



RACE WEEKLY INTELLIGENCE BRIEF

RETINA

AI · CYBERSECURITY · CLOUD · LEADERSHIP

THIS WEEK'S COVER:
Frugal AI for India:
When the Constraint Becomes the Strategy

ALSO INSIDE
The World in 5 · Three Forces ·
The Frugal AI Playbook · Metric of the Week · The Compute
Signal · Fun Fact · CANDID · Book Shelf ·
Research Radar · RACE Pulse

FROM THE EDITOR'S DESK

For nine issues, RETINA tracked who owns the AI stack — the agents, the supply chain, the power bill — and, last week, the proof gap. This week we invert the question. The world's AI story in 2026 is one of abundance: \$2.52 trillion in projected spend, trillion-parameter models, gigawatt data centres. India's opportunity is the opposite. It is not to win the abundance race. It is to make abundance unnecessary.

Frugal AI - smaller models, cheaper compute, offline hardware, vernacular-first design - is not a consolation prize for a country that cannot afford frontier scale. It is emerging as the more defensible strategy, and the numbers this week explain why. When inference costs collapse a hundredfold, when a 30-billion-parameter model beats a frontier one on Indian-language tasks, and when the government is renting H100 and 200s at a fraction of cloud rates, the constraint stops being a constraint. It becomes a design advantage.

This issue maps the hardware and operating-cost economics honestly, names the three forces reshaping the build-versus-buy decision, and gives leadership teams a five-move playbook for building AI that survives a CFO's scrutiny - in rupees, not slideware.

**SIGNAL**

THE AI PRODUCTIVITY PARADOX – BY THE NUMBERS

₹65

per GPU-hour — IndiaAI Mission's subsidised rate for high-end GPUs, roughly one-third of commercial cloud (₹300–600 on AWS/Azure)

600×

collapse in LLM inference price in under six years: \$60 per million tokens (GPT-3, Nov 2021) to about \$0.06 by March 2026

100×

reduction in cost per million queries when routine work is routed to a small model instead of a frontier LLM

50–70%

share of an Indian data-centre's operating cost that is electricity alone — the decisive frugal-AI variable

<https://staging2.pmgatishakti.gov.in/IndiaAICompute/pricelist>

**THE WORLD IN 5**

NEW · GLOBAL HEADLINES A CXO SHOULD BE ABLE TO REFERENCE

Five signals from the past fortnight, worth a note at your next leadership table.

- 1. COMPUTE:** The IndiaAI Mission has onboarded 38,000+ GPUs - nearly quadruple its original 10,000 target - offered to startups, researchers and academia at subsidised rates as low as 65 per GPU-hour, against 300–600 on commercial cloud. A further 20,000 units are in the pipeline under “AI Mission 2.0.”
- 2. MODELS:** Twelve Indian organisations are now building sovereign foundation models under the mission - Sarvam, BharatGen, Gnani, Soket, Avataar and others - spanning 2B to 1T parameters, most on efficiency-first Mixture-of-Experts designs. Sarvam and BharatGen weights are open on Hugging Face and AIKosh.
- 3. ECONOMICS:** Gartner projects that by 2030, inference on a trillion-parameter model will cost providers over 90% less than in 2025 - up to 100× more cost-efficient than 2022-era models - driven by inference-specialised silicon and edge deployment.
- 4. ENERGY:** India's data centres draw about 1.2 GW today; that could rise to 13.5 GW as AI scales. AI racks now pull 30–100 kW each versus 7–10 kW for traditional racks, and power is 50–70% of operating cost - which is why energy, not the GPU invoice, decides the frugal winner.
- 5. FUNDING:** Sarvam AI became a unicorn in June 2026 on a \$234M Series B - built with a core LLM team of about 15 people and 4,096 rented H100s. Live proof that frontier-class models no longer require a billion-dollar war chest.

