse Launch Cost: \$100/kg Manufacturing cost: \$300/kg	7%	9%	11%	13%	15%
Orbital DC	SpaceX Medium-Term Target	Aspirational	16%	20%	24%	28%	32%
Orbital DC	SpaceX Long-Term Target	Aspirational	26%	32%	38%	45%	51%

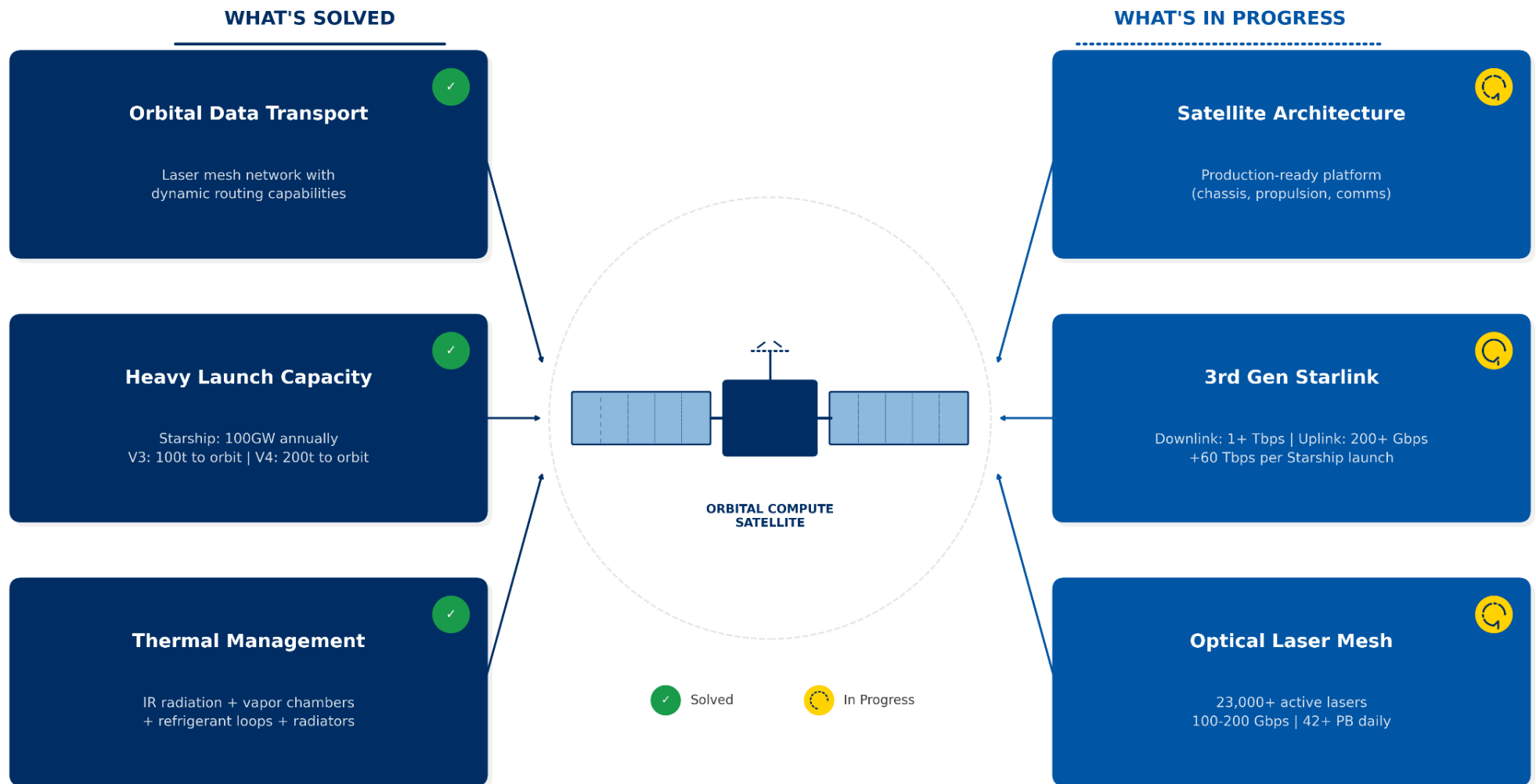
Source: Company reports, RBC Capital Markets estimates. All figures shown at \$50M/MW capex.

Payback Sensitivity

Payback - Based on EBIT (in years)							
			Annual Revenue (\$M/MW)				
			\$10.0	\$12.5	\$15.0	\$17.5	\$20.0
Terrestrial DC	Hyperscaler		11.6	9.2	7.7	6.6	5.8
Terrestrial DC	Generic Neocloud		33.4	26.7	22.3	19.1	16.7
Terrestrial DC	xAI	C1+C2 Current	8.9	7.1	5.9	5.1	4.4
Terrestrial DC	xAI	Steady-state	9.8	7.8	6.5	5.6	4.9
Orbital DC	SpaceX Near-Term Target	Mass: 36 kg/kW + Starship reuse Launch Cost: \$2,700/kg Manufacturing cost: \$600/kg	Negative Cash Flow	Negative Cash Flow	Negative Cash Flow	Negative Cash Flow	Negative Cash Flow
Orbital DC	SpaceX Near-Term Target	Mass: 36 kg/kW + Starship reuse Launch Cost: \$270/kg Manufacturing cost: \$600/kg	Negative Cash Flow	Negative Cash Flow	Negative Cash Flow	Negative Cash Flow	Negative Cash Flow
Orbital DC	SpaceX Medium-Term Target	Mass: 15 kg/kW + Starship reuse Launch Cost: \$270/kg Manufacturing cost: \$400/kg	31.0	24.8	20.7	17.7	15.5
Orbital DC	SpaceX Long-Term Target	Mass: 10 kg/kW + Starship reuse Launch Cost: \$100/kg Manufacturing cost: \$300/kg	13.7	10.9	9.1	7.8	6.8
Orbital DC	SpaceX Medium-Term Target	Aspirational	6.3	5.0	4.2	3.6	3.1
Orbital DC	SpaceX Long-Term Target	Aspirational	3.9	3.1	2.6	2.2	2.0

Source: Company reports, RBC Capital Markets estimates. All figures shown at \$50M/MW capex.

The crossover of orbital compute



The LLM landscape



Capital
Markets

Verticalization vs LLM labs- Points of leverage vs challenges

HYPERSCALERS		LLM COMPANIES	
LEVERAGE <i>Competitive Advantages</i> • Capital & proprietary components • Enterprise distribution • Access to power		• Algorithmic velocity • Architectural agility • No tech debt / legacy product baggage	
CHALLENGES <i>Key Risks</i> • Innovator's dilemma / slower research adaptation • Infrastructure depreciation		• Existential compute dependency • Persistent dilution • Likely low/zero switching costs for APIs	
MONETIZATION <i>Revenue Engine (Current)</i> • Own & competitor compute • SaaS / productivity suites		• Model as a service / APIs • Subscriptions • Agentic enterprise apps / orchestration	

Cross Currents Across the LLM Product Stack

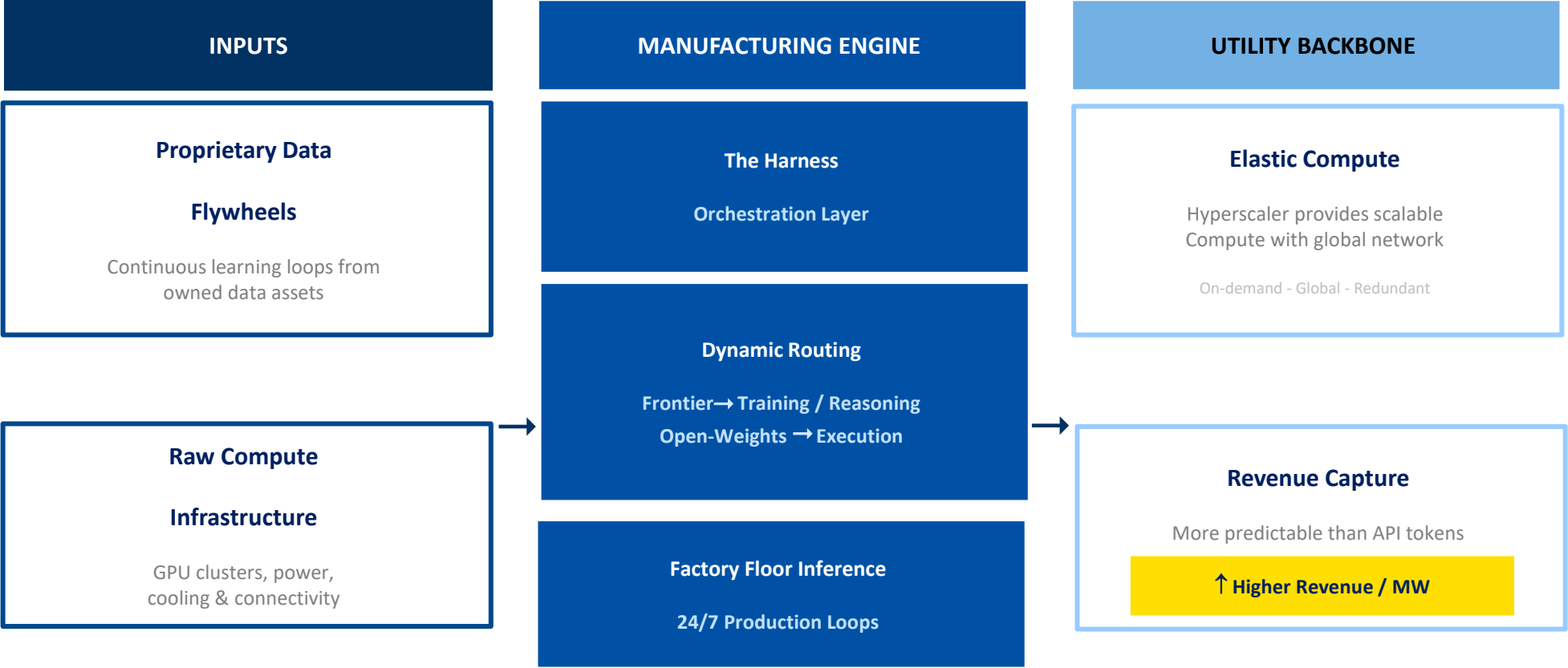
- Here are some observations on the cross currents across the product/service landscape:
 - Convergence toward full-stack:** All four providers now offer a flagship LLM, consumer chatbot, enterprise API, voice/audio AI, coding tools, and agentic capabilities — these categories have become table stakes and are no longer differentiators.
 - The divergence is in physical extension:** Google and xAI are pushing AI into physical form factors (wearables, humanoid robots) while OpenAI is building toward it (device expected late 2026/early 2027) and Anthropic has shown zero appetite for hardware of any kind.
 - Video generation has flipped:** OpenAI exited video (Sora shutdown, March 2026) while xAI entered it (Grok Imagine Video 1.5, June 2026) — a reversal in less than six months that reflects the compute-intensity tradeoffs labs are now making in real time.
 - Custom silicon is a two-player game (for now):** Only Google (TPU) and OpenAI (Jalapeño) have shipped proprietary chips — Anthropic relies entirely on others' infrastructure, and xAI runs on Nvidia Blackwell without a custom chip announced.
 - Anthropic's absence is the story:** No image generation, no video, no search, no advertising, no hardware, no chips — yet it commands the highest valuation (\$965B) and ~\$47B ARR, suggesting the market is pricing enterprise depth over product breadth.
 - OpenAI is the only lab monetizing attention directly:** ChatGPT Ads (launched early 2026) makes OpenAI the sole provider operating its own advertising platform — a revenue model none of the other three have pursued.

Product / Service Category	OpenAI (GPT)	Anthropic (Claude)	Google (Gemini)	xAI (Grok)
Flagship LLM (text)	✓	✓	✓	✓
Consumer Chatbot App	✓	✓	✓	✓
Enterprise API Platform	✓	✓	✓	✓
Image Generation	✓		✓	✓
Video Generation			✓	✓
Voice/Audio AI	✓	✓	✓	✓
AI-Powered Search	✓		✓	✓
Coding Assistant/IDE	✓	✓	✓	✓
Agentic AI / Computer Use	✓	✓	✓	✓
AI-Powered Advertising	✓			
Custom AI Hardware (Chips)	✓		✓	
Wearable AI Devices	✓		✓	
Humanoid Robots			✓	✓

Source: Company reports, Bigdata, RBC Capital Markets

We're in a transition from API consumers to intelligence manufacturers

The Shift from API Consumers to intelligence Manufacturers



LLM optimization for lowest cost of compute, memory and latency

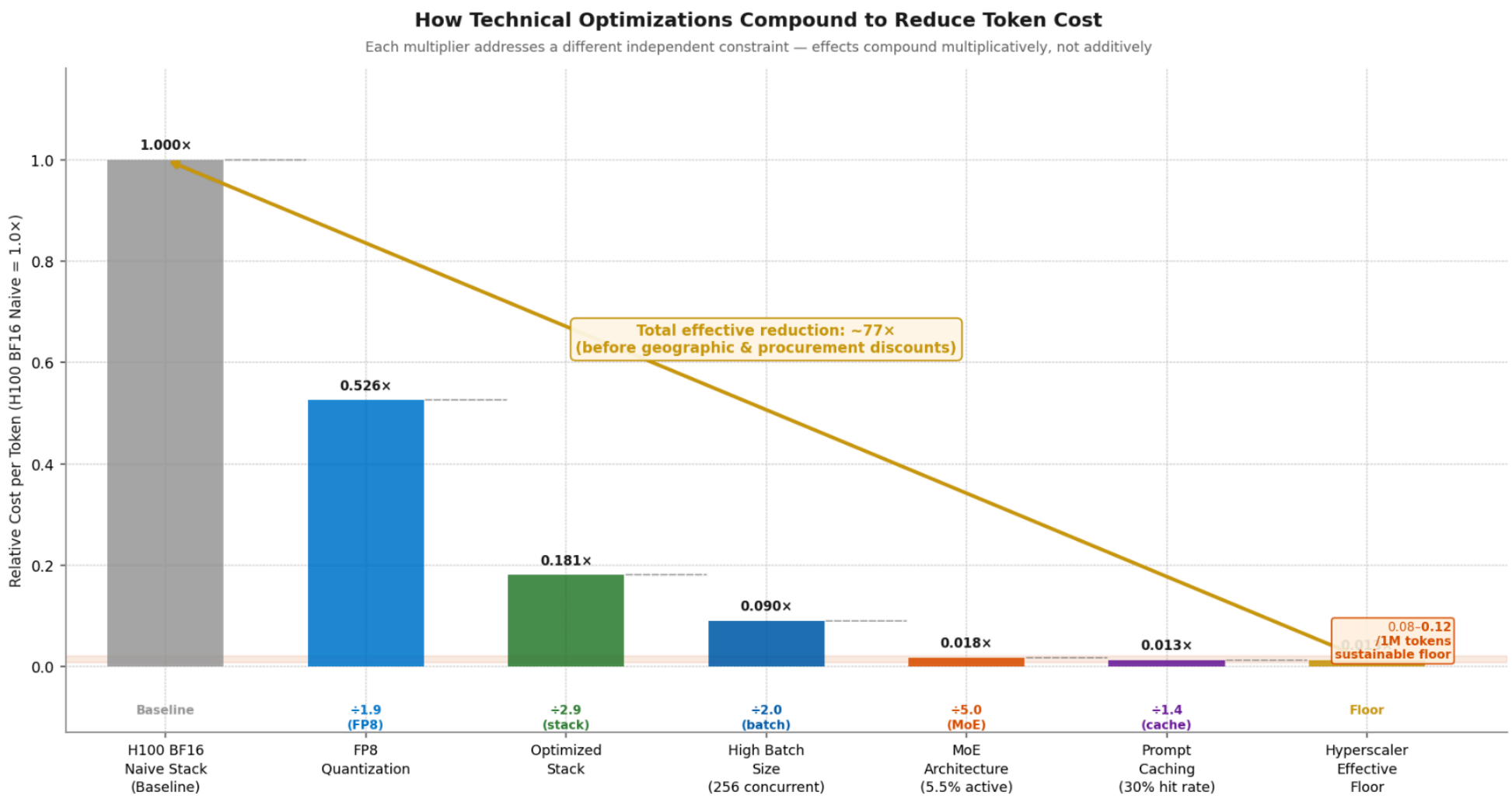
- With model training costs ranging from a few million \$s for small models (single billion parameter) to several hundred million dollars for multi-trillion parameter frontier models, there's a pipeline of interconnected techniques that collectively aim to maximize hardware performance with the lowest compute costs, memory footprints and lowest latency.
- **Core Efficiency Techniques**
 - **Model Distillation:** A smaller "student" model learns to mimic a larger "teacher" model, achieving 95-97% of the teacher's performance at 3-5x smaller size and faster inference.
 - **Quantization:** Reducing numerical precision (e.g., 32-bit to 8-bit or 4-bit) to shrink model size 4-8x and speed up inference 2-4x with minimal accuracy loss.
 - **Mixture of Experts (MoE):** Architecture where only a subset of specialized "expert" networks activate per input, enabling massive parameter counts (56B+) while only computing with a fraction (7B active).
 - **Data Parallelism:** Distributing training across multiple GPUs, where each processes different data batches simultaneously, then synchronizes to update the model.
 - **Inference Optimization**
 - **Pre-fill vs Decode:** Pre-fill processes the input prompt in parallel (fast); decode generates output one token at a time sequentially (slow bottleneck).
 - **KV Cache:** Stores previously computed attention keys and values to avoid recalculating them for each new token, reducing decode-phase computation by 10-100x.
 - **Self-Summarization:** Model periodically condenses its conversation context into a summary to stay within memory limits while retaining key information from long conversations.
- **Reinforcement Learning**
 - **RLHF/RLAIF:** Training technique using human (RLHF) or AI (RLAIF) feedback on model outputs to align behavior toward helpfulness, harmlessness, and following instructions effectively.

4 buckets of model usage: Frontier, mid-tier, open-weights & local

- **Frontier/premium models** (Claude Fable 5, GPT 5.5)
 - Reserved for architectural tier tasks, complex strategic planning, multi-file refactoring
 - High-cost \$5-30+ per million justified by reliability
- **Mid-tier** (Gemini 3.5 Flash, GPT 4.1)
 - Probably could do 60-80% of enterprise workloads longer term with higher reliability from techniques requiring primary sources (avoiding hallucinations)
 - Priced for higher throughput \$0.20-\$4.50 per million tokens depending on capability
- **Open-weights** (Llama 4, Qwen 3.7, KimiK 2.7 Code)
 - Require maintenance, i.e., engineering resources but allows for private/ring-fenced deployments
 - Intended for high-volume, pre-established tasks (often done by frontiers), also helpful for high security in case of private API or local
 - API model (3rd hosting): pay infrastructure provider for hosting
 - Cost range \$0.10-\$0.60/M tokens for input, \$0.25-\$1.50/M for output
- **Locally run (most often leveraging open-source/open-weights):** Tends toward prosumer/SMB small-scale tasks, zero latency but static models which need updating.
 - Cost range \$2k to tens of thousands of \$s up-front for hardware & engineers to update models periodically, maintain uptimes, load balancing, etc.
 - For many, we believe locally run models' free marginal token cost likely does not outweigh ongoing maintenance challenges combined with initially higher capex

How Technical Optimizations Compound to Reduce Token Cost

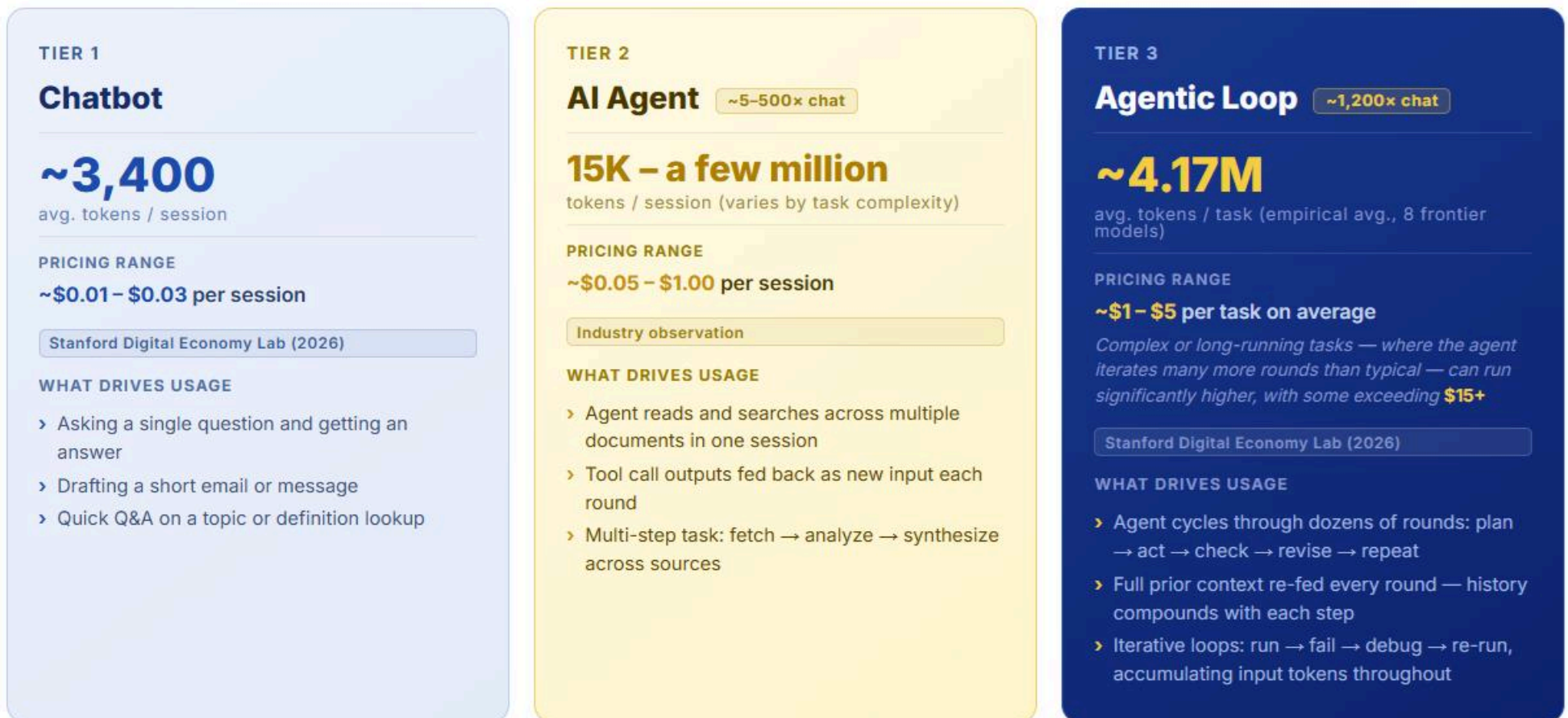
- xAI captures all five compounding optimization layers — FP8/FP4 quantization, optimized SW stack, high batching, MoE architecture (Grok), and prompt caching — placing it at or near the \$0.08–\$0.12/1M token floor achievable by hyperscale MoE providers, rather than the \$0.15–\$0.20 dense-model floor applicable to less optimized deployments.



Source: RBC Capital Markets estimates, company reports, MLcommons

Token consumption grows exponentially with agentic loop adoption

- Agentic loops carry significantly higher utility potential vs. chat or simple agents but per-task cost scales quickly given exponentially higher token usage.
- Enterprises are becoming more shrewd with model choices (Meta, Uber).
- Reinforces infrastructure demand with even the most efficient open-source models still requiring state of the art compute.



How big of a risk is the rising availability of cheaper models (open-source, local models, SLMs, etc.)?

■ Current Market Pricing

- Frontier Models: \$5-30/million tokens | Open-Weight Models: \$0.05-\$2/million tokens

■ Key Market Trends

- Replaced frontier model token costs can drop up to 90% | Third-party reports show agentic-fueled token consumption increased

■ Primary Risk Factors

- Companies must deploy ring-fenced models through cloud partners or owned infrastructure to protect core IP
- Cannot expose data through public endpoints or create single vendor dependency as third-party model routing providers proliferate

■ Secondary Considerations

- Jevons Paradox: Lower compute costs drive broader AI exploration and higher ROI potential
- Frontier models often deliver lower total cost of completion despite higher token prices
- Owning physical infrastructure enables model-agnostic approaches, though ring-fencing open-weight models remains inefficient except for largest enterprises

■ Competitive Positioning Strategy

- SpaceX/Grok Approach: Price software and Grok model access below competing frontier providers' retail rates. Ring-fenced open weights face limitations in information freshness and utilization efficiency.

HIGH



LOW

HIGH

Claude Fable 5 (w/ fallback)

Premium costs — highest per-token pricing, maximum capability

Superior reasoning — best-in-class complex tasks, nuanced understanding

Reliability — fallback mechanism ensures service continuity

Best for — mission-critical, complex analysis, highest accuracy

MEDIUM

Gemini 3.1 Pro Preview

Moderate costs — cost-effective balance of price and quality

Better reasoning — improved context retention, accurate outputs

Versatility — analysis to creative work, wide task range

Best for — business apps, content creation, moderate complexity

LOW

gpt-oss-20B (high)

Lower costs — significantly cheaper per token or API call

Adequate — good for straightforward queries, basic generation

Trade-offs — slower speed, less nuance, shorter context

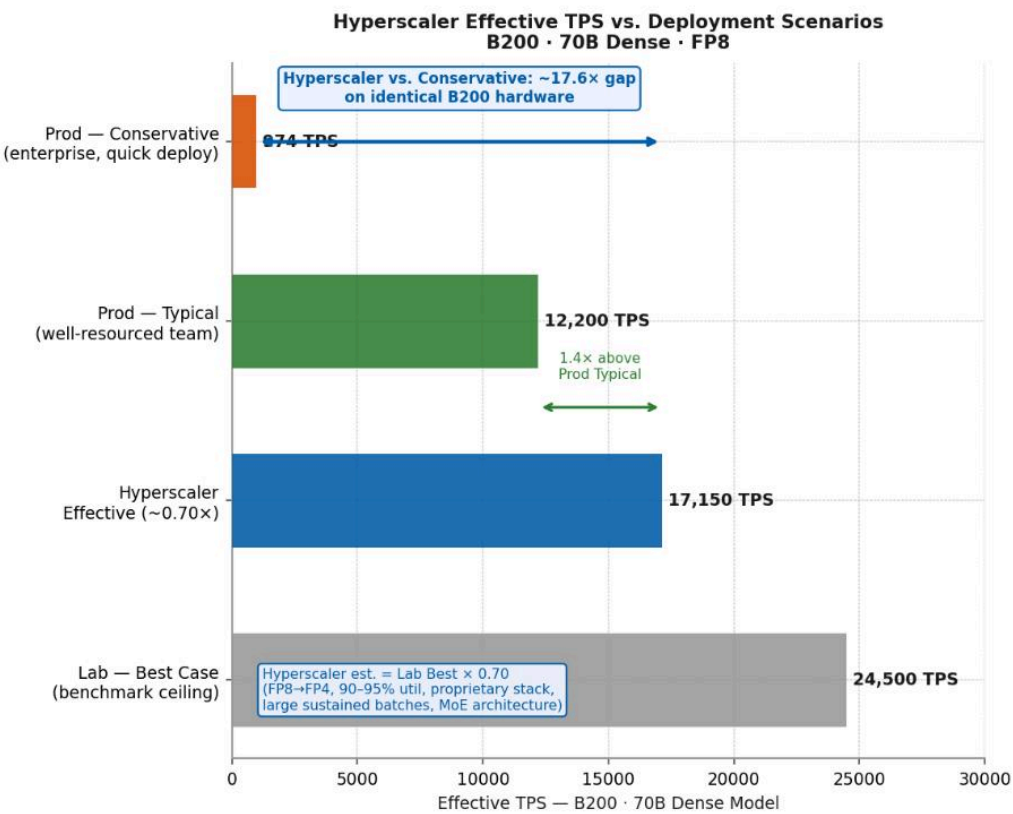
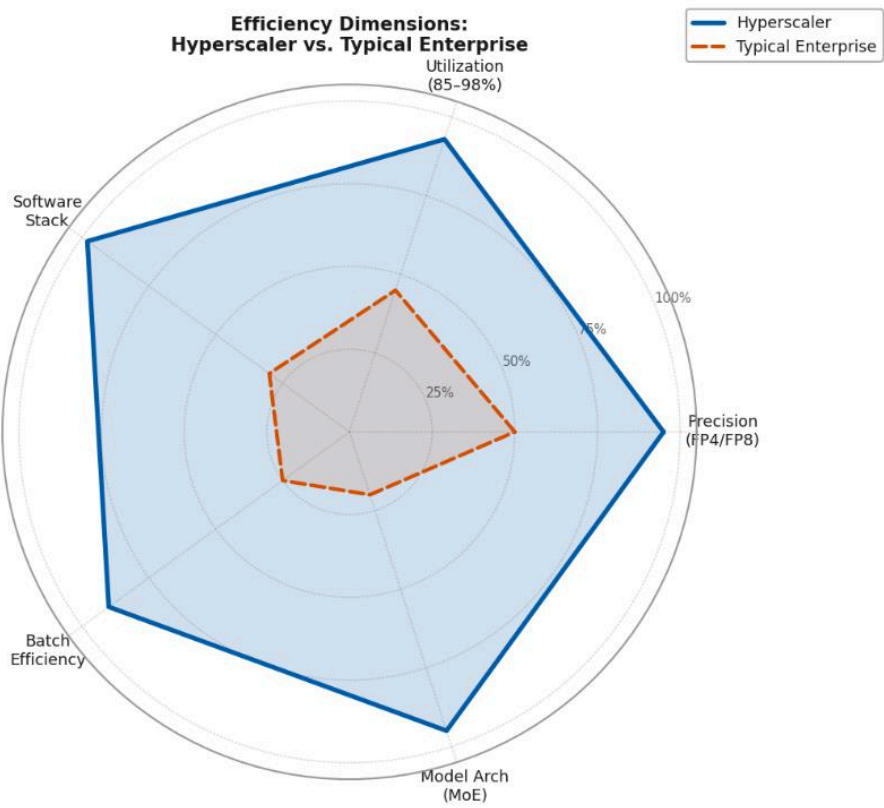
Best for — bulk processing, simple classification, chatbots

Hyperscaler vs. Enterprise: The Efficiency Gap

- xAI is explicitly grouped with Anthropic and OpenAI in the hyperscaler operating band on all five radar dimensions: FP8/FP4 precision, 85–98% utilization, optimized proprietary SW stack, large sustained batch sizes, and MoE architecture. The bar chart quantifies this as ~17,150 TPS effective on B200 70B — 17.6× above a conservative enterprise deployment on identical hardware. xAI is on the favorable side of every dimension shown.

