

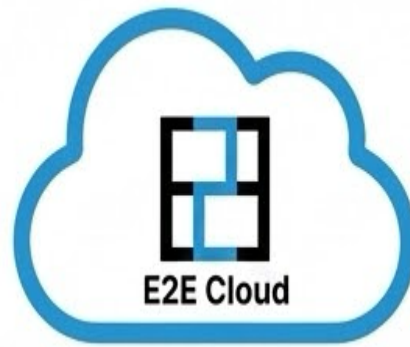


TABLE OF CONTENTS

1.	About the Company.....	2
2.	Industry Outlook.....	7
3.	Explaining the Business Model.....	12
4.	Investment Rationale	14
5.	Key Trackables.....	19
6.	Financials.....	21

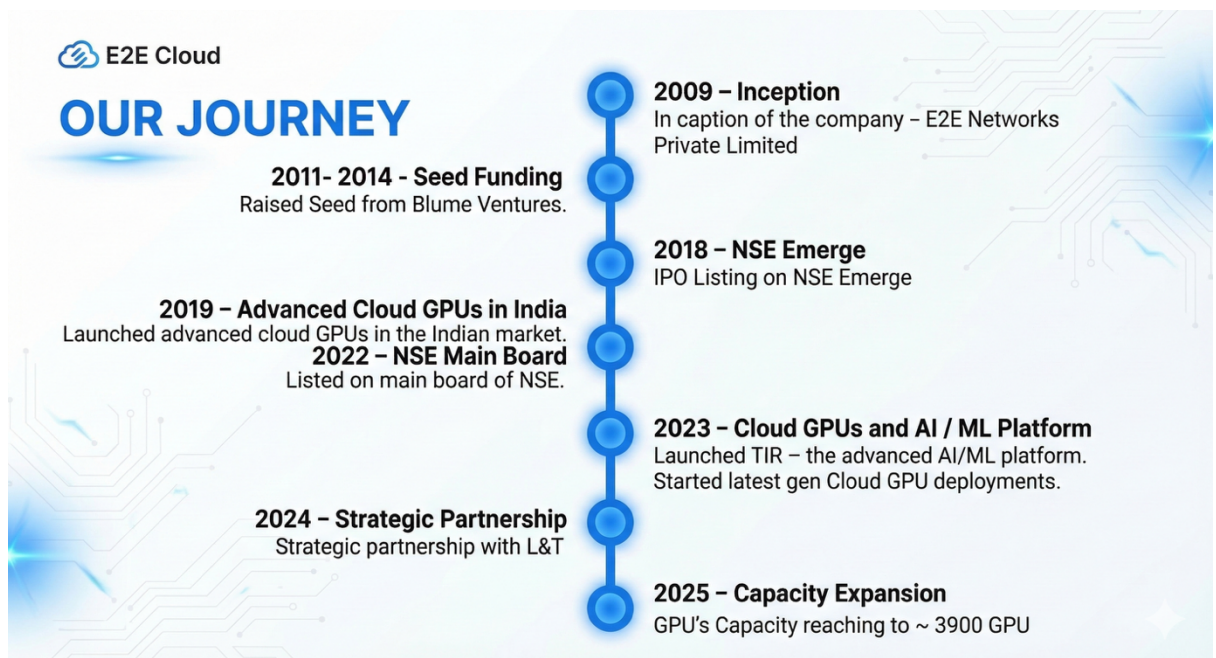
About the Company

Company Analysis: The company that brought H100s to India first

E2E Networks was founded in 2009 by Tarun Dua, a computer engineering graduate from REC Kurukshetra with over two decades of experience in open-source infrastructure, Linux, and cloud computing. The company started as a traditional cloud services provider but pivoted decisively toward GPU-first AI infrastructure beginning in 2019. Early seed funding came from Blume Ventures (2011–2014), and the company listed on NSE Emerge in 2018.

The transformative moment came in November 2023, when E2E became the first Indian cloud provider to deploy NVIDIA H100 GPUs. This first-mover advantage positioned the company as a credible Indian alternative for AI workloads that had previously defaulted to hyperscalers. By 2024, the fleet expanded to include H200 GPUs, and on January 9, 2026, E2E procured its first batch of 1,024 NVIDIA Blackwell B200 GPUs at the L&T Vyoma Data Center in Chennai making it among the earliest Blackwell deployers in India.

Journey of the Company



Platform Architecture & Capabilities

E2E operates a full-stack AI cloud ecosystem built on open-source and in-house software:

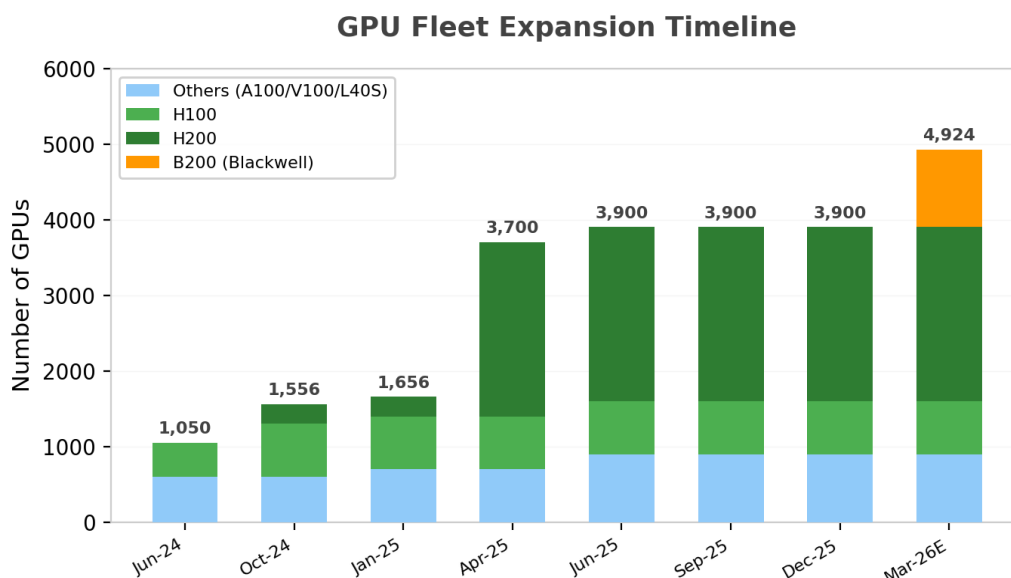
- **TIR GenAI Platform:** End-to-end AI infrastructure supporting the complete ML lifecycle development, training, deployment, and inference with pre-configured environments and NVIDIA Enterprise software integration.
- **Cloud GPUs:** NVIDIA B200, H200, H100, A100, L40S, and V100 available on-demand, spot, and reserved instances.
- **Linux Cloud:** Self-service deployment via control panel and API
- **Sovereign Cloud Platform:** Launched in 2025, designed for minimal foreign software dependencies targeting national data centers, government, and enterprises with compliance requirements
- **Supporting infrastructure:** Object storage, CDN, load balancers, firewalls, VPC, DBaaS, DNS/rDNS, and continuous data protection.

The company holds SOC2, ISO 27001, ISO 27017, ISO 27018, and PCI DSS certifications and is MeitY-empanelled for government cloud procurement.