Source: <https://www.eetimes.com/> · <https://www.medianama.com/> · <https://www.gartner.com/en> · <https://www.etvbharat.com/> · https://en.wikipedia.org/wiki/Sarvam_AI



HYPERFOCAL

WHAT BROKE THIS WEEK

Frugal Stopped Being a Compromise and Became a Thesis

At the 5th India Global Innovation Connect (IGIC 2026), the dominant message from investors and policymakers was blunt: India should skip the foundation-model sprint and double down on frugal innovation and vertical AI. Capria Ventures' Will Poole argued the opportunity is not to chase foundation models but to build vertically focused solutions that solve real problems at scale and can be exported globally. The summit put the base to do it on the table; over \$20 billion committed to AI, roughly 2,200 Global Capability Centres, and 2.5 million professionals, and returned repeatedly to Claude Smadja's framing that India's strength "has always been its ability to achieve more with less."

The signal for enterprises: the question is no longer "which frontier model do we license?" It is "what is the smallest, cheapest system that clears the task — and can we own it?"

RETINA TAKE

The abundance narrative was written for balance sheets that can absorb a \$100-million training run and a gigawatt power contract. India cannot — and, increasingly, does not need to. The entire cost curve of AI is bending toward the frugal player. The constraint is becoming the moat.



PANORAMA

THREE FORCES — AND THE MOVE EACH ONE DEMANDS

FORCE 1

The Hardware Bill Is the Real Barrier - and India Just Lowered It

The binding constraint on Indian AI has never been talent; it has been the price of compute. A single H100 hour on the commercial cloud runs 300–600. The IndiaAI Mission's subsidised commons - 38,000+ GPUs at 65–150 per hour, with eligible projects getting a further 40% off - resets that maths. A training run that costs 1,00,000 on AWS can land near 60,000 through the mission.

The Move: Before signing another cloud GPU contract, check eligibility for IndiaAI compute. For fine-tuning and inference at scale, not frontier training, the cost delta is the difference between "we can try" and "we can't compete."

FORCE 2

Operating Expense, Not CapEx, Decides the Frugal Winner

The GPU sticker price is the headline; the power bill is the story. Electricity is 50–70% of an Indian data-centre's operating cost, and AI racks draw 30–100 kW versus 7–10 kW for traditional ones. Industrial tariffs range from 6.90/kWh in Karnataka to 11.5/kWh in Maharashtra; captive renewables cut that to 3–4/kWh. The frugal advantage compounds at the meter, not the invoice.

The Move: Model total cost of ownership; energy, cooling, opex, before compute, not after. A workload that runs on a small model on-prem or at the edge can be cheaper by an order of magnitude once the power line is honest.

**PANORAMA**

THREE FORCES — AND THE MOVE EACH ONE DEMANDS

FORCE

3**The Model Doesn't Need to Be Big — It Needs to Be Right**

LLM inference prices fell up to 600× in six years; small language models cut cost per million queries by over 100× versus frontier models. The winning production pattern in 2026 routes roughly 80% of predictable queries to a small model and escalates only the hard 20% — cutting inference spend by about 70%. Sarvam-105B, a Mixture-of-Experts model activating around 10B of 105B parameters per token, outperforms Gemini 2.5 Flash on Indian-language tasks. Bigger is no longer the safe default.

The Move: Architect for routing, not for a single flagship model. Match each task to the smallest model that clears the bar and reserve frontier calls for genuine complexity. For Indian-language workloads, benchmark Sarvam or BharatGen before defaulting to GPT or Gemini.

**PERSPECTIVES**

DEEP TECH — ANALYSIS BY THE RACE TEAM

Jugaad Is Not a Shortcut - It Is a Cost Model

By **Dr J. B Simha**, Chief Mentor - AI, REVA Academy for Corporate Excellence, REVA Univeristy

When I ask executive cohorts at RACE how they chose their AI model, the answer is almost always the same: they picked the most capable one on the market, never the most cost-effective one that would do the job. That instinct, to reflexively reach for the frontier, is exactly the approach that leads to bloated tech stacks, eroding margins, and permanent vendor lock-in. It is a strategic blind spot that enterprises cannot afford, and increasingly do not need.