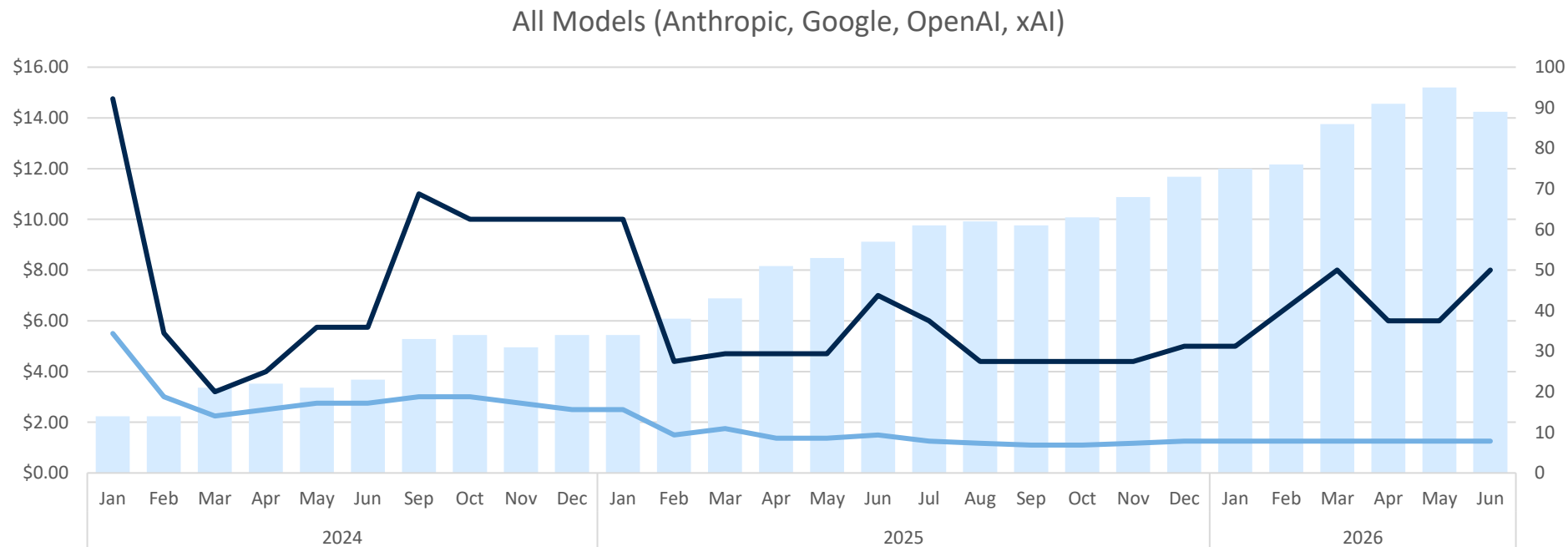
Hyperscaler vs. Enterprise: The Efficiency Gap

Hyperscalers operate closer to Lab — Best Case than typical enterprise deployments across all five efficiency dimensions



Source: RBC Capital Markets, company reports, MLcommons

LLM Pricing Trends (Anthropic, Google, OpenAI, xAI)



- Median output token price of all models for the major model providers (Anthropic, Google, OpenAI, xAI) declined throughout 2024-2025 before increasing during 2026 – this is largely due to an increase in the quantity of frontier models (higher \$/M tokens) paired with the deprecation of older, cheaper mid-tier and lightweight models.
- Median input token price of all models has declined steadily from 2024-2026, from ~\$3.00/M tokens in September 2024 to \$1.25/M tokens (~60%) due to hardware efficiency gains, inference optimization, and scale.
- The quantity of active models has also increased, as providers have iterated over their lightweight, mid-tier, and frontier models, and also developed specialized models (coding, instruct variants, image).

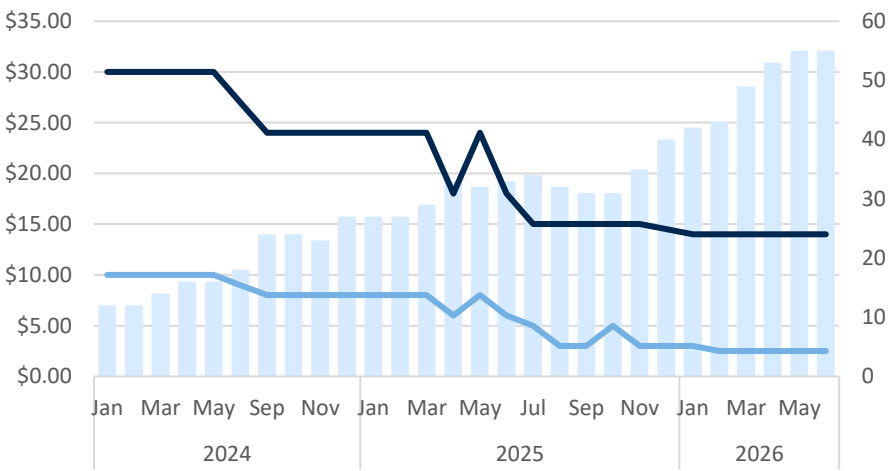
Unique Models Output Token Price (\$/M tokens) Input Token Price (\$/M tokens)

Note: Model providers considered are Anthropic, Google, OpenAI, xAI. Source: OpenRouter, RBC Capital Markets, RBC Elements

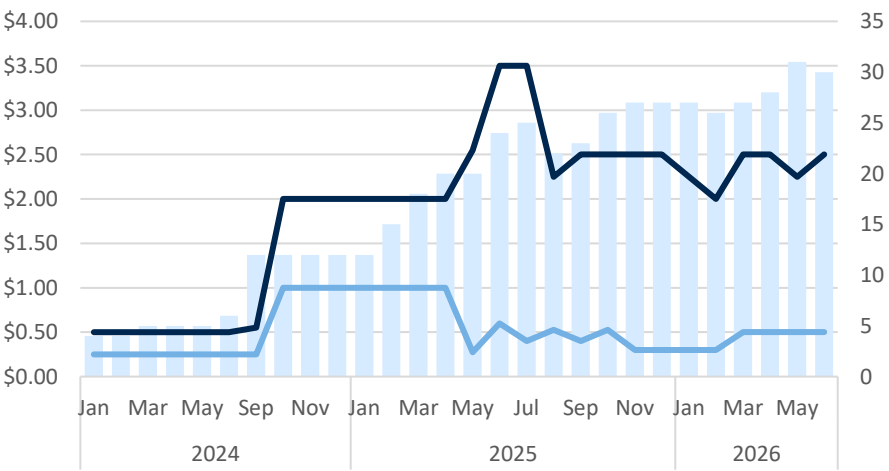
LLM Pricing Trends (Anthropic, Google, OpenAI, xAI)

- Frontier models in by the major model providers in aggregate have seen steady median output token and median input token pricing erosion driven by both hardware/architecture/inference gains and competition for enterprise business. Pricing has remained steady throughout 2026 at \$2.50/M input tokens and \$14.00/M output tokens, though new models (ex. Mythos) (rather than new variants) stand to disrupt this trend.
- In contrast, mid-tier models have increased in output token pricing while remaining flat in input token pricing – the trend is due to increased capabilities of mid-tier output tokens (thinking/modalities), while input token processing has become commoditized.
- Lightweight model pricing exhibits a race to the bottom, as these models are designed to handle minimal complexity. The increase in output token price in May 2026 is from the deprecation of cheaper models.

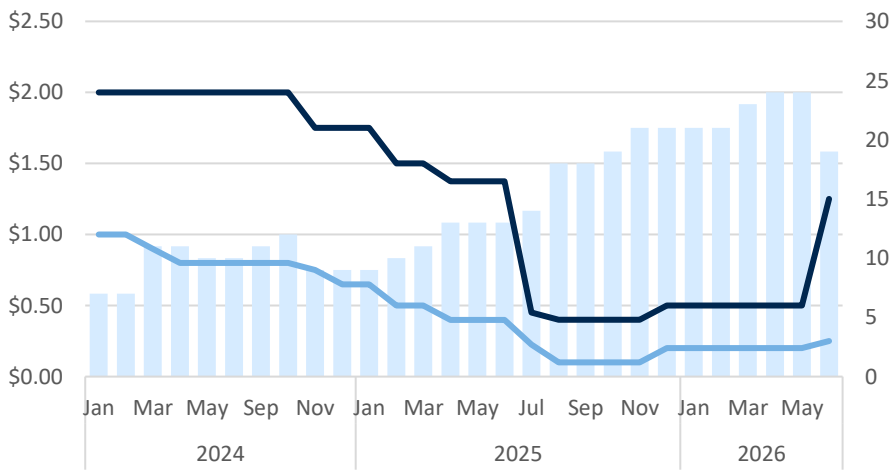
Frontier Models



Mid-Tier Models



Lightweight Models



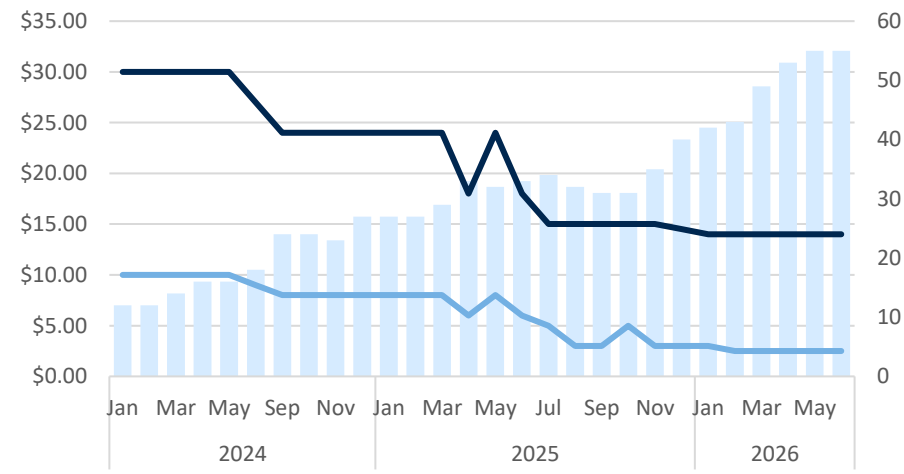
Unique Models Output Token Price (\$/M tokens) Input Token Price (\$/M tokens)

Source: OpenRouter, RBC Capital Markets, RBC Elements

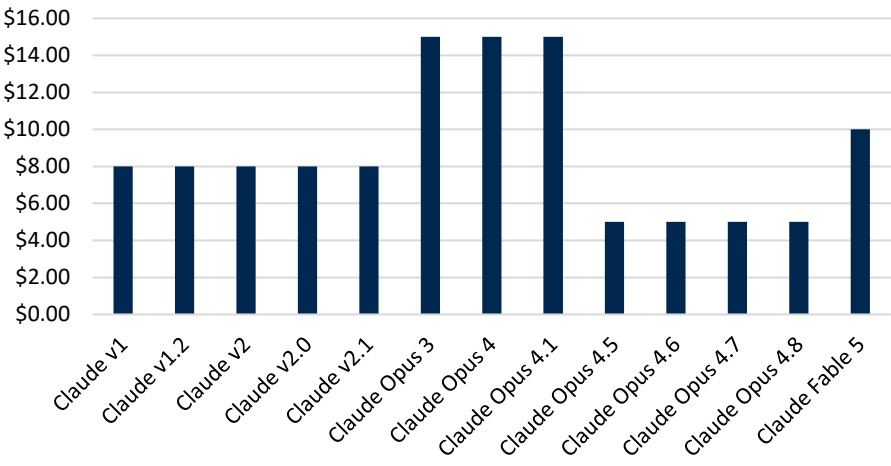
LLM Pricing Trends (Anthropic)

- **Early Claude models started at mid-tier pricing.** The initial Claude v.x series was priced at \$8/\$24 (input/output per million tokens), in line with the broader market at the time.
- **Opus introduced a premium tier.** With Claude 3 Opus, Anthropic moved upmarket. Opus 3, 4, and 4.1 were priced at \$15/\$75—nearly 3x the output price of earlier models.
- **Anthropic then cut Opus pricing to drive adoption.** Opus 4.5 through 4.8 dropped to \$5/\$25, an 80% reduction in output pricing. The move signaled a push for enterprise volume over per-token margin.
- **Fable 5 re-established premium pricing.** At \$10/\$50, Fable 5 doubled the price of late-generation Opus, reflecting its stronger coding, reasoning, and agentic capabilities.

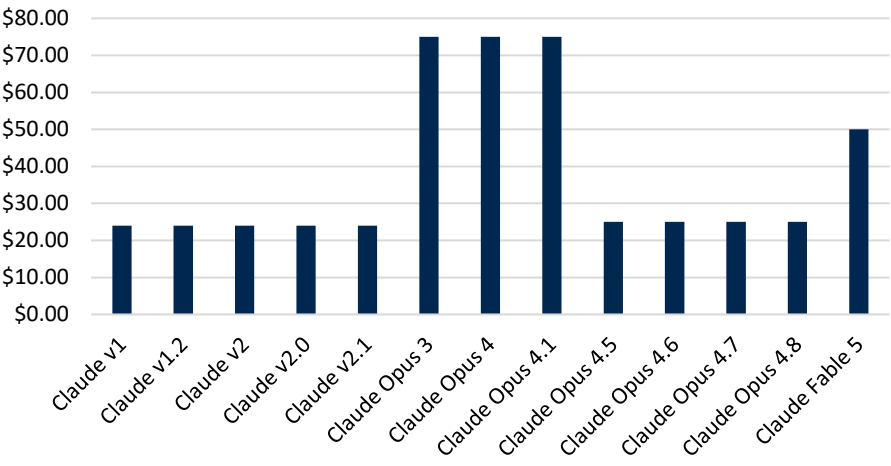
Frontier Models



Input Token Price



Output Token Price

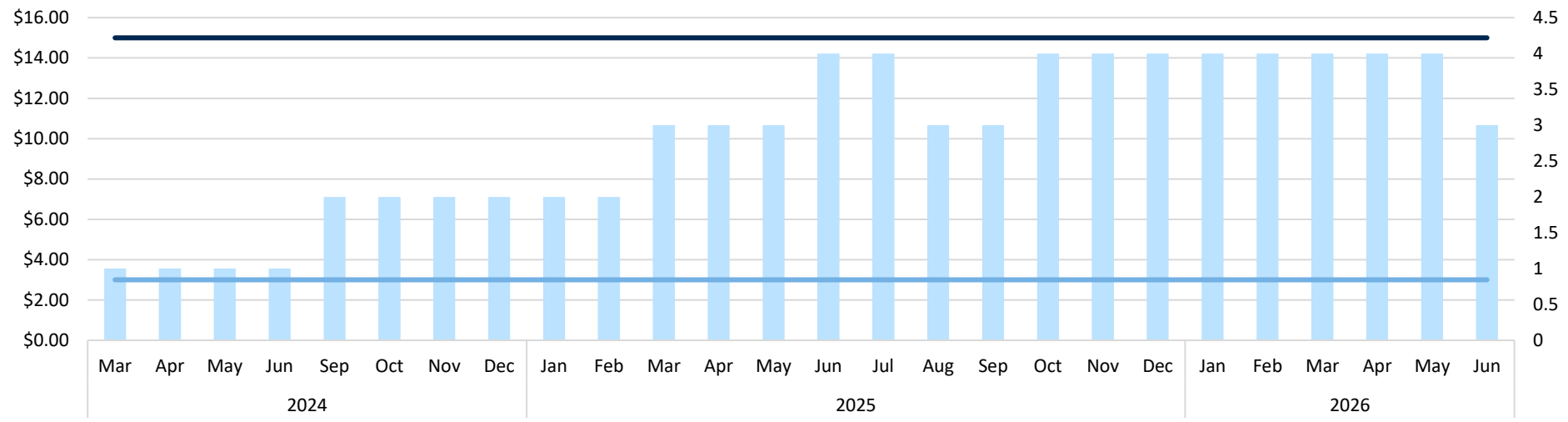


Unique Models Output Token Price (\$/M tokens) Input Token Price (\$/M tokens)

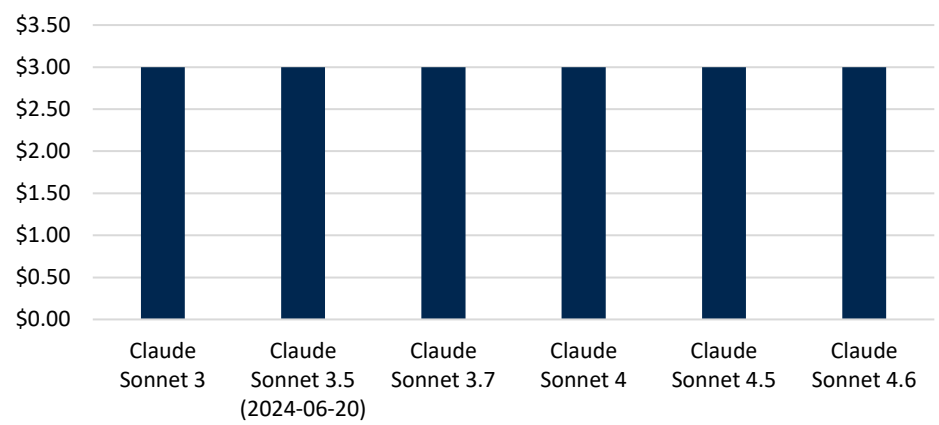
Source: OpenRouter, RBC Capital Markets, RBC Elements

LLM Pricing Trends (Anthropic)

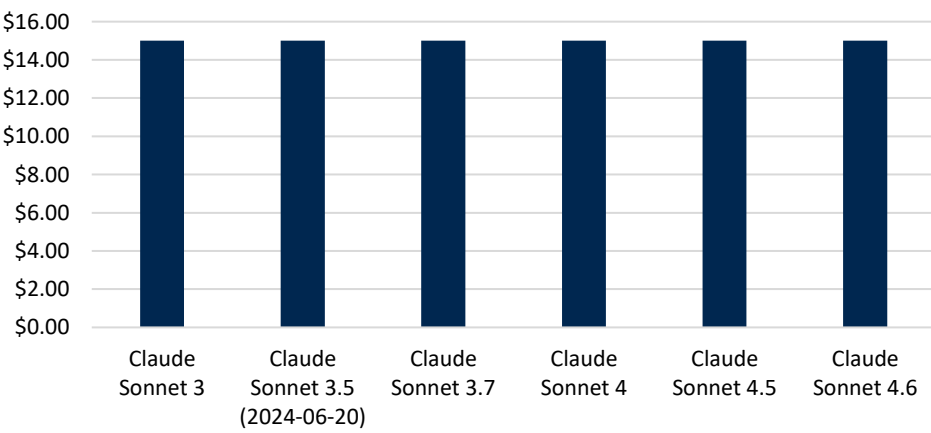
Mid-Tier Models



Input Token Price



Output Token Price

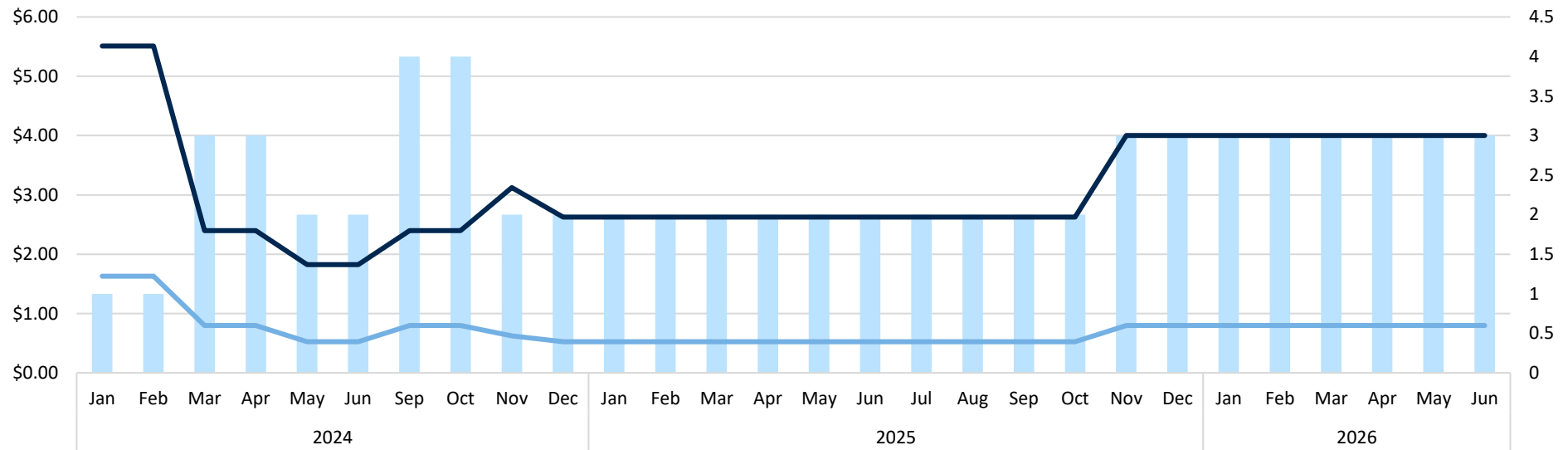


Unique Models Output Token Price (\$/M tokens) Input Token Price (\$/M tokens)

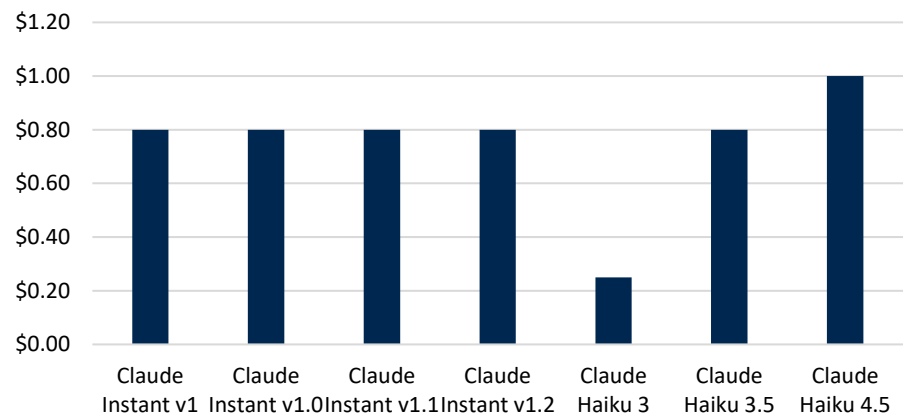
Source: OpenRouter, RBC Capital Markets, RBC Elements

LLM Pricing Trends (Anthropic)

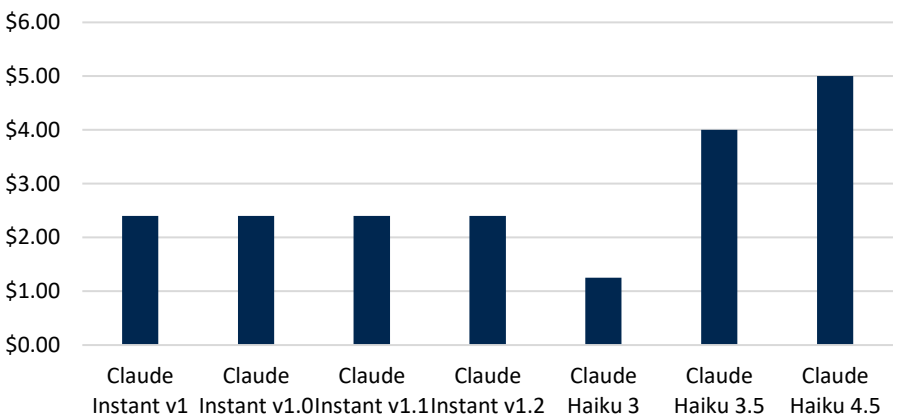
Lightweight Models



Input Token Price



Output Token Price



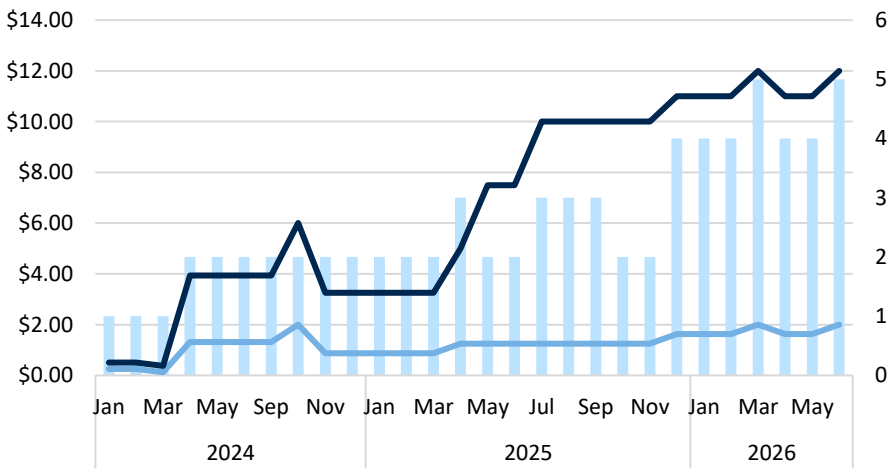
Unique Models Output Token Price (\$/M tokens) Input Token Price (\$/M tokens)

Source: OpenRouter, RBC Capital Markets, RBC Elements

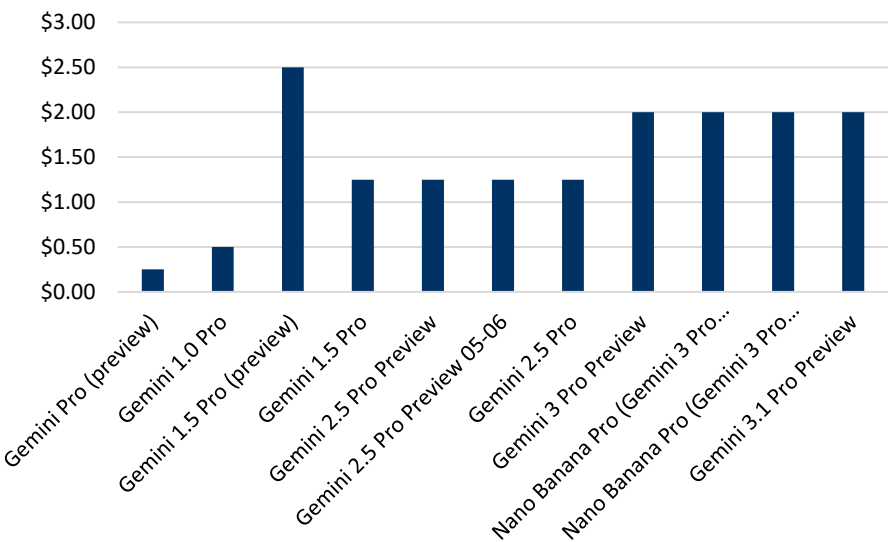
LLM Pricing Trends (Google)

- **Output prices rose while input stayed flat.** Google's frontier output pricing climbed from \$0.50 to \$12 over the Gemini Pro lineage, while input hovered around \$1–2.
- **Higher prices reflect higher capability.** Each Pro generation brought stronger reasoning, expanded multimodal support, and massive context windows (1M → 2M tokens). Google priced in the value.
- **Google raises prices where peers cut them.** Unlike OpenAI and Anthropic, Google has increased pricing with each Gemini Pro iteration, capturing a premium for meaningfully better performance.

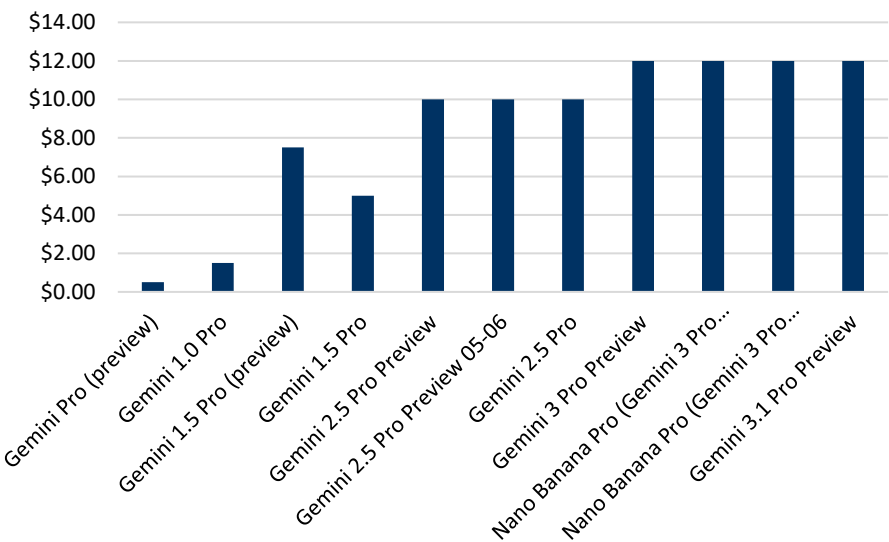
Frontier Models



Input Token Price



Output Token Price

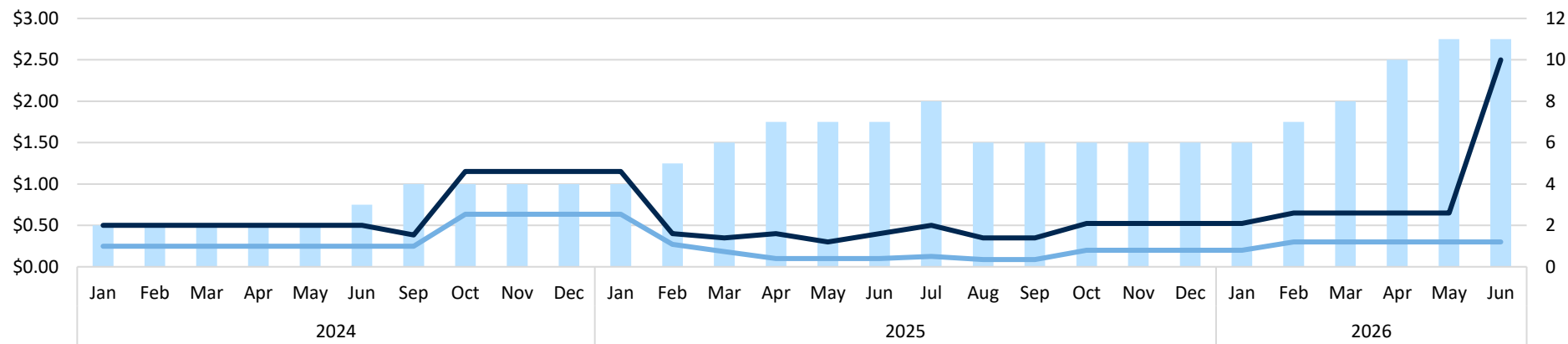


Unique Models Output Token Price (\$/M tokens) Input Token Price (\$/M tokens)

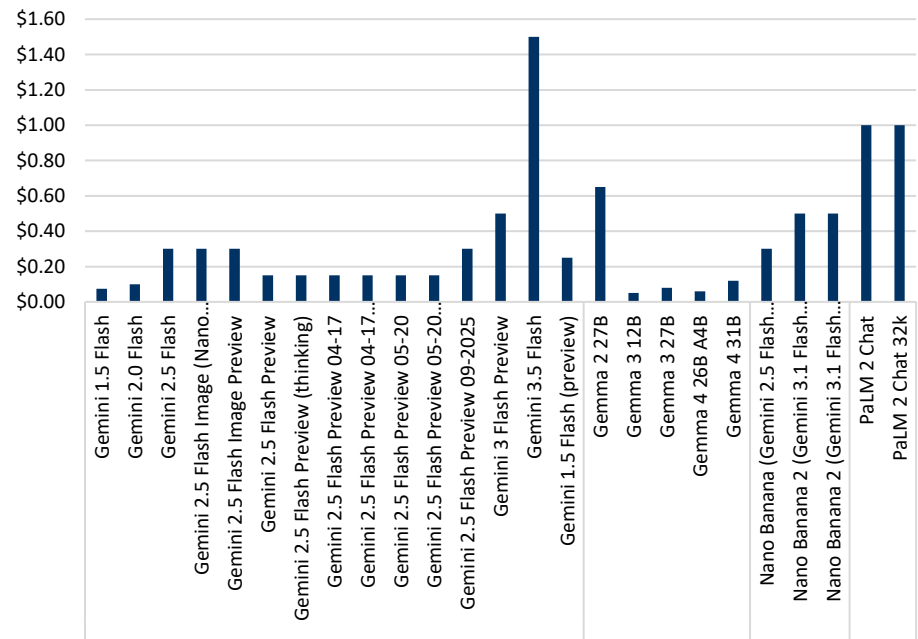
Source: OpenRouter, RBC Capital Markets, RBC Elements