Management Team

Tarun Dua (Founder, Chairman & Managing Director) drives technical and strategic direction. His re-appointment through January 2031 signals long-term founder commitment. Srishti Baweja (Whole-Time Director, former CFO) brings 20 years in finance and compliance. Nitin Jain serves as CFO, with prior experience at Bata Malaysia. Mohammed Imran (CTO) leads cloud and networking architecture, while Bakshish Dutta (Chief Business Officer) handles enterprise sales with 21+ years of B2B experience. The board includes two L&T nominees Seema Ambastha (CEO of L&T-Cloudfiniti, 33 years in IT/data center management) and Prashant Jain providing direct enterprise connectivity.

GPU fleet and data centers



As of December 2025, E2E operates 3,900 GPUs: approximately 700 H100s, 2,300 H200s, and 900 others (A100, V100, L40S). The 1,024 Blackwell B200 cluster (184 TB GPU RAM) procured in January 2026 is expected to bring the fleet to 5,000 GPUs by March 2026. Management has indicated plans to potentially scale to 2,048–4,096 Blackwell units. Data centers are located in Chennai (L&T Vyoma Data Center primary facility hosting the Blackwell cluster), Delhi NCR, and Mumbai (being wound down after a reliability incident). The Chennai facility, operated by L&T-Cloudfiniti, provides colocation with enterprise-grade power and cooling infrastructure.

Subsidiaries and strategic acquisitions

Jarvis Labs acquisition was completed in December 2025. Jarvis Labs, a Coimbatore-based GPU cloud startup focused on AI and deep learning, brings intellectual property, hardware, domain expertise, and a customer base that positions E2E for global self-serve customer acquisition.

E2E maintains strategic technology partnerships with NVIDIA (Preferred Cloud Partner), Intel, AMD, HPE, Dell Technologies, and Microsoft.

L&T Partnership: From Equity Investor to Infrastructure Scaling Partner

The E2E–L&T relationship has evolved rapidly through three distinct phases, each deepening the strategic lock-in:

Phase 1 - Equity Investment (November 2024): Larsen & Toubro acquired a 21% stake in E2E Networks for ₹1,407 crore, structured through a preferential equity allotment of ~30 lakh shares at ₹3,622/share plus a secondary acquisition of ~12 lakh shares at ₹2,750/share from founders. The deal included a software license agreement, reseller agreement, and colocation agreement.

Phase 2 - Enterprise Channel Activation (January 2026): A ₹8.49 crore, one-year GPU services contract was signed — the first commercial order flowing through the L&T channel. Management confirmed on the Q3 FY26 concall that early enterprise conversions have begun, with similar L&T-type orders in active pipeline.

Phase 3 - MOU for GPU Cloud Infrastructure Scaling (February 20, 2026): E2E and Larsen & Toubro-Vyoma (Data Center & Cloud Services business of L&T) signed a Memorandum of Understanding to partner in scaling GPU cloud infrastructure. Under the MOU, GPU infrastructure procured, owned, or deployed by L&T-Vyoma will be integrated into E2E's proprietary cloud and orchestration platform — powered by E2E's Technology Integration & Runtime (TIR) platform. The collaboration focuses on technical integration, commercial structuring, pricing frameworks, and operational alignment for scalable GPU infrastructure delivery.

This is a structurally significant development. It means E2E is not just a tenant in L&T's data centers it is becoming the software and commercialization layer for L&T-Vyoma's entire GPU cloud capacity. In effect, L&T brings the physical infrastructure and capital, while E2E brings the AI-native cloud platform, observability, workload intelligence, and customer relationships. This "asset-light scaling" model allows E2E to grow its addressable GPU fleet without proportionally growing its own balance sheet L&T owns the GPUs, E2E monetizes them through its platform.

The MOU is currently a framework agreement and subject to definitive agreements. However, when read alongside L&T's announcement at the India AI Impact Summit (February 18, 2026) of building gigawatt-scale NVIDIA AI factory infrastructure — including 30 MW expansion in Chennai and a new 40 MW facility in Mumbai — the scale of GPU capacity that could flow through E2E's TIR platform over the next 2–3 years is potentially transformative relative to E2E's current ~5,000 GPU fleet.

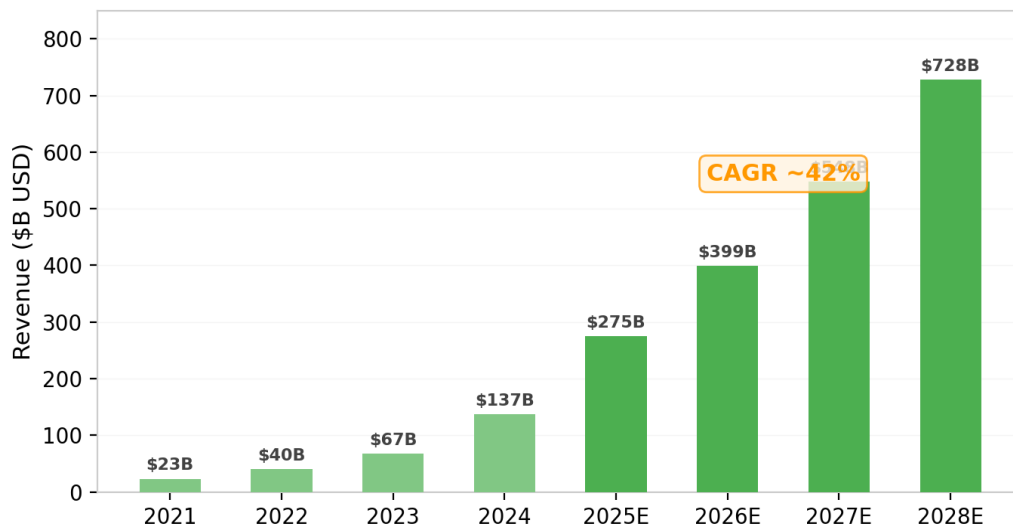
Customer mix and business segments

E2E serves a large base of registered customers across multiple segments. AI startups contribute the majority of revenue, followed by educational institutions and colleges, research organizations, and the remaining share coming from enterprises and government clients. Notable past customers include Zomato, CarDekho, Cars24, Healthkart, and 1mg. The IndiaAI Mission has introduced government-backed AI model developers (Gnani.ai, GANAI) as significant new revenue contributors

Industry Outlook

A \$200+ billion global AI infrastructure wave, and India is just getting started