The word jugaad is often misunderstood in technology circles as a mere shortcut or a temporary compromise. It is neither. At its core, jugaad is a sophisticated cost model. It is disciplined engineering under constraints: delivering the exact required outcome with the absolute minimum expenditure of computation, energy, infrastructure, and capital, without ever compromising on quality, reliability, or security.

India has rewritten the rules of innovation using this exact framework before. The Jaipur Foot, Aravind Eye Care, and a Mars orbiter

delivered for less than the budget of a Hollywood film were not successful despite our constraints; they were successful because the architecture was designed around those constraints from day one.

The global AI economy has finally caught up to this philosophy. Inference costs have plummeted over the past few years. Smaller, domain-specific models increasingly match or outperform massive frontier models on narrowly defined enterprise tasks. The availability of high-performance open-weight models, efficient hardware, and advanced optimisation techniques have shifted the competitive optimum away from "biggest" and toward precision engineering. Using a frontier model for a routine enterprise workflow isn't superior tech leadership; it is poor asset allocation. The discipline that matters now is computational precision—deploying the leanest model required to solve the problem reliably.

This is where FrugalAI evolves from a cost-saving exercise into a core strategic



PERSPECTIVES

DEEP TECH — ANALYSIS BY THE RACE TEAM

framework. It optimises the entire AI lifecycle, from model selection and deployment to governance, maintenance, and future upgrades. Every unnecessary parameter, redundant API call, or unoptimized GPU hour represents a recurring financial leak and operational vulnerability. Eliminating them is not an act of austerity; it is elite engineering.

For leadership teams, this shift introduces a critical geopolitical and operational advantage. A frontier-first strategy forces an enterprise to import its dependencies: foreign models, foreign clouds, foreign silicon, and a power bill it cannot control. Conversely, a frugal strategy—anchored by smaller models, open weights, edge deployment, and localised design—keeps the data, the unit economics, and the intellectual property firmly within your control. It enables organizations to build AI that is affordable enough to scale beyond restricted pilots and resilient enough to operate where connectivity and infrastructure are highly variable.

The breast-screening tool running at the point

of care, the edge sprayer that identifies weeds without a cloud round-trip, the analytics engine running securely inside an SME, or the learning application cached offline on a shared television are not scaled-down, inferior versions of Western AI. They are precision systems designed from realities outward.

The question for a leadership team is no longer, "Are we using the most powerful model?" It is, "Are we using the most profitable model for the job, do we own it, and can we sustain it?"

Organizations that internalize this shift will do more than just spend a fraction of what their competitors spend. They will deploy AI faster, adapt to market changes more fluidly, and insulate themselves from vendor dependencies. In the coming decade, structural advantage will belong not to those who build the heaviest AI systems, but to those who master the most efficient ones. FrugalAI is how we transform resource-efficient innovation into global leadership in the age of artificial intelligence.

RACE PERSPECTIVE

The AI advantage in India will not go to the enterprise with the biggest model. It will go to the one that treats compute, power and sovereignty as a single cost equation, and solves it with the smallest system that clears the task.



FUN FACT

NEW · SOMETHING TO REPEAT AT DINNER

In the 1960s, India built its space programme by carrying rocket components to the launch pad on bicycles and bullock carts. Decades later, the same agency put a spacecraft into Mars orbit on its first attempt for about \$74 million - less than the budget of the Hollywood film Gravity. The pattern India is now applying to AI is not new. It has a name, a heritage, and a track record: do more with radically less, and let the constraint do the designing.



Dr. Madhan Kumar Srinivasan, Ph.D.

Founder and CEO, Imaginext

<https://www.linkedin.com/in/drmadhaniimcal/>

Q1

Most enterprises still default to the most capable model available. In your experience, what does an honest total-cost-of-ownership picture - compute, energy, integration - actually reveal once you look past the sticker price?

