

VOLUME 6.0 | 2024-25

AINA

AI & ANALYTICS MAGAZINE

AI IN CULTURE HERITAGE

SWARM INTELLIGENCE

AI IN SPACE

NEUROSymbolic AI

DEEPSEEK

INTELLIGENT CHAINS



ANIKET THAKUR
Editor-In-Chief



ANKAMREDDI PRAVEEN
Lead Designer



LENKA HARI SRIJAA
Creative Head



MRINAL JHA
Content Architect



CHIMALA ROHITH
Features Editor

From The Team

Dear Readers,

In an era where artificial intelligence is no longer just a technological advancement but a societal force, it becomes more important than ever to foster spaces where ideas, innovation, and introspection can co-exist. Understanding AI is no longer optional. It is essential. And just as AI reshapes our world, it is equally important to shape how we think about it.

This magazine is one such space, crafted with care, curiosity, and conviction by the students of the PGDBA program. In the AINA 6.0 issue, we explore the extraordinary, from how AI has touched Nobel Prize-winning science, redefined the preservation of cultural heritage, and is now unlocking the secrets of the cosmos. We unpack deep concepts like self-supervised learning and neuro-symbolic AI, while also grounding ourselves in real-world applications, from supply chains to ethical dilemmas around deepfakes. Our interviews with educators, researchers, and AI startup founders bring fresh, authentic voices to the fore, of people not just working in AI, but shaping its future.

We thank the team of AINA 5.0 for their constant guidance, and we believe this will be their true sign-off. We are also grateful to the faculty and administration of all three institutes, IIM Calcutta, IIT Kharagpur, and ISI Kolkata, for their unwavering support. This magazine is also a celebration of learning in motion. It reflects not only where the field is heading, but also where we, as emerging professionals and thinkers, position ourselves in this landscape. The aim is not just to inform, but to spark questions, challenge assumptions, and invite dialogue. We hope you find these pages as enriching to read as they were exciting to create.

Here is to exploring, learning, and questioning. Happy reading!

Warm regards,
The AINA Team

From the Desk of PGDBA Chairperson, IIM Calcutta

Post Graduate Diploma in Business Analytics (PGDBA) is a full-time residential program in business analytics offered jointly by the Indian Institute of Management Calcutta (IIMC), the Indian Statistical Institute Kolkata (ISI-K), and the Indian Institute of Technology Kharagpur (IIT-KGP). Through a highly competitive selection

process and a good mix of experienced professionals and fresh graduates, the PGDBA program has an impressive batch and alum base catering to the increasing global demand for business data scientists.



The PGDBA program's uniqueness stems from its integration of business, statistics, and technology and an extensive internship. Our students study various statistical and machine learning theories for analytics at ISI-K, technological aspects of analytics at IIT-KGP, and applications of analytics in the functional areas of management at IIMC. They are primed for success due to the rigorous and demanding learning at

the three institutes through theory, case studies, labs, and business simulations. Naturally, such rich exposure helps students build critical thinking ability, develop cutting-edge skills, and provide data-driven solutions to management problems using appropriate data science tools.

What truly sets this program apart is its unique six-month industry internship that equips participants with hands-on experience and makes them wholly industry-ready for an upward and onward career in business analytics. Indeed, the fusion of business, technology, and statistics, along with the long industry internship, gives our PGDBA program students a head start in the professional world.

I am confident that the PGDBA program will continue to evolve and expand, meeting the increasing demand for business data professionals worldwide and significantly contributing to the industry.



Sudhir S. Jaiswall
PGDBA Chairperson

From the Desk of PGDBA Coordinator, IIT Kharagpur

The Post Graduate Diploma in Business Analytics (PGDBA) program, a unique association between three of India's premier institutions belonging to the Eastern zone- the Indian Institute of Management Calcutta (IIMC), the Indian Statistical Institute (ISI), and the Indian Institute of Technology (IIT) Kharagpur, provides an amalgamation of business, statistics and technology related concepts in its course curriculum to the students. This blend ensures that the program distinctly stands out among analytics courses offered by several peer institutes across the country.



While the first three semesters of the program help the students acquire theoretical knowledge through rigorous classroom training, the fourth semester equips them with practical experience through an industry internship exposure of six-months duration. This not only prepares the students to handle the real-life problems of the business world through a data-driven approach but also helps them get ready to tackle the challenges that they may face professionally.

The program encourages the bright minds of the nation, selected through a rigorous intake process, to stay abreast with the latest advancements in the field of analytics and take it up as their career. I wish the students good luck in their journey through the course of this program and hope they will be able to meet the increasing demands of data scientists worldwide.

Sangeeta Sahney
PGDBA Coordinator
IIT Kharagpur



From the Desk of PGDBA Coordinator, ISI Kolkata

Post Graduate Diploma in Business Analytics (PGDBA) is a tri-institutional full-time residential two years program conducted by three premier institutes (IIM Calcutta, ISI and IIT Kharagpur) in our country. This program presents a unique opportunity for students to receive world-class education and training in the field of business analytics. From these three institutes, students can get the benefit to learn interdisciplinary subjects, like statistical and machine learning theories for analytics from ISI, some



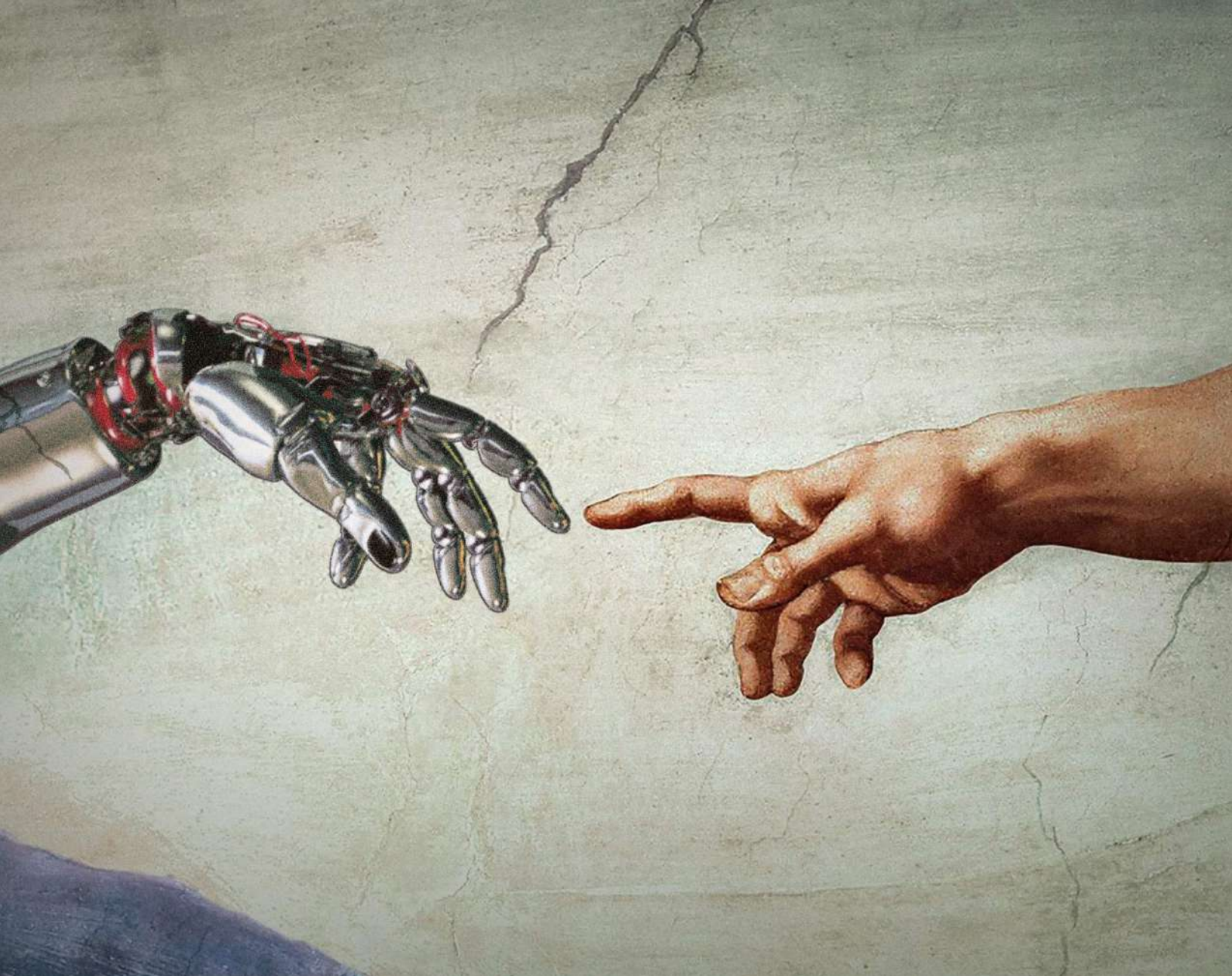
technological training on analytics from IIT Kharagpur, and application of analytics from IIM Calcutta. A major highlight of the PGDBA is the six-month industry internship, where students apply their knowledge to real analytics projects. This hands-on experience not only strengthens their technical and analytical skills but also provides critical industry exposure, preparing them to take on leadership roles in analytics-driven organizations.