LLM Pricing Trends (Google)

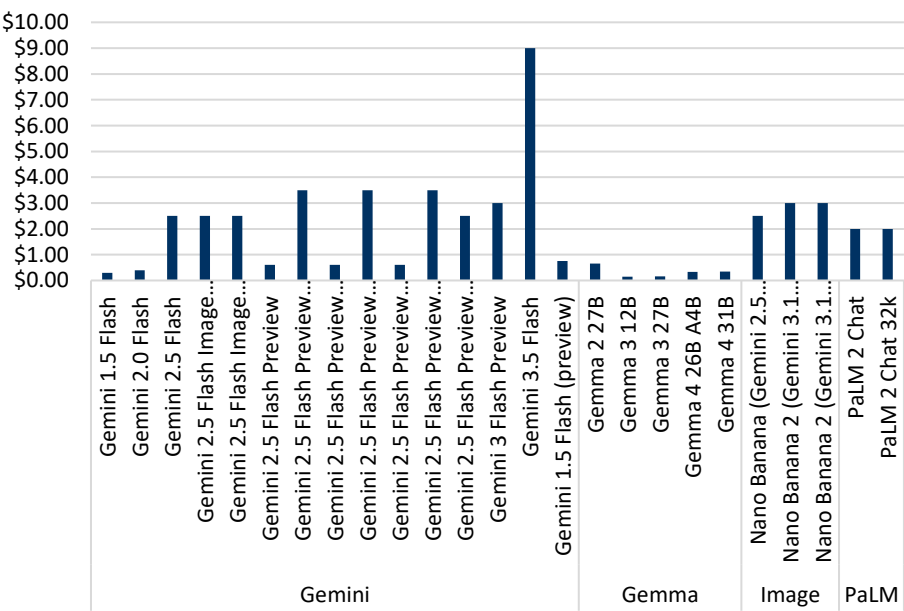
Mid-Tier Models



Input Token Price



Output Token Price

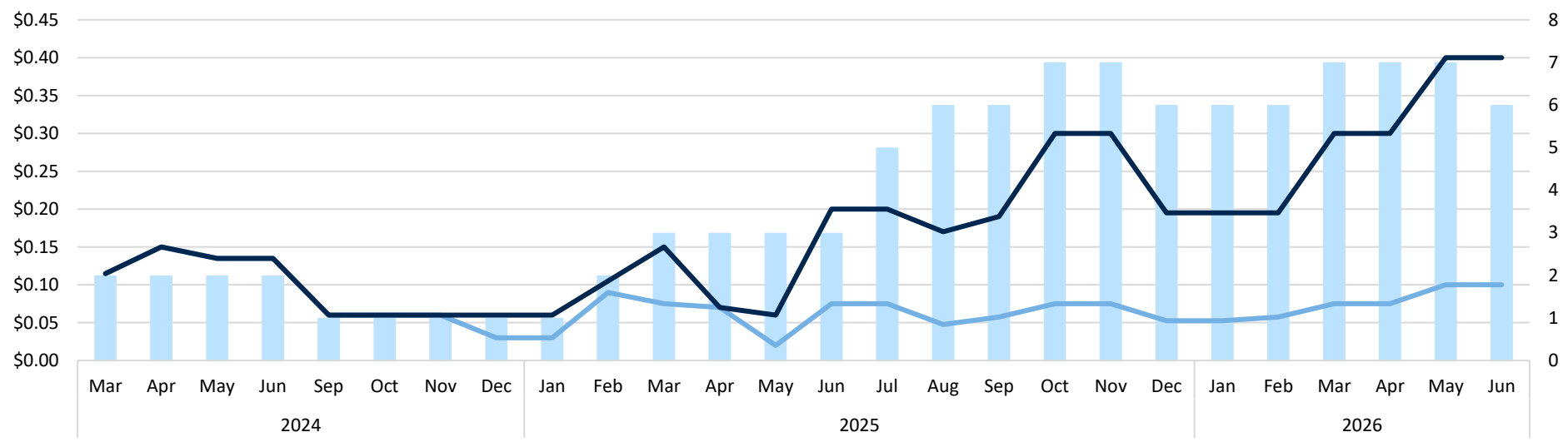


Unique Models Output Token Price (\$/M tokens) Input Token Price (\$/M tokens)

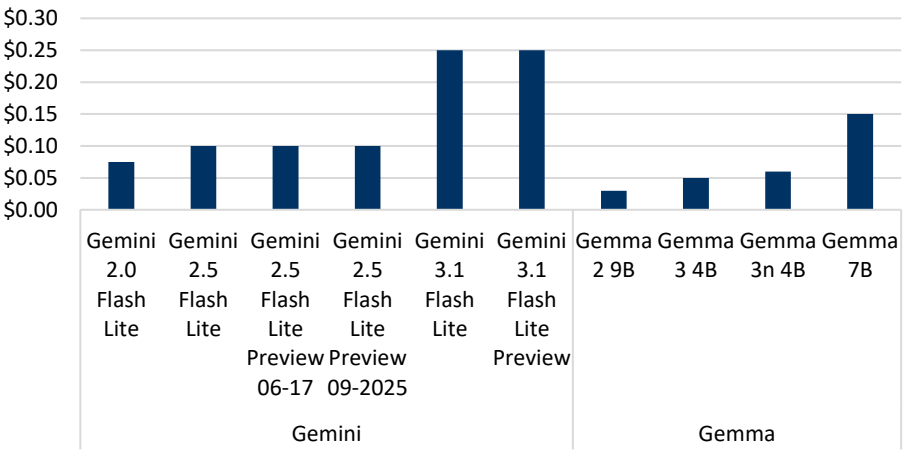
Source: OpenRouter, RBC Capital Markets, RBC Elements

LLM Pricing Trends (Google)

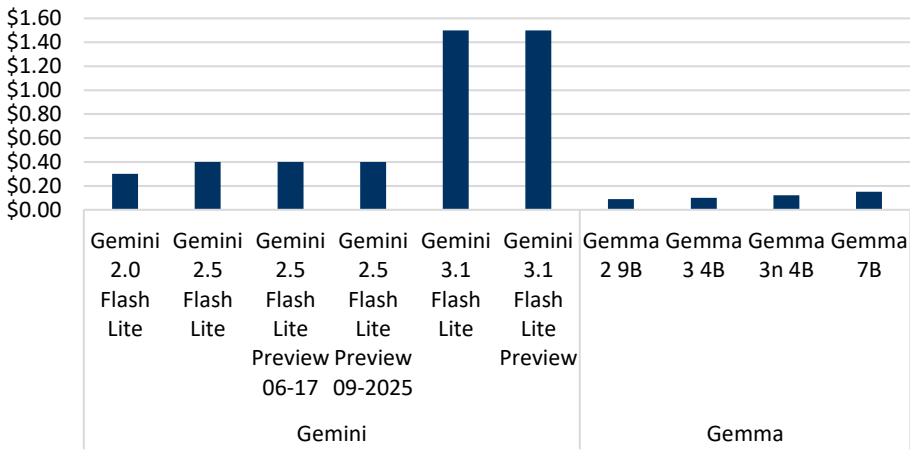
Lightweight Models



Input Token Price



Output Token Price

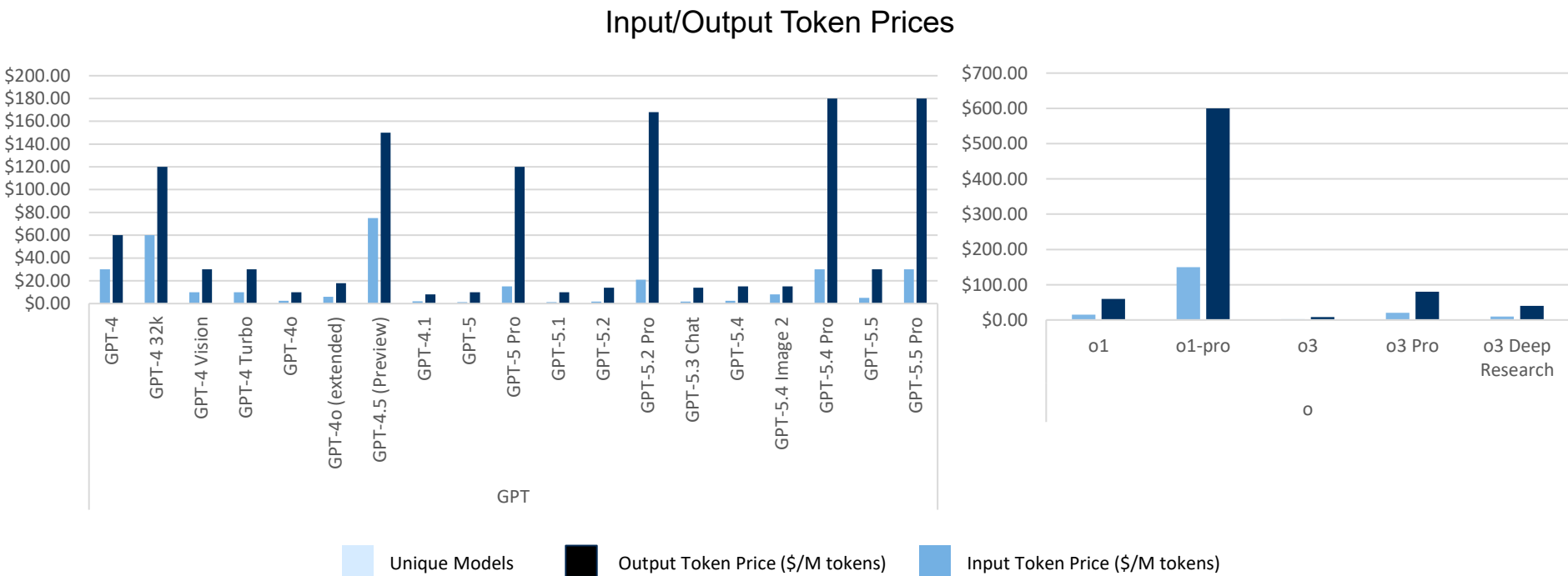
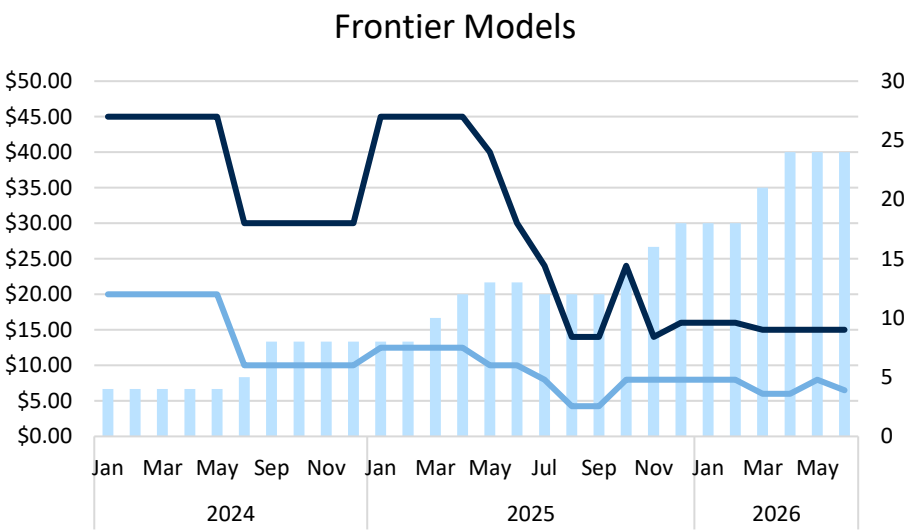


Unique Models Output Token Price (\$/M tokens) Input Token Price (\$/M tokens)

Source: OpenRouter, RBC Capital Markets, RBC Elements

LLM Pricing Trends (OpenAI)

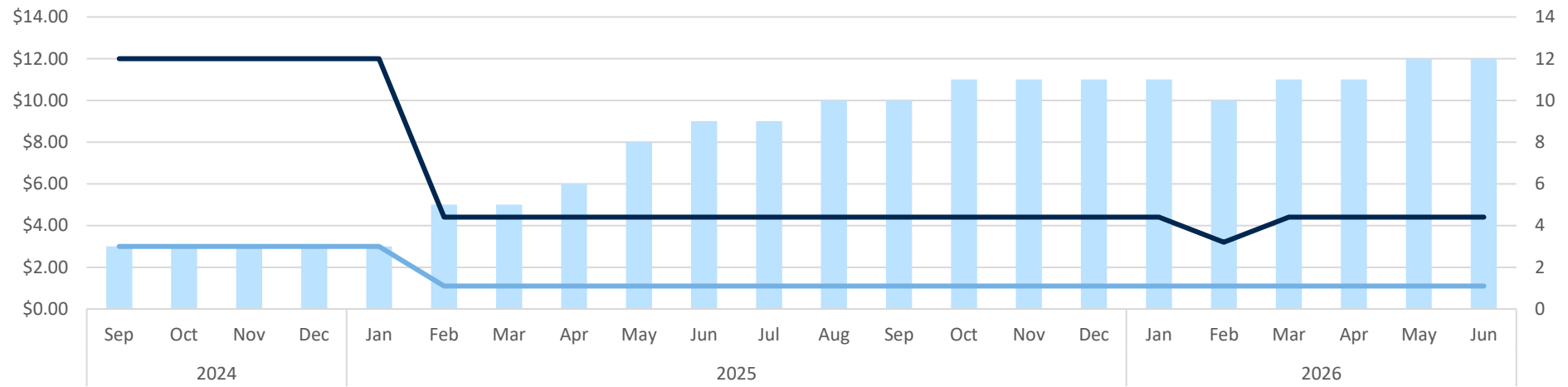
- **New models get cheaper, old models stay frozen.** Input prices dropped ~96% from GPT-4 (\$30) to GPT-5 (\$1.25) over two years. Only two models received mid-life price cuts: GPT-4o (50% input, 33% output) and o3 (80% both).
- **Output tokens getting relatively more expensive.** Output/input ratio went from 2x (GPT-4) to 6–8x (GPT-5 era). Input is cheap; generation is where they capture margin.
- **Pro tiers command huge premiums.** Base GPT-5 is \$1.25 input; GPT-5 Pro is \$15 (12x). Same pattern with o1 vs o1-pro (10x). This reflects commodity pricing for volume, and premium pricing for power users.



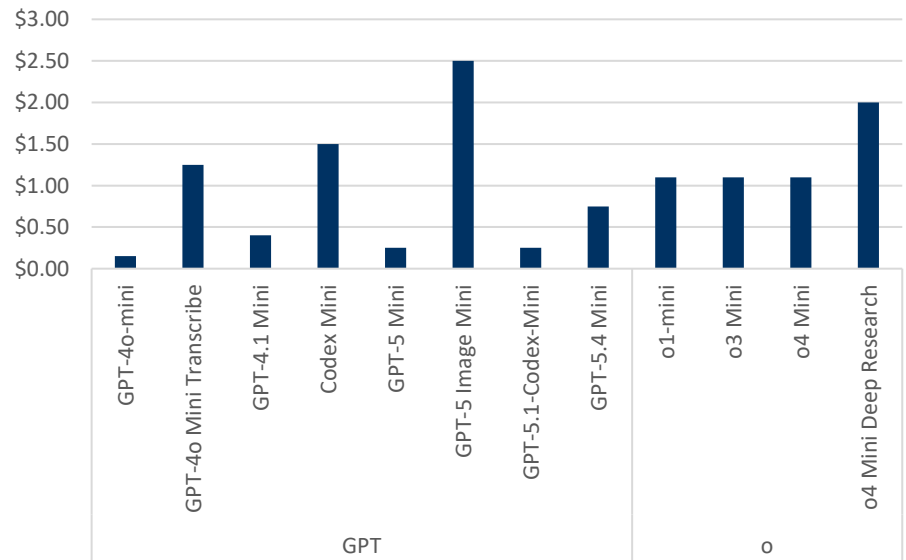
Source: OpenRouter, RBC Capital Markets, RBC Elements

LLM Pricing Trends (OpenAI)

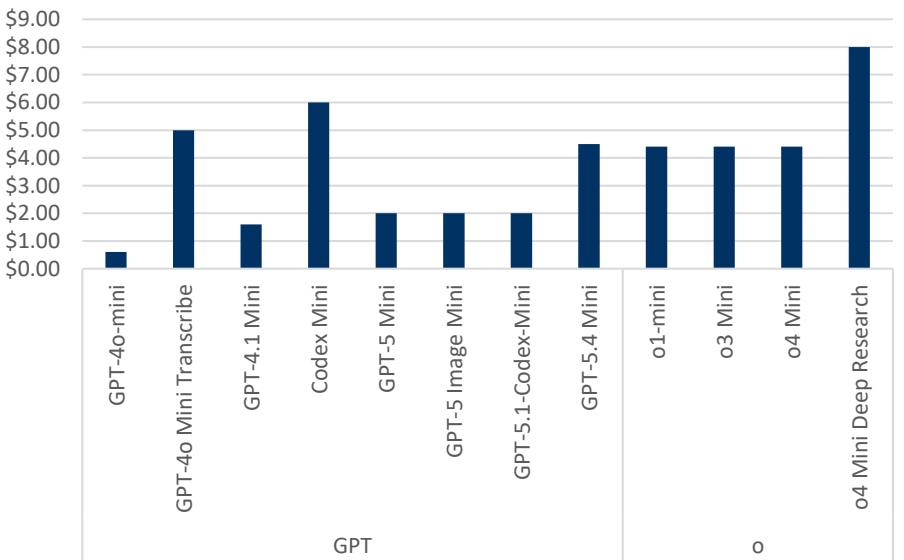
Mid-Tier Models



Input Token Price



Output Token Price

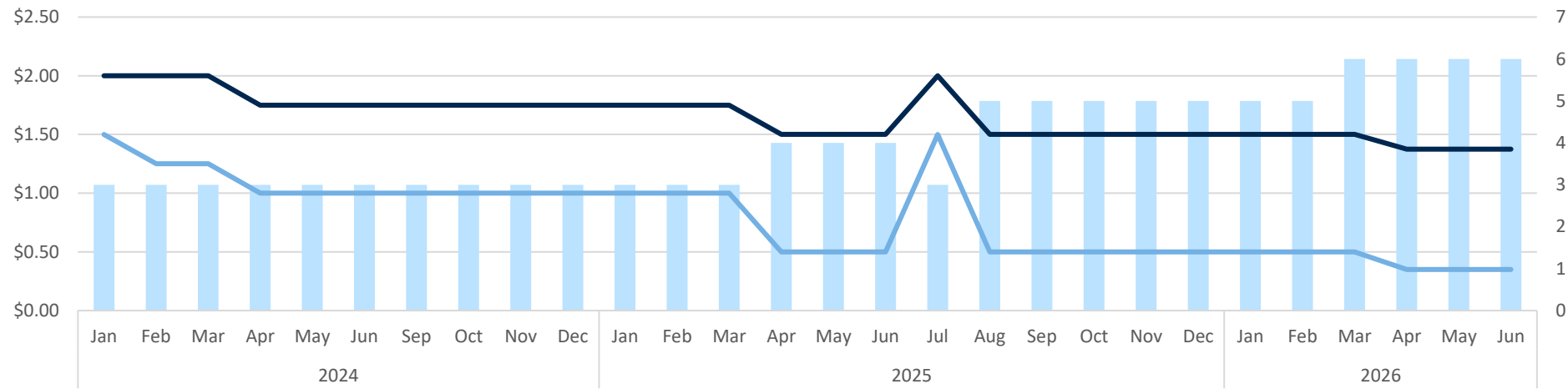


Unique Models Output Token Price (\$/M tokens) Input Token Price (\$/M tokens)

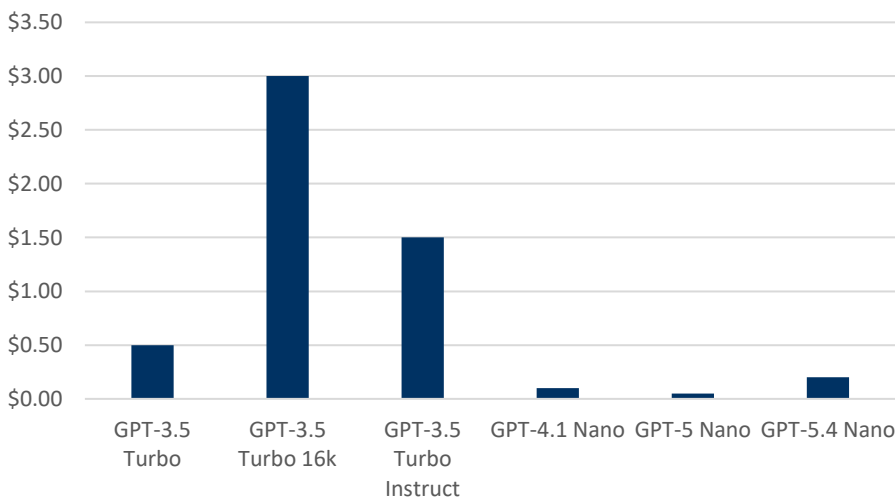
Source: OpenRouter, RBC Capital Markets, RBC Elements

LLM Pricing Trends (OpenAI)

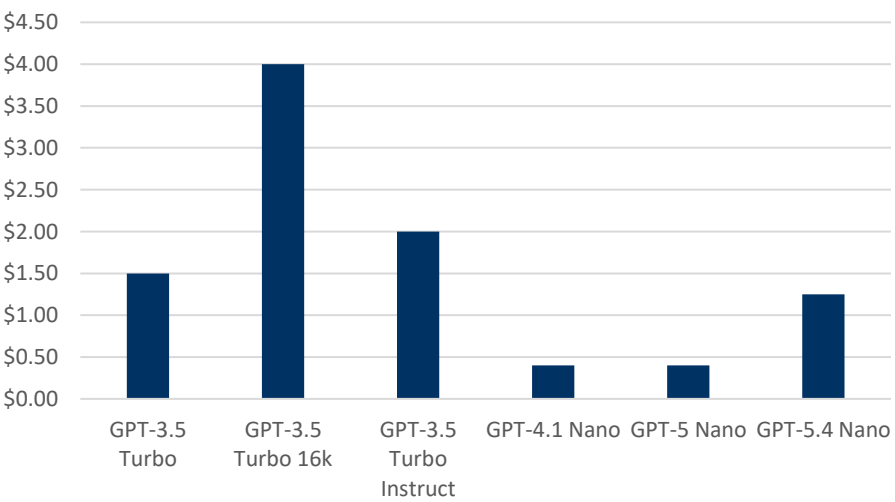
Lightweight Models



Input Token Price



Output Token Price



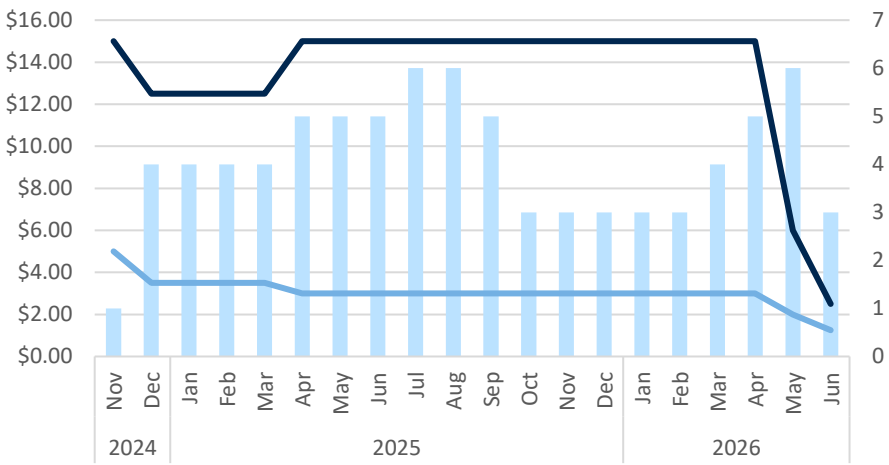
Unique Models Output Token Price (\$/M tokens) Input Token Price (\$/M tokens)

Source: OpenRouter, RBC Capital Markets, RBC Elements

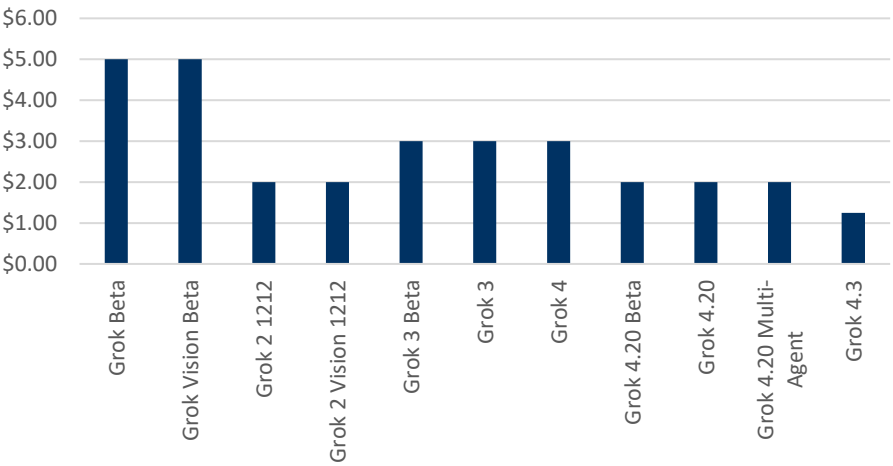
LLM Pricing Trends (xAI)

- **Aggressive late entrant.** xAI held pricing flat through Grok 3/4, then undercut the market with Grok 4.20 and 4.3. At \$1.25/\$2.50, xAI is now the cheapest frontier provider.
- **Output deflation, not inflation.** Output pricing dropped from \$15 to \$2.50 (83%) in under a year, compressing the output/input ratio from 5x to 2x.

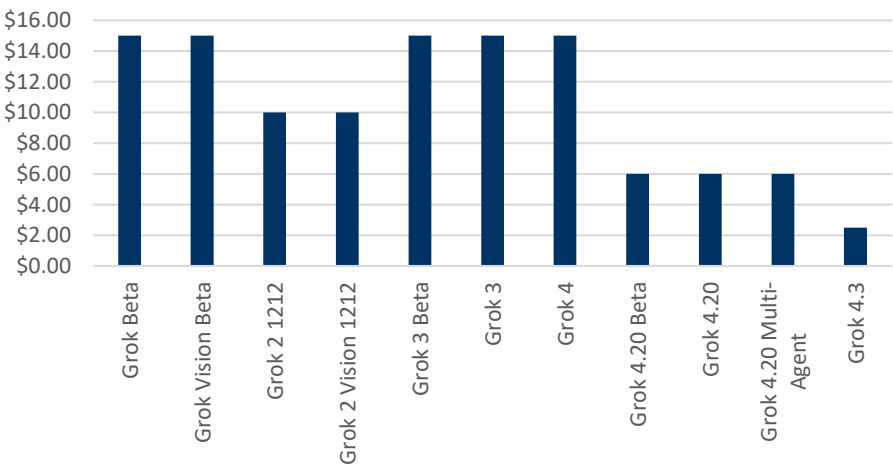
Frontier Models



Input Token Price



Output Token Price

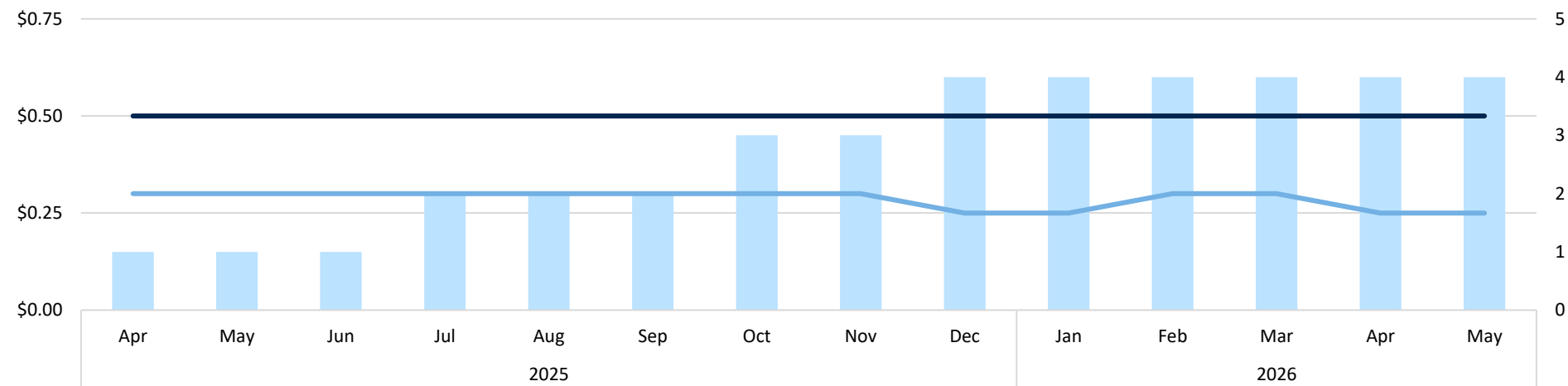


Unique Models Output Token Price (\$/M tokens) Input Token Price (\$/M tokens)

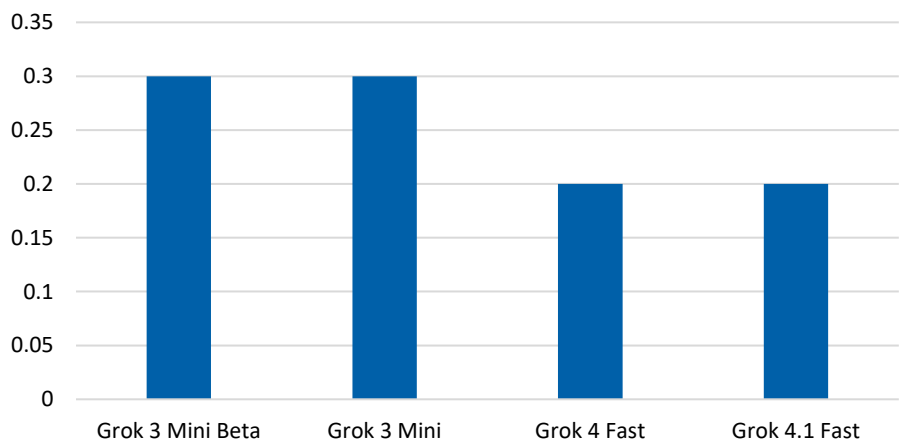
Source: OpenRouter, RBC Capital Markets, RBC Elements

LLM Pricing Trends (xAI)