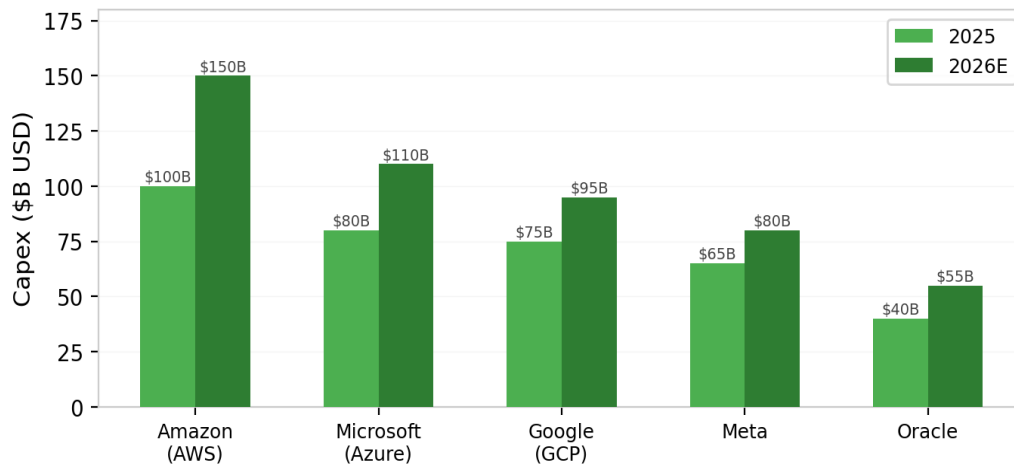
Global Generative AI Revenue Forecast



The global AI infrastructure market encompassing GPU servers, networking, cooling, and supporting data center hardware is estimated at \$80–160 billion in 2025 depending on scope definition. McKinsey projects AI data center spending alone reaching \$934 billion by 2030 (31.6% CAGR), with 70% of global data center demand driven by AI workloads by decade's end. The more narrowly defined GPU-as-a-Service (GPUaaS) market stands at \$8.2 billion in 2025, projected to reach \$26.6 billion by 2030 (26.5% CAGR per MarketsandMarkets), with some estimates stretching to \$50 billion by 2032.

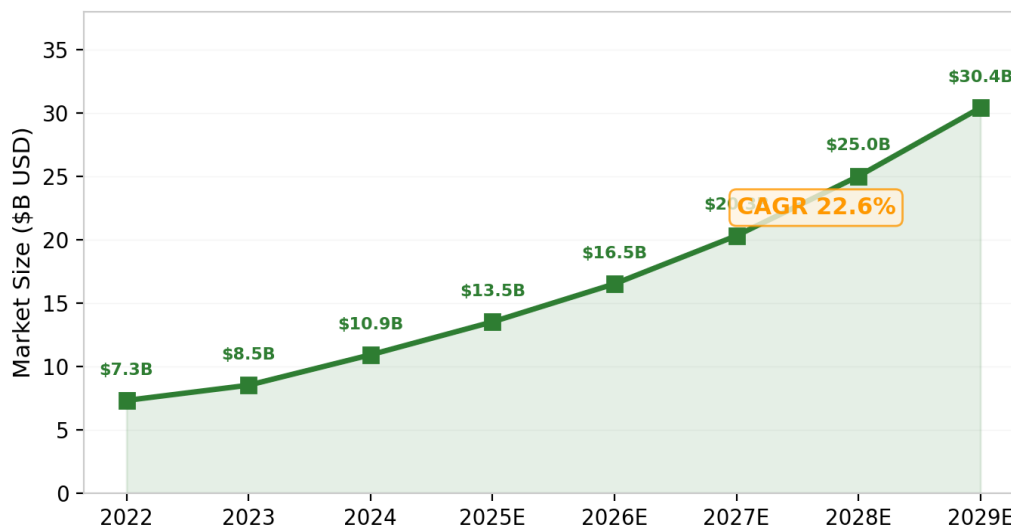
Hyperscaler capital expenditure tells the most compelling story. Combined 2026 capex for the five largest cloud providers (Amazon, Microsoft, Google, Meta, Oracle) is projected at \$660–690 billion up 60% from 2025 with approximately 75% directed at AI infrastructure. Goldman Sachs estimates hyperscaler capex for 2025–2027 at \$1.15 trillion, more than double the \$477 billion spent in 2022–2024.

Hyperscaler AI Infrastructure Capex



India's cloud and AI opportunity

India Public Cloud Services Market



India's public cloud services market reached \$10.9 billion in 2024 and is projected to hit \$30.4 billion by 2029 (22.6% CAGR per IDC's July 2025 update). The critical statistic for the E2E thesis: India generates ~20% of global data but hosts only 3% of global data center capacity. Cloud penetration at 1.5% of the global market versus 6.6% of global GDP represents a massive structural underinvestment.

The generative AI market is projected to reach \$1.3 trillion by 2032 (42% CAGR per Bloomberg Intelligence), with AI infrastructure-as-a-service contributing \$247 billion of that figure. India receives only 1.12% of global AI investment despite its economic weight a gap that government policy and hyperscaler commitments are now aggressively addressing.

Hyperscaler investments reshape India's AI landscape

Since October 2025, the three largest hyperscalers have committed \$67.5 billion to Indian infrastructure: Microsoft (\$17.5B over 4 years), Google (\$15B over 5 years including a gigawatt-scale hub in Visakhapatnam), and Amazon (\$35B+ through 2030). Reliance is building a 1–3 GW data center in Gujarat with NVIDIA Blackwell technology. India's data center capacity is expected to grow from ~1.5 GW in 2025 to 8–9 GW by 2030 a 5x increase.

GPU demand-supply dynamics

The NVIDIA controls 80–92% of the AI data center GPU market. The company's data center revenue reached \$170 billion in FY2026 (ending January 2026), with total revenue of \$215.9 billion up 65% YoY. Blackwell GPUs remain sold out through mid-2026, with a backlog of 3.6 million units from major cloud providers. Key supply bottlenecks include TSMC's CoWoS advanced packaging capacity, HBM memory production, and power availability at customer sites.

The technology roadmap is accelerating. Blackwell (B200/GB200) shipped in volume from early 2025. Blackwell Ultra (B300) began shipping in H2 2025 with 1.5x performance improvement. Rubin (R200) is expected in H2 2026 with 3.3x the inference performance of Blackwell and HBM4 memory. NVIDIA now operates on an annual cadence a new GPU generation every year, a new architecture every two years.

IndiaAI Mission: the government demand catalyst

The IndiaAI Mission, approved by the Cabinet on March 7, 2024, carries a ₹10,372 crore (\$1.25 billion) budget over five years, with ₹4,563 crore allocated specifically to compute capacity. Administered by MeitY through the IndiaAI Independent Business Division, the mission has seven pillars spanning compute, innovation centers, applications, skills, startup financing, datasets, and AI safety.

As of the India AI Impact Summit in February 2026, 38,000+ GPUs have been provisioned under the mission, with 20,000 more announced for deployment in the coming weeks. Round 1 empanelment (January 2025) qualified 10 bidders offering 18,693 GPUs; Round 2 (May 2025) added 15,916 more. A new bidding round for 12,000–15,000 Blackwell B100/B200 GPUs was announced in January 2026.

E2E Networks' IndiaAI position: The company offered 1,353 GPUs (700 H100, 256 H200, 100 each A100/L40S/L4) and was the only provider to make 100% of empanelled GPUs available on day one. E2E secured a ₹177 crore contract to provide GPU resources to Gnani.ai (voice AI model development) and an additional ₹88 crore contract totaling ₹265 crore in IndiaAI orders.

Yotta Data Services is the largest contributor to IndiaAI with 9,216 GPUs (8,192 H100s + 1,024 L40S), representing roughly half of the mission's initial compute capacity. Other empaneled providers include NxtGen (1,088 GPUs), Jio Platforms (312 GPUs, slow deployment), CtrlS, Tata Communications, and AWS managed service providers.