At ISI, we take pride in our rich legacy of excellence in education and research. Renowned globally for our contributions to statistics, mathematics, and computer science, we bring deep academic rigor to the PGDBA curriculum. Our faculty members are leading experts in mathematics, statistics, computer science and other interdisciplinary area of researches and students will get a high-quality learning experience from our faculty members. We actively encourage students to engage in physical activities, social initiatives, and interactions with faculty and researchers, fostering both personal and intellectual growth.

We are confident that the PGDBA will be a transformative journey for our students. I look forward to witnessing their growth, contributions, and success in the dynamic world of business analytics.

Dibakar Ghosh
PGDBA Coordinator
ISI Kolkata





Disclaimer & Contact

The contents of this magazine have been prepared with due attention to ensure they are plagiarism-free to the best extent possible. The numbers and data quoted are referenced from publicly available online sources and may vary over time. All external images used have been duly credited in the reference sections, and their license status has been verified before inclusion.

Several visuals and text elements in the magazine have been generated using AI tools, including large language models (LLMs) such as ChatGPT and image generators like Playground AI. As per current international norms, purely AI-generated content typically does not carry traditional copyright protection and is considered free for reuse unless explicitly restricted.

The views and opinions expressed in the articles are solely those of the authors and do not necessarily reflect the views of the PGDBA program or its partner institutes.

For feedback, suggestions, or queries, please contact the PGDBA magazine team at: pgdba.aina@email.iimcal.ac.in

Table of Contents

01 When AI earned Nobel Glory

2024 Nobel Prize recognizes AI breakthrough

07 Interview: Gaurav Aggarwal

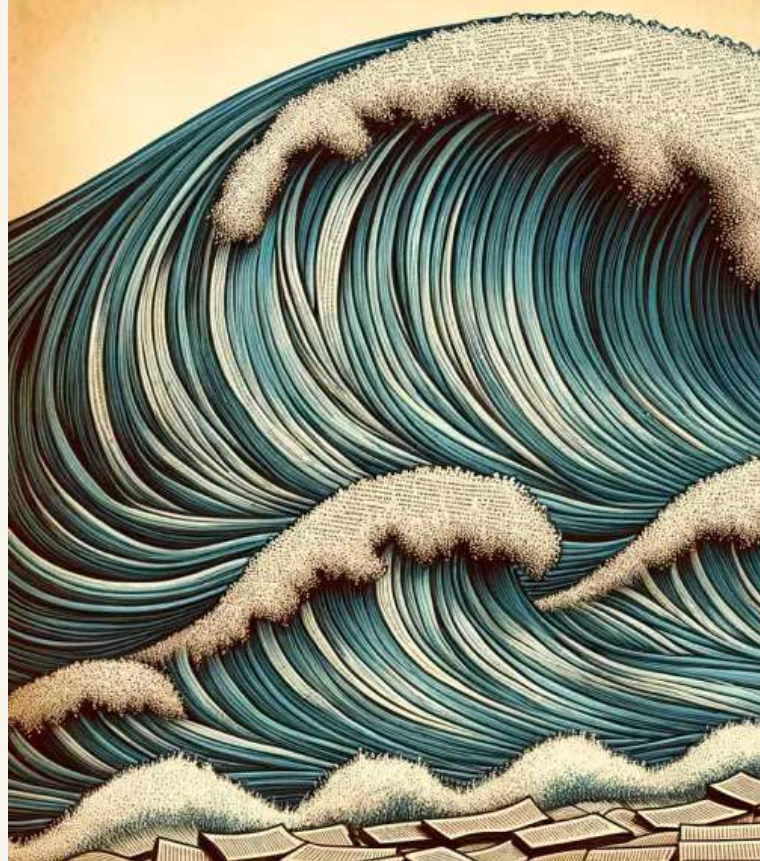
Co-founder, educator, and AI industry veteran

11 Learning the Art of Prompting

Mastering prompts to guide AI better

21 AI in Cultural Heritage

Preserving history using AI



31 Self- Supervised Learning

Training AI without labeled data

43 Neurosymbolic AI

Blending logic and deep learning

51 Book Recommendations

Essential reads in AI and data

54 Crosswords

Fun AI-themed puzzle

55 Swarm Intelligence

Collective behavior powering AI systems



65 AI in Space

AI's role in space exploration

75 Interview : Panocle Labs

Startup journey shaping AI innovations

81 Deepseek - A Deep Dive into Its Cutting-Edge Features

Breakdown of Deepseek's unique capabilities

87 Major Happenings

Latest breakthroughs and AI trends

89 How to make the world a better place for AI

Ethics and responsibility in AI's future

97 Interview: Nitish Singh

Educator shares journey and tips

105 Intelligent Chains

Exploring blockchain + AI integration

113 Trilytics 2024

115 References



When AI Earned Nobel Glory

- Aniket Thakur



On October 8th, 2024, When the Nobel Committee announced that the 2024 Physics Prize would be given to John Hopfield (Princeton University) and Geoffrey Hinton (University of Toronto), the world did a double-take. Social media buzzed with confusion: “Why award a physics prize for... artificial intelligence?” Critics argued AI belonged in computer science, not alongside quantum pioneers and cosmic explorers. But the Nobel Committee saw deeper. This year’s prize wasn’t just about algorithms, it honored a 50-year quest to decode the physics of intelligence itself. From the chaotic magnetism of spin-glasses to the brain’s synaptic whispers, the winners revealed a startling truth: AI’s “magic” is rooted in the laws of nature.