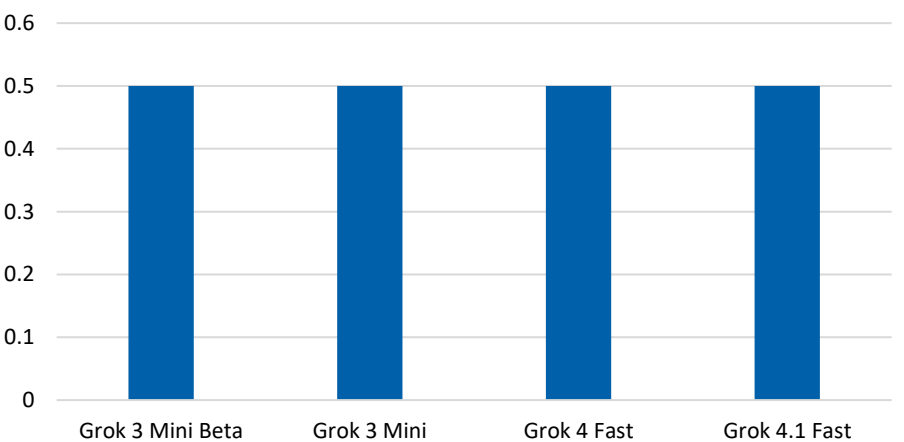
Lightweight Models



Input Token Price



Output Token Price

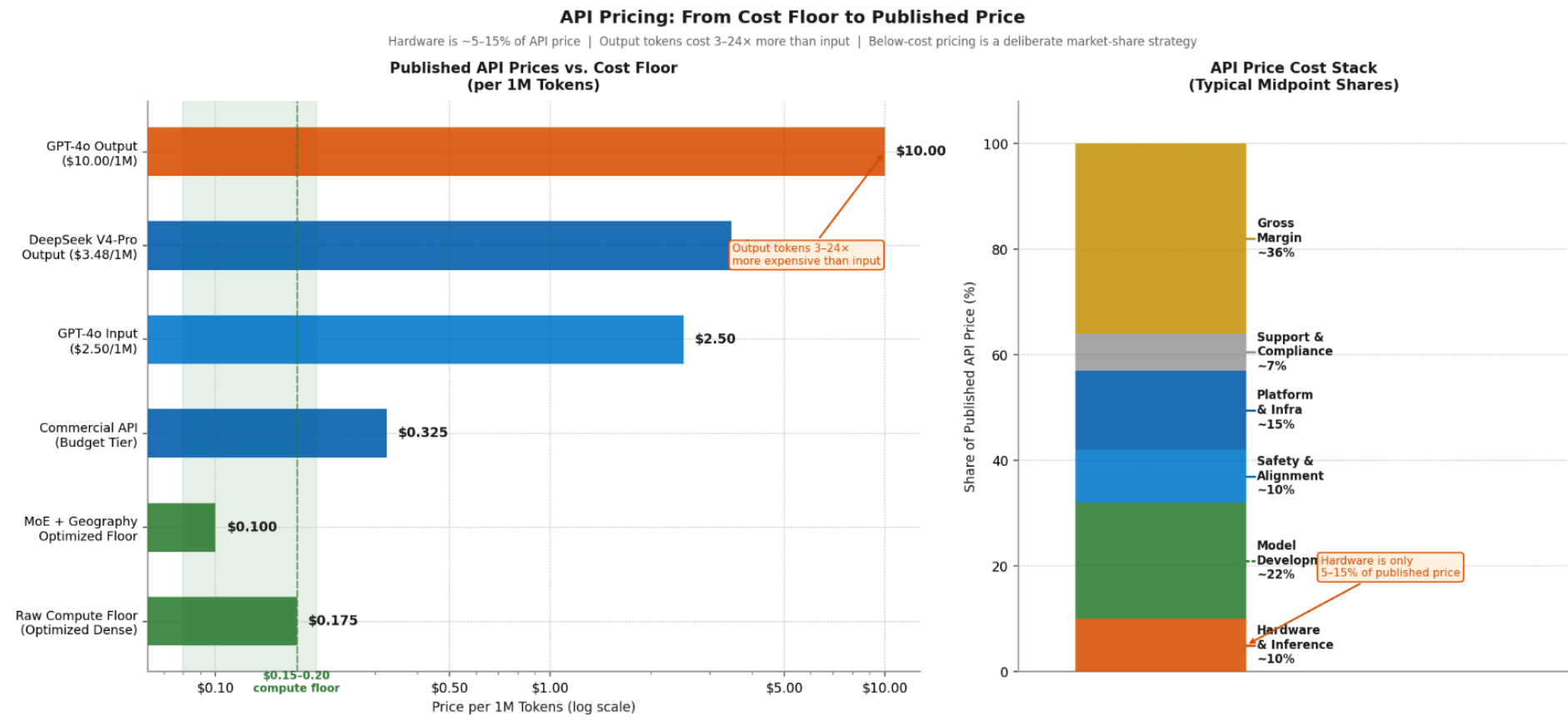


Unique Models Output Token Price (\$/M tokens) Input Token Price (\$/M tokens)

Source: OpenRouter, RBC Capital Markets, RBC Elements

API Pricing: From Cost Floor to Published Price

- We estimate xAI's underlying compute cost floor — at approximately \$0.03–\$0.034/1M tokens (adjusted, C1+C2 mix) — sits well below the \$0.15–\$0.20 raw compute floor applicable to less optimized dense-model deployments. GPU hardware represents only ~10% of a published frontier API price; xAI's COGS advantage of 23% below the H100 industry average operates on that 10% slice, providing a meaningful but not dominant contribution to published API pricing relative to the broader cost stack.



Source: RBC Capital Markets, company reports, MLcommons, Spheron

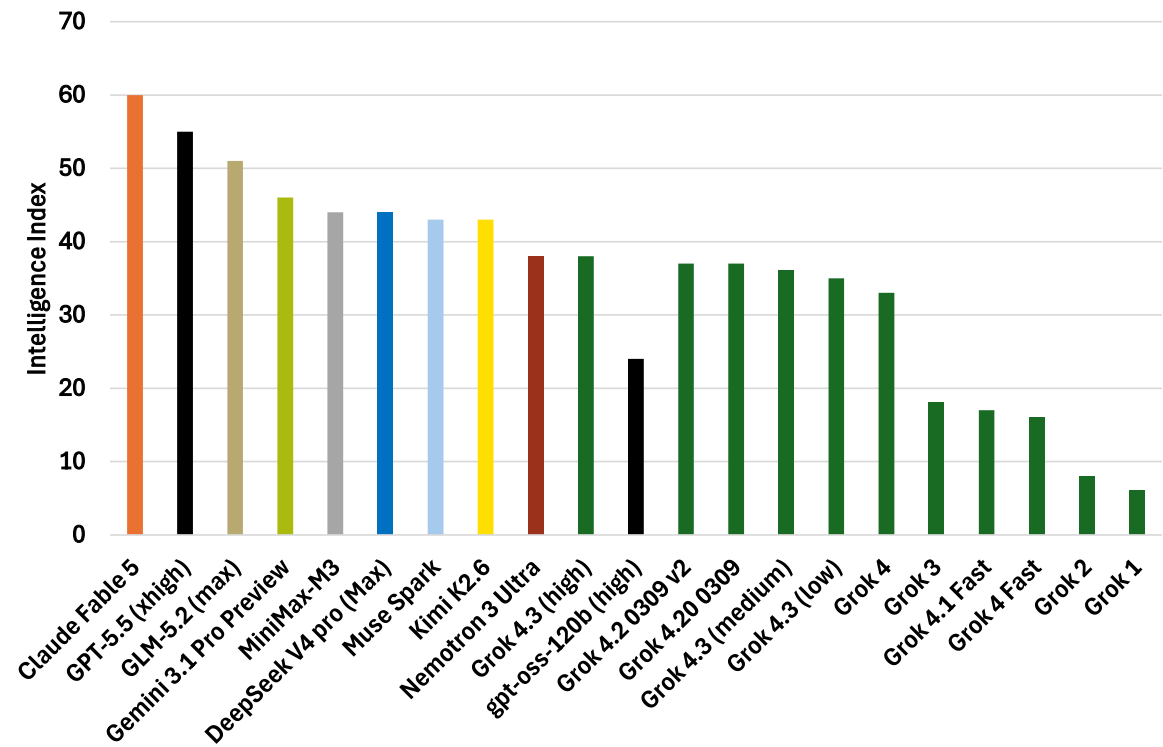
Training Grok – Still room to go

- Grok 4.3 currently trails industry leaders with an intelligence index score of 38 compared to the market-leading score of 60, representing a significant competitive gap.
- **Next Generation - Grok 5**
 - Architecture: 6-10T parameter Mixture of Experts (MoE)
 - Context Window: Native 1.5M tokens
 - Infrastructure: 550K NVIDIA Blackwell GPUs | ~1GW power draw

Grok release timeline w/ parameters & intelligence index

Models	Release Date	Parameters (B)	Intelligence Index
Grok-0	Never released	33	N/A
Grok-1	Nov-23	314	6
Grok-1.5	Mar-24	"Multi-billion"	N/A
Grok-2 & Mini	Aug-24	270	8
Grok-3	Feb-25	3,000	18
Grok-4	Jul-25	1,700	33
Grok-4.2	Feb-26	500	37
Grok-4.3	Apr-26	500	38
Grok-5	May-Jun 2026	6,000-10,000	N/A

Intelligence Index for Grok & others in the Industry

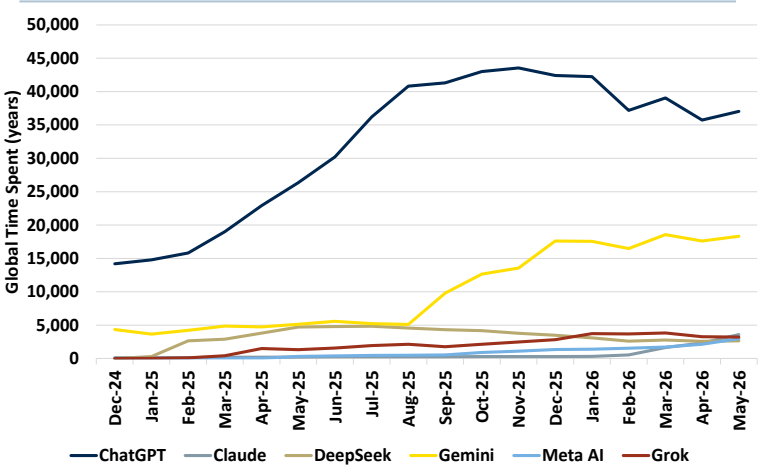


Source: Company reports, Artificial Analysis, RBC Capital Markets

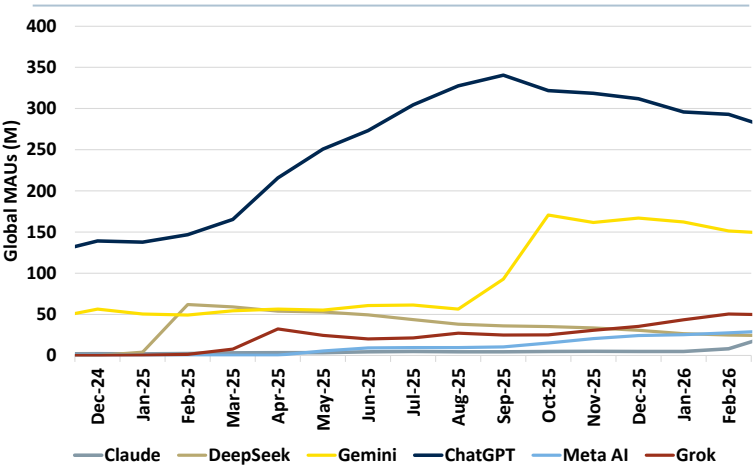
Grok's engagement growth creates a self-reinforcing model improvement cycle

Grok's engagement metrics (e.g., time spent, MAUs) are trending positively but remain dependent on continued model quality improvements and broader awareness. Stronger models drive longer sessions and repeat usage, which in turn generates more training signal — but this flywheel requires parallel investment in consumer marketing and developer outreach to expand the user base feeding it.

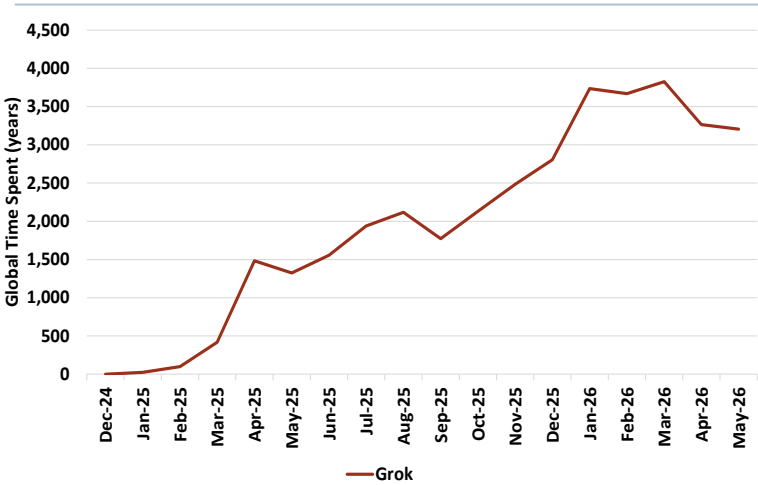
Global Mobile Time Spent steadily increases for Grok



Global Mobile MAUs passed Perplexity in Jan '26



Steady increases in Global Time spent via web for Grok



Source: Apptopia, Similarweb, RBC Capital Markets

Intelligence & Ads



Capital
Markets

Flow of Funds

Regardless of whether the winning architecture is vertically integrated, disaggregated, or router-based, the compute owner consistently captures stable economics. App and model layers face wide margin variance depending on workflow stickiness and pricing power, but the hyperscaler gets paid in every scenario.

Scenario	Chain	Example	App Op Margin	Model Op Margin	Infra Op Margin	Leverage	Key Risk
1. Vertically Integrated	User → Integrated Co.	Google (Gemini), xAI (Grok), Apple	50%	(embedded)	(embedded)	Integrated company	Regulatory/antitrust on bundling
2A. Disaggregated — Thin Wrapper	User → App → Model API → Hyperscaler	Notion AI/Jasper → OpenAI → Azure	45%	35%	35%	Hyperscaler (capacity gate)	Model price war (GPT-5 Nano @ \$0.05); app margin compressed by 40% COGS
2B. Disaggregated — Workflow-Rich	User → App → Model API → Hyperscaler	ServiceNow/SAP → OpenAI → Azure	55%	35%	35%	App (workflow lock-in)	Model cost only 10% of COGS; workflow stickiness insulates pricing
3. Lab/App on 3P Infra	User → Lab/App → Hyperscaler	Anthropic → AWS/GCP	45%	(combined)	35%	Hyperscaler (allocation)	Forward-integration by hyperscaler; concentration on single cloud
4. Model-Agnostic Router	User → Router/App → Multi-Model → Infra	Databricks, enterprise orchestration	67%	54%	35%	Router (workflow lock-in)	Race to zero on model pricing; margin depends on stickiness not access
5. VPC / Ring-Fenced	User → Hyperscaler (managed) → Model (licensed)	Azure OpenAI Svc, Bedrock	42%	50%	42%	Hyperscaler (owns customer + compliance)	Enterprise locked by 6-12mo security review cycle; compliance premium erosion
6. On-Prem / Sovereign	User → Enterprise → HW OEM + OSS Model	Dell/Lenovo + Llama	48%	N/A (open-weight = free)	18%	Enterprise (bears utilization risk)	Needs >60-70% GPU utilization; model freshness degrades without update pipeline
7. Agentic Multi-Model	User → Agent Platform → Model(s) → Hyperscaler	Claude Code, Devin, Cursor	-338%	28%	35%	Model provider (token volume explosion accrues to them)	Cost-revenue scissors: COGS = 420% of revenue; usage grows multiplicatively but flat subscriptions can't keep pace
8. Inference Marketplace	User → App Builder → Inference Platform	Harvey/Abridge → Together AI/Fireworks	81%	N/A (open-weight = free)	28%	App builder (near-zero COGS = unconstrained pricing)	Inference layer commoditizes; platform competes on price in race to bottom; 72% COGS on thin margins

Source: RBC Capital Markets

Company x Scenario Positioning Matrix

Hypercalers (e.g., Microsoft, Amazon, and Google) score as dominant across the majority of deployment scenarios and at minimum participate in the rest — a structural breadth no other player replicates. Standalone model labs and inference platforms are confined to a narrow subset of scenarios at best. We believe choosing compute owners is effectively scenario-agnostic: they seemingly win regardless of whether the market evolves toward vertical integration, agentic workflows, or inference commoditization.

Dominant	Major	Participating
★	●	○

	Arrangements / Architectures / Configurations / Deployments							
	Sc.1 Vertically Integrated	Sc.2 Disaggregated	Sc.3 Lab on 3P Infra	Sc.4 Router	Sc.5 VPC	Sc.6 On-Prem	Sc.7 Agentic	Sc.8 Inference Marketplace
Microsoft	★	★	○	●	★		★	○
Amazon/AWS	★	○	★	○	★		○	●
Google/Alphabet	★	○	●	○	●		○	●
OpenAI		★	●		●		★	
Anthropic		★	★		●		★	
Meta	★							★
NVIDIA	○	○	○	○	○	★	○	○
Together AI				★				★
Databricks				★			○	
Oracle		●			●	●		
CoreWeave		●	○					○
xAI	●		○	○	○		○	

Source: RBC Capital Markets

Infrastructure Moat Determinants

The highest-defensibility moats (e.g., energy access, custom silicon, workflow tie-in) — are levers that hyperscalers have spent years building and that require billions in capital and multi-year development cycles to replicate. Long-dated power contracts, proprietary chip stacks, and identity/productivity lock-in among others collectively create a compounding cost advantage that widens over time.

	Determinant	Mechanism	Cost Impact	Defensibility	Margin Impact	Rationale	Who Benefits	Key Companies
1	Energy Access	Below-market LT contracts (nuclear, hydro); behind-the-meter generation	Lowers variable cost 20-40% vs. grid	5	5	Long-dated PPAs take 5-10 years to replicate and the supply of nuclear/hydro sites is physically finite. 20-40% variable cost advantage compounds over every GPU-hour and widens as grid demand outstrips supply.	Hyperscalers w/ PPA portfolios	MSFT, AMZN, GOOG, CoreWeave
2	Scale Without Caps	Permitted capacity, grid interconnection queue, water rights, zoning	Removes revenue ceiling; avoids turning away demand	4	5	Determines whether revenue can grow without turning away customers during the 2025-27 capacity crunch. Land banks and grid queue positions are hard to replicate quickly but not permanent monopolies — new sites eventually get permitted.	AWS, GOOG (large land banks), xAI	AWS, GOOG, xAI
3	Data Freshness Layer	Real-time data ingestion (X firehose, enterprise data lakes) for RAG	Justifies 2-3x premium on enriched inference	3	5	Justifies a 2-3x pricing premium (revenue lever, not cost lever) but doesn't reduce \$/token. Defensibility is maximum because proprietary data sources (X firehose, LinkedIn Graph, enterprise data lakes) are exclusive and non-replicable.	xAI ; MSFT (LinkedIn/Graph); Snowflake/Databricks	xAI , MSFT, Snowflake
4	Physical Stack (custom silicon)	Custom silicon (TPU, Trainium, Maia), networking, co-designed storage	30-50% lower \$/token vs. renting NVDA GPUs	5	5	30-50% lower \$/token vs. renting NVIDIA GPUs — the largest controllable cost lever. Requires \$1B+ and 3–5-year development cycles to replicate; internal chip→compiler→model→serving optimization loop compounds across generations.	GOOG (TPU v6), AMZN (Trainium2/Inferentia 3), MSFT (Maia 100), xAI (TeraFab)	GOOG (TPU), AMZN (Trainium), MSFT (Maia), xAI (TeraFab)

Source: RBC Capital Markets

Infrastructure Moat Determinants (cont.)

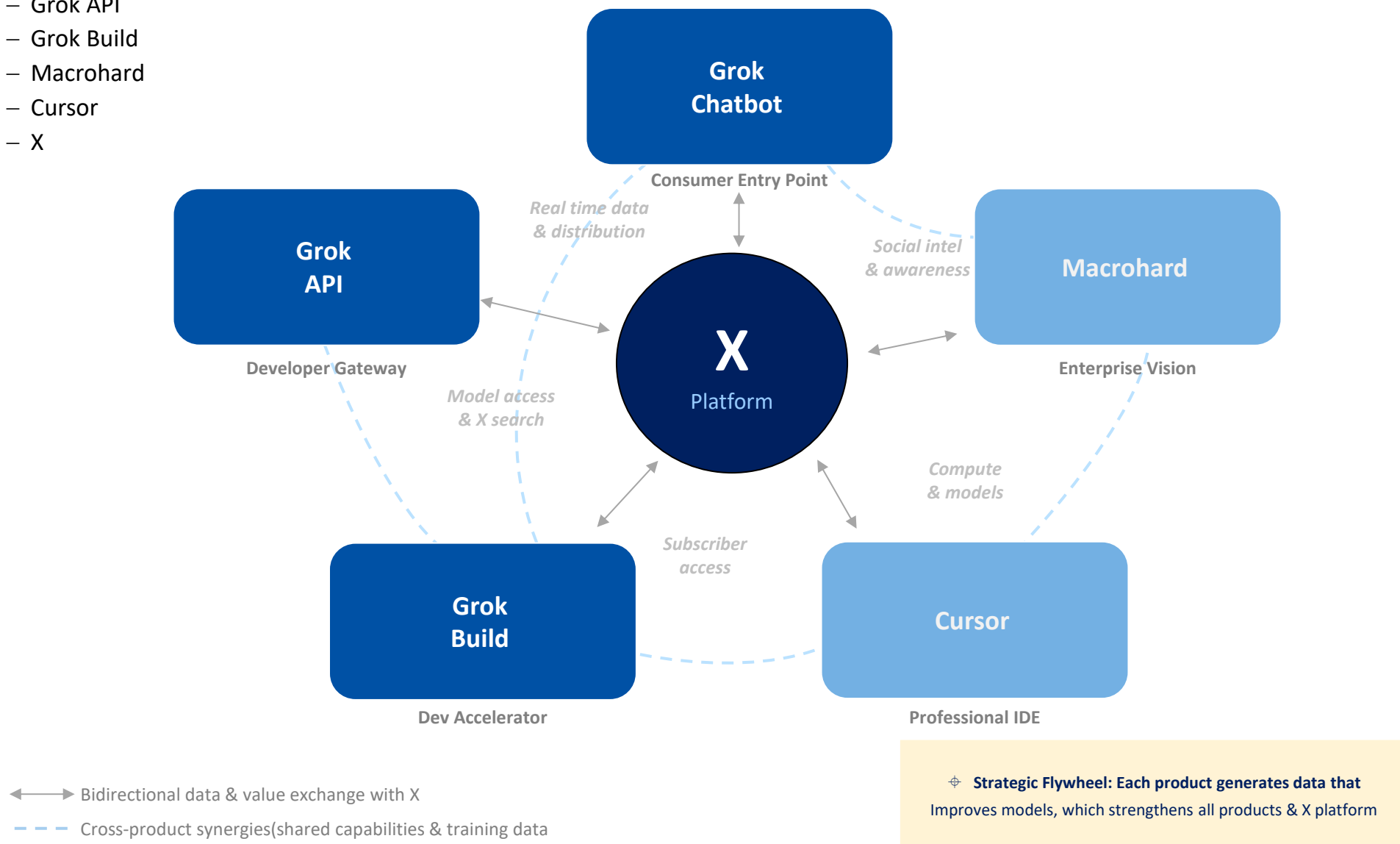
	Determinant	Mechanism	Cost Impact	Defensibility	Margin Impact	Rationale	Who Benefits	Key Companies
5	Workflow Tie-In	Identity (Entra), productivity (M365), DevOps (GitHub), ERP integration	Increases switching cost → pricing power; upsell surface	5	5	Doesn't reduce \$/token but makes pricing power unconstrained — AI embedded in M365/GitHub/Oracle ERP can't be comparison-shopped independently. Displacing these relationships requires multi-year full-stack migration, not weeks of API integration.	MSFT (M365+GitHub+Entra), GOOG (Workspace), Oracle (ERP)	MSFT (M365+GitHub), GOOG (Workspace), Oracle (ERP)
6	Committed Spend Contracts	Multi-yr commits (MACCs, EDPs) w/ AI consumption credits bundled	De-risks capex recovery; customer pre-pays floor utilization	3	3	De-risks capex recovery and guarantees floor utilization but doesn't lower unit cost of delivering compute. Every hyperscaler can offer identical contract structures — the moat is proportional to book size (MSFT \$368B), not structural.	All hyperscalers; larger contractual book = more leverage, xAI	All hyperscalers; MSFT, xAI
7	Open-Weight Hosting Speed	Proprietary serving stack (FlashAttention 4, kernel optimization, speculative decoding) on dedicated GPU fleet tuned for open-weight models; delivers 3-10x cost improvement vs. general-purpose hyperscaler inference.	Cuts \$/token up to 10x vs. hyperscaler general-purpose inference; enables floor pricing at lower thresholds	2	2	Up to 10x \$/token savings vs. general-purpose cloud, but inference optimization is replicable (FlashAttention is open-source, hardware vendors enable all players equally). Operates at thin 28% margin with commoditization risk as open-weight models consolidate.	Specialized inference platforms with proprietary serving stacks and research-to-production pipelines; benefits any app builder using open-weight models	Together AI, Fireworks, Baseten
8	Model Distribution / User Base	Installed base of users on proprietary platform (ChatGPT 300M+ users, Claude.ai, Gemini) creates default consumption channel for inference revenue	Lowers CAC to near-zero for incremental token revenue; amortizes training cost over massive query volume	4	4	Self-reinforcing flywheel (users → feedback → better model → more users) guarantees token volume that amortizes training costs. Takes years to build — no shortcut to hundreds of millions of MAUs.	Model labs with consumer-facing products	OpenAI (ChatGPT), Google (Gemini), Anthropic (Claude.ai), xAI (Grok)
9	Inference Cost Curve / GPU Utilization	Sustained >70% GPU utilization through workload mixing (training + inference), spot/preemptible scheduling, and multi-tenant orchestration	Directly determines whether infra margin is 28% (low utilization) or 42% (high utilization); 15pp margin swing from 60% → 80% utilization	4	3	The single most impactful operational variable for infra margins. Requires diverse workload mixing and deep scheduling infrastructure, but any scaled provider can eventually achieve it — an execution advantage correlated with scale, not a permanent structural moat.	Hyperscalers with diverse workload base; providers mixing training + inference	AMZN (AWS), MSFT (Azure), GOOG (GCP), CoreWeave, xAI

Source: RBC Capital Markets

One Flywheel, Compounding Returns: Data → Models → Stronger Products

- When thinking about Space X AI Solution & Infrastructure there are 6 main products:

- Grok chatbot
- Grok API
- Grok Build
- Macrohard
- Cursor
- X



Organizational size will ultimately determine each company's enterprise AI strategy

- With hundreds of companies emerging across infrastructure, model, middleware and applications, each organization's AI strategy will vary based on its size, resources, data privacy and compliance requirements.
- **SMBs**
 - Primary plan: Subscription, API
 - Inference pattern: sporadic (human driven)
 - Tokenomics: Commodity API pricing
 - Beneficiaries: Software companies, model labs, hyperscalers
- **Mid-tier companies**
 - Primary plan: hybrid (managed API or virtual cloud)
 - Inference pattern: more predictable/periodic workflows
 - Tokenomics: Low API pricing through distilled/tuned models
 - Beneficiaries: Deeply integrated software, model labs, hyperscalers
- **Large organizations**
 - Primary: hybrid (On-prem + cloud for excess/training)
 - Inference pattern: near constant from agent loops + frontier training & bursts
 - Tokenomics: Amortized hardware cost, infrastructure maintenance and hyperscaler consumption
 - Beneficiaries: On-prem hardware, hyperscalers and to lesser degree, specialized model builders
- **Frontier Models cost: \$5-30/million tokens**
 - Open-Weight Models: \$0.05-\$2/million tokens
 - Replaced frontier model token costs can drop up to 90+%
 - Token consumption grows by thousands of percent when prices drop like this
- **Primary Risk Factors**
 - Companies must deploy ring-fenced models through cloud partners or owned infrastructure to protect core IP
 - Cannot expose data through public endpoints or create single vendor dependency as third-party model routing providers proliferate
- **Secondary Considerations: Jevons Paradox: Lower compute costs drive broader AI exploration and higher ROI potential**
 - Frontier models often deliver lower total cost of completion despite higher token prices
 - Owning physical infrastructure enables model-agnostic approaches, though ring-fencing open-weight models remains inefficient except for largest enterprises
- **Competitive Positioning Strategy**
 - SpaceX/Grok Approach: Price software and Grok model access below competing frontier providers' retail rates. Ring-fenced open weights face limitations in information freshness and utilization efficiency

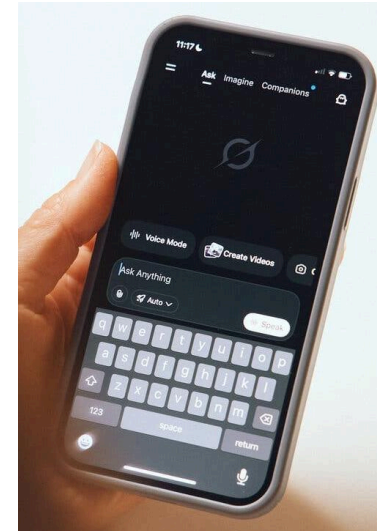
Harness defensibility

- **What defines a harness?**
 - An interface layer that connects an AI model to a specific environment or workflow:
 - **Provides context** about environment (files, structure, dependencies)
 - **Translates** user intent into model prompts
 - **Executes** model outputs (edits, commands, actions)
 - **Manages feedback loop** between model and environment
- **Examples in Practice**
 - **Cursor** – Harness for code development (wraps virtually any model into a full IDE)
 - **Cowork** – Agentic harness for knowledge workers (desktop AI assistant)
 - **GitHub Copilot** – Harness for OpenAI Codex (integrates into VS Code, manages context, applies suggestions)
- **Determinants of a plausible harness moat**
 - Integration depth – Deep embedding with IDEs, build systems, version control creates high switching costs
 - Context management – Sophisticated understanding of codebases, project structure, and workflows that models alone can't replicate
 - User experience – Learned interaction patterns and workflow optimizations that become habitual
 - Data flywheel – Collection of usage data improves suggestions over time (with permission)
 - Ecosystem lock-in – Extensions, configurations, team setups built around specific harnesses
- **The model is commoditized; the harness captures value**

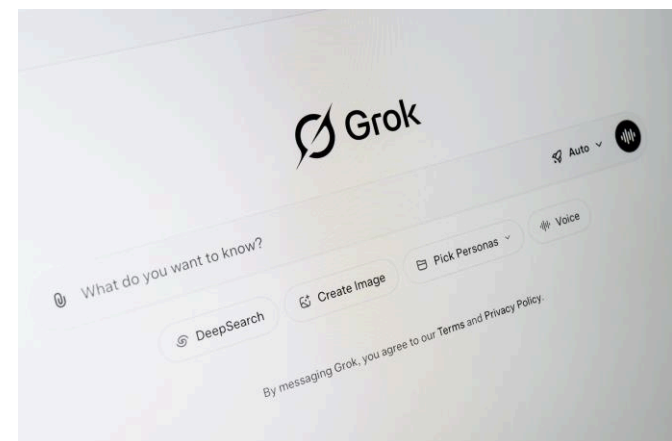
Grok Chatbot – xAI's initial foray into the consumer