Budget execution has lagged: only ₹800 crore of the ₹2,000 crore FY26 budget estimate was utilized (40% utilization rate), and FY27 allocation was cut to ₹1,000 crore suggesting the mission's pace, while directionally transformative, faces bureaucratic and execution friction.

Sovereign AI and data localization tailwinds

The sovereign AI theme nations building domestic AI infrastructure for strategic autonomy is a \$250 billion global shift. India's approach emphasizes "application-led sovereignty" through multilingual models, voice-first interfaces, and Digital Public Infrastructure. The RBI's payment data localization mandate (2018) requires all payment system data to be stored exclusively in India. The Digital Personal Data Protection Act (2023) enables sectoral regulators to impose data localization requirements. SEBI mandates Indian hosting for all regulated entities. These regulations create a structural floor for domestic cloud demand that hyperscalers operating from foreign jurisdictions cannot fully serve.

India's Frontier AI Model Ecosystem: Building Intelligence, Not Just Infrastructure

The NVIDIA blog reveals a rapidly maturing ecosystem of Indian AI-native companies building sovereign frontier models on domestic infrastructure. Key model builders include:

- **BharatGen** - A government-backed sovereign AI initiative developing a 17B-parameter mixture-of-experts model for public services, agriculture, security, and cultural preservation
- **Sarvam.ai** - Open-sourcing its Sarvam-3 series across 3B, 30B, and 100B parameter sizes for 22 Indic languages, trained on NVIDIA H100 GPUs through cloud partners including Yotta
- **NPCI** - India's payment backbone is exploring training FiMi, a financial model for India, using NVIDIA Nemotron — underscoring that even critical financial infrastructure is moving toward sovereign AI
- **Zoho** - Advancing its Zia LLM platform with proprietary models on NVIDIA Blackwell and Hopper, integrated across its global SaaS applications
- **Tech Mahindra** - Building an 8B-parameter foundation model for Indian languages and dialects, targeting educational applications
- **CoRover.ai** - Already deployed with Indian Railways, supporting ~10,000 concurrent users and 5,000+ daily ticket bookings

The critical investment insight: every one of these model builders requires GPU cloud infrastructure for training and inference. As India's sovereign AI model ecosystem matures from R&D to production deployment, the demand for domestic GPU cloud capacity will compound not from a single government program, but from a growing flywheel of model builders, enterprises, and government agencies. E2E Networks, as one of only three NVIDIA-named cloud partners, sits directly in this value chain.

NVIDIA Cloud Partners Powering India's AI Infrastructure

At the India AI Impact Summit (February 17, 2026), NVIDIA formally announced collaborations with three next-generation cloud providers to deliver advanced AI factories for India's sovereign compute needs. This is a landmark moment it places E2E Networks alongside Yotta and L&T as one of only three named NVIDIA cloud partners for India's AI infrastructure buildout.

Yotta is deploying over 20,000 NVIDIA Blackwell Ultra GPUs across its Shakti Cloud platform, with campuses in Navi Mumbai and Greater Noida delivering GPU-dense, high-bandwidth AI cloud services on a pay-per-use model.

Larsen & Toubro (L&T) is building sovereign, gigawatt-scale NVIDIA AI factory infrastructure, with initial expansions in Chennai to 30 megawatts and a new 40-megawatt facility in Mumbai powering sovereign cloud workloads and hyperscale deployments.

E2E Networks is building an NVIDIA Blackwell GPU cluster on its TIR platform, hosted at the L&T Vyoma Data Center in Chennai. The TIR cloud compute platform features NVIDIA HGX B200 systems and NVIDIA Enterprise software, along with NVIDIA Nemotron open models to accelerate sovereign development across agentic AI, healthcare, finance, manufacturing, and agriculture.

Additionally, Netweb Technologies is launching its Tyrone Camarero AI Supercomputing systems built on the NVIDIA Grace Blackwell architecture (GB200 NVL4) manufactured in India under the "Make in India" mission.

This public endorsement from NVIDIA validates E2E's positioning as a credible sovereign AI infrastructure provider and de-risks concerns around GPU supply allocation.

Explaining the Business Model

How the GPU-as-a-Service economics work

E2E operates a GPU-as-a-Service model with three pricing tiers. On-demand pricing (pay- per-hour, no commitment) for H100 SXM GPUs is ₹249/hour (\$2.93), for H200 it is ₹300/hour, and for B200 it is ₹430/hour (\$5.05). Spot instances offer significant discounts H100 at ₹70/hour (72% discount), H200 at ₹88/hour. Reserved/committed instances (3+month commitments of hundreds of GPUs) provide further savings, with H100 committed pricing at ₹155.90/hour (\$1.83). Management has indicated international Blackwell pricing of \$3–4/hour on longer-term contracts.

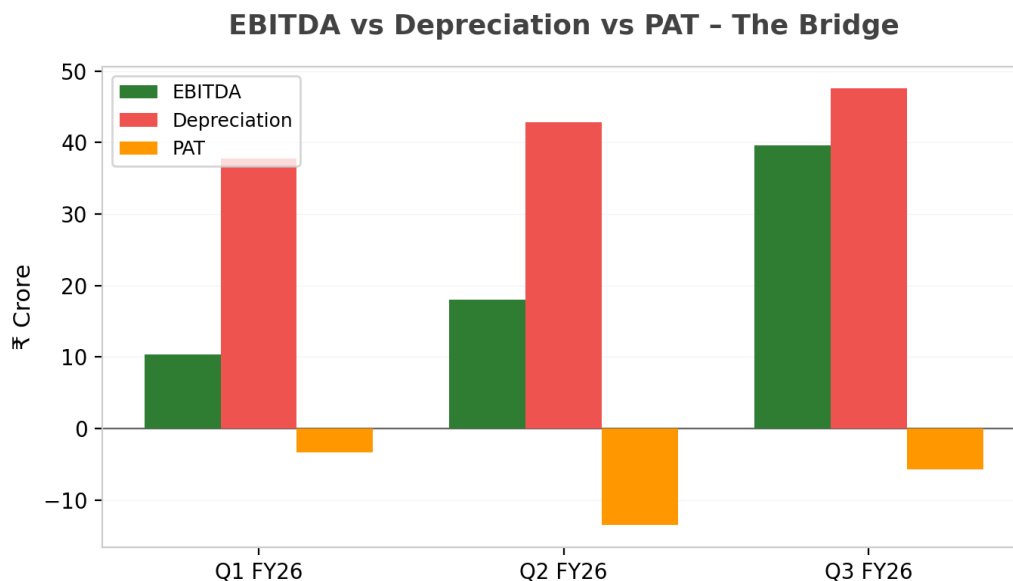
Pricing benchmarks versus competitors

E2E's pricing is almost 2-2.5 times cheaper than hyperscalers in India (AWS India H100: ₹595+/hr; Azure: ₹535+/hr; GCP: ₹815+/hr) and competitive with global neocloud players (Lambda: \$2.49–2.99/hr; CoreWeave: \$4.76/hr on-demand). The IndiaAI Mission's average L1 rate of ₹115.85/hr represents aggressive government-discovered pricing that benefits E2E through volume.