Physics Nobel

The Neural Network Revolution

Our brain learns by rewiring itself. When we memorize a fact or master a skill, the tiny connections between brain cells, known as neurons, either strengthen or weaken through a process called synaptic plasticity. In 1949, psychologist Donald Hebb famously summed this up: *“Neurons that fire together, wire together.”*

But for decades, scientists struggled to turn this idea into machines that learn like humans.

Enter **John Hopfield**, a physicist turned AI pioneer. In 1982, he made a surprising connection: the brain’s learning process resembles how atoms behave in a weird magnetic material called a spin-glass and gave ‘Hopfield network’, an associative memory structure that can store and reconstruct information.

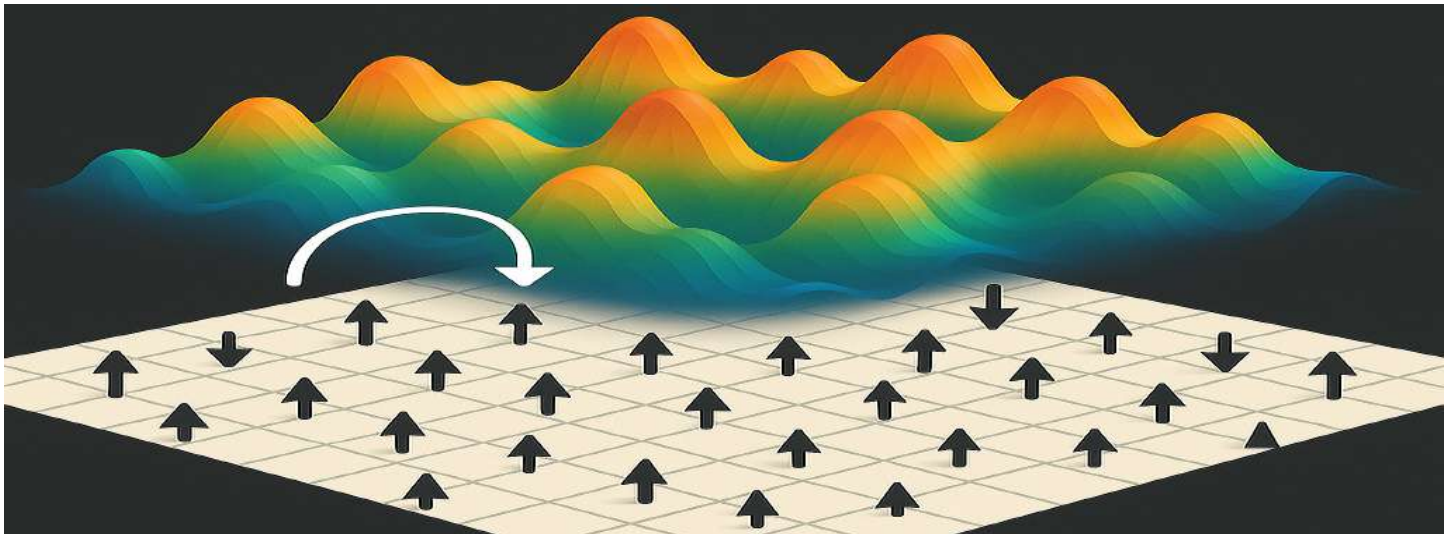
neurons could mimic this? Just as spin-glasses stabilize into memory-like states, a neural network could “remember” patterns (like images or words) by balancing connections between neurons.

Hopfield designed a neural network that:

1. Stores memories like a spin-glass: Inputs (e.g., a photo) become stable states in the network.

2. Recall them by association: Show it a blurry image, and the network “snaps” to the closest memory (like finishing a friend’s sentence).

So, Hopfield proposed a model, **an artificial neural network** built from binary neurons and weighted connections, specifically designed to store and retrieve memories, similar to how a physical system stores energy. His model assigned each memory (e.g., a face or word) to a stable “energy minimum”, a state the network naturally “settles into” when presented with input, much like a ball rolling to the bottom of a valley.



Here’s a simple analogy:

1. Spin-glasses are metals with randomly scattered magnetic atoms. These atoms “fight” over which direction to point (like teammates arguing over strategies).

2. The material settles into many stable “memory” states, each reflecting a compromise between competing forces.

Hopfield realized: What if a network of artificial

Hopfield’s network used an energy function to guide its neurons into configurations matching stored memories. These energy valleys, known as dynamical attractors, formed the basis for associative recall.

This was more than just clever engineering; it was physics reshaping the very concept of intelligence. By framing learning as energy minimization, Hopfield proved machines could replicate a core human trait: **associative memory**.

It was great at recognising patterns but wasn't enough to build AI systems that could predict or generate new information. That's where Geoffrey Hinton came in.

In 1983, **Geoffrey Hinton** tackled a fundamental challenge: How can neural networks learn without explicit instructions? His answer was the Boltzmann machine, a model inspired by physics. Named after **Ludwig Boltzmann**, the 19th-century physicist who linked energy and entropy, this system introduced neurons that updated their states probabilistically.

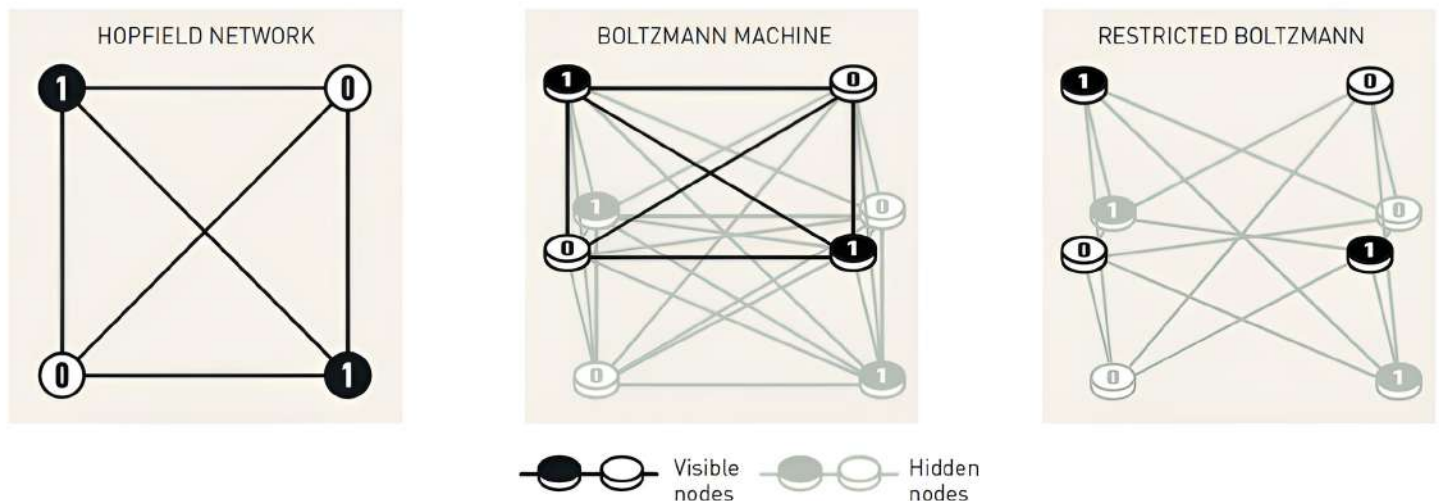
Hinton's genius lay in two radical additions:

1. **“Hidden” Neurons:** Unlike Hopfield's visible-only neurons, Boltzmann machines included layers of hidden units. These layers autonomously detected patterns (like edges in images or grammar in text), acting as the network's intuition.

data resembling its training set by mimicking the Boltzmann distribution, a statistical rule governing how particles settle in heated materials.