- The Grok chatbot is an AI chatbot developed by xAI that offers web-based access and a standalone app.
- The most recent flagship model is Grok 4.3 (released Jun '26) with the first model for public rollout released in Nov '23.
- **Core capabilities:**
 - Real-time X platform integration: Direct access to public posts and live conversations on X
 - Breaking news awareness: Stays current with trending topics and time-sensitive information
 - Image generation: AI-powered visual content creation
 - Code writing & debugging: Advanced software development assistance
 - Voice interaction: Natural, hands-free conversation features
 - Multimodal functionality: Supports text, image, code, and voice inputs/outputs
- **Key Differentiators:**
 - Unique X integration: Real-time access to social media data unavailable to competing chatbots
 - Social media intelligence: Analyzes trending topics, sentiment, and breaking news as it happens
 - Conversational personality: More casual and humorous tone compared to competitors
 - Current awareness: Maintains up-to-the-minute knowledge through X platform connection
- **Access Tiers & Pricing:**
 - Free Tier: Limited daily queries for testing core functionality
 - SuperGrok: \$30/month with enhanced access and higher usage limits
 - Heavy Tier: \$300/month for professional users requiring maximum capacity
 - X Platform Bundles:
 - Basic: \$8/month (includes limited Grok access)
 - Premium+: From \$40/month (includes expanded Grok features)

Grok Mobile App



Grok Web-Based Access

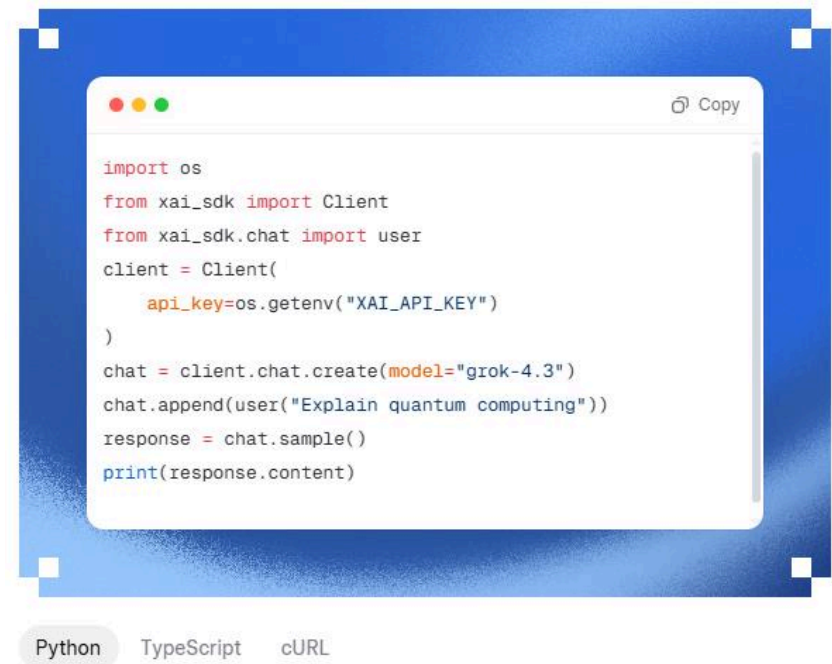


Grok API – Scaling intelligence through infrastructure

- The Grok API, released Oct '24, is a developer toolkit that enables integration of Grok AI models into custom applications, which enables developers to build AI-powered features without managing underlying infrastructure.
- **Core capabilities:**
 - Three types of API:
 - Voice API: Natural language voice interaction and speech processing capabilities
 - Imagine API: AI-powered image generation and visual content creation
 - Live Search API: Real-time semantic search across web, X platform, and news sources
 - Other capabilities:
 - Model Access: Direct integration with xAI's flagship Grok models
 - Real-time Intelligence: Up-to-the-minute information retrieval and analysis
 - Multimodal Processing: Combined text, voice, image, and search functionalities
- **Key Differentiators:**
 - Real-time Search Integration: Unique live data access from X platform unavailable in competing APIs
 - Current Information Advantage: Semantic search across multiple real-time sources (web, social, news)
- **Access Tiers & Pricing**
 - Pay-as-you-go Model: Token-based pricing with no monthly minimums
 - Grok 4.3 (Full Reasoning):
 - Input: \$1.25 per 1M tokens
 - Output: \$2.50 per 1M tokens
 - Non-Reasoning Variant (for lighter tasks):
 - Input: \$0.20 per 1M tokens
 - Output: \$0.50 per 1M tokens

Source: Company reports, RBC Capital Markets

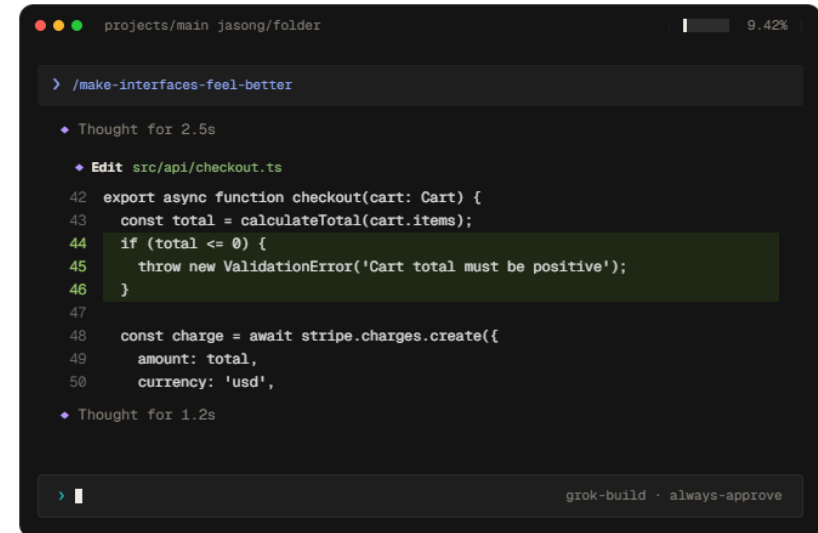
Grok Mobile App



Grok Build – xAI’s move from product to platform

- The Grok Build, released May '26, is an AI-powered coding agent designed for automated software development that assists developers with intelligent code generation and development planning tools.
- **Core capabilities:**
 - Plan-first mode: Drafts structured development plans before writing any code
 - Multi-agent architecture: Main agent delegates tasks to specialized subagents
 - Automated code generation: Produces code based on approved plans and specifications
 - Task decomposition: Breaks complex projects into manageable subtasks
- **Key Differentiators:**
 - Structured planning approach: Review development strategy before implementation begins
 - Multi-agent delegation system: Specialized subagents handle specific coding tasks efficiently
 - Concurrent processing: Multiple subagents execute tasks in parallel for improved speed
- **Access Tiers & Pricing**
 - SuperGrok Subscribers: Full access included in \$30/month subscription
 - X Premium Plus Subscribers: Access bundled with Premium Plus plan
 - API Key option: Pay-per-token pricing model for programmatic access
 - Flexible payment models: Choose between subscription or usage-based pricing

Grok Build- Skills



The screenshot shows a terminal window with the prompt `> /make-interfaces-feel-better`. It displays a thought process for 2.5s and then an edit to `src/api/checkout.ts`. The code defines an async function `checkout` that calculates the total of cart items and throws a `ValidationError` if the total is not positive. It also creates a Stripe charge. A second thought process for 1.2s is shown below the code.

```
> /make-interfaces-feel-better

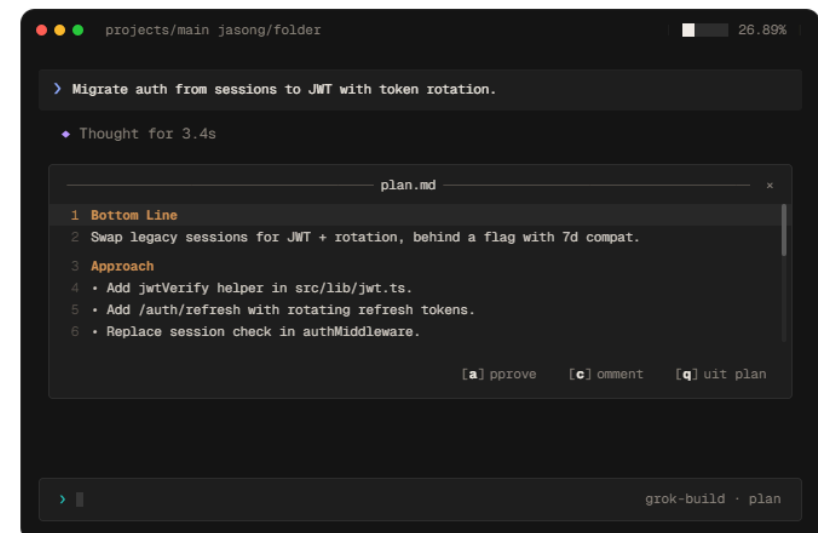
Thought for 2.5s

Edit src/api/checkout.ts
42 export async function checkout(cart: Cart) {
43   const total = calculateTotal(cart.items);
44   if (total <= 0) {
45     throw new ValidationError('Cart total must be positive');
46   }
47
48   const charge = await stripe.charges.create({
49     amount: total,
50     currency: 'usd',
51   });
52 }

Thought for 1.2s

grok-build · always-approve
```

Grok Build- Plan



The screenshot shows a terminal window with the prompt `> Migrate auth from sessions to JWT with token rotation.`. It displays a thought process for 3.4s and then a plan in `plan.md`. The plan includes a bottom line, a task to swap legacy sessions for JWT with rotation, and an approach with three steps: adding a `jwtVerify` helper, adding `auth/refresh` with rotating refresh tokens, and replacing the session check in `authMiddleware`. At the bottom, there are buttons for `[a] pprove`, `[e] comment`, and `[q] uit plan`.

```
> Migrate auth from sessions to JWT with token rotation.

Thought for 3.4s

plan.md
1 Bottom Line
2 Swap legacy sessions for JWT + rotation, behind a flag with 7d compat.
3 Approach
4 • Add jwtVerify helper in src/lib/jwt.ts.
5 • Add /auth/refresh with rotating refresh tokens.
6 • Replace session check in authMiddleware.

[a] pprove [e] comment [q] uit plan

grok-build · plan
```

Macrohard – The next-gen workplace platform

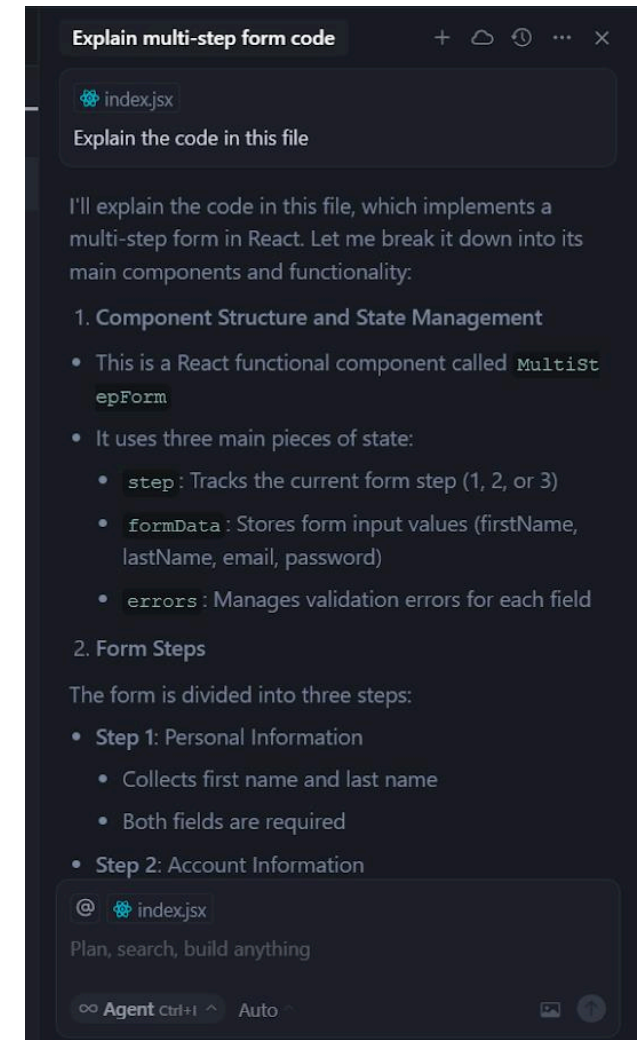
- **Macrohard is an AI-powered agentic platform forming the next-generation digital workforce automation system currently in active development through a strategic partnership with Tesla.**
- **Core capabilities:**
 - End-to-End Digital Workflow Automation: Manages complete processes from coding to management
 - Intelligent Co-Pilot Functionality: Goes beyond traditional automation as active work partner
 - Emulation and Augmentation: Replicates and enhances how humans use computers across all work types
 - Comprehensive Work Coverage: Handles diverse tasks spanning technical to strategic functions
- **Key Differentiators:**
 - Tesla Partnership Advantages: Access to cutting-edge physical AI capabilities and processors
 - Frontier AI Models: Combines state-of-the-art AI with Tesla’s specialized technology
 - Vertical Integration Benefits: Leverages Tesla’s integrated stack for superior performance and cost efficiency
 - Intelligent Co-Pilot vs. Task Executor: Active reasoning partner rather than simple automation tool
 - Holistic Work Approach: Addresses entire spectrum of work activities, not isolated tasks
- **Access Tiers & Pricing**
 - Not available yet

Cursor – Capturing developers at the point of work

- **Cursor is an AI-native development environment (IDE) for professional developers featuring embedded LLM agents and workflows.**
- **Strategic rationale:**
 - Vertical Integration Extension: Completes stack from compute infrastructure → models → applications
 - Compute Capacity Leverage: Utilizes SpaceX's significant infrastructure for competitive advantage
 - First enterprise application exposure for xAI—shifts value capture from infrastructure to high-margin developer workflows
 - Developer activity (code generation, debugging, acceptance/rejection) creates continuous training signal for Grok
- **Core capabilities:**
 - LLM Agent Integration: Embedded AI agents that understand and assist with coding tasks
 - AI-Powered Workflows: Intelligent automation of common development patterns and processes
 - Code Generation and Assistance: Real-time code suggestions, completions, and problem-solving
 - Professional Developer Focus: Enterprise-grade features for serious software development teams
- **Key Differentiators:**
 - In-house frontier model: specifically optimized for code generation
 - AI-Native IDE Functionality: Development environment designed specifically around AI capabilities
 - AI-Powered Workflows: Intelligent automation of common development patterns and processes
 - Code Generation and Assistance: Real-time code suggestions, completions, and problem-solving
 - Professional Developer Focus: Enterprise-grade features for serious software development teams
- **Access Tiers & Pricing:**
 - Hobby (Free) with limited agent requests
 - Pro (\$20/month)
 - Pro+ (\$60/month)
 - Ultra (\$200/month)
- **Vertical Integration Advantage**
 - Only competitors with full-stack control: Google
 - Direct compute access: Wholesale token pricing + ability to use excess capacity for Grok training (precedent: Composer saved \$100M+ via self-hosted RL tuning)
 - Peak efficiency: Route inference at lower token costs than any retail-compute-dependent provider



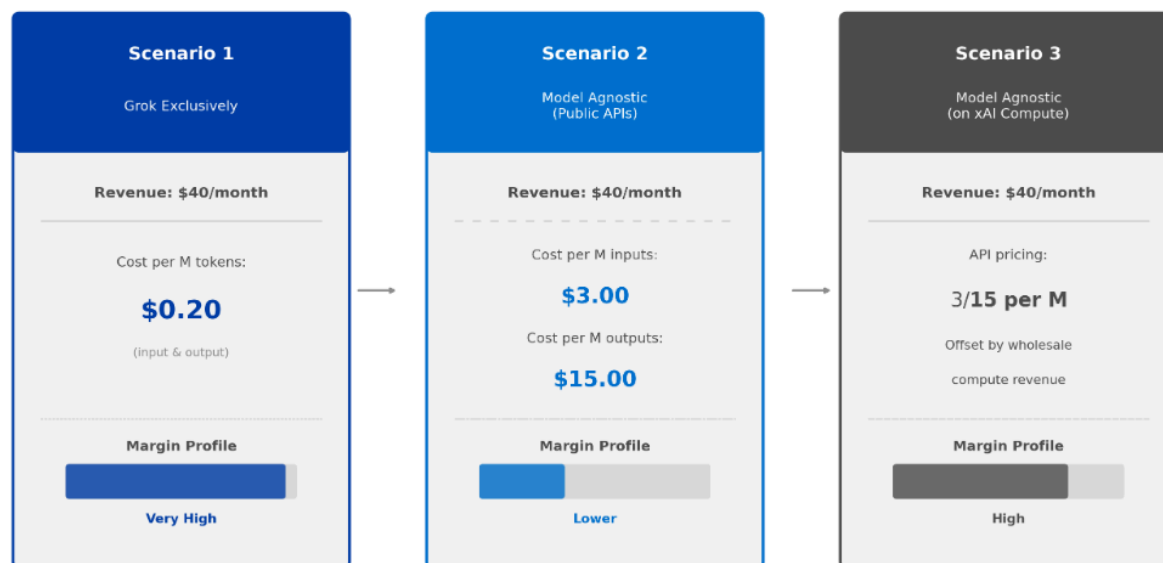
Cursor AI Code Editor



SpaceX + Cursor is a production engine

- An app provider runs the orchestration layer: users pay subscription or usage-based fee & app manages agentic execution loop
 - Multi-step agentic loops can be several multiples individual token costs – lots of margin sensitivity
- 3 approaches to model allocation app providers can offer:
 - Retail: Lab frontier models: necessary to create new workflows/tasks but expensive both initially and as workflows grow
 - Wholesale/bring your own key: Volume discounts from model providers or user inputs own API, which reduces the moat
 - Open source: App provider hosts open-source models on rented cloud compute; better margins than frontier APIs but still a price taker from compute owners
- What Cursor brings that's so important for having preserved the economics even when they didn't have claim to infrastructure
 - Speculative edits: only spend tokens focused on necessary coding changes, not entire coding file
 - Model routing: being embedded in the IDE created understanding of easy tasks vs. hard
 - Enterprise vs. pro account arbitrage: charge more for enterprise seats where persistence & latency requirements have proven lower

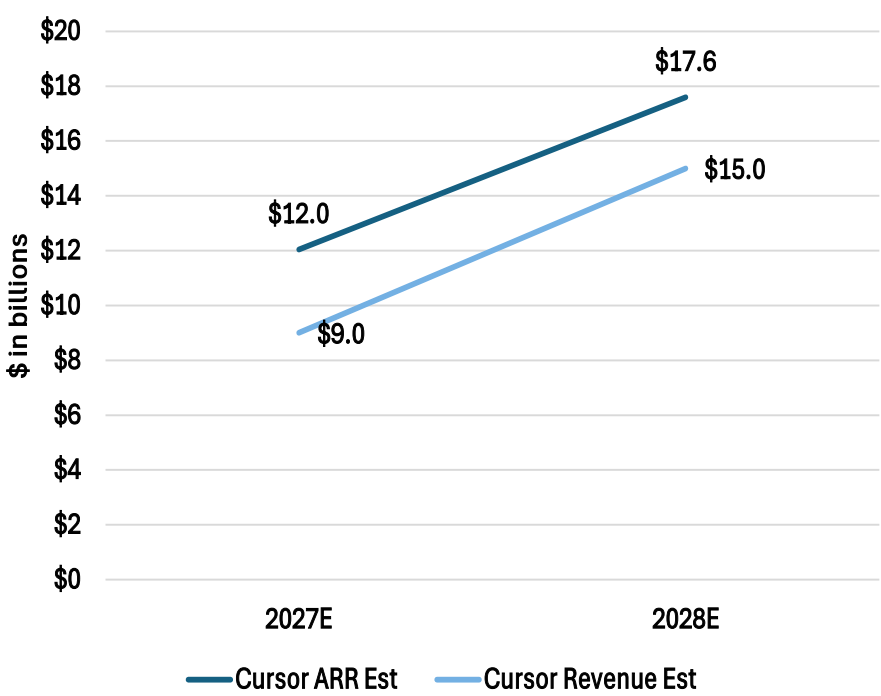
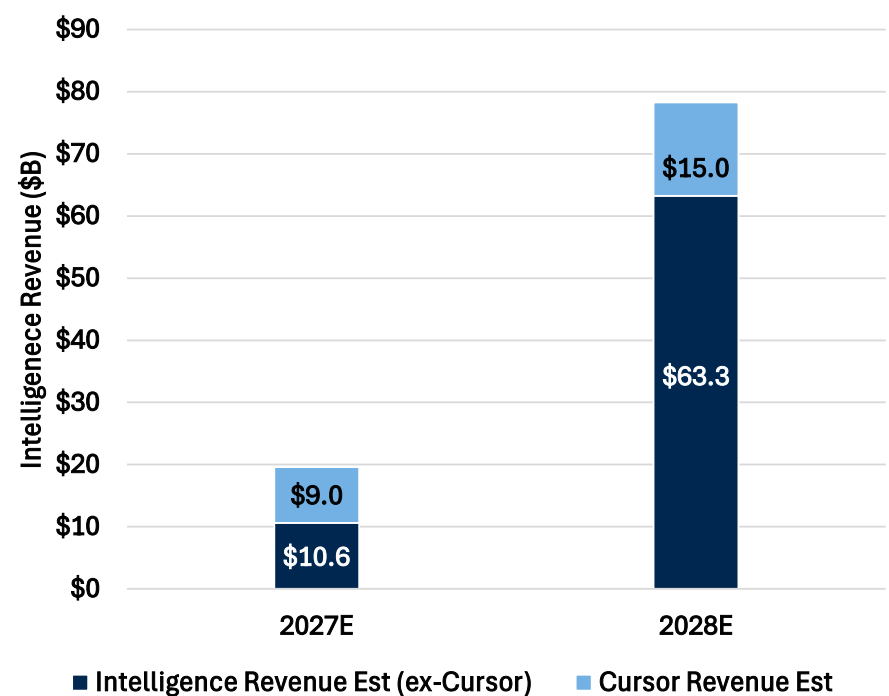
\$40 Subscription scenarios



Source: Company reports, RBC Capital Markets

Cursor revenue impact estimate post-deal close

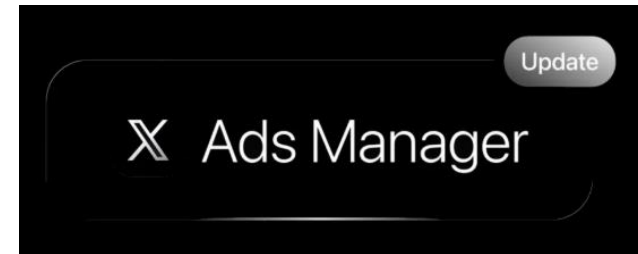
- **Revenue impact:** We estimate the addition of Cursor could potentially drive \$9B and \$15B in '27 and '28
 - We have not included this in our estimates as we await further guidance from the company, but we'd expect at least half in '27 to be incremental while this may already be largely embedded in '28.



Source: Company reports, RBC Capital Markets estimates

X Ad Stack Rebuild: adding AI

- **X Ad Stack Rebuild – Phased Rollout (April 2026)**
- **Key Features:**
 - Modern AI-powered retrieval and ranking system
 - Enhanced marketer controls for targeted campaign creation
 - Integration with xAI following 2025 merger
 - Iterative deployment model (no single launch date)
- **Technical Improvements:**
 - Retrieval: AI determines which impressions and users enter the ad pool
 - Ranking: Advanced logic prioritizes candidates within auction and pacing systems
- **Early reviews from X Ads Manager rollout:**
 - "Sales" objective (for website conversions) now shows "Coming Soon" label and is unavailable
 - Audience targeting drastically reduced—only "Follower Look-alike" remains accessible
 - Previously available options (Keywords, Conversations, etc.) have been removed



- **Three New Capabilities (Now Available- as of 6.17.26):**
 - Google Tag Manager Integration
 - No-code setup for X Pixel and Conversion API (CAPI)
 - Eliminates need for developer support during campaign configuration
 - Consolidated Developer Tools
 - Single hub in Events Manager for all Conversion API resources
 - Centralized access to tracking data and management tools
 - Real-Time Conversion Diagnostics Dashboard
 - Live monitoring of Pixel and CAPI event health
 - Immediate identification and resolution of tracking issues

Recently launched (April) X Ad Stack Rebuild: What signals does X have vs. Meta?

- X’s ad estimates out to 2031 imply reaching ARPU levels that approach META’s ~\$50+
- No SMID-cap ad platform has ever come close to these levels. Why? Largely, engagement & behavioral-generated signal challenges vs. Facebook & Instagram
- Tale of Two Approaches

Meta: Behavioral Prediction Engine
"Who should see this ad?"

POWERED BY

Andromeda — nearly a decade of development, publicly discussed 2023

ARCHITECTURE

Candidate retrieval layer scans billions of ads, retrieves most relevant subset before final auction

SIGNALS

Cross-platform engagement (FB, IG, WhatsApp), Reels, shopping, pixel events, Conversion API, on-platform commerce

ADVANTAGE
Closed-loop feedback — knows what happened after the click

X: Real-Time Intent Detection
"What is this user interested in right now?"

POWERED BY

Grok/xAI models — early in transition from traditional to AI-driven delivery

ARCHITECTURE

Contextual and semantic signals determine ad placement

SIGNALS

Real-time public conversations, trending topics, community notes processed via Grok/xAI

ADVANTAGE
Captures intent in the moment vs. historical behavior

- The Viability Question: Can X’s Signals Compete?

X's Structural Challenges

- **Scale gap:** Smaller user base limits training data volume vs. Meta's billions across multiple platforms
- **Model maturity:** Grok/xAI is early-stage; Andromeda has a decade of refinement
- **Feedback loop:** Limited commerce ecosystem = weaker post-click conversion data
- **Signal diversity:** Single platform vs. Meta's cross-app data enrichment

X's Potential Edge

- **Intent advantage:** "Interest right now" signals could outperform historical predictions for certain verticals
- **Unique data asset:** Real-time conversational data difficult to replicate elsewhere

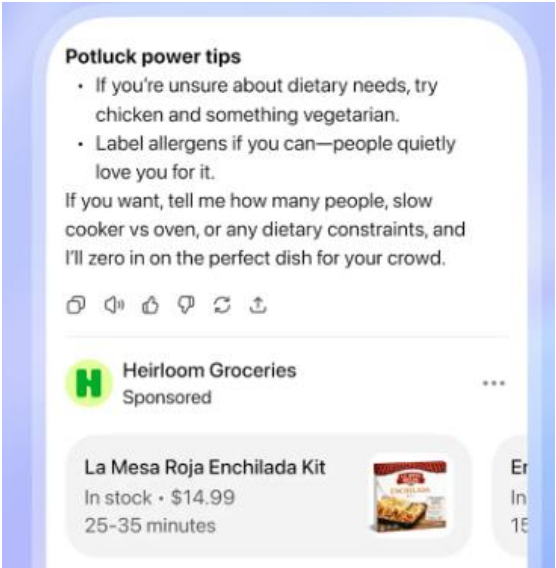
Source: Net Influencer, TechCrunch, Reddit, X, RBC Capital Markets

ChatGPT Ads: How is it going?

- Launched Feb '26 and going from an estimated \$2.4-2.5B and reaching a \$100B target by 2030E

Ad Format & Placement	Geographic Reach & Partners
Where: Sponsored box below organic answers Frequency: 84% of queries show no ads Privacy: No conversation data shared with advertisers; users control personalization Who: Free & Go tier users only	Live Markets: US, Canada, UK, Australia, New Zealand, Japan, South Korea Coming Soon: Brazil, Mexico, India Partners: Criteo, StackAdapt, Adobe, LiveRamp, CallRail, major agency holding companies

Current ChatGPT Ad Format



ChatGPT Ad Partners:



Source: OpenAI, Company reports, RBC Capital Markets

ChatGPT Ads: Scaling the Conversational Commerce Opportunity

- LLM search even more contextual given significantly longer query lengths.
- ChatGPT's challenge is attribution but building all the right tools like CAPI, conversion pixel and advertiser API.

- **Platform Updates**
 - New 'Conversions' Objective: ChatGPT Ads Manager now includes a dedicated conversions objective, allowing advertisers to optimize campaigns specifically for conversion actions rather than just awareness or engagement
 - Self-Serve Ads Manager Beta: OpenAI has launched a beta version of their self-serve Ads Manager interface, replacing the previous spreadsheet-based setup process and making campaign creation more intuitive and accessible
 - Enhanced Budget Controls: Advertisers can now set average daily budgets with weekly pacing allocations to smooth spend over time, plus the option to switch campaigns to lifetime budgets for greater flexibility in campaign management
- **New Ad Formats**
 - Dedicated E-Commerce Formats: New ad formats designed specifically for e-commerce advertisers, supporting both portrait and landscape image orientations to accommodate different creative assets and placements
 - Richer Creative Options: Advertisers can now use larger, more prominent images alongside personalized call-to-action buttons including options like "Book Now," "Sign Up," "Shop Now," and other action-oriented CTAs that drive specific user behaviors
 - Multi-Ad Placements (Testing): OpenAI is currently testing a new format where multiple relevant ads from different advertisers are grouped together in a single placement, potentially increasing visibility while maintaining user experience

Conversational Commerce vs. the Click Economy

- **ChatGPT Ads vs. Google Search Ads.** The below table highlights how ChatGPT's ad model departs from Google's across every structural dimension — from context-based targeting and single-slot inventory to a relevance-weighted auction—creating a fundamentally different UX for both users and advertisers. The key gating factor remains whether OpenAI can expose enough measurement & attribution depth to prove ChatGPT isn't just doing upper-funnel work that Google then takes credit for downstream.