The depreciation challenge driving near-term PAT losses

GPU depreciation is the single most important P&L dynamic for E2E's near-term financials. The company depreciates GPU servers over 6 years on a straight-line basis in line with CoreWeave (6 years), Google (6 years), and Oracle (6 years), though more aggressive than Nebius (4 years, moving to 5) and Amazon (5 years, shortened from 6 in February 2025).

In Q3 FY26, depreciation surged to ₹47.6 crore up 167% YoY from ₹17.8 crore as ₹1,400+ crore of GPU capex entered the books. This single line item turned a ₹39.6 crore EBITDA quarter into a ₹5.7 crore net loss. Management characterizes this as "accounting- led" rather than operational distress.



The book life versus economic useful life debate is consequential. Bulls argue older GPUs cascade through a value chain from frontier training (years 1–2) to real-time inference (years 3–4) to batch processing (years 5–6). CoreWeave's CEO notes that A100 GPUs launched in 2020 remain "fully booked" at 95% of original rental rates. Bears, including Michael Burry, estimate \$176 billion of understated depreciation across hyperscalers for 2026–2028, arguing true economic life is 2–3 years given NVIDIA's annual release cadence.

For E2E, if the depreciation schedule were shortened to 4 years (Nebius standard), annual depreciation would increase approximately 50%, deepening near-term losses. Conversely, the 6-year schedule risks carrying assets at book values that exceed economic reality if rental rates continue declining.

Capex and unit economics

GPU Blackwell B200 unit economics form the core near-term value creation story. E2E is investing ₹600–650 crore for 1,024 B200 GPUs (implying ~₹60–63 lakh per GPU including full infrastructure servers, InfiniBand networking, storage, cooling). At ₹430/hour on-demand pricing and targeted utilization of 70–80%, the 1,024-unit cluster carries an annual revenue potential of ₹250 crore implying a payback period of approximately 2.5–3 years.

Incremental EBITDA margins on deployed GPUs run at 75–80%, as operating costs (primarily power, cooling, and staff) represent only 20% of revenue per cluster. Power consumption per B200 GPU is approximately 1,000W (14kW per 8-GPU system), with India's relatively lower electricity costs providing a structural advantage.

Capital structure and funding

E2E has raised ₹1,485 crore through preferential equity issues (₹406 crore in Q2 FY25, ₹1,079 crore from L&T in Q3 FY25) and an additional ₹107 crore via QIP in February 2026 at ₹2,500/share. Debt facilities total ₹611 crore across term loans (Axis Bank, HDFC at 7.75– 8.6%) and lease facilities, with only ₹154 crore outstanding as of December 2025. The company maintains a conservative debt-to-equity ratio with substantial undrawn capacity (₹350 crore undrawn term loan + ₹100 crore facility).

MRR progression and revenue visibility

MRR has scaled from ₹5.1 crore (June 2022) to ₹28 crore (December 2025) a 5.5x increase in 3.5 years. The trajectory accelerated sharply in Q3 FY26 as IndiaAI Mission contracts commenced deployment. Management targets ₹35–40 crore MRR by March 2026, implying ₹420–480 crore annualized revenue run rate. The Blackwell cluster, once live, could add ₹20+ crore of monthly revenue, providing a clear bridge to the target.

Investment Rationale

Why the investment case is compelling and what could go wrong

IndiaAI Mission as structural demand catalyst

The ₹265 crore in confirmed IndiaAI contracts provide near-term revenue visibility, but the larger opportunity lies in the mission's scaling trajectory. With 38,000+ GPUs provisioned and 20,000 more coming, the ₹4,563 crore compute allocation represents a multi-year revenue pool for empaneled providers. E2E's track record as the only provider to deploy 100% of empanelled capacity on day one positions it well for the upcoming 12,000–15,000 Blackwell GPU procurement round.

The shift from training to inference workloads inference is projected to reach two-thirds of all AI compute by 2026 (Deloitte) means E2E's existing H100/H200 fleet can transition to serving production inference workloads, extending economic utility while newer Blackwell GPUs handle frontier training.

Blackwell economics are highly attractive

The 1,024 B200 cluster's economics ~₹250 crore ARR potential on ₹600–650 crore capex represent the strongest unit economics in E2E's history. At 70% utilization and \$3–4/hour international pricing, the Blackwell cluster generates 35–45% ROIC with sub-3-year payback. This compares favorably to hyperscaler

returns, which are diluted by massive overhead and infrastructure costs that E2E avoids through colocation at L&T facilities.

Metric	Value
Capex (1,024 B200 GPUs)	₹600–650 Crore
ARR Potential (at target utilization)	~₹250 Crore
Incremental EBITDA Margin	75–80%
Payback Period	~2.5–3 years
Revenue Contribution Start	Q1 FY27

Gnani.ai: The E2E Customer Reference Case - From NVIDIA's Own Blog

Gnani.ai, E2E's largest IndiaAI Mission customer (₹177 crore contract), was featured prominently in NVIDIA's official India AI Impact Summit blog post as a key frontier model developer. The validation is significant:

Gnani.ai offers enterprises a multilingual agentic AI platform for voice and text customer interactions. The company is building a 14-billion-parameter speech-to-speech model using NVIDIA Nemotron Speech models, datasets, and NeMo libraries deployed through NVIDIA Cloud Partner E2E Networks with plans to expand to a 32-billion-parameter model. By fine-tuning NVIDIA's Nemotron Speech model for Indic languages, Gnani has achieved a 15x reduction in inference costs, enabling the company to scale to support more than 10 million calls per day for customers in telecom, banking, and hospitality.

This is a powerful proof point for several reasons: it demonstrates that E2E's infrastructure is capable of supporting frontier model training at the 14B–32B parameter scale; it shows that NVIDIA itself publicly names E2E as the cloud partner powering this workload; and it validates the training-to-inference migration thesis — Gnani's 15x inference cost reduction means these workloads will persist on E2E's H100/H200 fleet long after training completes.

Software differentiation and customer stickiness

The TIR GenAI Platform provides differentiation beyond raw compute. Pre-configured environments, NVIDIA Enterprise software integration, and NVIDIA Nemotron open model support create switching costs that commodity GPU rental lacks. The Jarvis Labs acquisition adds global self-serve customer acquisition capability. The L&T enterprise channel opens regulated industries BFSI, healthcare, government where data sovereignty requirements make Indian-hosted infrastructure mandatory.

The hyperscaler competition reality

The risk that AWS, Azure, and GCP could capture India's sovereign AI demand deserves careful analysis. Three factors limit hyperscaler dominance in this specific segment:

- **Data localization compliance** - RBI mandates exclusive India storage for payment data; SEBI requires Indian hosting for regulated entities. While hyperscalers operate Indian regions, their control planes, management tools, and metadata often traverse global infrastructure, creating compliance gaps that domestic providers avoid.
- **Pricing structure** - E2E's H100 at ₹249/hour is 50–60% cheaper than hyperscaler India pricing. Even accounting for hyperscaler price cuts (AWS reduced H100 pricing by ~44% in June 2025), the structural cost advantage of a lean Indian operator without hyperscaler overhead persists.
- **Government procurement preference** - IndiaAI Mission's empanelment process favors MeitY-registered domestic providers. E2E's L1 bidder status and fastest-deployment record create institutional trust that hyperscalers struggle to replicate in government contexts

However, hyperscaler risks are real. Microsoft's \$17.5 billion India commitment, Google's \$15 billion, and Amazon's \$35 billion will dramatically increase Indian GPU capacity. If hyperscalers compete aggressively on price in India's enterprise segment, E2E's margin advantage narrows.