A

The sticker price is the smallest line item. Once you factor in API costs at production scale, retry/token overhead from prompt engineering, integration into existing enterprise stacks, and the latency tax on user experience, frontier models often cost 3-5x what the per-token price suggests.

At Imaginext, building Cortexa (our patent intelligence platform) taught us this directly: running every document-parsing and classification task through a frontier model would have made unit economics unworkable. The real TCO question isn't "which model is smartest"; it's "which (cheapest) model is smartest enough for this specific task, at this volume, at this latency."

Q2

Where have you seen a smaller, on-prem or edge model outperform a frontier one on the metric that mattered, cost, latency, or data sovereignty?

A

We are increasingly seeing this in enterprise AI deployments. In products like Cortexa, we adopt a multi-model orchestration approach rather than relying on a single frontier model. For repetitive enterprise workflows—document classification, procurement automation, compliance validation, internal knowledge search, or structured reasoning, smaller domain-specialised models often deliver comparable business outcomes with significantly lower latency and cost.

On-prem or edge deployments become even more valuable in regulated sectors such as banking, healthcare, manufacturing, and government, where data sovereignty, privacy, and predictable response times are more important than achieving the absolute highest benchmark scores.

The winning strategy is rarely "largest model everywhere"; it is "right-sized intelligence deployed where it creates maximum value."



Q2

If a board asked for the one question they should put to their team every quarter about AI cost discipline, what would it be?

A

If nobody on the team can answer that with data, you don't have an AI cost strategy; you have a subscription.

Final thoughts about Frugal AI

According to me, Frugal AI is not about using cheaper models. It is about engineering intelligent systems that know when to use a large model, when to use a small model, and when not to invoke AI at all. Sustainable AI will be defined by intelligent orchestration, not computational extravagance.

Note: CANDID features direct contributions from CXO leaders. To contribute, contact: retina@reva.edu.in



THE FRUGAL AI PLAYBOOK

TURNING THE PROOF
GAP INTO AN ADVANTAGE

Five moves, ordered from fastest to most strategic, that any leadership team can act on.

Right-size the model to the task

MOVE 1

Inventory your AI workloads and ask, for each, the smallest model that clears the bar. Route roughly 80% of predictable queries to a small or on-prem model and reserve frontier APIs for the genuinely hard 20%. Teams doing this report about 70% lower inference spend.

The Move: Stop paying frontier prices for routine work. A router plus a small model is the highest-ROI change most enterprises can make this quarter.

Cost the power, not just the chip

MOVE 2

Build total cost of ownership around energy, cooling and opex, 50–70% of the real bill, before compute. Where workloads are steady, on-prem or edge with captive renewables (3–4/kWh) can beat cloud by an order of magnitude.

The Move: Put the electricity line at the top of the AI business case, not in a footnote. The frugal winner is decided at the meter.

Use the sovereign compute commons

MOVE 3

Check IndiaAI Mission eligibility before your next cloud contract. Subsidised GPUs at 65–150 per hour, with up to 40% additional subsidy for national-priority projects, change the build-versus-buy maths for fine-tuning and scaled inference.

The Move: Treat subsidised compute as a real procurement option. The approval friction is worth it on any run longer than a pilot.



THE FRUGAL AI PLAYBOOK

NEW · TURNING THE PROOF GAP INTO AN ADVANTAGE

Design vernacular- and offline-first

MOVE 4

India's edge is linguistic and rural. Build for 22 languages, low bandwidth and intermittent connectivity from the start. Benchmark Sarvam, BharatGen and Bhashini for Indian-language tasks before defaulting to a global model.

The Move: Make Indian-language and offline capability a design requirement, not a later port. It is where the domestic advantage actually lives.