Dimension	Google Search Ads	ChatGPT Ads
Who Sees Ads	All users	Free & Go tiers only; no ads on Plus, Pro, Business, or users <18
Ad Format	Text ads, shopping ads, display ads above/beside results	Single unit below conversation: advertiser name, favicon, headline, description, landing page, image
Targeting	Keyword match bidding (exact, broad, phrase)	Context/intent based (topics/keywords describing relevant conversations); not exact-match
Auction	Quality Score × bid (keyword relevance, landing page quality, expected CTR)	Relevance-weighted second-price auction
Pricing	CPC, CPM, CPA options available	CPM (Reach) or CPC (Clicks); recommended starting CPC: \$3–\$5
Ad Density	Multiple ads per results page	One ad per response
Measurement	Conversions, click-through rate, view-through, cross-device	Ads Manager Beta: impressions, clicks, CTR, CPC, CPM, conversions; UTM parameters supported
Brand Safety	Advertiser exclusions + third-party verification (IAS, DV)	OpenAI policy: ads only near safe/appropriate chats; no third-party verification announced
Answer Independence	Ads displayed separately from organic results	OpenAI states ads do not influence answers; users can disable personalization

Source: Company reports, RBC Capital Markets

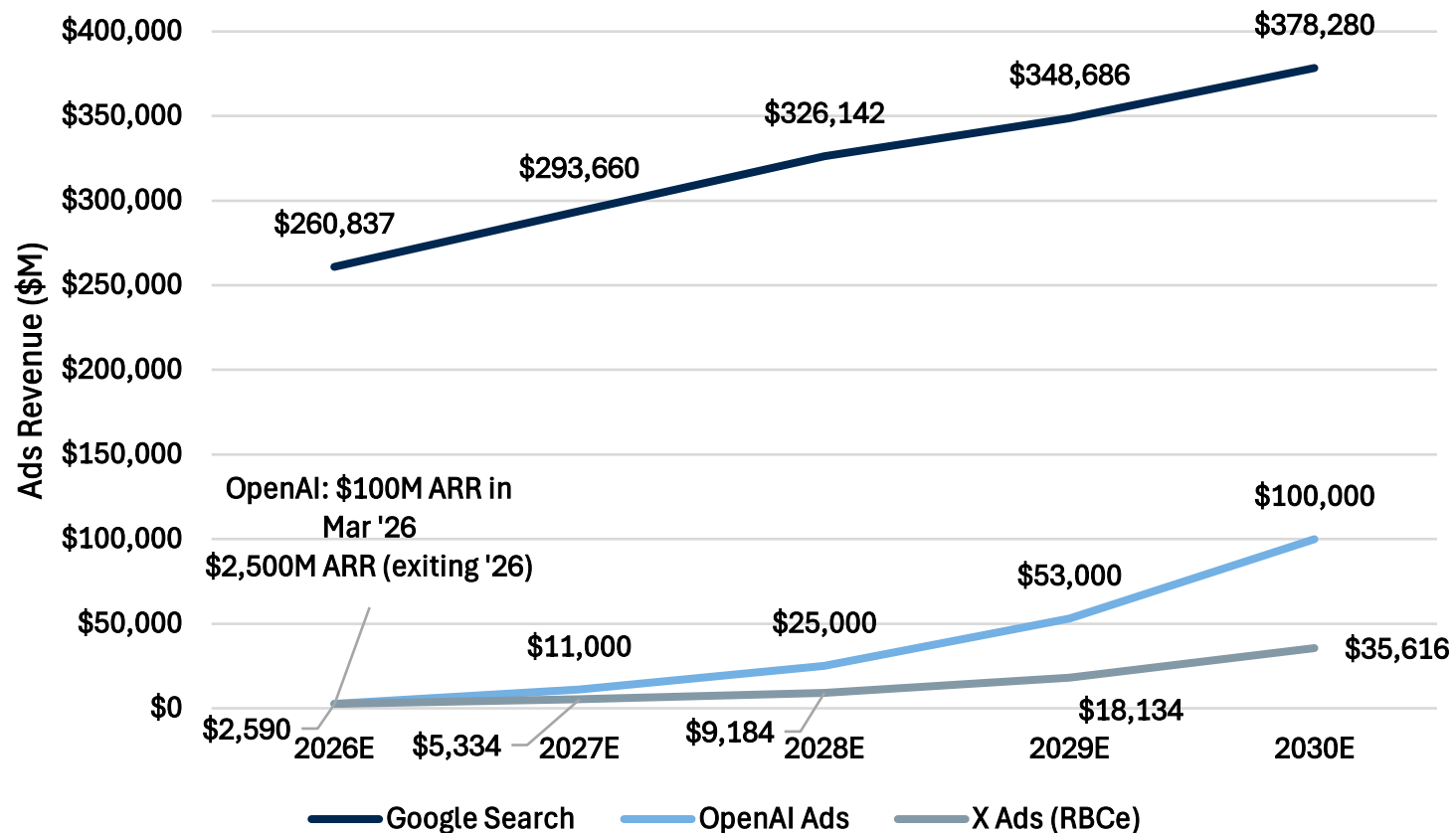
Where LLMs Struggle With Ads

- ChatGPT Ads are a relevance-first, context-driven, single-slot format that compresses the entire purchase funnel into one conversation — but the model faces fundamental interface challenges: limited inventory (one ad per response), nascent measurement, and the inherent unpredictability of multi-turn dialogue.

Challenge	Detail
Conversation Unpredictability	A user's topic can shift mid-thread after an ad has already been served against earlier context.
Measurement Gap	The measurement stack is nascent. Ads Manager Beta covers foundational metrics (impressions, clicks, CTR, conversions) but advertisers cannot yet track whether a ChatGPT interaction influenced a conversion that occurred later or on a different device.
Lower Click-Out Incentive	LLMs answer within the interface; users have less structural reason to click outbound to an advertiser's landing page.
No Exact-Match Control	Context hints guide matching but do not guarantee delivery in specific conversations. Advertisers cannot ensure their ad appears against a particular topic or query.
Brand Safety Verification	OpenAI enforces internal safeguards. 3p verification providers (IAS, DoubleVerify) not announced yet.
Limited Inventory	One ad slot per response. Combined with a smaller query base relative to traditional search, total available impressions are lower.

High Ambitions for LLM ads

- Ads have proven incremental for GOOGL where despite AIO having driven some impression headwinds, incremental searches from generative AI's capability more than offsets.
- ChatGPT's user base at over 1 billion MAUs along with rising context should allow for revenues to scale quickly, though we are reserving judgment on how achievable medium/long-term targets appear.
- Grok's sub-125M MAU levels will need to grow meaningfully to have traction for meaningful ad spend.
 - Management not pointing to that near term & nothing baked into the model.



Source: FactSet, Reuters, RBC Capital Markets estimates

SpaceX and Tesla Collaboration



Capital
Markets

SpaceX–Tesla Partnership Development & Transactions

Partnership Formation

- SpaceX and Tesla established an initial relationship through a series of limited but commercially successful engagements prior to 2026.
- In January 2026, Tesla committed to invest in xAI; upon SpaceX's subsequent acquisition of xAI, Tesla's investment converted into an equity interest in SpaceX.
- Tesla and xAI continue to build on their collaborative relationship by evaluating future strategic opportunities between the companies.
- Certain projects, including Macrohard and Terafab, remain in early stages; SpaceX and Tesla have not finalized financial terms, intellectual property rights, or the ultimate term of collaboration.
- Mr. Musk and other affiliated businesses may compete with SpaceX for investment or business opportunities and may direct opportunities to other businesses in which they have invested.

Commercial Purchases from Tesla

- Q1 2026: SpaceX/xAI purchased \$34 million of Megapack products from Tesla, recorded in PP&E.
- FY 2025: SpaceX/xAI purchased \$506 million of Megapack products and \$131 million of Cybertrucks (at MSRP) from Tesla, recorded in PP&E.
- FY 2024: SpaceX/xAI purchased \$191 million of Megapack products from Tesla, recorded in PP&E.
- Total xAI Megapack purchases from Tesla amount to approximately \$1 billion, representing approximately 4.5% of Tesla's total energy segment revenue of \$12.8 billion in 2025.

Deployment at xAI (Colossus I, Memphis)

- xAI has deployed 168 Megapack units at Colossus I, targeting approximately 1 GWh of buffering capacity.

Strategic Initiatives (Summary)

- **Macrohard (Digital Optimus):** Joint Tesla/xAI software-based AI agent system announced March 11, 2026, designed to autonomously execute computer-based tasks by pairing Tesla's AI4 real-time processing with xAI's Grok LLM reasoning layer. Tesla could deploy millions of computing units at Supercharger locations utilizing approximately 7 GW of available grid power.
- **Terafab:** Joint AI chip manufacturing initiative vertically integrating lithography mask design, logic/memory chip fabrication, and advanced packaging in a single facility. Long-term goal: one terawatt of compute per year. Two chip types planned — terrestrial edge/inference (Tesla Optimus robots and vehicles) and space-optimized (SpaceX orbital compute).

Strategic Collaboration Detail — Macrohard (Digital Optimus) & Terafab

Macrohard / Digital Optimus

- Digital Optimus is distinct from the physical Optimus humanoid robot; it is a software-based AI agent that processes real-time screen video (past five seconds) and executes keyboard and mouse actions.
- Two-component architecture: Digital Optimus serves as the real-time action execution layer; Grok LLM serves as the high-level reasoning and navigation layer directing actions.
- Runs on Tesla's AI4 hardware paired with a targeted allocation of xAI's Nvidia GPU equipment, described by Musk as relatively limited in scale.
- Tesla plans to convert Supercharger locations into decentralized edge-computing hubs for AI inference during off-peak travel times.
- Tesla has filed a modular datacenter trademark (MEGAPOD), expected to house AI server cabinets at Supercharger stalls.
- SpaceX expects Macrohard to benefit from running on both state-of-the-art processors and next-generation Tesla processors.

Terafab

- Intended to be the world's largest chip manufacturing facility, alleviating potential future chip shortages and optimizing compute performance.
- End-to-end single-facility approach enables rapid testing, iteration, and faster scaling of chip design and manufacturing.
- Speed and cost advantages from vertical integration support efficient scaling toward one terawatt of compute per year.
- SpaceX expects to continue sourcing compute hardware from third-party suppliers; Terafab is complementary, extending control to the foundational chip layer.
- SpaceX's stated view: key constraints to AI growth are physical (chip manufacturing, datacenter infrastructure, power generation), and SpaceX is positioned with control over the full physical stack.

Financial Overview of AI



Capital
Markets

Financial Overview: AI Segment

We estimate overall revenue to grow at a 104% CAGR between '27 and '31.

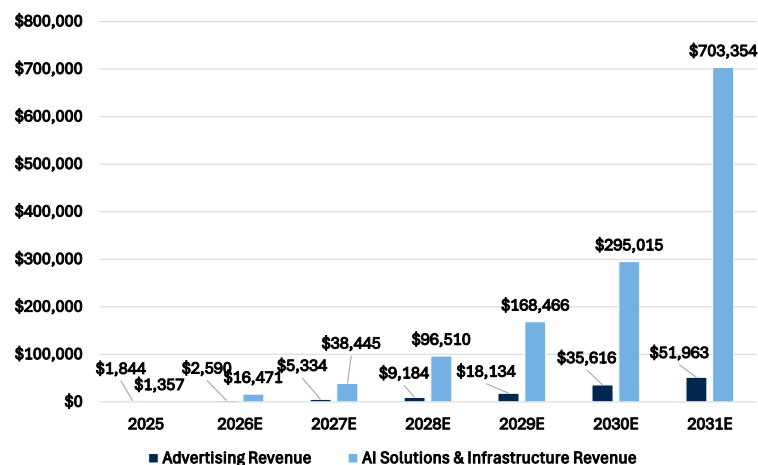
■ Infrastructure & Solutions (I&S)

- Expected to be considerably larger than the Advertising segment, driven primarily by Software, which we anticipate will grow significantly in 2027 and beyond as the company integrates Cursor and scales its enterprise platform (Grok API, Build, browser & computer agents, Macrohard, etc.).
- Subscription growth remains solid as the Grok platform ramps, though longer term we expect the consumption model (token-based pricing) to become the more meaningful growth driver.

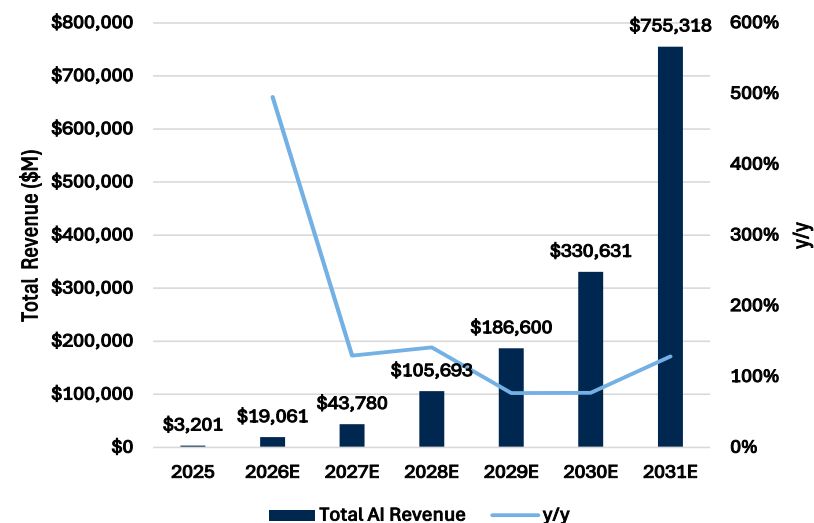
■ Advertising

- The X platform re-build, combined with a growing user base and improving engagement, is expected to drive an inflection in ARPU. We are somewhat skeptical relative to management's forecasts; however, we believe the company's access to its own scaled compute should fuel structurally better content recommendation, creation, and campaign performance over time versus other walled-garden platforms (ex-META, GOOGL).

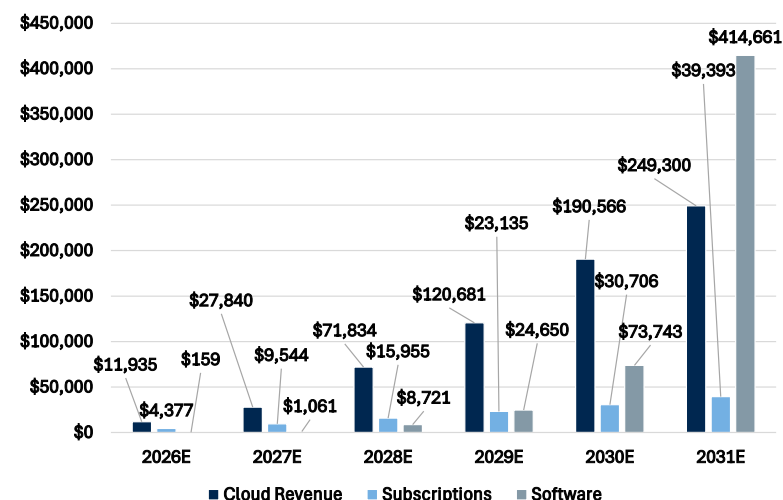
Total Revenue (By Segments) (\$M)



Total Revenue (\$M)



AI Solutions & Infrastructure Revenue (\$M)

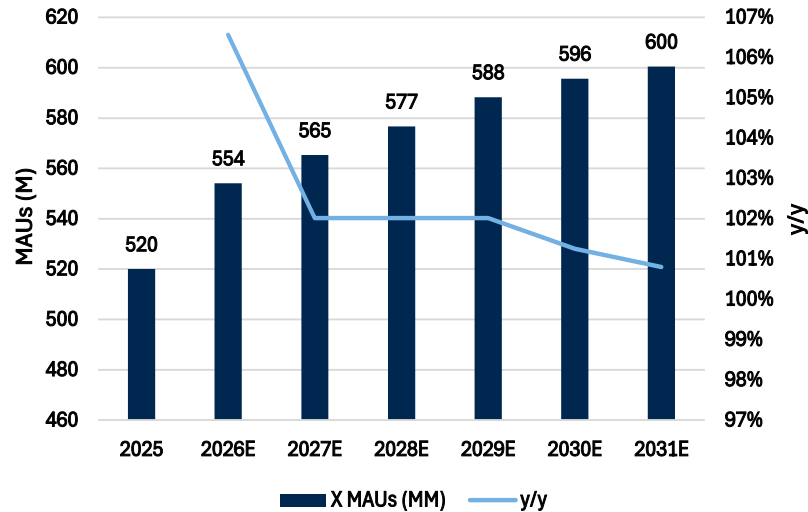


Source: Company reports, RBC Capital Markets estimates

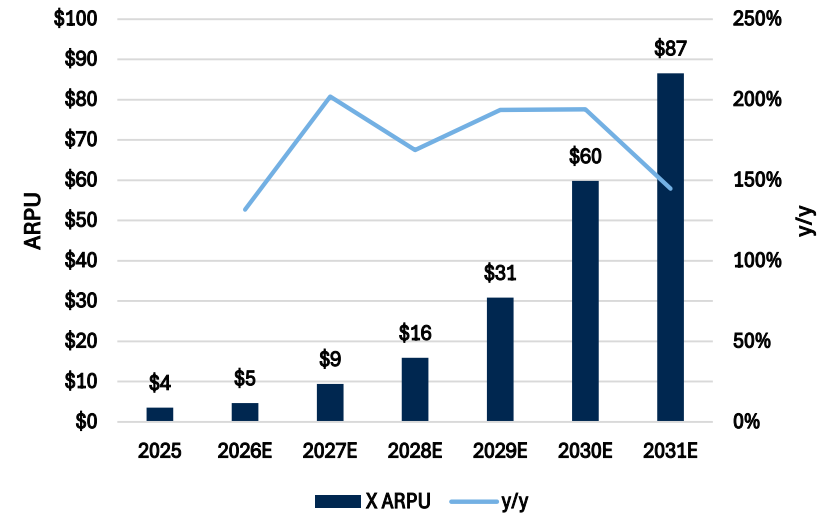
Financial Overview: AI Segment – KPIs

We expect Grok engagement KPIs such as subscribers and MAUs to ramp meaningfully in 2028 before moderating through 2031, while ARPU declines from ~\$28 to sub-\$20 over the same period as the user base expands into lower-monetization tiers.

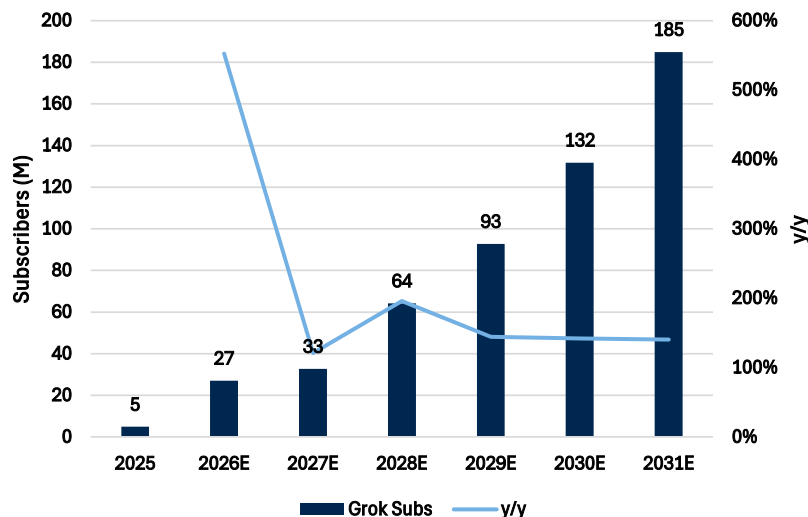
X MAUs (M)



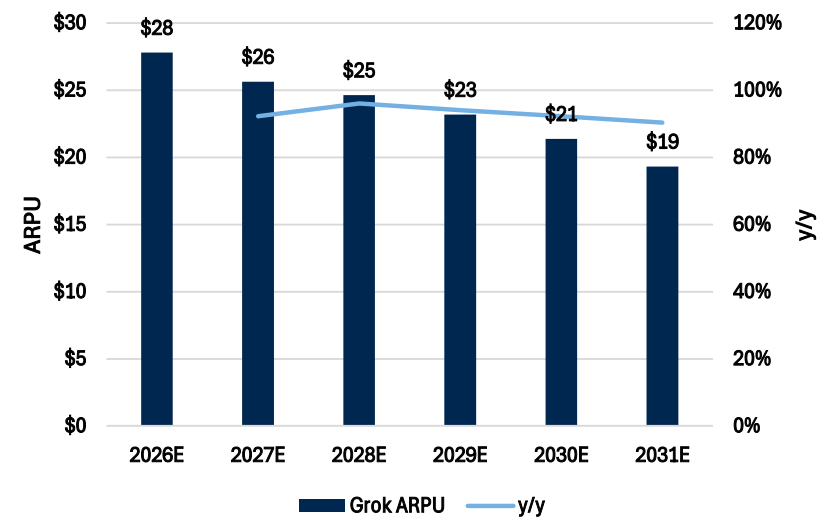
X ARPU



Grok Subscribers (M)



Grok ARPU

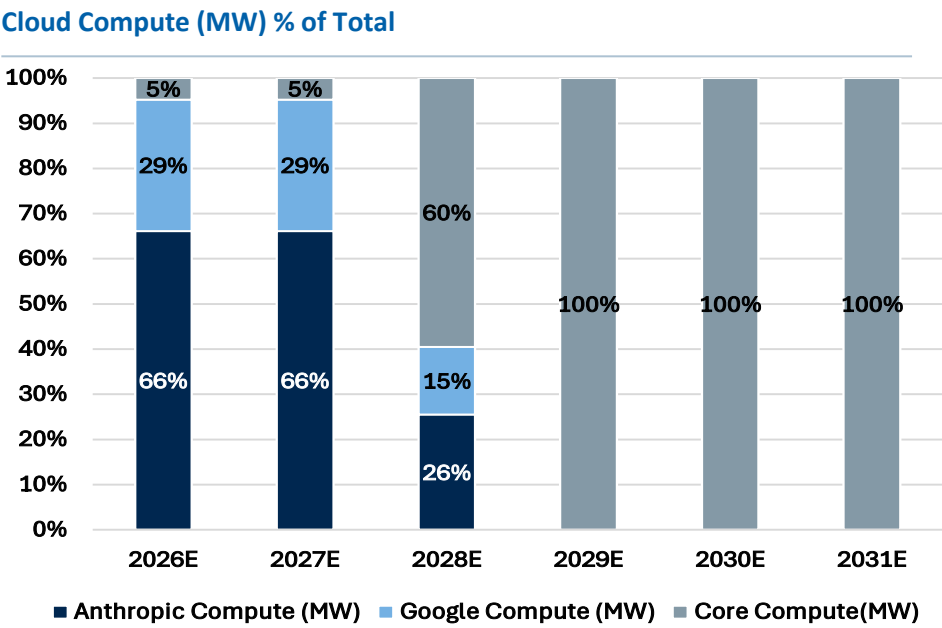
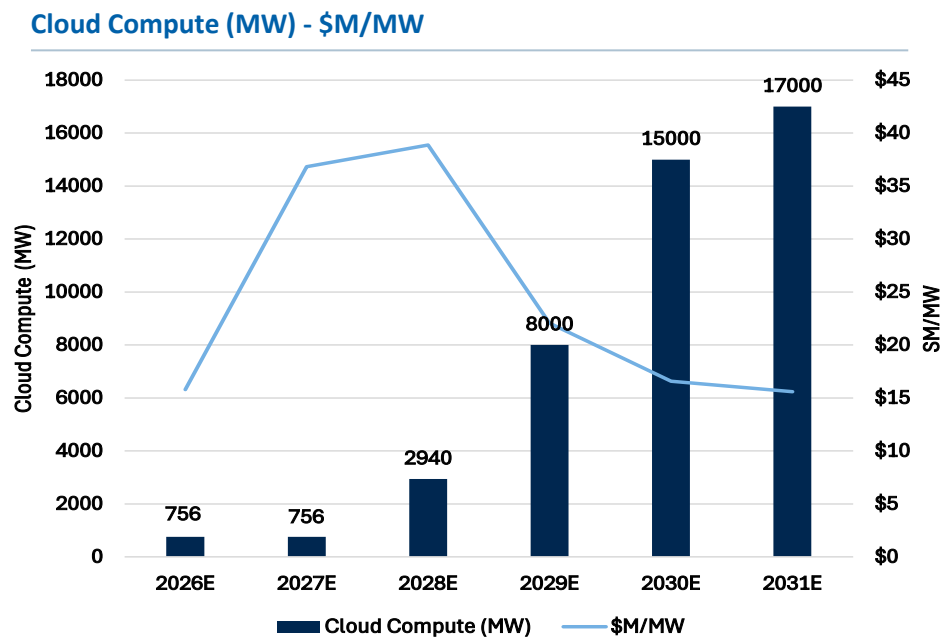


Source: Company reports, RBC Capital Markets estimates

Financial Overview: AI Segment – KPIs

We estimate overall neocloud compute to grow at a 118% CAGR between '27 and '31.

- We expect cloud compute revenue to scale rapidly beginning in 2028, growing at a CAGR above 100% through 2031, driven by increasing capacity utilization and new customer additions that more than offset the expected termination of the Google and Anthropic deals in 2Q29, with further upside supported by the ramp of orbital compute from 2029 onward.

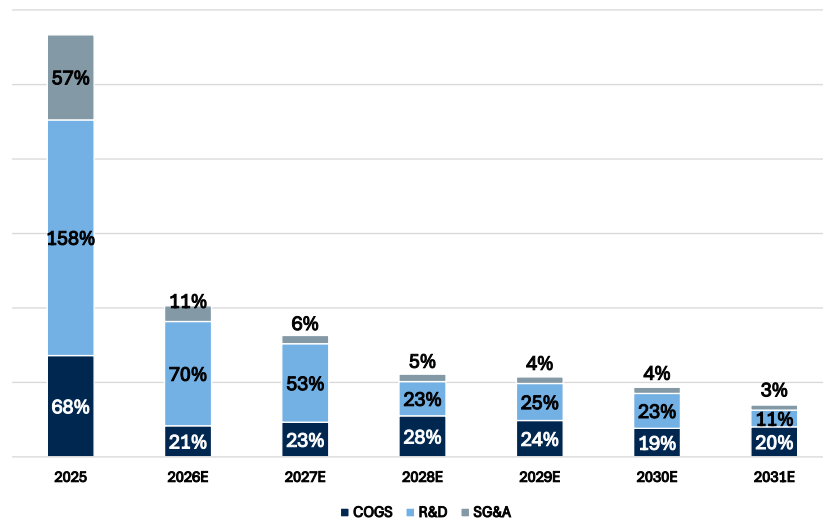


Source: RBC Capital Markets estimates

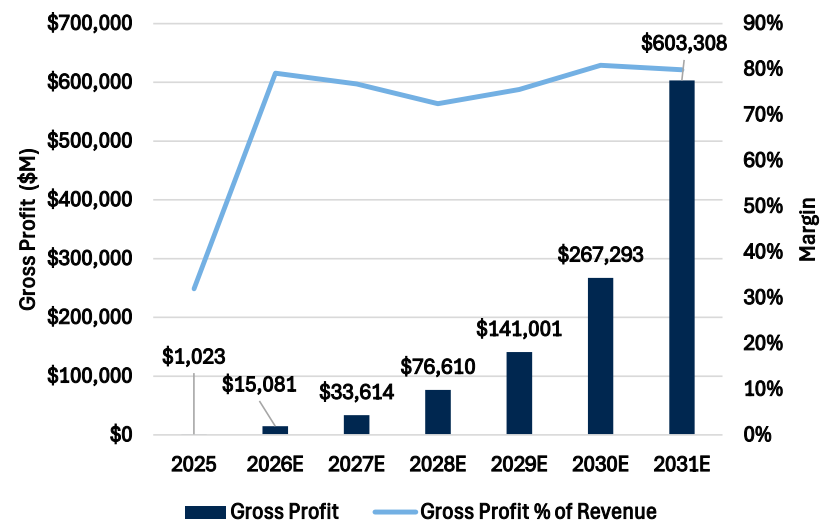
Financial Overview: AI Segment – Profitability

We forecast EBITDA margins expanding from ~54% in 2026 to ~90% by 2031, driven by operating leverage as R&D declines from 70% to 11% of revenue and SG&A from 11% to 3%, with D&A as a percentage of EBITDA compressing from 87% to 27% over the same period.

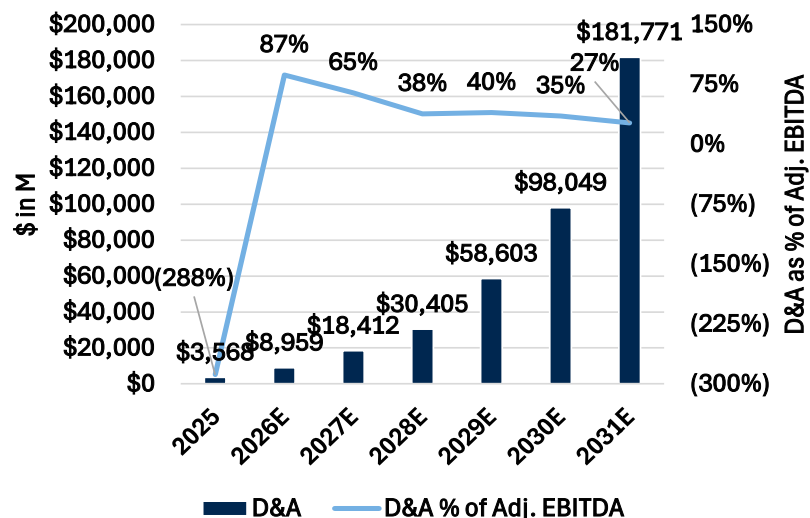
COGS & Opex % of Revenue



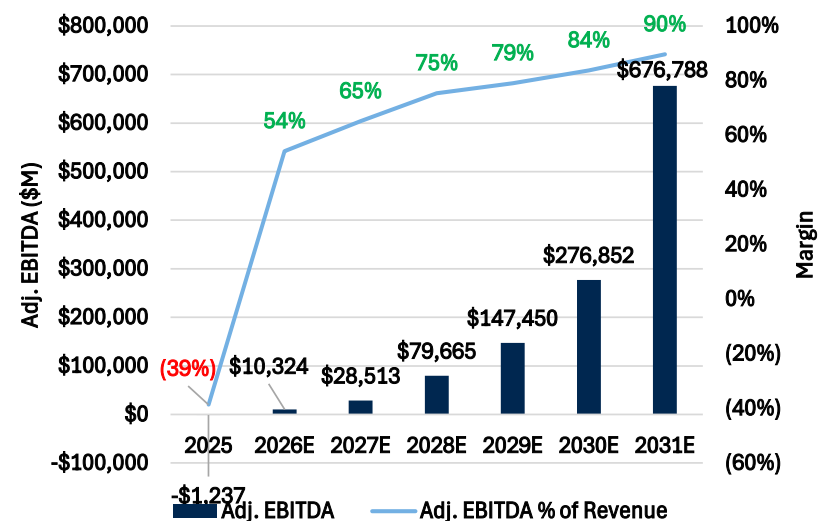
Gross Profit (\$M)



D&A (\$M)



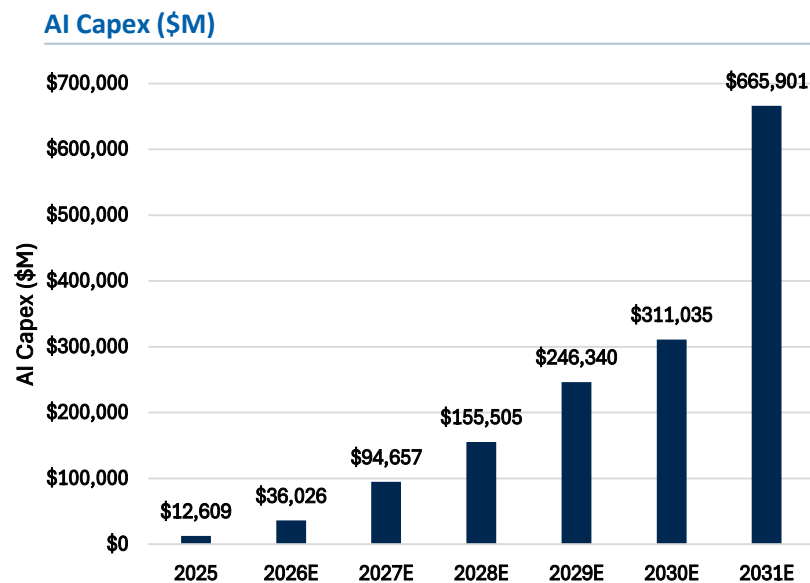
Adjusted EBITDA (\$M)



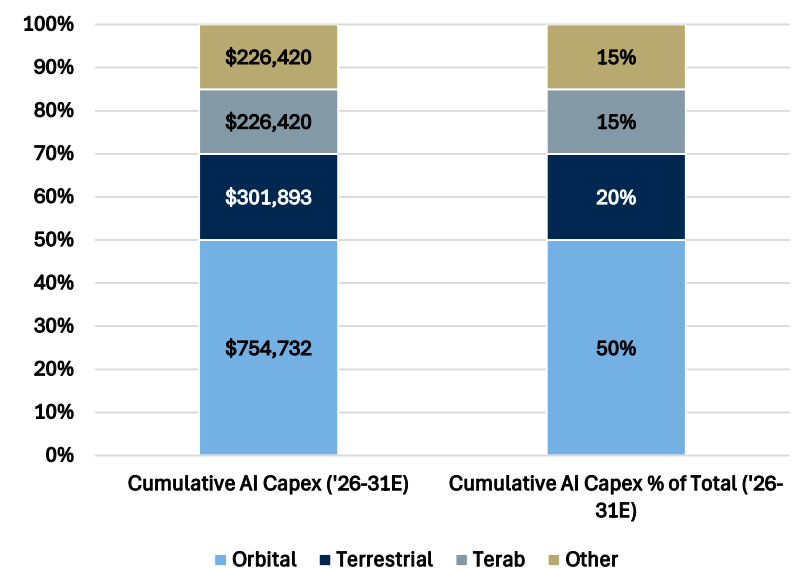
Source: Company reports, RBC Capital Markets estimates

Financial Overview

- **Similar to the hyperscalers, we expect xAI to spend more than \$1 trillion of capex looking out to 2031**
 - Initially, we expect it will largely be around terrestrial infrastructure, but Terafab, solar fab and orbital will likely also be drivers in '27 and beyond.
 - We expect the '26-'31 mix of capex to be the following:
 - 20% terrestrial datacenters
 - 15% Terafab
 - 50% orbital AI
 - 15% other
 - We model AI solutions on \$M/MW, but it becomes somewhat disconnected in the latter part of the forecast.