The L&T-Vyoma MOU: Asset-Light Scaling as a Business Model Shift

The February 20, 2026 MOU with L&T-Vyoma is arguably the most underappreciated development in E2E's recent trajectory. To understand why, consider the business model implication:

Until now, E2E's growth has been capital-constrained every GPU added to the fleet required E2E's own balance sheet (equity raises or debt). The company raised ₹1,485 crore in preferential issues and ₹107 crore via QIP specifically to fund GPU procurement. This created a cycle of: raise capital → buy GPUs → deploy → generate revenue → raise more capital.

The L&T-Vyoma MOU introduces a fundamentally different model: L&T-Vyoma procures, owns, and deploys the GPU infrastructure using L&T's vastly larger balance sheet (L&T's consolidated net worth exceeds ₹90,000 crore). E2E integrates this capacity into its TIR platform and commercializes it earning platform fees, management fees, or revenue-share without bearing the full capex burden.

If this model scales as intended, it transforms E2E from a capital-intensive GPU infrastructure company into a platform company that monetizes software and orchestration capabilities on top of third-party hardware. The financial implications are significant: higher ROCE (less capital employed for same revenue), lower equity dilution risk, and potentially faster fleet scaling than E2E could achieve independently.

The key risk: this remains an MOU, not a definitive agreement. The commercial terms particularly the revenue-share or fee structure have not been disclosed. If the economics tilt too heavily toward L&T (the infrastructure owner), E2E's margin on these workloads could be thinner than on self-owned GPUs. Investors should track the signing of definitive agreements and the first commercial deployments under this framework as critical milestones.

How E2E compares to global and Indian GPU cloud peers

The global neocloud landscape

Metric	CoreWeave	Nebius	Lambda	Crusoe	E2E Networks
Revenue	\$5.13B	~\$500M	~\$500M	~\$1B est.	~₹183 Cr (\$21 Mn TTM)
2026 Guidance	\$12–13B	\$3.0–3.4B	\$1B+ est.	~\$2B est.	Growing
Valuation	\$37B	\$23.8B	\$4–5B	\$10B	₹4,800 Cr(\$570 Mn)
GPU Fleet	250,000+	1,00,000 planned	(30-50k Est.)	Scaling(20-30k est.)	~3,900 (→5,000)
EBITDA Margin	57% (Q4)	Negative company level	~50%	N/A	56.6% (Q3 FY26)

Metric	CoreWeave	Nebius	Lambda	Crusoe	E2E Networks
Net Income	Loss	Loss	Loss	Loss	Loss ₹5.7 Cr (\$.64Mn Q3)
Total Debt	\$29.82B	\$4.89 B	~\$775M	Significant	₹154 Cr (\$17.5 Mn)
GPU	6 years	4→5 years	5 years	N/A	6 years

E2E's dramatically lower debt (₹154 Cr (\$17.5 Mn) vs CoreWeave's \$29.82B), near-breakeven operations (vs CoreWeave's \$1.2B net loss), and India-specific moats (data sovereignty, government empanelment, INR billing) justify a distinct valuation framework.

Indian competitive landscape

Yotta Data Services is the dominant Indian GPU infrastructure player, operating 16,000+ H100 GPUs and deploying 20,736 Blackwell Ultra (B300) GPUs in a \$2 billion investment expected live by August 2026. Yotta has committed \$1 billion+ to NVIDIA's DGX Cloud program and claims 60–70% of India's GPU capacity. An IPO is expected within 12 months. E2E's advantages versus Yotta are its public listing (transparency, liquidity), self-serve developer-friendly model with transparent pricing, and lean cost structure. Yotta's advantages are overwhelming in scale, deeper NVIDIA relationship (Exemplar Cloud Partner), and Hiranandani Group's capital depth.

The India AI Impact Summit blog establishes a clear hierarchy among Indian GPU cloud providers. NVIDIA names exactly three cloud partners for India's AI infrastructure: Yotta (scale leader with 20,000+ Blackwell Ultra GPUs), L&T (infrastructure partner building gigawatt-scale AI factories), and E2E Networks (platform partner with Blackwell cluster on TIR platform at L&T Vyoma Chennai).

Notably absent from this list: Jio Cloud Platform, CtrlS, NxtGen, and all hyperscaler India regions. This is not accidental — NVIDIA's blog reflects commercial partnerships with active deployment timelines, not aspirational announcements. E2E's inclusion alongside a \$20B+ conglomerate (L&T) and a Hiranandani-backed infrastructure giant (Yotta) validates its technical credibility at a level that its ₹4,800 crore market cap might not immediately suggest.

The blog also names Netweb Technologies as a hardware manufacturing partner (GB200 NVL4 systems under "Make in India"), positioning Netweb as a complementary rather than competing player in the value chain — Netweb builds the servers, E2E runs them as cloud infrastructure.

Jio Cloud Platform represents the largest long-term threat. Reliance's partnership with NVIDIA for end-to-end AI infrastructure, plans for a 1 GW data center in Gujarat, and track record of aggressive "Jio-style" pricing disruption could compress margins across the Indian GPU cloud market. However, Jio had not deployed its empanelled GPUs as of mid-2025, and no commercial GPU cloud product is at scale making this a medium-to-long-term rather than immediate threat.

CtrlS and NxtGen are IndiaAI Mission participants but lack E2E's public market presence and AI-first software stack.

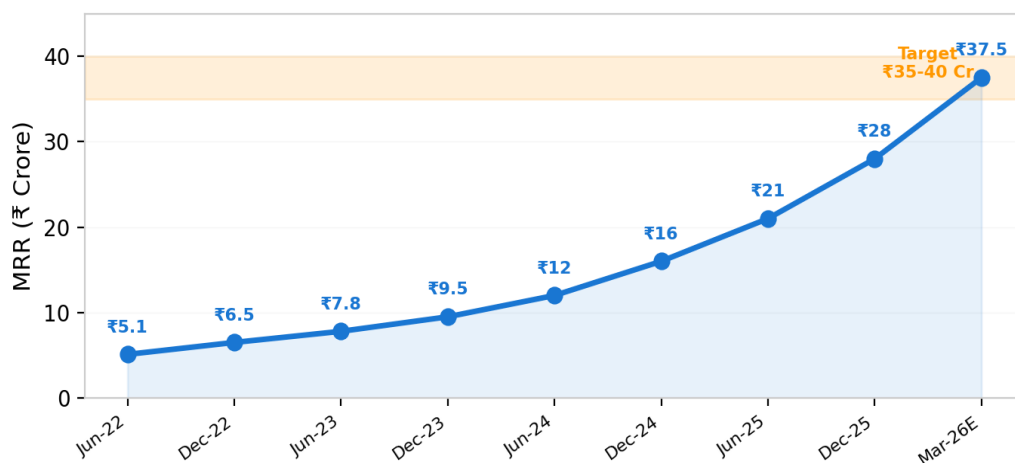
Key Trackables

Key metrics to track quarter by quarter

GPU utilization rates stood at 60–65% in December 2025, with management targeting 70% by Q4 FY26 and 80–90% by FY27. Every 10 percentage-point improvement in utilization on the existing fleet translates to approximately ₹3–4 crore in incremental monthly revenue.