These hidden layers act like a “subconscious,” helping computers not just recognise things but also make predictions. For example, instead of just identifying how the cat would look like in a photo, a computer could now guess what the cat might look like in a completely different scene.

Critics called it impractical. Yet Hinton's work laid the groundwork for **unsupervised learning**, where AI learns from raw data, not labeled examples. Hinton's later work on **backpropagation**, which is the algorithm that enables networks to learn from their mistakes, would transform these physics-inspired models into practical tools.



© Johan Jarnestadt / The Royal Swedish Academy of Science

2. **Controlled Chaos:** Instead of rigidly updating neurons, Hinton introduced randomness, like heating a metal. Just as atoms jostle unpredictably under heat, his neurons “flipped” states probabilistically. This stochasticity allows the network to explore solutions in a wild manner before gradually settling into the best one, a process known as simulated annealing.

His model didn't just memorize patterns; it generated them. The machine could create new

“AI didn't “win” a physics prize. Physics helped invent AI”

The Physics Prize honored the why behind thinking machines, but the Chemistry Prize answers the what next. Imagine AI not just finishing your sentences, but folding proteins to cure diseases or crafting molecules to clean oceans. The bridge between these Nobels isn't just silicon, it's the audacity to see science as a conversation, not a competition.



Chemistry Nobel

The Fold Revolution

How AI Cracked Life's Origami Code

When the 2024 Nobel Prize in Chemistry was announced, it wasn't awarded for a new element or reaction—it celebrated a breakthrough in solving life's oldest puzzle: How do proteins fold? For decades, this question haunted biologists. Proteins—the molecular machines driving every heartbeat, immune response, and thought—are born as linear chains of amino acids. Within milliseconds, they twist into intricate 3D shapes that dictate their function. Misfolded proteins cause diseases like Alzheimer's and Parkinson's.

The 2024 Nobel Prize in Chemistry has been awarded to **David Baker (University of Washington)** “for computational protein design” and to **Demis Hassabis and John M. Jumper (Google DeepMind)** “for protein structure prediction”.

Hassabis and Jumper have developed a groundbreaking series of artificial intelligence models to address the decades-long structural

biology problem of how to predict the complex 3D structures of proteins solely from their linear amino acid sequences, a breakthrough that many believe will redefine biology as profoundly as the invention of the microscope did centuries ago.

Baker has dedicated his entire scientific career to designing and constructing novel proteins that are not, and even cannot, be found anywhere in nature.

The “Protein Folding Problem”

In 1972, Nobel laureate **Christian Anfinsen** theorized that a protein’s amino acid sequence alone determines its 3D structure. But solving it manually was impossible. With even small proteins having more possible shapes than atoms in the universe, brute-force computation collapsed under its own weight. For 50 years, scientists relied on painstaking lab experiments, often taking months to map a single protein.

Demis Hassabis & John Jumper

Cracking the Code with AlphaFold

Hassabis, a neuroscientist turned AI entrepreneur, co-founded **DeepMind**, known for training AI to master chess and Go. In 2018, his team entered the CASP competition, a biennial “Olympics” of protein prediction, with AlphaFold, an AI model that achieved 60% accuracy. Impressive, but not groundbreaking.



The breakthrough came in 2020 with AlphaFold2, developed alongside physicist John Jumper. This system combined:

1. Evolutionary Insights:

AlphaFold2 detected co-evolved amino acid pairs that mutate together to preserve structural stability by analyzing millions of protein sequences across species.

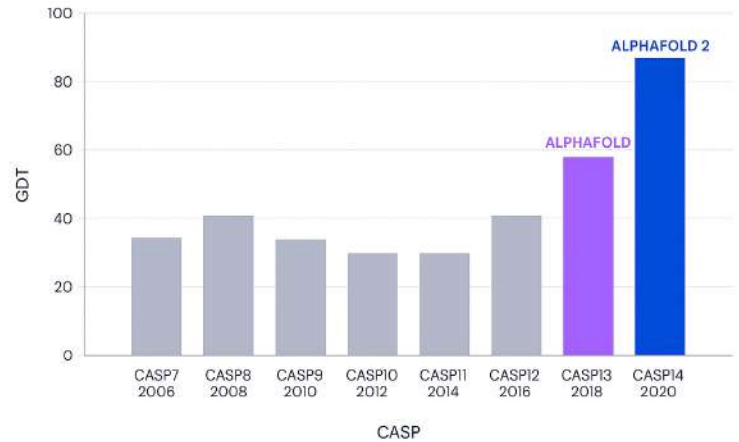
2. Geometric Deep Learning:

Neural networks predicted atomic distances and bond angles, assembling structures like a 3D puzzle guided by physics-based rules (e.g., avoiding atomic clashes).

3. Confidence Scoring:

Each prediction included a per-residue accuracy score, flagging uncertain regions for lab validation. At the 2020 CASP competition, AlphaFold2 stunned scientists with **92.4%** accuracy, rivaling

Median Free-Modelling Accuracy



experimental methods like cryo-electron microscopy. By 2021, DeepMind had predicted **98.5%** of human proteins, releasing them in a public database of 200+ million predicted protein structures, including the entire human proteome, now freely available to scientists.

David Baker: Designing Proteins from Scratch

In early 2020, as the world scrambled for ways to stop a novel coronavirus, Baker’s lab pivoted fast. Within weeks, they designed a synthetic protein shaped like a tripod that latched onto the SARS-CoV-2 virus’s spike and blocked its entry into cells. It was a feat not of evolution, but invention, one enabled by decades of patient, behind-the-scenes work at the interface of biology, physics, and computer science. This was no longer theory or academic promise; it was a glimpse into the future of medicine where life-saving molecules are not discovered but crafted.

While AlphaFold2 decoded nature’s designs, **David Baker’s Rosetta@AI** flipped the script: **Why not create entirely new proteins?** His team combined DeepMind’s neural networks with physics simulations to engineer proteins unseen in nature, an achievement the Nobel Committee



described as “*an almost impossible feat.*”

Instead of entering amino acid sequences to get protein structures out, Baker’s team inverted the workflow. They began with a desired 3D structure, say, a molecule shaped to bind cancer cells or break down plastic, and used AI to generate amino acid sequences that would fold into that exact shape. This reverse-engineering breakthrough turned Rosetta into a design tool, enabling scientists to create entirely new proteins never seen in nature.

Why Their Work Matters

The Nobel Committee captured it best: “That we can now so easily visualise the structure of these small molecular machines is mind boggling; it allows us to better understand how life functions, including why some diseases develop, how antibiotic resistance occurs or why some microbes can decompose plastic. The ability to create proteins that are loaded with new functions is just as astounding. This can lead to new nanomaterials, targeted pharmaceuticals, more rapid development of vaccines, minimal sensors and a greener chemical industry.