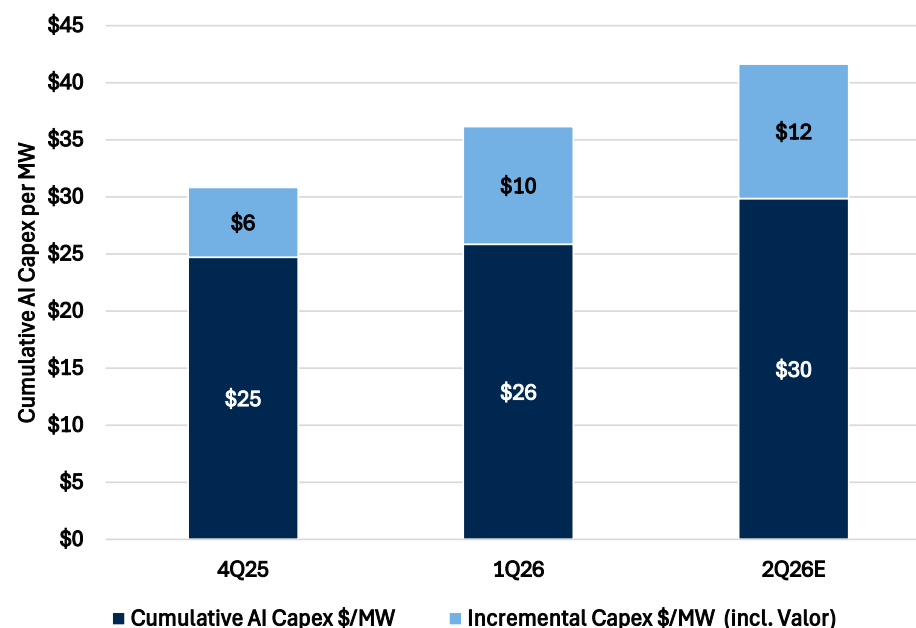
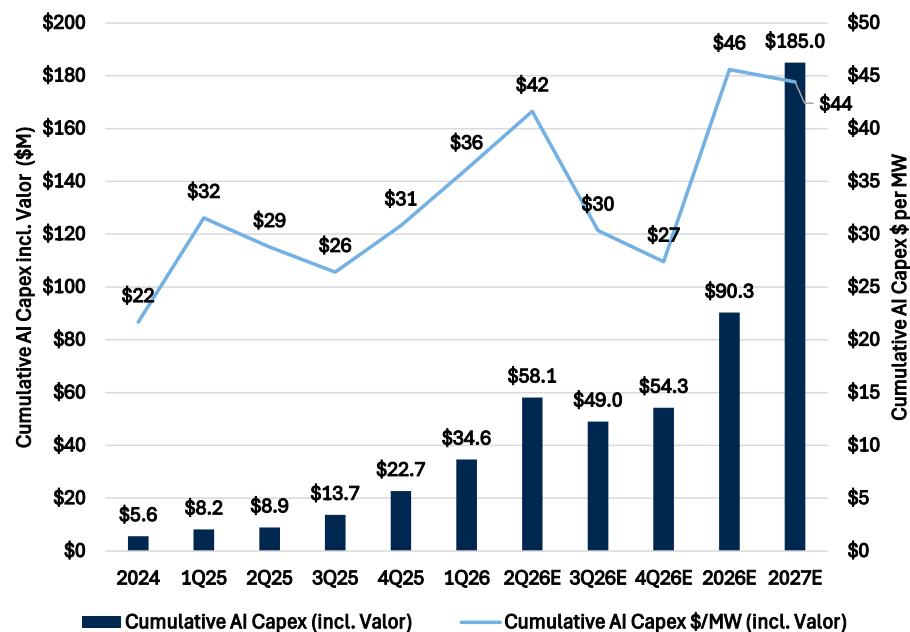


Cumulative AI Capex % of Total '26E-31E



Getting comfortable with alternative financing for AI infrastructure

- **xAI executed 3 infrastructure financing transactions treated as failed sale leaseback**
 - Treated as related-party debt (not a true sale leaseback) due to xAI retaining control of the purchased GPUs
 - ~\$16.5B in aggregate gross capex financed across three transactions:
 - Valor Transaction I (Oct 2025): ~\$4.5 billion
 - Valor Transaction II (Jan 2026): ~\$5.4 billion
 - Valor Transaction III (Apr 2026): ~\$6.6 billion
- **Including these Valor transactions, we estimate a cumulative capex of ~\$35B in 1Q26 and ~\$58B in 2Q26**
 - We estimate over \$7B of construction-in-progress (as of 1Q26) relates to Colossus 3
 - Removing that from ~\$35B in cumulative AI capex at \$36M/MW (as of 1Q26) implies a \$/MW of ~\$30M versus ~\$26M/MW implied solely from reported capex ('24-1Q26)



Source: Company reports, RBC Capital Markets estimates

AI Capex ROIC

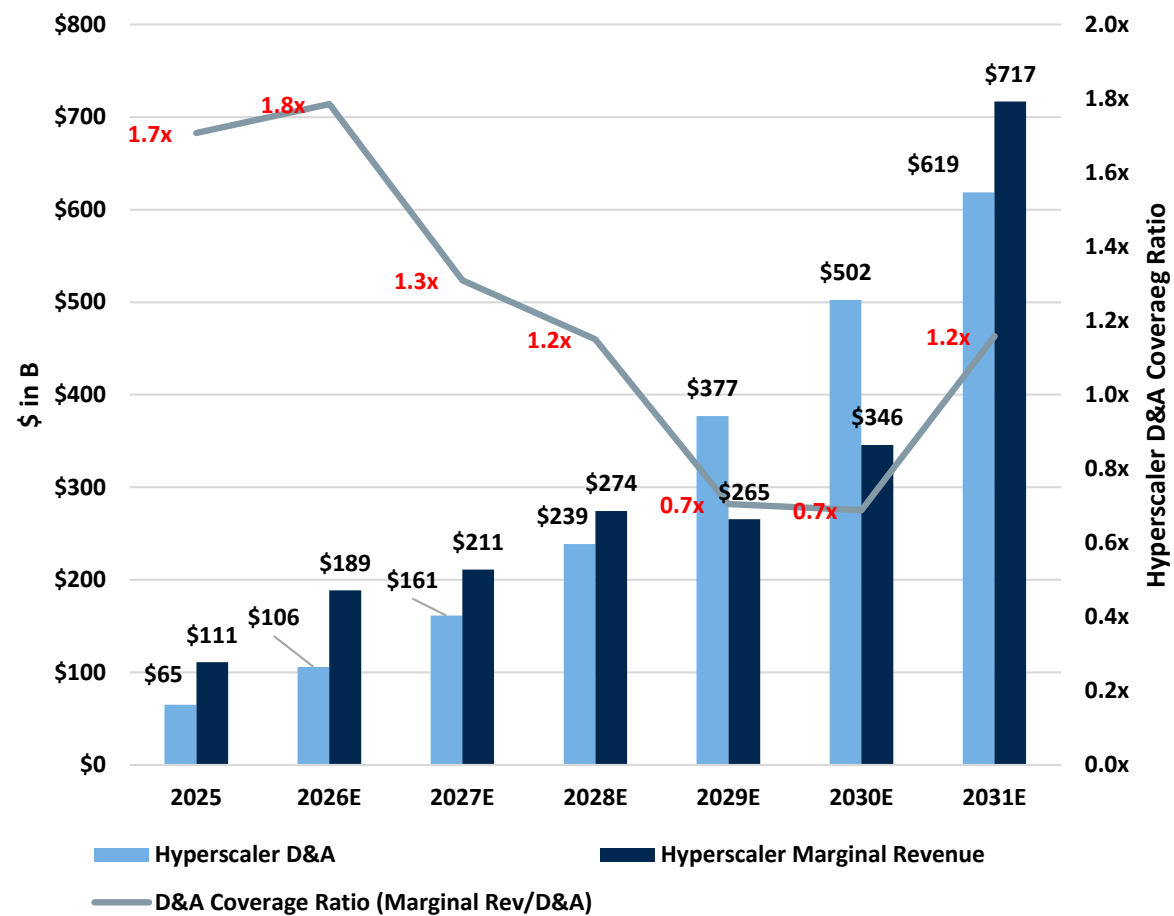
- Our ROIC framework suggests net returns expanding from ~15% in '28E to ~35% by '31E as AI Solutions & Infrastructure revenue approaches ~\$703B and EBIT margins reach ~65%. We expect Orbital compute infrastructure coming online in '28 to reduce per-token inference costs, with COGS compressing from 28% to 20% of revenue over the period. Incremental ROIC reaches ~53% by '31E, implying meaningful operating leverage as revenue scales against a moderating capex growth trajectory.

Metrics (\$M)	2025	2026E	2027E	2028E	2029E	2030E	2031E
AI Solutions & Infrastructure Revenue	\$1,357	\$16,471	\$38,445	\$96,510	\$168,466	\$295,015	\$703,354
COGS % of Revenue	68%	21%	23%	28%	24%	19%	20%
COGS	\$923	\$3,459	\$8,842	\$27,023	\$40,432	\$56,053	\$140,671
Gross Profit	\$434	\$13,012	\$29,603	\$69,487	\$128,034	\$238,962	\$562,684
Opex % of Revenue	231%	81%	58%	28%	29%	28%	15%
Opex	\$3,135	\$13,342	\$22,298	\$27,023	\$48,855	\$81,129	\$105,503
AI EBIT	-\$2,700	-\$329	\$7,305	\$42,464	\$79,179	\$157,833	\$457,180
AI EBIT Margin	(199.0%)	(2.0%)	19.0%	44.0%	47.0%	53.5%	65.0%
Tax Rate	15%	15%	15%	15%	15%	15%	15%
AI NOPAT	-\$2,700	-\$329	\$6,209	\$36,095	\$67,302	\$134,158	\$388,603
Annual AI Capex	\$12,609	\$36,026	\$94,657	\$155,505	\$246,340	\$311,035	\$665,901
Cumulative AI Capex (Gross)	\$18,242	\$54,268	\$148,924	\$304,430	\$550,769	\$861,805	\$1,527,706
Annual Depreciation	\$3,568	\$8,959	\$18,412	\$30,405	\$58,603	\$98,049	\$181,771
Cumulative Depreciation	\$5,247	\$14,206	\$32,617	\$63,022	\$121,625	\$219,674	\$401,446
Cumulative AI Capex (Net)	\$12,995	\$40,062	\$116,307	\$241,407	\$429,144	\$642,130	\$1,126,260
AI ROIC:							
AI ROIC (Net Capex)	(20.8%)	(0.8%)	5.3%	15.0%	15.7%	20.9%	34.5%
AI ROIC (Gross Capex)	(14.8%)	(0.6%)	4.2%	11.9%	12.2%	15.6%	25.4%
Incremental AI ROIC (Net Capex)		8.8%	8.6%	23.9%	16.6%	31.4%	52.6%
Incremental AI ROIC (Gross Capex)		6.6%	6.9%	19.2%	12.7%	21.5%	38.2%
Per MW:							
AI Rev	\$1.8	\$8.3	\$9.2	\$11.7	\$9.5	\$8.0	\$8.1
AI EBIT	-\$3.7	-\$0.2	\$1.8	\$5.2	\$4.5	\$4.3	\$5.3
AI NOPAT	-\$3.7	-\$0.2	\$1.5	\$4.4	\$3.8	\$3.7	\$4.5
Invested Capital (Capex, net)	\$17.6	\$20.2	\$27.9	\$29.3	\$24.3	\$17.5	\$13.0
Invested Capital (Capex, gross)	\$24.7	\$27.4	\$35.8	\$36.9	\$31.2	\$23.5	\$17.6

Source: RBC Capital Markets estimates

Tracking Marginal Revenue Relative to Depreciation Growth

D&A Coverage Ratio for hyperscalers (e.g., GOOGL, META, AMZN AWS, SPCX) is estimated to decline from 1.8x in '26E to a trough of 0.7x in '29E–'30E, suggesting that the pace of capex-driven depreciation could materially outstrip incremental revenue generation during peak investment years. Based on current estimates, this compression would imply a period where marginal returns on invested capital are under pressure, as each dollar of new D&A is supported by less than a dollar of marginal revenue. The ratio's estimated recovery to 1.2x by 2031E suggests revenue growth may eventually begin to catch up as earlier capacity investments are monetized, though coverage is expected to remain well below the ~1.7–1.8x levels estimated for '25–'26E.



Source: FactSet, Visible Alpha, RBC Capital Markets estimates



Space Exploration Technologies "SpaceX" (SPCK)	FY'23	FY'24	FY'25	1Q'26	2Q'26	3Q'26	4Q'26	FY'26E	1Q'27E	2Q'27E	3Q'27E	4Q'27E	FY'27E	FY'28E	FY'29E	FY'30E	FY'31E
Income Statement (\$M)	2023	2024	2025	Mar-26	Jun-26	Sep-26	Dec-26	2026	Mar-27	Jun-27	Sep-27	Dec-27	2027	2028	2029	2030	2031
Adj. EBITDA	\$3,821	\$5,350	\$6,584	\$1,127	\$2,043	\$6,372	\$11,139	\$20,681	\$11,275	\$11,595	\$12,361	\$13,566	\$48,797	\$110,863	\$192,474	\$340,411	\$765,473
% Change, Y/Y	-	40.0%	23.1%	(34.9%)	68.6%	193.9%	655.7%	214.1%	900.5%	467.5%	94.0%	21.8%	135.9%	127.2%	73.6%	76.9%	124.9%
Segment revenues																	
Space	3,557	3,796	4,086	619	829	1,610	1,675	4,733	1,916	1,980	1,766	1,597	7,259	7,470	7,679	8,070	10,205
Connectivity	3,869	7,599	11,387	3,257	4,036	4,788	6,047	18,127	5,538	6,609	7,523	9,166	28,836	46,277	69,609	100,940	144,229
AI	2,961	2,620	3,201	818	1,907	5,858	10,479	19,061	10,372	10,809	11,023	11,575	43,780	105,693	186,600	330,631	755,318
Total Segment Revenues	10,387	14,015	18,674	4,694	6,772	12,255	18,201	41,922	17,826	19,398	20,312	22,338	79,874	159,440	263,888	439,641	909,751
Cost of Sales	(6,110)	(7,996)	(9,451)	(2,388)	(2,900)	(4,061)	(5,365)	(14,713)	(5,395)	(6,505)	(6,995)	(7,660)	(26,555)	(53,533)	(80,634)	(111,306)	(218,035)
Gross Profit	4,277	6,019	9,223	2,306	3,872	8,194	12,836	27,208	12,431	12,893	13,316	14,679	53,320	105,907	183,254	328,335	691,716
Research & development	(2,105)	(3,464)	(8,643)	(3,514)	(4,665)	(5,267)	(5,506)	(18,952)	(5,341)	(6,272)	(7,262)	(7,841)	(26,717)	(28,310)	(50,839)	(82,649)	(91,611)
Selling, general & administrative expenses	(1,665)	(1,813)	(2,644)	(746)	(885)	(944)	(1,040)	(3,615)	(995)	(1,135)	(1,193)	(1,306)	(4,628)	(8,443)	(12,772)	(19,749)	(34,647)
Restructuring credits / (charges)	(237)	(213)	(487)	11	-	-	-	11	-	-	-	-	-	-	-	-	-
Impairment & other charges	(3,775)	(63)	(38)	-	-	-	-	-	-	-	-	-	-	-	-	-	-
Operating Income	(3,505)	466	(2,589)	(1,943)	(1,678)	1,983	6,290	4,652	6,095	5,486	4,861	5,532	21,975	69,155	119,643	225,937	565,458
Operating Margin	(33.7%)	3.3%	(13.9%)	(41.4%)	(24.8%)	16.2%	34.6%	11.1%	34.2%	28.3%	23.9%	24.8%	27.5%	43.4%	45.3%	51.4%	62.2%
Interest expense, net	(1,444)	(1,209)	(1,453)	(451)	(633)	(223)	(269)	(1,576)	(344)	(553)	(982)	(1,378)	(3,257)	(6,198)	(12,108)	(17,954)	(19,976)
Other expense, net	(42)	985	(177)	(1,876)	-	-	-	(1,876)	-	-	-	-	-	-	-	-	-
Pretax Income	(4,991)	242	(4,219)	(4,270)	(2,310)	1,760	6,021	1,200	5,752	4,933	3,879	4,154	18,718	62,957	107,535	207,982	545,482
Income tax (benefit)	363	549	(718)	(6)	-	-	-	(59)	(11)	(23)	(26)	(39)	(99)	14,202	(4,537)	(36,071)	(96,092)
Net income, continuing ops	(4,628)	791	(4,937)	(4,276)	(2,310)	1,760	5,967	1,141	5,741	4,910	3,853	4,114	18,619	77,159	102,998	171,911	449,390
Net income attributable to non-controlling interest	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
Net income, continuing ops	(4,628)	791	(4,937)	(4,276)	(2,310)	1,760	5,967	1,141	5,741	4,910	3,853	4,114	18,619	77,159	102,998	171,911	449,390
Net income from discontinued ops	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
Net Income	(4,628)	791	(4,937)	(4,276)	(2,310)	1,760	5,967	1,141	5,741	4,910	3,853	4,114	18,619	77,159	102,998	171,911	449,390
Non-GAAP adjustments	4,691	1,060	2,472	628	743	771	765	2,906	687	835	893	890	3,305	3,746	4,132	4,372	4,684
Non-GAAP Tax adjustments	-	-	(828)	(1)	-	-	7	6	1	4	6	8	20	(890)	174	758	825
Adj. Net Income	(4,628)	791	(3,293)	(3,649)	(1,568)	2,530	6,739	4,053	6,429	5,749	4,752	5,013	21,943	80,015	107,304	177,041	454,899
Diluted Earnings Per Share	(\$3.38)	(\$0.03)	(\$1.69)	(\$0.36)	(\$0.19)	\$0.13	\$0.46	\$0.09	\$0.44	\$0.38	\$0.29	\$0.31	\$1.42	\$5.90	\$7.87	\$13.13	\$34.29
Non-GAAP adjustments	\$1.70	\$0.11	\$0.56	\$0.05	\$0.06	\$0.06	\$0.06	\$0.23	\$0.05	\$0.06	\$0.07	\$0.07	\$0.25	\$0.22	\$0.33	\$0.39	\$0.42
Adj. Diluted Earnings Per Share	(\$1.68)	\$0.08	(\$1.13)	(\$0.31)	(\$0.13)	\$0.19	\$0.52	\$0.32	\$0.49	\$0.44	\$0.36	\$0.38	\$1.68	\$6.12	\$8.20	\$13.52	\$34.71
% Change, Y/Y	-	(104.7%)	(1516.6%)	-	-	-	-	(128.8%)	(258.9%)	(437.8%)	87.7%	(25.7%)	417.4%	264.4%	34.0%	64.9%	156.8%
Seasonality (% of FY)	0.0%	0.0%	0.0%	(95.4%)	(40.1%)	59.7%	158.9%	83.1%	29.3%	26.2%	21.7%	22.8%	100.0%	100.0%	100.0%	100.0%	100.0%
Average diluted shares	2,759	9,956	2,926	11,790	12,042	13,066	13,068	12,492	13,070	13,072	13,074	13,076	13,073	13,081	13,089	13,097	13,105
% Change, Y/Y	-	(70.6%)	-	-	-	-	-	326.9%	10.9%	8.6%	0.1%	0.1%	4.7%	0.1%	0.1%	0.1%	0.1%
Net income	(4,628)	791	(4,937)	(4,276)	(2,310)	1,760	5,967	1,141	5,741	4,910	3,853	4,114	18,619	77,159	102,998	171,911	449,390
EBITDA	(870)	4,290	4,112	499	1,301	5,602	10,374	17,775	10,588	10,761	11,468	12,675	45,492	107,118	188,342	336,040	760,789
Adjusted EBITDA	3,821	5,350	6,584	1,127	2,043	6,372	11,139	20,681	11,275	11,595	12,361	13,566	48,797	110,863	192,474	340,411	765,473
Adj. EBITDA Margin	36.8%	38.2%	35.3%	24.0%	30.2%	52.0%	61.2%	49.3%	63.3%	59.8%	60.9%	60.7%	61.1%	69.5%	72.9%	77.4%	84.1%
Net Income	(4,628)	791	(4,937)	(4,276)	(2,310)	1,760	5,967	1,141	5,741	4,910	3,853	4,114	18,619	77,159	102,998	171,911	449,390
Restructuring, impairment & other	4,012	276	525	(11)	-	-	-	(11)	-	-	-	-	-	-	-	-	-
Stock-based comp, other	679	784	1,947	639	743	771	765	2,917	687	835	893	890	3,305	3,746	4,132	4,372	4,684
Adj. Net Income, pre-tax adjustments	63	1,851	(2,465)	(3,648)	(1,568)	2,530	6,732	4,047	6,428	5,745	4,746	5,005	21,924	80,904	107,130	176,283	454,074
Tax adjustments	341	(2,405)	(828)	(1)	-	-	7	6	1	4	6	8	20	(890)	174	758	825
Adj. Net Income	404	(554)	(3,293)	(3,649)	(1,568)	2,530	6,739	4,053	6,429	5,749	4,752	5,013	21,943	80,015	107,304	177,041	454,899
Margin Analysis:																	
Gross margin	41.2%	42.9%	49.4%	49.1%	57.2%	66.9%	70.5%	64.9%	69.7%	66.5%	65.6%	65.7%	66.8%	66.4%	69.4%	74.7%	76.0%
SG&A Margin	(16.0%)	(12.9%)	(14.2%)	(15.9%)	(13.1%)	(7.7%)	(5.7%)	(8.6%)	(5.6%)	(5.9%)	(5.9%)	(5.8%)	(5.8%)	(5.3%)	(4.8%)	(4.5%)	(3.8%)
Operating Margin	(33.7%)	3.3%	(13.9%)	(41.4%)	(24.8%)	16.2%	34.6%	11.1%	34.2%	28.3%	23.9%	24.8%	27.5%	43.4%	45.3%	51.4%	62.2%
EBITDA Margin	(8.4%)	30.6%	22.0%	10.6%	19.2%	45.7%	57.0%	42.4%	59.4%	55.5%	56.5%	56.7%	57.0%	67.2%	71.4%	76.4%	83.6%
Adj. EBITDA Margin	36.8%	38.2%	35.3%	24.0%	30.2%	52.0%	61.2%	49.3%	63.3%	59.8%	60.9%	60.7%	61.1%	69.5%	72.9%	77.4%	84.1%
Pretax margin	(48.1%)	1.7%	(22.6%)	(91.0%)	(34.1%)	14.4%	33.1%	2.9%	32.3%	25.4%	19.1%	18.6%	23.4%	39.5%	40.8%	47.3%	60.0%
Net income margin	(44.6%)	5.6%	(26.4%)	(91.1%)	(34.1%)	14.4%	32.8%	2.7%	32.2%	25.3%	19.0%	18.4%	23.3%	48.4%	39.0%	39.1%	49.4%
Effective tax rate	(7.3%)	226.9%	17.0%	0.1%	0.0%	0.0%	(0.9%)	(5.0%)	(0.2%)	(0.5%)	(0.7%)	(1.0%)	(0.5%)	22.6%	(4.2%)	(17.3%)	(17.6%)

Source: Company reports, RBC Capital Markets estimates



Space Exploration Technologies "SpaceX" (SPCX)	FY'24	FY'25	FY'26E	FY'27E	FY'28E	FY'29E	FY'30E	FY'31E
Balance Sheet (\$M)	2024	2025	2026	2027	2028	2029	2030	2031
Free cash flow	(5,387)	(13,834)	(32,113)	(83,838)	(77,809)	(84,832)	(23,825)	44,368
% Change, Y/Y	-	156.8%	132.1%	161.1%	(7.2%)	9.0%	(71.9%)	(286.2%)
Assets								
Cash and cash equivalents	11,385	24,747	69,405	59,820	95,578	104,250	127,356	159,500
Short-term marketable securities	800	-	7,823	7,823	7,823	7,823	7,823	7,823
Receivables, net	1,052	1,579	2,201	5,550	11,631	19,141	30,993	43,246
Inventory, net	2,003	2,416	4,950	4,903	3,794	3,457	3,326	3,030
Prepaid expenses and other current assets	868	2,210	2,500	2,376	2,011	2,030	2,024	2,099
Total Current Assets	16,108	30,952	86,879	80,473	120,837	136,701	171,522	215,698
Property, plant and equipment, net	21,147	42,602	82,266	186,225	325,506	523,179	743,705	1,234,607
Finance Leas ROU assets	1,686	1,260	1,182	1,182	1,182	1,182	1,182	1,182
Other assets	17,425	17,124	17,088	17,088	17,088	17,088	17,088	17,088
Total Assets	56,366	91,938	187,415	284,968	464,613	678,149	933,497	1,468,575
Liabilities & Shareholder's Equity								
Accounts payable	4,413	11,792	9,370	12,144	14,450	23,219	33,166	57,967
Deferred revenue	10,179	12,116	26,611	25,320	20,188	20,236	23,339	34,571
Other short-term liabilities	1,508	2,569	8,154	7,384	5,672	5,532	6,156	17,209
Total current liabilities	16,100	26,477	44,135	44,847	40,310	48,987	62,661	109,747
Long-term debt	12,262	21,659	33,800	111,468	232,992	329,000	374,952	354,276
Finance lease liabilities	1,531	1,237	1,154	1,154	1,154	1,154	1,154	1,154
Deferred tax liabilities / (asset)	(725)	(113)	8	84	(14,305)	(9,702)	13,201	71,755
Other long-term liabilities	1,394	1,353	1,293	1,293	1,293	1,293	1,293	1,293
Total liabilities	30,562	50,613	80,390	158,846	261,444	370,732	453,260	538,225
Convertible preferred stock	20,941	38,752	-	-	-	-	-	-
Additional paid-in capital	35,868	37,710	140,626	140,626	140,626	140,626	140,626	140,626
Retained earnings	(31,005)	(35,137)	(32,441)	(12,170)	66,302	170,880	345,039	796,333
Accumulated other comprehensive income	-	-	-	-	-	-	-	-
Other	-	-	-	-	-	-	-	-
Total Shareholder's Equity	25,804	41,325	107,024	126,122	203,170	307,417	480,236	930,350
Noncontrolling interest	-	-	-	-	-	-	-	-
Total stockholders' equity	25,804	41,325	107,024	126,122	203,170	307,417	480,236	930,350
Total Liabilities & Shareholder's Equity	56,366	91,938	187,415	284,968	464,613	678,149	933,497	1,468,575
Key Metrics								
Returns Analysis								
Return on invested capital	4.2%	(9.8%)	1.1%	9.8%	22.9%	19.2%	23.1%	42.0%
Return on equity	6.1%	(14.7%)	1.5%	16.0%	46.9%	40.3%	43.7%	63.7%
Return on assets	(2.1%)	(3.2%)	3.1%	9.2%	22.0%	19.4%	22.5%	37.2%
Leverage Summary								
Debt/EBITDA	2.3x	3.3x	1.6x	2.3x	2.1x	1.7x	1.1x	0.5x
Net Debt / EBITDA	0.2x	(0.5x)	(1.7x)	1.1x	1.2x	1.2x	0.7x	0.3x
Total debt/total capital	32.2%	34.4%	24.0%	46.9%	53.4%	51.7%	43.8%	27.6%
Net debt /Total cap	2.3%	(4.9%)	(25.3%)	21.7%	31.5%	35.3%	29.0%	15.2%
Interest coverage (EBITDA basis)	4.4x	4.5x	13.1x	15.0x	17.9x	15.9x	19.0x	38.3x

Source: Company reports, RBC Capital Markets estimates



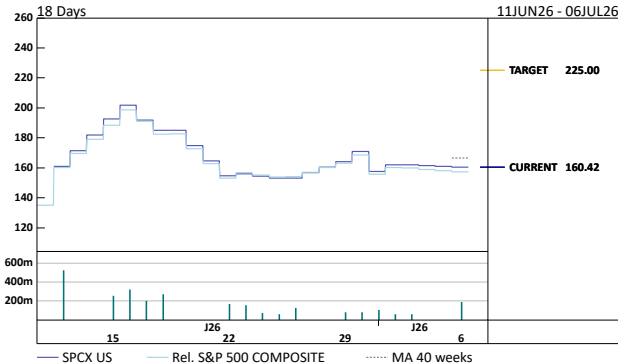
Space Exploration Technologies "SpaceX" (SPCX)	FY'23	FY'24	FY'25	FY'26E	FY'27E	FY'28E	FY'29E	FY'30E	FY'31E
Statement of Cash Flows (\$M)	2023	2024	2025	2026	2027	2028	2029	2030	2031
Free cash flow	105.0	(5,387)	(13,834)	(32,113)	(83,838)	(77,809)	(84,832)	(23,825)	44,368
% Change, Y/Y	-	(5230.5%)	156.8%	132.1%	161.1%	(7.2%)	9.0%	(71.9%)	(286.2%)
Net income	(4,628)	791	(4,937)	1,141	18,619	77,159	102,998	171,911	449,390
Depreciation and amortization	2,635	3,824	6,701	13,123	23,518	37,963	68,700	110,103	195,331
Stock-based compensation	679	784	1,947	2,917	3,305	3,746	4,132	4,372	4,684
Impairment and other charges	3,811	63	38	-	-	-	-	-	-
Non-cash interest income and other	-	(1,359)	947	998	(703)	(1,031)	(330)	(1,340)	(1,180)
Other, net	214	-	-	-	-	-	-	-	-
Accounts receivable	345	(347)	(543)	(586)	(3,349)	(6,081)	(7,510)	(11,852)	(12,253)
Inventories	(72)	(309)	(413)	(2,746)	47	1,110	337	131	296
Prepaid expenses and other current assets	41	(328)	(673)	(938)	124	365	(19)	6	(75)
Accounts payable	220	472	709	(1,160)	2,774	2,306	8,769	9,947	24,801
Deferred revenue	1,695	1,876	1,929	-	-	-	-	-	-
Other liabilities	(420)	309	1,080	3,118	(695)	(16,100)	4,463	23,527	69,606
Working capital requirements	1,809	1,673	2,089	(2,312)	(1,099)	(18,401)	6,040	21,758	82,375
Other	-	-	-	-	-	-	-	-	-
Cash flow from operations	4,520	5,776	6,785	15,868	43,638	99,435	181,540	306,804	730,600
Purchase of property, plant & equipment	(4,415)	(11,163)	(20,619)	(47,980)	(127,477)	(177,244)	(266,372)	(330,629)	(686,232)
Proceeds from the sale of assets and businesses	-	-	-	-	-	-	-	-	-
Cash paid for acquisitions, net of cash	-	-	-	-	-	-	-	-	-
Other	(452)	367	1,044	(7,812)	-	-	-	-	-
Cash flow from investing	(4,867)	(10,796)	(19,575)	(55,792)	(127,477)	(177,244)	(266,372)	(330,629)	(686,232)
Proceeds from issuance of debt	-	(77)	9,161	8,649	78,115	121,524	96,008	45,952	-
Repayments on debt	-	-	(30)	(737)	(447)	-	-	-	(20,676)
Payment of debt financing costs	(112)	-	-	-	-	-	-	-	-
Finance lease payments	-	(154)	(295)	(373)	(470)	(393)	-	-	-
Public offering proceeds / equity issuance	774	12,061	17,514	63,472	-	-	-	-	-
Repurchase of stock	(170)	-	-	(579)	(1,653)	(2,432)	(2,552)	(2,124)	(2,780)
Other	(70)	-	-	-	-	-	-	-	-
Cash flow from financing	422	11,830	26,350	70,431	75,545	118,699	93,456	43,828	(23,456)
FX Adjustment	(2)	1	63	36	-	-	-	-	-
Net change in cash	73	6,811	13,623	30,543	(8,293)	40,890	8,624	20,003	20,912
Beginning cash, cash equivalents and restricted cash	4,617	-	11,501	25,124	55,667	47,373	88,263	96,887	116,890
Ending cash	4,690	11,501	25,124	55,667	47,373	88,263	96,887	116,890	137,802
Less: Restricted cash, included in other assets	-	116	-	-	-	-	-	-	-
Cash and cash equivalents	4,690.0	11,385	25,124	55,667	47,373	88,263	96,887	116,890	137,802
Free cash flow	105	(5,387)	(13,834)	(32,113)	(83,838)	(77,809)	(84,832)	(23,825)	44,368
Cash from operations	4,520	5,776	6,785	15,868	43,638	99,435	181,540	306,804	730,600
Capex	(4,415)	(11,163)	(20,619)	(47,980)	(127,477)	(177,244)	(266,372)	(330,629)	(686,232)
Free cash flow	105	(5,387)	(13,834)	(32,113)	(83,838)	(77,809)	(84,832)	(23,825)	44,368
FCF conversion	(2%)	(681%)	280%	(2815%)	(450%)	(101%)	(82%)	(14%)	10%
Seasonality (% of FY)	100.0%	100.0%	100.0%	100.0%	100.0%	100.0%	100.0%	100.0%	100.0%
Free Cash Flow / Share	\$ 0.04	\$ (0.54)	\$ (4.73)	\$ (2.57)	\$ (6.41)	\$ (5.95)	\$ (6.48)	\$ (1.82)	\$ 3.39
Weighted average diluted shares	2,759.0	9,956.0	2,926.0	12,491.7	13,073.0	13,081.0	13,089.0	13,097.0	13,105.0

Source: Company reports, RBC Capital Markets estimates



Target/Upside/Downside Scenarios

SpaceX



Source: Bloomberg and RBC Capital Markets estimates for Target

Valuation

Our \$225 price target is based on applying a weighted average 15.2x EV/EBITDA multiple to our 2029E EBITDA \$192.5B. We value the company on a SOTP basis, applying a 25x multiple to Space, 15x to Connectivity, and 15x to AI, all at their respective FY29 EBITDA estimates. We view the company as very well-positioned to execute against its strategy, even if ambitious. We believe our valuation framework applies a multiple comparable to other A&D peers, but is slightly ahead from a Connectivity and AI perspective. We view our price target as supportive of our Outperform rating.