Deploy at the edge where the decision is made

MOVE 5

For diagnostics, agriculture, inspection and safety, on-device inference removes the cloud round-trip, the bandwidth cost and the data-sovereignty risk in one move. Niramai's radiation-free breast screening and on-device crop-spraying vision are already doing this in low-connectivity settings.

The Move: Treat the edge as the default and the cloud as the exception for real-time, sensor-driven work. It is cheaper, faster and more private at once.

Source: <https://techdodo.in/> · <https://www.imarcengineering.com/> · <https://www.orfonline.org/> · <https://agrospectrumindia.com/>



METRIC OF THE WEEK

ONE NUMBER,
PROPERLY UNDERSTOOD

₹65 per GPU-hour

What it is: The IndiaAI Mission's subsidised rate for high-end GPU access — against roughly 300–600 per hour on AWS or Azure for comparable H100 compute. With the mission's up-to-40% subsidy for national-priority projects, effective rates fall under 100.

Why it matters now: Compute price has been the single largest barrier to Indian AI. At one-third of cloud cost, it stops being the reason a project dies in committee. The catch

is friction — access is gated by an application and approval process, not self-serve provisioning — but for training runs and scaled inference, the saving is real and large.

How leaders use it: Treat subsidised compute as a line item in every AI business case, not an afterthought. The gap between 65 and 600 is the gap between a pilot you can afford to run and one you cannot.

⚡ THE COMPUTE SIGNAL

INDIA'S FRUGAL-AI NUMBERS

The infrastructure data that will shape your AI strategy whether you track it or not.

38,000+ - GPUs onboarded under the IndiaAI Mission, about 4× the original target, with 20,000 more in the pipeline.

12 - Indian organisations funded to build sovereign foundation models, spanning 2B to 1T parameters.

\$234M - Sarvam AI's June 2026 Series B, reaching unicorn status on a roughly 15-person core LLM team.

548M - rural internet users in India (57% of the total), with 41% already using AI-enabled features.

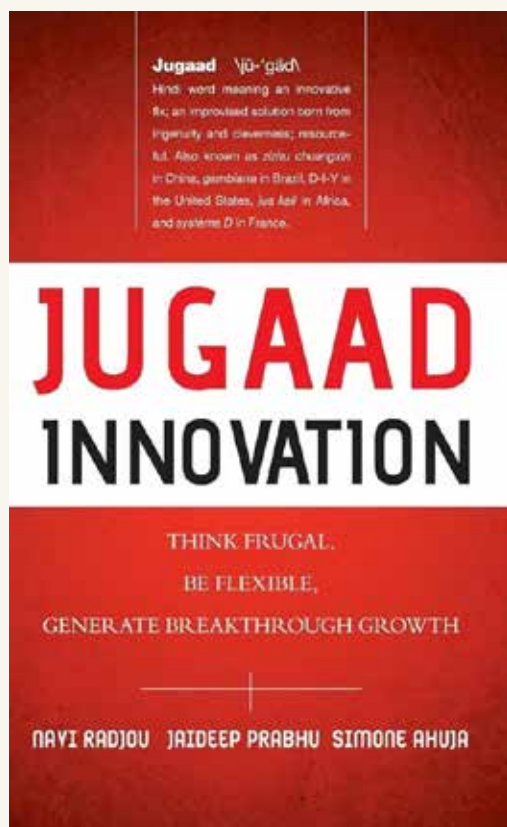
Sub-\$50 - hardware cost of Sarvam's Soliga voice models running fully offline on Raspberry Pi.

Source:

https://en.wikipedia.org/wiki/Sarvam_AI

📚 BOOK SHELF

RETINA'S READ OF THE WEEK



RETINA RECOMMENDS · THIS WEEK'S THEME

Jugaad Innovation: Think Frugal, Be Flexible, Generate Breakthrough Growth

Navi Radjou, Jaideep Prabhu & Simone Ahuja · Jossey-Bass, 2012

The book that named frugal innovation for a global business audience and traced it to its Indian roots. Its central argument is that breakthrough growth in resource-scarce markets comes not from bigger R&D budgets but from a different starting point: improvise under constraint, keep it simple, do more with less, and design for the customer who has been left out.