Conclusion

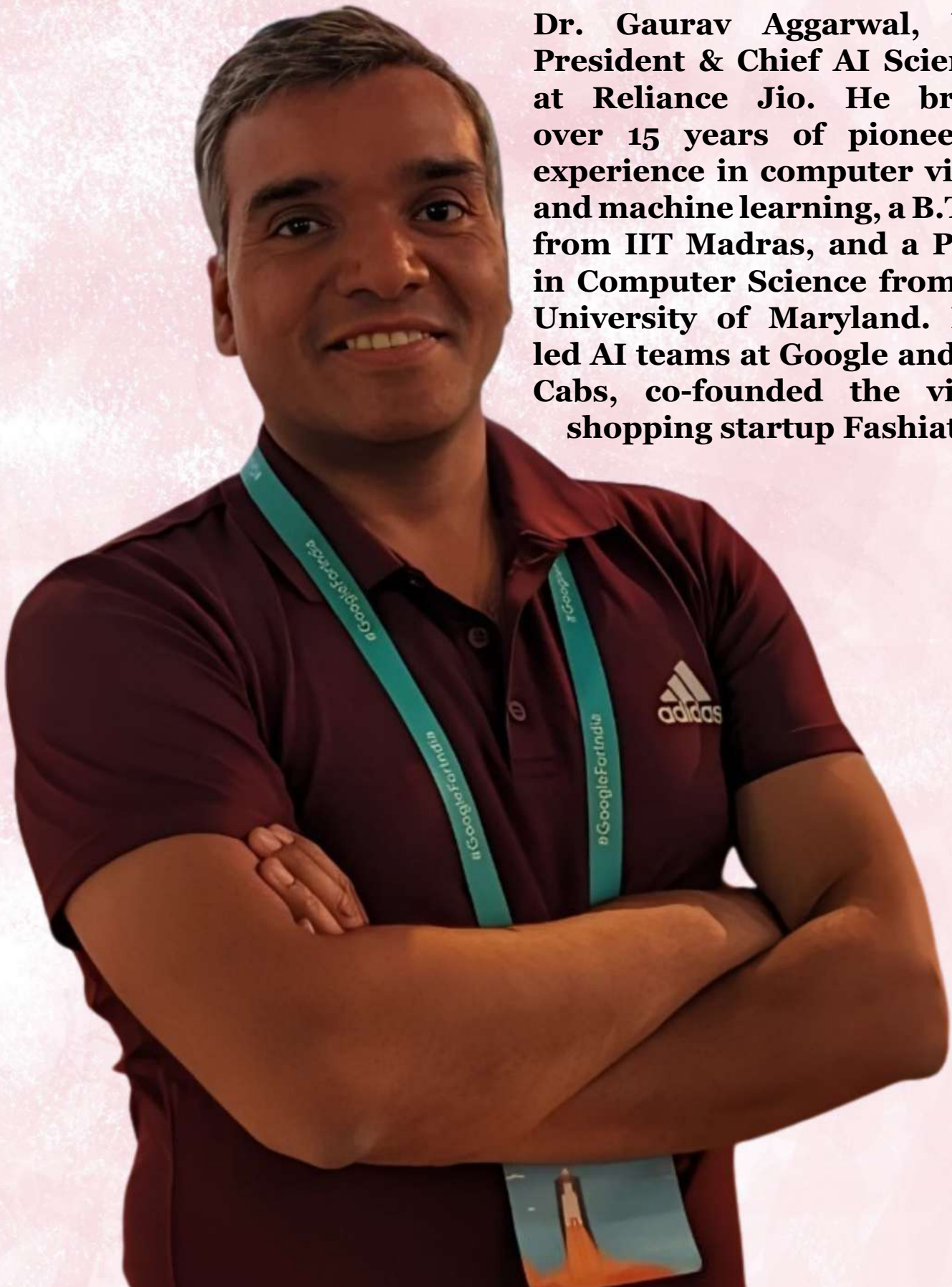
The 2024 Nobel Prizes in Physics and Chemistry mark a pivotal moment in science, a recognition that artificial intelligence is not merely a tool, but a bridge between disciplines. The Nobels remind us that the greatest breakthroughs emerge not from silos, but from the audacity to let physics converse with biology, and computation with creativity. In the words of the Nobel Committee, we are now “*redesigning life’s toolkit,*” blending curiosity with purpose.

This moment signals more than just scientific progress, it’s a shift in how we think about knowledge itself. The boundaries between discovery and invention are fading. We are no longer just uncovering nature’s secrets; we are beginning to shape them. In labs where atoms are arranged with precision and algorithms learn to predict life’s architecture, a new kind of science is taking root, one that is collaborative, bold, and increasingly humane. The future won’t be written by any one field alone, but by the willingness to imagine what they can build together.

An AI Journey: From Academia to Industry and Back to the Roots

Gaurav Aggarwal

Dr. Gaurav Aggarwal, Vice President & Chief AI Scientist at Reliance Jio. He brings over 15 years of pioneering experience in computer vision and machine learning, a B.Tech from IIT Madras, and a Ph.D. in Computer Science from the University of Maryland. He's led AI teams at Google and Ola Cabs, co-founded the visual shopping startup Fashiate.



AINA: We've followed your incredible journey in AI and Computer Vision. But looking back to the beginning: What truly motivated you to pursue a PhD in this field at a time when AI wasn't the "hot" topic it is today? And what kept you committed throughout—was there a moment, an idea, or a larger vision that anchored you?

I started working on neural networks around the year 2000, long before they became mainstream. What really anchored my interest was the influence of Professor Bayya Yegnanarayana at IIT Madras, who passionately believed in neural networks even when they weren't widely accepted and inspired many of us to think deeply about AI when it wasn't a popular field. His discipline, work ethic, and conviction left a strong impression on me. Coming from a Tier-3 town in Uttar Pradesh, even getting into IIT was a significant achievement for me. A PhD wasn't part of the plan, I had good job offers and was ready to enter the industry. But Professor Bayya Yegnanarayana strongly encouraged me to apply. His belief in my potential and his mentorship were key reasons I took that leap. I ended up applying to just one university, more out of respect than ambition. He wrote a strong recommendation, and that single application changed the course of my life. Looking back, it wasn't a single moment, but a combination of early exposure, inspiring mentors, and the desire to keep learning that kept me going.

AINA: After a distinguished career in the U.S., including roles in academia and at organizations like Google Research, you made the pivotal decision to return to India and get involved in the Indian AI

landscape. First, What were the reasons for this transition, and if you can tell how has your vision for India's AI ecosystem evolved since your return?

The desire to return was always there—it wasn't driven by a single event but a deep-rooted feeling. I've always believed in staying connected to my roots, and while my career in the US was rewarding, the idea of coming back never left me. My parents had always been hesitant about me going abroad, and somewhere there was an unspoken promise to return.

What truly triggered the move was a personal loss, my father-in-law passed away while we were in the US. That moment made us reflect deeply. We realized that as our parents aged, we didn't want to be thousands of miles away. It wasn't just about being helpful, it was about being present.

Once we made the decision, things fell into place. We found the right opportunities and moved back. It's been over 13 years now. Since then, I've seen India's AI ecosystem grow tremendously. There's still a long way to go in terms of research culture and infrastructure, but the momentum and talent here are undeniable. Being part of this journey—to help shape and support it- was a key reason I returned.