Upside scenario

Our upside scenario is \$260/share. We believe that the company could realize faster revenue growth and better EBITDA expansion with a faster-than-estimated Starship launch cadence. We view this scenario as likely to pull forward adoption and usage of both Starlink and the company's AI compute capacity. In this scenario, we expect total sales to hit ~\$276BB, which we believe could generate ~\$204B in EBITDA (74% margins). A 16.5x multiple on the higher EBITDA implies a per-share value of \$260.

Downside scenario

Our downside scenario is \$95/share. We believe that a slower Starship launch cadence could push revenue growth in the Connectivity and AI segments to the right. In this scenario, we model revenue growth of just ~\$132B and EBITDA of ~\$90B in FY29E. We then apply a lower multiple of 13.5x on the lower EBITDA framework to arrive at our \$95 downside scenario.

Investment summary

Space Exploration Technologies Corp. (SPCX) is an integrated space technology company operating across three segments — Space (launch), Connectivity (Starlink), and AI (Grok/compute) — targeting a combined TAM of ~\$28.5T. The company controls 80%+ of global mass to orbit, operates ~9,600 Starlink satellites serving 10.3M+ subscribers across 164 countries, and following its acquisition of xAI, has deployed gigawatt-scale AI compute infrastructure (Colossus I & II) anchored by landmark cloud contracts with Anthropic (\$1.25B/month) and Google (\$920M/month). We would be buyers of SPCX for the following reasons: (1) a structural and potentially durable cost-of-compute advantage in terrestrial datacenters, with a longer-term orbital compute opportunity that could dwarf anything achievable on the ground; (2) Starlink's highly scalable, capital-light connectivity flywheel generating 60%+ EBITDA margins with a \$24B signed revenue backlog through 2028; and (3) a CEO with unparalleled access to capital, engineering talent, and a target-rich distribution channel of hundreds of millions of users — positioning SPCX to emerge as a leading AI intelligence provider. We rate SPCX Outperform with a \$225 price target.

Risks to rating and price target

- **Starship Development Risk:** Delays or failures in Starship's development could impair SpaceX's launch, Connectivity, and AI growth strategies.
- **Launch Failures & Operational Hazards:** Mission failures, launch delays, and uninsured asset losses pose ongoing material risks to operations and reputation.
- **Regulatory & Licensing Risk:** Failure to obtain or renew FAA, FCC, or international licenses could materially disrupt operations across all segments.
- **Space Environment & Collision Risk:** Orbital debris, radiation, and LEO congestion expose satellites and launch vehicles to malfunction or total loss.
- **AI Segment Integration & Execution Risk:** The recently acquired xAI segment remains early-stage, loss-generating, and subject to significant integration and execution uncertainty.
- **Key Person Dependency — Elon Musk:** Loss or reduced involvement of Elon Musk could significantly disrupt strategy, culture, and investor confidence.
- **Concentrated Voting Control & Governance Risk:** Musk controls ~82% of voting power post-IPO, limiting Class A shareholder influence over corporate matters.
- **AI Infrastructure & Capital Intensity Risk:** Scaling AI operations requires massive, ongoing capital investment in power, GPUs, and datacenter infrastructure.



Key fundamental questions

Our view

Space - Can SpaceX achieve its desired Starship launch cadence — and is it realistic?

This is the foundational question for the entire SPCX thesis. Starship's ability to scale is the prerequisite for both the Connectivity and AI segments — without it, neither orbital compute nor V3 Starlink constellation upgrades are achievable. The company has outlined a framework requiring 3,500–7,800 cumulative Starship launches to reach 50 GW of orbital compute. We view this path as ambitious but achievable, provided the company continues to improve reusability and production throughput. We believe sustained launch cadence improvement post-2028 is the single most important stock catalyst to monitor.

Connectivity - Does Starlink's cost structure allow it to remain competitively priced against terrestrial broadband providers?

Starlink's ~5-year satellite replacement cycle creates a recurring capex burden that structurally disadvantages it versus fiber and cable operators with 15–25-year asset lives. However, we believe a combination of V3 satellite cost reductions, CPE equipment subsidies, and the high-margin Government & Enterprise revenue mix can meaningfully improve economic gross margins over time. Broadband subscribers nearly doubled in 2025, demonstrating Starlink can compete on price — but sustaining that pricing power as the constellation matures and ARPUs compress remains the key watch item, in our view.

Connectivity - What is the viable — and value-creating — growth path for Starlink Mobility?

The Direct-to-Device "dead zone" coverage model is elegant and capital-light, but RBC estimates it addresses no more than 5% of the Mobility TAM. Expanding meaningfully into the U.S. mobile market requires either an MVNO arrangement with one of the Big 3 carriers or significant terrestrial spectrum investment — both of which introduce operational complexity and capital intensity that could pressure ROIC. RBC's preferred path is MVNO over building a fourth wireless network, with IoT, V2X, and UAV connectivity representing the more compelling medium-term growth levers beyond 2028.

AI - Is the rise of cheaper open-source and open-weight models a risk to SPCX's compute business — or an accelerant?

We view this as an accelerant, not a risk. Token cost elasticity is proving explosive: the OpenRouter market grew 7x in nine months (6.5T to 44T tokens/month), and Anthropic's usage surged 1,711% even as prices rose 67%. Cheaper models expand the surface area of what organizations can automate, accelerating the shift to always-on agentic workloads that ultimately require more compute, not less. Crucially, enterprises that control the compute layer — as SPCX does — gain model-agnostic flexibility and durable positioning regardless of which model wins.

AI - Can cost of compute, Grok, and Cursor enable SPCX to become a leading AI intelligence provider?

SPCX today is the laggard among the four major AI labs, with Grok's intelligence index at 38 versus a market-leading 60. However, we believe the company possesses all the necessary building blocks: owned compute (Grokus cluster, no hyperscaler margin paid), a rich proprietary data flywheel from X and Tesla, and — pending close — the \$60B Cursor acquisition, which would embed Grok as the default model in a widely used developer IDE. Combined with distribution across hundreds of millions of X and Tesla users, we believe the path to competitive intelligence provider status is achievable, though execution risk on model quality and enterprise integration remains high.



Key sustainability questions

This section is intended to highlight key sustainability discussion points relevant to this company, as well as our views on the outlook. Both the questions we highlight and our responses will evolve over time as the dialogue between management, analysts and investors continues to advance. We welcome any feedback on the topics.

Our view

What are the most material sustainability issues facing this company?

We believe SpaceX faces several material sustainability issues. On the environmental side, the company's launch facilities, datacenters, and manufacturing operations are subject to broad environmental regulation, and growing scrutiny of rocket emissions in Earth's upper atmosphere could introduce new constraints on launch cadence. Energy intensity is also a pressing concern — SpaceX's AI datacenters rely heavily on natural gas and large volumes of water for cooling, creating both operational and reputational risk as the AI segment scales. Space debris and orbital congestion represent a sustainability issue unique to SpaceX's business model; as the operator of ~75% of all active satellites in Low-Earth Orbit, stricter debris mitigation regulations could increase costs and constrain constellation expansion. On the positive side, Starlink's provision of connectivity to underserved communities, disaster relief operations, and remote healthcare and education represents a meaningful social benefit and will likely anchor the company's broader ESG narrative going forward.

Does the company integrate sustainability considerations into its strategy?

As a recent IPO, we do not believe sustainability has been a significant focus for the company, and is not directly discussed within its prospectus. However, we believe the company will look to implement long-term carbon reduction efforts similar to its other large-scale industrial peers. The re-usability of rockets and reduction in cost/launch is a key focus for the company and investors.

What is diversity like at board / management level?

SpaceX has an 8-member Board of Directors. We note that 4 of the 8 are considered independent directors. We believe the Board of Directors has a broad diversity of experience and leadership. Moreover, Elon Musk holds approximately 82.3% of the total voting power of outstanding common stock, and with his ~92% ownership of the Class B common stock, he is entitled to elect 51% of the total number of authorized directors. The concentration of power also qualifies SpaceX as a controlled company.



Company description

Space Exploration Technologies Corp. (SPCX) is the only company building the fully integrated hardware and software infrastructure of the future across three converging domains: Space, Connectivity, and Artificial Intelligence (AI). Founded in 2002 by Elon Musk and headquartered in Starbase, Texas, SpaceX designs, manufactures, launches, and operates the world's most advanced rockets and spacecraft, transporting astronauts, satellites, and payloads with an industry-leading mission success rate exceeding 99% using its Falcon rocket family. Since 2021, the company has accounted for more than 80% of all mass launched to orbit globally each year.

SpaceX's Connectivity segment operates Starlink, a high-speed, low-latency global broadband network powered by approximately 9,600 satellites in Low-Earth Orbit, serving over 10 million subscribers across 164 countries as of March 2026. Its AI segment — anchored by xAI, acquired in early 2026 — develops Grok, a frontier AI model, alongside cutting-edge AI compute infrastructure, including the COLOSSUS datacenter in Memphis, Tennessee.

Guided by a mission to make life multiplanetary and extend the reach of human consciousness, SpaceX combines extreme vertical integration, first-principles engineering, and a deeply mission-driven culture to create what it believes to be an unprecedented, self-reinforcing engine of innovation and value creation.

Required disclosures

Non-U.S. analyst disclosure

One or more research analysts involved in the preparation of this report (i) may not be registered/qualified as research analysts with the NYSE and/or FINRA and (ii) may not be associated persons of the RBC Capital Markets, LLC and therefore may not be subject to FINRA Rule 2241 restrictions on communications with a subject company, public appearances and trading securities held by a research analyst account.

Conflicts disclosures

The analyst(s) responsible for preparing this research report received compensation that is based upon various factors, including total revenues of the member companies of RBC Capital Markets and its affiliates, a portion of which are or have been generated by investment banking activities of the member companies of RBC Capital Markets and its affiliates.

With regard to the MAR investment recommendation requirements in relation to relevant securities, a member company of Royal Bank of Canada, together with its affiliates, may have a net long or short financial interest in excess of 0.5% of the total issued share capital of the entities mentioned in the investment recommendation. Information relating to this is available upon request from your RBC investment advisor or institutional salesperson.

Please note that current conflicts disclosures may differ from those as of the publication date on, and as set forth in, this report. To access current conflicts disclosures, clients should refer to <https://www.rbccm.com/GLDisclosure/PublicWeb/DisclosureLookup.aspx?entityId=1> or send a request to RBC CM Research Publishing, P.O. Box 50, 200 Bay Street, Royal Bank Plaza, 29th Floor, South Tower, Toronto, Ontario M5J 2W7.

A member company of RBC Capital Markets or one of its affiliates received compensation for investment banking services from Space Exploration Technologies Corp. in the past 12 months.

A member company of RBC Capital Markets or one of its affiliates managed or co-managed a public offering of securities for Space Exploration Technologies Corp. in the past 12 months.

A member company of RBC Capital Markets or one of its affiliates received compensation for products or services other than investment banking services from Space Exploration Technologies Corp. during the past 12 months. During this time, a member company of RBC Capital Markets or one of its affiliates provided non-securities services to Space Exploration Technologies Corp..

RBC Capital Markets has provided Space Exploration Technologies Corp. with non-securities services in the past 12 months.

RBC Capital Markets, LLC makes a market in the securities of Space Exploration Technologies Corp..

RBC Capital Markets is currently providing Space Exploration Technologies Corp. with investment banking services.

A member company of RBC Capital Markets or one of its affiliates expects to receive or intends to seek compensation for investment banking services from Space Exploration Technologies Corp. in the next three months.



Explanation of RBC Capital Markets Equity rating system

An analyst's 'sector' is the universe of companies for which the analyst provides research coverage. Accordingly, the rating assigned to a particular stock represents solely the analyst's view of how that stock will perform over the next 12 months relative to the analyst's sector average.

Ratings

Outperform (O): Expected to materially outperform sector average over 12 months.

Sector Perform (SP): Returns expected to be in line with sector average over 12 months.

Underperform (U): Returns expected to be materially below sector average over 12 months.

Restricted (R): RBC policy precludes certain types of communications, including an investment recommendation, when RBC is acting as an advisor in certain merger or other strategic transactions and in certain other circumstances.

Not Rated (NR): The rating, price targets and estimates have been removed due to applicable legal, regulatory or policy constraints which may include when RBC Capital Markets is acting in an advisory capacity involving the company.

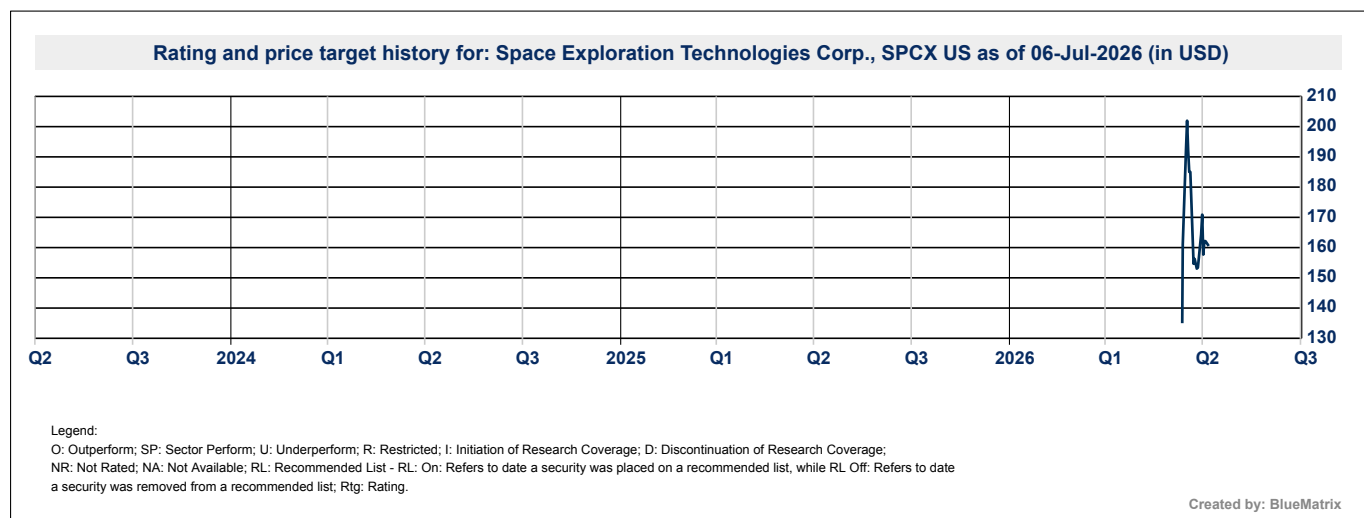
Risk Rating

The **Speculative** risk rating reflects a security's lower level of financial or operating predictability, illiquid share trading volumes, high balance sheet leverage, or limited operating history that result in a higher expectation of financial and/or stock price volatility.

Distribution of ratings

For the purpose of ratings distributions, regulatory rules require member firms to assign ratings to one of three rating categories - Buy, Hold/Neutral, or Sell - regardless of a firm's own rating categories. Although RBC Capital Markets' ratings of Outperform (O), Sector Perform (SP), and Underperform (U) most closely correspond to Buy, Hold/Neutral and Sell, respectively, the meanings are not the same because our ratings are determined on a relative basis.

Distribution of ratings				
RBC Capital Markets, Equity Research				
As of 30-Jun-2026				
Rating	Count	Percent	Investment Banking	
			Serv./Past 12 Mos.	
			Count	Percent
BUY [Outperform]	927	57.94	290	31.28
HOLD [Sector Perform]	626	39.12	161	25.72
SELL [Underperform]	47	2.94	3	6.38



References to a Recommended List in the recommendation history chart may include one or more recommended lists or model portfolios maintained by RBC Wealth Management or one of its affiliates. RBC Wealth Management recommended lists include the Guided Portfolio: Prime Income (RL 6), the Guided Portfolio: Dividend Growth (RL 8), the Guided Portfolio: ADR (RL 10), and



the Guided Portfolio: All Cap Growth (RL 12). The abbreviation 'RL On' means the date a security was placed on a Recommended List. The abbreviation 'RL Off' means the date a security was removed from a Recommended List. As of April 3, 2023, U.S. RBC Wealth Management's quarterly reports will serve as the primary communication for its models and will highlight any changes to the model made during the quarter.

Equity valuation and risks

For valuation methods used to determine, and risks that may impede achievement of, price targets for covered companies, please see the most recent company-specific research report at www.rbcinsightresearch.com or send a request to RBC Capital Markets Research Publishing, P.O. Box 50, 200 Bay Street, Royal Bank Plaza, 29th Floor, South Tower, Toronto, Ontario M5J 2W7.

Space Exploration Technologies Corp.

Valuation

Our \$225 price target is based on applying a weighted average 15.2x EV/EBITDA multiple to our 2029E EBITDA \$192.5B. We value the company on a SOTP basis, applying a 25x multiple to Space, 15x to Connectivity, and 15x to AI, all at their respective FY29 EBITDA estimates. We view the company as very well-positioned to execute against its strategy, even if ambitious. We believe our valuation framework applies a multiple comparable to other A&D peers, but is slightly ahead from a Connectivity and AI perspective. We view our price target as supportive of our Outperform rating.

Risks to rating and price target

- **Starship Development Risk:** Delays or failures in Starship's development could impair SpaceX's launch, Connectivity, and AI growth strategies.
- **Launch Failures & Operational Hazards:** Mission failures, launch delays, and uninsured asset losses pose ongoing material risks to operations and reputation.
- **Regulatory & Licensing Risk:** Failure to obtain or renew FAA, FCC, or international licenses could materially disrupt operations across all segments.
- **Space Environment & Collision Risk:** Orbital debris, radiation, and LEO congestion expose satellites and launch vehicles to malfunction or total loss.
- **AI Segment Integration & Execution Risk:** The recently acquired xAI segment remains early-stage, loss-generating, and subject to significant integration and execution uncertainty.
- **Key Person Dependency — Elon Musk:** Loss or reduced involvement of Elon Musk could significantly disrupt strategy, culture, and investor confidence.
- **Concentrated Voting Control & Governance Risk:** Musk controls ~82% of voting power post-IPO, limiting Class A shareholder influence over corporate matters.
- **AI Infrastructure & Capital Intensity Risk:** Scaling AI operations requires massive, ongoing capital investment in power, GPUs, and datacenter infrastructure.

Conflicts policy

RBC Capital Markets Policy for Managing Conflicts of Interest in Relation to Investment Research is available from us on request. To access our current policy, clients should refer to

<https://www.rbccm.com/global/file-414164.pdf>

or send a request to RBC Capital Markets Research Publishing, P.O. Box 50, 200 Bay Street, Royal Bank Plaza, 29th Floor, South Tower, Toronto, Ontario M5J 2W7. We reserve the right to amend or supplement this policy at any time.

Dissemination of research

RBC Capital Markets endeavors to make all reasonable efforts to provide research content simultaneously to all eligible clients, having regard to local time zones in overseas jurisdictions. RBC Capital Markets provides eligible clients with access to Research Reports on the Firm's proprietary INSIGHT website, via email and via third-party vendors. Please contact your investment advisor or institutional salesperson for more information regarding RBC Capital Markets' research.

For a list of all recommendations on the company that were disseminated during the prior 12-month period, please click on the following link: <https://rbcnew.bluematrix.com/sellside/MAR.action>

The 12 month history of Quick Takes can be viewed at [RBC Insight](#).

Analyst certification



All of the views expressed in this report accurately reflect the personal views of the responsible analyst(s) about any and all of the subject securities or issuers. No part of the compensation of the responsible analyst(s) named herein is, or will be, directly or indirectly, related to the specific recommendations or views expressed by the responsible analyst(s) in this report.

Third-party disclaimers

The Global Industry Classification Standard ("GICS") was developed by and is the exclusive property and a service mark of MSCI Inc. ("MSCI") and Standard & Poor's Financial Services LLC ("S&P") and is licensed for use by RBC. Neither MSCI, S&P, nor any other party involved in making or compiling the GICS or any GICS classifications makes any express or implied warranties or representations with respect to such standard or classification (or the results to be obtained by the use thereof), and all such parties hereby expressly disclaim all warranties of originality, accuracy, completeness, merchantability and fitness for a particular purpose with respect to any of such standard or classification. Without limiting any of the foregoing, in no event shall MSCI, S&P, any of their affiliates or any third party involved in making or compiling the GICS or any GICS classifications have any liability for any direct, indirect, special, punitive, consequential or any other damages (including lost profits) even if notified of the possibility of such damages.

RBC Capital Markets disclaims all warranties of originality, accuracy, completeness, merchantability or fitness for a particular purpose with respect to any statements made to the media or via social media that are in turn quoted in this report, or otherwise reproduced graphically for informational purposes.

Disclaimer

RBC Capital Markets is the business name used by certain branches and subsidiaries of the Royal Bank of Canada, including RBC Dominion Securities Inc., RBC Capital Markets, LLC, RBC Europe Limited, RBC Capital Markets (Europe) GmbH, Royal Bank of Canada, Hong Kong Branch, Royal Bank of Canada, Singapore Branch and Royal Bank of Canada, Sydney Branch. The information contained in this report has been compiled by RBC Capital Markets from sources believed to be reliable, but no representation or warranty, express or implied, is made by Royal Bank of Canada, RBC Capital Markets, its affiliates or any other person as to its accuracy, completeness or correctness. All opinions and estimates contained in this report constitute RBC Capital Markets' judgement as of the date of this report, are subject to change without notice and are provided in good faith but without legal responsibility. Nothing in this report constitutes legal, accounting or tax advice or individually tailored investment advice. This material is prepared for general circulation to clients and has been prepared without regard to the individual financial circumstances and objectives of persons who receive it. The investments or services contained in this report may not be suitable for you and it is recommended that you consult an independent investment advisor if you are in doubt about the suitability of such investments or services. This report is not an offer to sell or a solicitation of an offer to buy any securities. Past performance is not a guide to future performance, future returns are not guaranteed, and a loss of original capital may occur. RBC Capital Markets research analyst compensation is based in part on the overall profitability of RBC Capital Markets, which includes profits attributable to investment banking revenues. Every province in Canada, state in the U.S., and most countries throughout the world have their own laws regulating the types of securities and other investment products which may be offered to their residents, as well as the process for doing so. As a result, the securities discussed in this report may not be eligible for sale in some jurisdictions. RBC Capital Markets may be restricted from publishing research reports, from time to time, due to regulatory restrictions and/ or internal compliance policies. If this is the case, the latest published research reports available to clients may not reflect recent material changes in the applicable industry and/or applicable subject companies. RBC Capital Markets research reports are current only as of the date set forth on the research reports. This report is not, and under no circumstances should be construed as, a solicitation to act as securities broker or dealer in any jurisdiction by any person or company that is not legally permitted to carry on the business of a securities broker or dealer in that jurisdiction. To the full extent permitted by law neither RBC Capital Markets nor any of its affiliates, nor any other person, accepts any liability whatsoever for any direct, indirect or consequential loss arising from, or in connection with, any use of this report or the information contained herein. No matter contained in this document may be reproduced or copied by any means without the prior written consent of RBC Capital Markets in each instance.

Additional information is available on request.

To U.S. Residents:

This publication has been approved by RBC Capital Markets, LLC (member FINRA, NYSE, SIPC), which is a U.S. registered broker-dealer and which accepts responsibility for this report and its dissemination in the United States. Any U.S. recipient of this report that is not a registered broker-dealer or a bank acting in a broker or dealer capacity and that wishes further information regarding, or to effect any transaction in, any of the securities discussed in this report, should contact and place orders with RBC Capital Markets, LLC.

To Canadian Residents:

This publication has been approved by RBC Dominion Securities Inc.(member CICO). Any Canadian recipient of this report that is not a Designated Institution in Ontario, an Accredited Investor in British Columbia or Alberta or a Sophisticated Purchaser in Quebec (or similar permitted purchaser in any other province) and that wishes further information regarding, or to effect any transaction in, any of the securities discussed in this report should contact and place orders with RBC Dominion Securities Inc., which, without in any way limiting the foregoing, accepts responsibility for this report and its dissemination in Canada.

To U.K. Residents:

This publication has been approved by RBC Europe Limited ('RBCEL') which is authorized by the Prudential Regulation Authority and regulated by the Financial Conduct Authority ('FCA') and the Prudential Regulation Authority, in connection with its distribution in the United Kingdom. This material is not for general distribution in the United Kingdom to retail clients, as defined under the rules of the FCA. RBCEL accepts responsibility for this report and its dissemination in the United Kingdom.

To EEA Residents:

This material is distributed in the EU by either RBCEL on an authorised cross-border basis, or by RBC Capital Markets (Europe) GmbH (RBC EG) which is authorised and regulated in Germany by the Bundesanstalt für Finanzdienstleistungsaufsicht (German Federal Financial Supervisory Authority) (BaFin).

To Persons Receiving This Advice in Australia:

This material has been distributed in Australia by Royal Bank of Canada, Sydney Branch (ABN 86 076 940 880, AFSL No. 246521). This material has been prepared for general circulation and does not take into account the objectives, financial situation or needs of any recipient. Accordingly, any recipient should, before acting on this material, consider the appropriateness of this material having regard to their objectives, financial situation and needs. If this material relates to the acquisition or possible acquisition of a particular financial product, a recipient in Australia should obtain any relevant disclosure document prepared in respect of that product and consider that document before making any decision about whether to acquire the product. This research report is not for retail investors as defined in section 761G of the Corporations Act.

To persons receiving this from Royal Bank of Canada, Hong Kong Branch:



This document is distributed in Hong Kong by Royal Bank of Canada, Hong Kong Branch which is regulated by the Hong Kong Monetary Authority and the Securities and Futures Commission. This document is not for distribution in Hong Kong, to investors who are not “professional investors”, as defined in the Securities and Futures Ordinance (Cap. 571 of Hong Kong) and any rules made under that Ordinance. This document has been prepared for general circulation and does not take into account the objectives, financial situation, or needs of any recipient. Past performance is not indicative of future performance. WARNING: The contents of this document have not been reviewed by any regulatory authority in Hong Kong. Investors are advised to exercise caution in relation to the investment. If you are in doubt about any of the contents of this document, you should obtain independent professional advice.

To persons receiving this from Royal Bank of Canada, Singapore Branch:

This publication is distributed in Singapore by the Royal Bank of Canada, Singapore Branch, a registered entity licensed by the Monetary Authority of Singapore. This publication is not for distribution in Singapore, to investors who are not “accredited investors” and “institutional investors”, as defined in the Securities and Futures Act 2001 of Singapore. This publication has been prepared for general circulation and does not take into account the objectives, financial situation, or needs of any recipient. You are advised to seek independent advice from a financial adviser before purchasing any product. If you do not obtain independent advice, you should consider whether the product is suitable for you. Past performance is not indicative of future performance. If you have any questions related to this publication, please contact the Royal Bank of Canada, Singapore Branch.

To Japanese Residents:

Unless otherwise exempted by Japanese law, this publication is distributed in Japan by or through RBC Capital Markets (Japan) Ltd. which is a Financial Instruments Firm registered with the Kanto Local Financial Bureau (Registered number 203) and a member of the Japan Securities Dealers Association (JSDA) and the Financial Futures Association of Japan (FFAJ).

® Registered trademark of Royal Bank of Canada. RBC Capital Markets is a trademark of Royal Bank of Canada. Used under license.

Copyright © RBC Capital Markets, LLC 2026 - Member SIPC

Copyright © RBC Dominion Securities Inc. 2026 - Member Canadian Investor Protection Fund

Copyright © RBC Europe Limited 2026

Copyright © Royal Bank of Canada 2026

All rights reserved