MRR progression toward ₹35–40 crore by March 2026 is the single most important near-term KPI. Achievement would validate the Blackwell deployment timeline and IndiaAI Mission revenue ramp. Management indicated in January 2026 that 70–80% of the March target had been achieved, suggesting ₹28 crore was a floor.

Monthly Recurring Revenue (MRR) Progression



IndiaAI Mission revenue execution the ₹265 crore in confirmed orders requires sustained deployment and utilization. The Gnani.ai contract (₹177 crore for 1.30 crore GPU hours over 360 days) became fully operational by mid-January 2026. Payment cycles have moved to monthly billing, improving working capital dynamics.

Blackwell go-live and revenue contribution the 1,024 B200 cluster must be operational before end-Q4 FY26. Revenue contribution starts Q1 FY27, with full annualized run-rate of ~₹250 crore expected by Q2–Q3 FY27.

L&T-Vyoma MOU progression conversion from MOU to definitive agreements, first commercial GPU deployments under the platform model, and disclosed revenue-share/fee economics. This is the single most important structural development to track, as successful execution could shift E2E's business model from capex-heavy GPU owner to asset-light platform operator.

Risks that could derail the thesis

GPU depreciation and technology obsolescence

The 6-year depreciation schedule assumes sustained economic utility that may not materialize. H100 rental rates have declined ~75% from peak (\$8–10/hour in 2023 to \$1.50– 3.00 in 2026). If B200 rates follow a similar compression as Rubin arrives in H2 2026, the economic useful life could prove shorter than the accounting life. Michael Burry's estimate of \$176 billion in understated hyperscaler depreciation underscores the industry-wide risk. E2E's relatively small asset base means any write-down would disproportionately impact financials.

Hyperscaler and Yotta scale advantage

Yotta's 16,000+ H100 fleet and planned 20,736 Blackwell Ultra deployment dwarfs E2E's ~5,000 GPU target. With \$67.5 billion in hyperscaler India commitments and Jio Cloud's potential entry, E2E faces the risk of being squeezed between well-capitalized competitors.

The IndiaAI Mission's L1 pricing at ₹115.85/hour is already 40% below market rates margins on government business are structurally lower than commercial pricing.

Capital intensity and funding risk

E2E requires continuous capital infusion to maintain competitive GPU fleet scale. The company has raised ~₹1,600 crore in equity and has ₹240 crore in undeployed preferential issue funds. Additional QIP capacity exists (₹893 crore of the ₹1,000 crore authorization unraised). Each GPU generation requires fresh capex Rubin procurement in H2 2026/early 2027 will demand another ₹500–1,000 crore.

IndiaAI Mission Execution & Budget Risk

Risk: IndiaAI Mission budget utilization stood at only 40% in FY26 (₹800 crore utilized of ₹2,000 crore allocation), and the FY27 budget was reduced to ₹1,000 crore half the prior year. Government procurement cycles are inherently lumpy, and if payment timelines stretch beyond monthly billing terms or if the mission's scope contracts, E2E's ₹265 crore in confirmed orders could face revenue recognition delays. Additionally, the mission's aggressive L1 pricing (₹115.85/hour average — 89% discount to listed rates) means government revenue carries structurally lower margins than commercial business.

Mitigant: The shift from quarterly to monthly billing (confirmed on Q3 FY26 concall) materially improves working capital visibility. The India AI Impact Summit in February 2026 saw fresh commitments of 20,000+ additional GPUs,

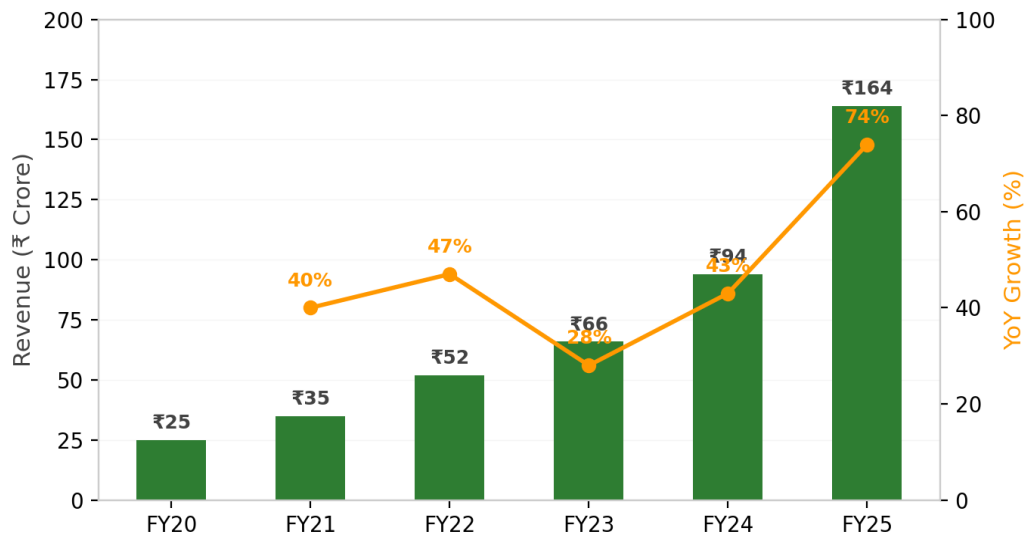
suggesting political will remains strong despite budget adjustments. E2E's Q3 FY26 revenue was largely non-IndiaAI, demonstrating the company's ability to capture commercial demand independently reducing single-program dependency. The upcoming 12,000–15,000 Blackwell GPU procurement round represents incremental upside beyond existing orders.