AINA: As someone who has experienced the evolution of computer vision first-hand — from the era of SIFT and HoG to CNNs and now transformers, which transitions or breakthroughs felt the most transformative or unexpected to you personally?

Transitions are natural in research. Transformations happen every

few years, often surprising even those of us in research. Right now, Transformers are everywhere, believe that in just a couple of years, we'll be moving beyond Transformers to something new and potentially even more powerful. The key is to understand that research is inherently iterative and uncertain. From the outside, these shifts might look like sudden disruptions, but for those of us immersed in the field, it's a natural and ongoing process. Many innovative ideas don't succeed, they're either wrong or ahead of their time—that's the nature of innovation. You dream, experiment, and often get it wrong. But occasionally, something sticks. You have to constantly dream big and be ready to pursue ideas when few believe in them. That's the essence of innovation.

AINA: Today, AI models have become easily accessible But this accessibility might also cause a detachment from foundational understanding. As someone who has worked on fundamental computer vision problems since the pre-deep learning era, what are some conceptual pitfalls you see in how today's practitioners use these models?

That's a very important question.

In my view, many young AI practitioners today often lack a solid mathematical and theoretical foundation. They might be able to get models running or fine-tune them superficially, but without understanding the core principles, they struggle to diagnose problems or push the boundaries with novel innovations.

To put it simply, many are like people who know how to change

a bulb but don't understand the wiring behind it. They can replace components, but they can't design or repair the entire system. True engineering requires knowing the wiring, its specifications, how it works, and how to build or improve it. Similarly, many practitioners treat AI models like magical black boxes. They believe the "magic" is in their manipulation of the model, but in reality, the power lies in the underlying model design and algorithms. Without understanding these fundamentals, they become mere operators, not innovators. This detachment from foundational knowledge limits progress. It's not enough to just apply models; to truly advance AI, we must deeply understand how and why they work. I see this gap clearly in the Indian ecosystem, and it's something I am passionate about changing, encouraging practitioners to master the basics, so they can contribute meaningfully to building the future of AI.

AINA: After returning to India, you transitioned into entrepreneurship by co-founding Fashiate, an AI-powered startup. What was the founding vision behind it? Was there a specific gap in the market you wanted to address, a personal passion that drove it, or an emerging opportunity in AI that made the timing right?

When I returned to India around 2013-14, startup culture was booming, and we were already deep into AI research at Yahoo Labs, using deep learning before it became mainstream. We saw an opportunity to apply AI to fashion, not because I'm passionate about fashion, but because there was a clear gap where AI could help.

Yahoo India was shutting down, so instead of job hunting, we decided to start our own company. We built everything from scratch, including our own deep learning models, and even gathered our own training data. Within a few months, Snapdeal and Flipkart wanted to acquire us, and we went with Snapdeal.

Back then, we were probably India's first deep learning startup to get acquired. Our product let users take a picture and find similar fashion items, something quite rare globally at that time.

Though Snapdeal later struggled and the product didn't last, I'm proud we executed and scaled such an innovative solution early on. It was truly ahead of its time.

AINA: "As Indian enterprises adopt AI—amid long ROI cycles and limited leadership familiarity—how can they build sustainable AI capabilities while securing top-management support?"

The biggest challenge in AI adoption isn't the technology or the ROI—it's the people. Change is hard, and humans naturally resist it, even when it's for the better. For AI to truly succeed, organizations need passionate change-makers, especially at the leadership level, who are committed to driving transformation.

Think of AI like a high-performance race car: no matter how powerful it is, if no one knows how to drive it or is willing to get behind the wheel, it will just sit idle. Similarly, without a willingness to rethink processes, reorganize teams, and build new skills, AI initiatives will struggle to deliver value.

Adopting AI means embracing cultural and structural shifts—not just purchasing new technology. The

true magic happens when people actively integrate AI into their daily work. Without that commitment, even the most advanced AI tools remain untapped potential.

AINA: As Jio's Chief AI Scientist, you're at the helm of AI strategy within one of India's largest and most diverse digital ecosystems. Could you share how AI and analytics are being leveraged across different layers of Jio—whether in customer experience, network optimization, new product development, or even shaping the company's long-term digital vision?

At Reliance, AI touches every part of our vast digital ecosystem, from improving customer experiences and optimizing our network to developing new products and businesses. Our focus isn't just on cutting costs or boosting efficiency, but on creating entirely new opportunities that didn't exist before.

For instance, powering our telecom networks consumes huge amounts of electricity. We use AI to reduce this cost by optimizing energy use across 3G, 4G, and 5G networks, a complex, location-specific challenge few outside India fully grasp.

AI also helps us understand customer pain points before they even reach out, enhancing service quality. In retail, whether online or offline, AI guides inventory, pricing, discounts, and personalized loyalty programs.

Logistics is another key area, where AI optimizes deliveries, saving time and cost.

The impact has been huge, our AI projects often deliver returns many times the investment in the teams behind them. Our goal is to increase this impact even more.

What excites me most is Reliance's leadership commitment to innovation and growth. Beyond business, there's a deep sense of responsibility to contribute to India's development. When I present plans, I frame them around "India's growth" rather than just Reliance, because our success is tied to the country's progress.

Being at Reliance means driving world-class AI innovation that's truly Indian, and that's what motivates me every day.

AINA: From your perspective, what are some emerging trends in AI and analytics—be it generative AI, edge AI, or multimodal learning—that you believe will significantly reshape how enterprises operate and make decisions in the coming years? And maybe if you can share how is Jio preparing for these shifts?

I believe AI will soon touch every part of our lives — not just in apps or phones, but in the very environment around us. The real transformation will come from the form factor. Today, AI is mostly accessed via smartphones or laptops, but in the future, we'll see AI embedded in wearables, glasses, voice-first devices, and other ambient tools that understand context, where you are, what you're doing, and how best to assist you.

Think of it like a co-pilot, but not just for coders. A co-pilot for every profession: factory workers, retail staff, logistics, farmers, even house help. That's when AI becomes truly democratized.

At Jio, we're preparing for this shift not just with cutting-edge models, but by building for India-first realities. Cost matters. Multilingual access matters. Simplicity matters. Just like Jio made data affordable

and widespread, we're working towards a similar "Jio moment" for AI — where every Indian, regardless of income or education, can benefit from it.

We're also deeply excited about real-time language translation. Imagine going to Chennai, speaking in Hindi, and being heard in Tamil — instantly. These kinds of breakthroughs don't just help enterprises work better; they bring people together. That's the deeper power of AI, and we want to lead that change.

AINA: Often the magic happens at the intersection of tech and domain. How important is domain knowledge when building impactful AI solutions? and how do you balance it with technical proficiency within industrial teams? What would you suggest to develop that in aspiring youth?

That's a great question. Modern AI models, especially large ones, reduce the need for deep domain knowledge because everything, whether it's text or images, gets converted into tokens. So, technically, data is just data.

But the real impact comes not from just building models it's about building products that solve real problems. And that's where product thinking and domain understanding become critical. You need to understand users deeply, sometimes even better than they understand themselves.

For example, most of India's blue-collar workers have smartphones, but many never use tools or to-do list apps. These tools weren't designed for them. That's the gap, and also the opportunity.