For any CXO building an AI strategy under real cost and infrastructure limits, this is the operating manual for the mindset the moment now demands. The frugal-AI wave is jugaad applied to compute — and this is the clearest account of why constraint, handled well, becomes an advantage rather than an apology. Read it before you approve your next frontier-model licence.

**EDITOR'S NOTE**

This week's Research Radar features three new IEEE publications spanning **foundational education, enterprise threat detection, and multilingual phishing defence**. Together, they reflect RACE's growing focus on applied AI systems designed for real classrooms, SOC environments, and linguistically diverse digital users. The structure follows RETINA's established research format: a concise technology summary followed by clear enterprise relevance.

EDTECH · COMPUTER VISION · KANNADA AI

Real-Time Detection and Correction of Handwritten Kannada Text Using YOLO*Soujanya Kelaginamani; Paramesh G; Kiran Kumar K V*

A real-time learning support system that identifies and corrects handwritten Kannada text produced by first-grade students. Using a custom dataset of children's handwriting, the researchers trained YOLOv8 and YOLOv9 models, with YOLOv9 achieving **99.4% mAP@50, 89.6% mAP@50–95**, and an inference time of just 3.5 milliseconds.

The system gives students immediate feedback on spelling errors while reducing repetitive correction work for teachers.

Enterprise Relevance: For schools and EdTech providers, this demonstrates how regional-language computer vision can enable personalised foundational learning at scale.

<https://ieeexplore.ieee.org/document/11559340>

CYBERSECURITY · LATERAL MOVEMENT · EXPLAINABLE AI

Machine-Learning Detection of Lateral Movement Using Windows Security Event Logs*M Shashank; Nishanth Krishna; G Paramesh; Kv Kiran Kumar*

This paper presents an operational machine-learning pipeline for detecting lateral movement within Windows-based enterprise environments. It uses **73 per-user behavioural features** covering authentication patterns, privilege changes, temporal behaviour, and movement across hosts.

The ensemble combines Isolation Forest, Random Forest, XGBoost, and LSTM models, producing a weighted risk score mapped to MITRE ATT&CK techniques. It achieved an **ROC-AUC of 0.98 and an F1-score of approximately 96.4%**, supported by a Streamlit interface and SHAP-based explanations.

Enterprise Relevance: For CISOs and SOC teams, the framework provides an explainable method for detecting credential misuse, privilege escalation, and suspicious cross-host activity while reducing low-value alerts.

<https://ieeexplore.ieee.org/document/11564332>



MULTILINGUAL AI · PHISHING DEFENCE · CONTINUAL LEARNING

Multilingual Phishing Detection Using XLM-RoBERTa and Continual Learning*S Sanjay; Nishant Krishna; G Paramesh; KV Kiran Kumar*

A multilingual phishing detection framework designed for **English, Hindi, Kannada, Tamil, Telugu, and Bengali**. The system uses XLM-RoBERTa to classify messages as phishing or safe and incorporates user feedback into a continual-learning pipeline.

Interactions are stored locally in SQLite to support retraining and transparency. The model achieved an **F1-score of 0.975** and an **ROC-AUC of 0.997**, demonstrating strong performance across regional and code-mixed communication.

Enterprise Relevance: For banks, telecom providers, government platforms, and digital-service companies, this offers a practical approach to protecting regional-language users from phishing and social-engineering attacks.

<https://ieeexplore.ieee.org/document/11564455>

Explore the full RACE research archive at <https://raceresearch.reva.edu.in>.

For collaboration, licensing or research-partnership enquiries: race@reva.edu.in

 RACE PULSE

EVENTS · COMMUNITY · WHAT'S COMING

NEW BATCH OPEN ON JULY 25TH 2026

Certified Agentic AI Engineer

Practical expertise in LLMs, AI agents, multi-agent systems, agentic workflows, governance and capstone-based implementation. Real-world use cases, deployable solutions.