Financials

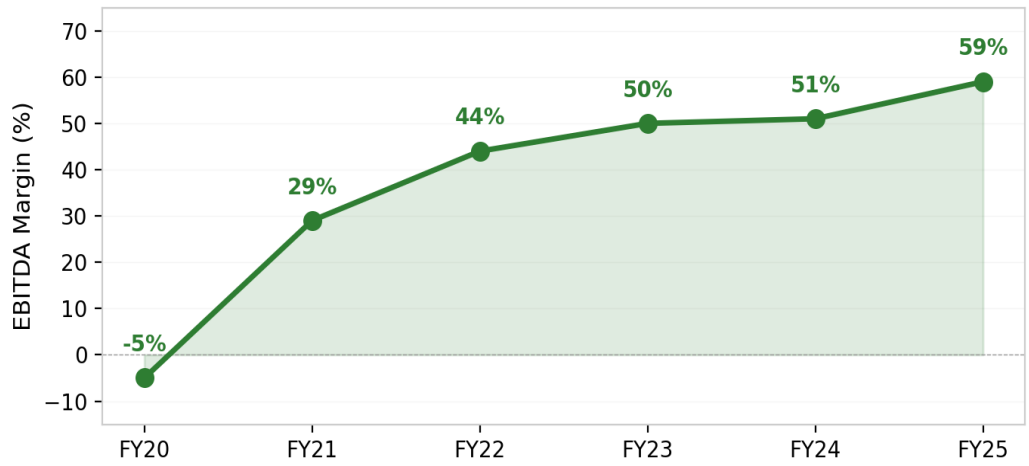
Income Statement

Income Statement (Cr)	FY23	FY24	FY25	TTM
Revenue from Operations	66.20	94.46	163.96	183.41
Service Cost	16.94	20.16	30.48	53.56
COGS	16.94	20.16	30.48	53.56
Gross Profit	49.26	74.30	133.48	129.85
Gross Margins	74.4%	78.7%	81.4%	70.8%
Employee Benefit Expenses	12.11	17.58	23.55	31.49
Other Expenses	4.09	8.78	13.27	16.86
Total Operating Expenses	33.14	46.52	67.30	101.91
Operating Profit (EBITDA)	33.06	47.94	96.66	81.50
Operating Margins	49.9%	50.8%	59.0%	44.4%
Finance cost	0.47	3.61	13.22	10.76
Depreciation	20.11	15.75	60.08	136.87
Total Expenses	53.73	65.88	140.60	249.54
PBT (excl. Other Income & Exceptional Items)	12.48	28.58	23.36	-66.13
Other Income	0.76	1.63	39.43	53.96
Exceptional Items	0.00	0.00	0.00	0.00
PBT	13.24	30.21	62.79	-12.17
Tax	3.33	8.35	15.30	-7.76
PAT	9.91	21.87	47.49	-8.39
PAT Margins	15.0%	23.2%	29.0%	-4.6%

E2E Networks - Revenue Trajectory



EBITDA Margin Expansion



Balance Sheet

Balance Sheet (Cr)	FY23	FY24	FY25	H1 FY26
Shareholders Funds	14.48	14.48	19.97	20.12
other equity	34.92	56.38	1,572.81	1,555.33
Total (1)	49.40	70.85	1,592.77	1,575.45
Avg				
Non Current Liabilities (2)				
Long Term Borrowings	0.21	88.50	6.10	103.40
Lease liabilities	3.06	29.40	43.84	35.44
Long term provisions	0.77	0.90	1.14	1.29
Deferred tax liabilities	0.89	8.41	23.40	16.74
Total Current Liabilities (3)	4.93	127.20	74.47	156.87
Short Term Borrowings	0.39	14.58	5.31	4.05
Lease liabilities	2.23	11.55	17.72	17.26
Trade Payables	2.50	6.04	7.15	10.62
Other Financial liabilities	5.66	16.18	875.74	6.77
Provisions	0.04	0.04	0.05	0.04
Current tax liabilities	0.00	0.00	0.00	0.00
deferred revenue	0.00	0.00	0.00	0.00
Other Current Liabilities	3.26	8.79	7.53	7.53
Total borrowings	0.60	103.08	11.41	107.45
TOTAL LIABILITIES (1+2+3)	68.41	255.24	2,580.74	1,778.58
Assets				
Non Current Assets (1)				
PPE	22.94	155.78	310.91	944.44
Right of use assets	5.55	42.38	63.49	53.21
Intangible Assets	13.51	12.22	14.93	13.48
Intangible under development	0.00	0.00	636.18	0.00
Other Intangible Assets	0.00	0.00	0.00	0.00
Defferd tax assets	0.00	0.00	0.00	0.00
Financial assets				
Other financial assets	0.00	3.87	1.65	1.94
Non Current tax assets	0.27	2.44	4.97	7.17
Other Non-Current Assets	0.00	0.00	0.00	0.00
Cash	16.31	7.77	463.68	350.50
Loans	0.00	0.00	0.00	0.00

bank balance	5.46	1.24	893.26	69.96
Contract Assets	0.00	0.00	0.00	0.00
investments	0.00	0.00	0.00	0.00
Trade Receivables	0.60	2.56	9.75	6.28
Other financial assets	2.78	3.66	3.49	3.97
Current tax assest	0.00	0.00	0.00	0.00
Other Current Assets	0.97	23.31	178.43	327.63
Total Current Assets (2)	26.13	38.54	1,548.61	758.34
TOTAL ASSETS (1+2)	68.41	255.24	2,580.74	1,778.58

Cash Flow

Cash Flow Statement (Cr)	FY23	FY24	FY25	H1 FY26
OPERATING ACTIVITIES				
PBT	13.24	30.21	62.79	-22.25
(+) Depreciation - PPE	18.51	10.28	41.61	59.37
(+) Depreciation - ROU	1.55	5.31	17.60	10.28
(+) Amortization	0.05	0.16	0.86	0.58
(-) Interest Income	-0.72	-1.48	-32.47	-20.10
(+) ESOP Expense	0.59	1.89	2.34	0.62
(+) Other Adjustments	0.46	3.98	13.18	3.77
Operating Profit before WC	33.67	50.36	105.93	32.27
Working Capital Changes	3.65	-5.33	-14.89	3.00
Cash Generated from Ops	37.32	45.03	91.04	35.27
(-) Taxes Paid	-1.79	-2.17	-2.57	-2.20
CFO	35.53	42.86	88.47	33.07
INVESTING ACTIVITIES				
(-) Capex	-18.94	-145.53	-125.93	-936.89
(+) Proceeds - Asset Sales	0.03	0.04	0.07	0.00
(+/-) Bank Deposits	-2.29	0.36	-889.40	823.74
(+) Interest Received	0.55	1.37	32.47	20.10
Other Investing Activities	0.00	0.00	0.00	-136.90
CFI	-20.65	-143.77	-982.80	-229.94
FINANCING ACTIVITIES				
(+) Equity Raised	0.00	0.00	1,472.83	1.19

(+) ESOP Proceeds	0.17	0.16	-0.12	-0.86
(+/-) Net Borrowings	-0.34	102.28	-91.47	96.11
(-) Lease Payments	-2.13	-7.68	-23.71	-12.03
(-) Interest Paid	-0.10	-2.40	-7.29	-0.73
Other Financing	-0.01	0.00	0.00	0.00
CFF	-2.41	92.36	1,350.24	83.68
NET CASH SUMMARY				
Net Change in Cash	12.47	-8.54	455.91	-113.18
Opening Cash	3.84	16.31	7.77	463.68
Closing Cash	16.31	7.77	463.68	350.50

Disclaimer:

Investments in the securities market are subject to market risks. Read all the related documents carefully before investing.

Registration granted by SEBI, membership of BSE, and certification from NISM in no way guarantee the performance of the intermediary or provide any assurance of returns to investors.

The securities quoted are for illustration and educational purpose only and are not recommendatory.

The Information is based on publicly available information and sources considered reliable; however, such information has not been independently verified. So, Niveshaay, its employees or its associates, shall not be liable for any direct or indirect loss arising from the use of or reliance on this material, and any such liability is explicitly disclaimed.

No representation is made as to the accuracy, completeness, reasonableness, or sufficiency of the information contained in this material, or, in the case of projections, as to their attainability or the assumptions on which they are based. Viewer (Spectators) are expected to conduct their own independent due diligence.

Niveshaay, its partners, employees, affiliates, and associated AIF schemes may have existing investments or exposures in securities mentioned herein and may also offer other advisory, consultancy, or fund management services that could potentially create conflicts of interest.