My advice to young people: yes, build strong technical skills, but also

focus on understanding user needs. Think about how people will use your solutions. AI is powerful, but it needs to be made accessible and relevant, especially in a country like India. That's where real innovation lies.

AINA: "As someone starting AI in 2025—with open-source models, online communities, and evolving tools—what path would you take today? And what roadmap would you suggest for students at the intersection of business, data, and tech who aim to shape AI's future?"

Great question. If I were starting today, I'd say: focus on the intersection of technology, business, and product thinking. You don't need to build new AI models from scratch, but you must understand what models can do & think creatively about how they can solve real-world problems.

Too often, we think small or play it safe. But real innovation comes from being bold and imaginative even "wild." Think of the iPad it wasn't a bigger phone or smaller laptop; it was a new idea altogether. That's what we need: wild ideas with purpose.

Don't let the "middle-class mindset" of playing it safe hold you back. Most wild ideas fail, but the few that succeed can change everything. Take risks. Build things that people don't yet know they need. Whether you're technologist, business leader, or product builder be yourself, believe in your vision, and don't wait for someone else to define your path. That's the mindset we need to truly shape the future of AI.

Learning the Art of Prompting

- Rohith Chandra

“The future of AI is not about replacing humans, it’s about augmenting human capabilities.”

As a Gen Z by birth and Gen X in mindset, whenever I sit to understand new things and gain knowledge, I find googling the star of the past and the cash cow of the present, tiresome, as it needs more effort but the results are decent. AI chatbots, the stars of the present time easy as it need minimal effort but the results are a question mark as there may be chances of bias, hallucination, etc... and apart from this there may be other solutions like books and consulting an expert, which are not always feasible in general for everyone.

So, going back to our top two solutions, though search engine is decent for basic queries, most of the search engines are in peak especially for fact checking purposes but they are not suited when



we are looking for advanced information like a programming code, brief summary of a topic, problem solving methodology and many such questions. Now, as the limelight shifts towards AI chatbots, even though it comes with its set of downside but it can give a human-like response with less effort, and these downsides can be bypassed or overcome.

As there is scope for development to get a solution from the chatbots that is nearest to our desired solution. It can be like selecting a suitable AI chatbot for the kind of the problem we are dealing with, involving in an effective interaction with the chatbot, which steers us to a perfect solution. So going forward we will be exploring the art of



Prompt Engineering, which helps the chatbot to effectively understand our problem and provide us with a better solution. Before diving into the topic, we need to take a step back. I want to ensure we are all on the same starting line of the knowledge race. So, I will be explaining some basic foundational concepts for better understanding of the upcoming topics. These basics are defined by AI chabots from the prompt methods we will learn in this article.

Large Language Models (LLMs)

Large Language Models (LLMs) are a fascinating breed of artificial intelligence. At their core, they are **sophisticated computer programs**

trained on colossal amounts of text and code. This extensive training allows them to understand, summarize, translate, predict, and generate remarkably human-like text. Think of them as versatile digital wordsmiths, capable of crafting articles, answering questions, and even writing various kinds of creative content.

But what exactly makes them “**large**”?

The “large” refers to the sheer scale of the neural networks underpinning these models and the massive, diverse datasets they’re trained on. This vastness enables them to capture intricate patterns and nuances in language, ultimately powering their wide-ranging and impressive abilities.

Unveiling the Magic: How LLMs Work

Large Language Models (LLMs) might seem like magic, but their impressive abilities are rooted in clever engineering and massive datasets. At their heart lies a complex neural network, a system inspired by the human brain, designed to process information. These networks are trained on vast quantities of text and code, learning the intricate patterns and relationships within language.

The key to their operation is prediction. Given a sequence of words, the LLM predicts the most likely next word. This process, repeated iteratively, allows the model to generate coherent and contextually relevant text.

Think of it like completing a sentence – you anticipate the next word based on what you’ve already read. LLMs do this on a massive scale, leveraging the patterns learned from their training data.

During training, the model’s predictions are compared to the actual text, and the network’s parameters are adjusted to improve its accuracy.

This process, known as “backpropagation,” refines the model’s ability to understand and generate language.

The larger the dataset and the more complex the network, the better the LLM becomes at its task. While they don’t “think” like humans, their ability to process and predict language makes them powerful tools for a wide range of applications.

Prompt: Guiding LLMs

Large Language Models (LLMs) are powerful, but their output depends heavily on the input they receive: the prompt. A prompt is essentially the instruction you give to the LLM, guiding its generation of text. Think of it as a chef taking a recipe – the clearer the recipe, the better the dish.

Effective prompts are key to unlocking the full potential of LLMs. They can range from simple

questions (“What is the capital of France?”) to complex instructions (“Write a short story about a robot learning to feel”). The more specific and detailed the prompt, the more likely the LLM will produce the desired output.

LLMs Breaking Down Language:

Before a Large Language Model (LLM) can understand a prompt, it must first break it down. This process, called tokenization, involves splitting the text into smaller units called tokens. These tokens can be words, subwords, or even individual characters. Think of it as dissecting a sentence into its constituent parts for closer examination.

Tokenization is crucial because LLMs process information at the token level. By breaking the prompt into manageable chunks, the model can analyze the relationships between these units and understand the overall meaning. Different LLMs might use different tokenization methods, impacting how they interpret the input.

Consider the prompt: “The quick brown fox jumps over the lazy dog.” A typical tokenizer might split this into tokens like “The,” “quick,” “brown,” “fox,” “jumps,” “over,” “the,” “lazy,” “dog,” and “.”. However, some tokenizers might break “brown” into “bro” and “wn” or even treat “jumps” as “jump” and “s.” This subword tokenization helps LLMs handle unfamiliar words and morphological variations. Each token is then converted into a numerical representation, an embedding, which the LLM understands. Different tokenizers and vocabularies exist, impacting the final embedding and thus, the LLM’s interpretation.

Prompt Engineering: Steering LLMs to Success

Large Language Models (LLMs) are powerful tools, but their true effectiveness hinges on how we interact with them. This is where prompt engineering comes in. It’s the art and science of thoughtfully crafting effective prompts to guide LLMs toward desired, meaningful outputs.

Think of it as training a highly intelligent, but sometimes unpredictable, assistant. Prompt engineering involves more than just asking a question. It requires understanding how LLMs process language and experimenting with different phrasing, structures, and techniques.

Clear instructions, specific examples, and well-defined constraints can significantly improve results. For complex tasks, breaking them down into smaller, manageable sub-prompts can be beneficial.

Prompt engineering is the bridge for effective human-AI communication. It's more than getting the right answer; it's about ensuring AI understands the context, the nuances, and the intent behind every query.

So, from the above section it is clear that Understanding tokenization is crucial for prompt

engineers, as it influences how the LLM “sees” the input and, consequently, its output, because a well-crafted prompt considers how the LLM will tokenize the text, ensuring the intended meaning is preserved. This knowledge helps users avoid ambiguity and guide the LLM towards the desired output.

Mastering prompt engineering is crucial for unlocking the true potential of LLMs. It's the key to generating high-quality content, automating tasks, and leveraging the power of AI for a wide range of applications. By carefully crafting our prompts, we can transform LLMs from general-purpose tools into specialized problem solvers.