Apply now:

<https://race.reva.edu.in/certified-agentic-ai-engineer>

UPCOMING · SEPTEMBER 2026 BATCH
APPLY BY 31 JULY

Applications open for all six RACE Master's & PG Programmes — Business Analytics (#1 India), Artificial Intelligence (#3 India), Cybersecurity (#5 India), and Cloud Architecture & Security. Weekend format. AICTE/UGC approved. Industry mentors from Amazon, Microsoft, AWS, EC-Council, Terralogic, IBM and more.

Apply now: race.reva.edu.in

COMING SOON

THE BIGGEST TECH CONCLAVE IN TOWN – RACE360 and Top Women Tech Awards

<https://race.reva.edu.in/race360>

<https://race.reva.edu.in/top-10-women-tech-leaders-india-awards>

RACE PROGRAMMES

REVA Academy for Corporate Excellence · REVA University · Bengaluru · AICTE / UGC Approved · Weekend Batches · Industry Mentors

Master's & PG Programmes

New Batch: September 2026 · 12–24 months · Weekend · For Working Professionals

#1 IN INDIA · ANALYTICS INDIA MAGAZINE ANALYTICS PG Diploma / M.Sc. in BUSINESS ANALYTICS AWS Academy · Microsoft Azure 12–24 Months	#3 IN INDIA · TECHGIG ARTIFICIAL INTELLIGENCE M.Tech. in ARTIFICIAL INTELLIGENCE AWS Academy · Microsoft Azure 24 Months	#3 IN INDIA · TECHGIG ARTIFICIAL INTELLIGENCE PG Diploma / M.Sc. in ARTIFICIAL INTELLIGENCE AWS Academy · Microsoft Azure 12–24 Months
#5 IN INDIA · ANALYTICS INDIA MAGAZINE CYBERSECURITY M.Tech. in CYBERSECURITY EC-Council · Microsoft Azure · Terralogic Inc. 24 Months	#5 IN INDIA · ANALYTICS INDIA MAGAZINE CYBERSECURITY PG Diploma / M.Sc. in CYBERSECURITY EC-Council · Microsoft Azure · Terralogic Inc. 12–24 Months	CLOUD PG Diploma / M.Sc. in CLOUD ARCHITECTURE & SECURITY AWS Academy · Microsoft Azure 12–24 Months

Short-Term Certification Programmes

100 Hours · Project-Driven · Customisable for Organisations

LEADERSHIP 8–12 weeks TECHNO-FUNCTIONAL LEADERSHIP Executive presence, decision-making, leading tech-driven teams. Built for mid-to-senior managers.	CYBERSECURITY 8–12 weeks CYBERSECURITY & SOC OPERATIONS Threat detection, incident response, SIEM/SOAR. Hands-on in RACE's live Security Operations Centre.
AI ENGINEER 8–12 weeks AI ENGINEERING CERTIFICATION Data pipelines, model training, LLMs, deployment. Real production projects every batch.	CLOUD 8–12 weeks CLOUD ARCHITECTURE & DEVSECOPS Multi-cloud, IaC, secure CI/CD on real AWS/Azure/GCP environments. Not demos.
AGENTIC AI 4–6 weeks AGENTIC AI SYSTEMS DESIGN Multi-agent systems using LangChain, LangGraph, CrewAI. Capstone: deployed autonomous agent.	DATA & AI 12 weeks DATA SCIENCE & BUSINESS ANALYTICS Statistics to ML to BI. Python, SQL, Power BI, Tableau on real business datasets.

Enroll with us · race@reva.edu · race.reva.edu.in

THIS WEEK, ASK YOUR LEADERSHIP TEAM

“What is the smallest, cheapest AI system that would clear this task - and do we own it?”

RETINA PUBLISHES EVERY SUNDAY · FORWARD TO A CXO YOU RESPECT