So, crafting good prompts is an Art, which basically involves understanding how LLMs interpret language and experimenting with different phrasing. Techniques like providing examples,

POV of a LLM

Good practice of prompt
A CAREful prompt

Context: Describe the situation

Ask: Request specific action

Rules: Provide constraints

Examples: Demonstrate what you want

There is clarity, easy to comprehend
and gives effective output.

Bad practice of prompt
When you treat it as a RACE

Rules:- Provide constraints

Ask:- Request specific action

Context:- Describe the situation

Examples:- Demonstrate what you want

There is no clarity in the prompt,
hard to comprehend and gives
ineffective output.



specifying the desired format, and setting constraints can significantly improve results. Mastering the art of the prompt is crucial for anyone looking to leverage the power of LLMs. And the final nail in the coffin of these basic-explanations section, the Art (with capital 'A') in the title refers to art (with smaller 'a') that is an important or finer version of the art. If you don't believe me then ~~google it~~ use AI chatbot to check it

So now we can move into the prompt. We start by Dissection of the prompts:

Elements of a Prompt

A prompt contains any of the following elements:

Instruction - a specific task or instruction you want the model to perform

Context - external information or additional context that can steer the model to better responses, this can also be like giving a persona to it, providing reference text.

Input Data - the input or question that we are interested to find a response for

Output Indicator - the type or format of the output.

Types of prompt:

Based on the type of the task we want we can use well-crafted prompts to perform different types of tasks like

- Text Summarization
- Information Extraction
- Question Answering
- Text Classification
- Conversation
- Code Generation
- Reasoning

Techniques in prompt engineering

These are for basic requirements, a user can use to make their prompts better.

• Role-playing

By making the model act as a specific entity, like a historian or a scientist, you can get tailored responses. For example, "As a nutritionist, evaluate the following diet plan" might yield a response grounded in nutritional science.

• Iterative refinement

Start with a broad prompt and gradually refine it based on the model's responses. This iterative process helps in honing the prompt to perfection.

• Feedback loops

Use the model's outputs to inform and adjust subsequent prompts. This dynamic interaction ensures that the model's responses align more closely with user expectations over time.

Advanced techniques

Now we look into more intricate strategies that require a deeper understanding of the model's behavior.

• Zero-shot prompting

In simple language you describe the task you want the model to perform without providing examples. The LLM then accesses its extensive pre-trained knowledge to generate a relevant response. It tests the model's ability to generalize and produce relevant outputs without relying on prior examples.

Eg: **Prompt:** Classify the animal based on its characteristics. Animal: This creature has eight legs, two pincers, segmented tail, and often eats insects.

Output: Scorpion.

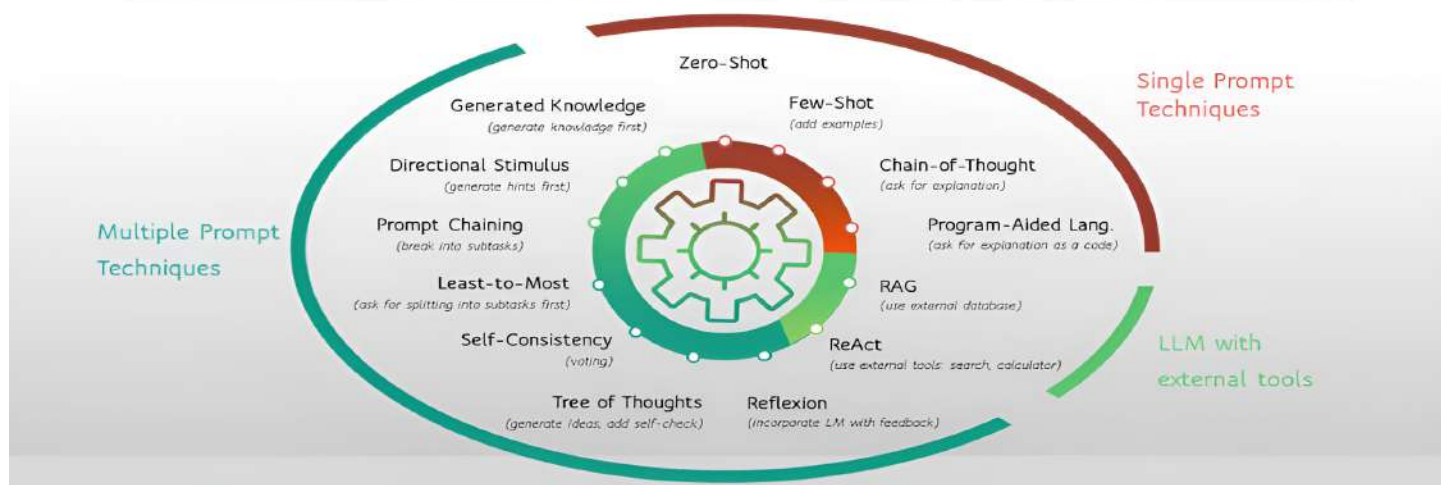
Effective prompt design is crucial for successful zero-shot prompting. It bridges the user's intent and the model's capabilities.

• Few-shot prompting/in-context learning

Here, the model is given a few examples (shots) to guide its response. By providing context or previous instances, the model can better understand and generate the desired output. For example, showing a model several examples of similar problems before asking it to solve a new one.

Prompt: Provide a solution to a logical question like tree: seed, ocean: river answer this one bird

Prompt Engineering Techniques



Output: bird : egg (or nest, or wings, or song - depending on the specific relationship being highlighted. Egg is probably the most direct analogy to seed/tree and river/ocean).

• Chain-of-Thought (CoT)

This advanced technique involves guiding the model through a series of reasoning steps. By breaking down a complex task into intermediate steps or “chains of reasoning,” the model can achieve better language understanding and more accurate outputs. It’s similar to guiding someone step-by-step through a complex math problem.

Eg: Prompt: Write a story about a dog.

Output: (Writes part of the story)

Prompt: Continue the story, but make it take a surprising turn.

Output : (Continues the story with a twist)

Prompt: Now, resolve the story in a heartwarming way.

Each user turn builds upon the previous interaction, creating a chain of prompts that guide the LLM’s narrative generation.

These techniques are also used to train the LLM by prompt.

Creating the perfect prompt often involves experimentation. Till now you might think what madness have I entered, don’t worry there is a

method to this madness which we will explore now.

Prompt Optimization Techniques

Now let’s dive into prompt optimization techniques, these fundamental techniques will help you improve your interactions with AI chatbots.

1. Be specific and direct

Avoid ambiguity. Vague or open-ended prompts often lead to unpredictable results because the AI might focus on aspects you’re not interested in or provide overly general information.

Example: You want to know about what causes climate change and its impact on global temperature and sea level.

Vague prompt: Tell me about climate change.

Output: Climate change refers to long-term shifts in global weather patterns and average temperatures. It’s primarily caused by human activities, especially the burning of fossil fuels. Effects include rising sea levels, more extreme weather events, and impacts on ecosystems and agriculture. Many scientists consider it one of the

most pressing issues facing our planet today.

Optimized prompt: What are the two main causes of climate change and their specific impacts on global temperature and sea levels?

Output:

1. Burning fossil fuels: Releases CO₂, raising global temperatures by 1.1°C since pre-industrial times. This contributes to sea level rise of about 3.3mm per year due to thermal expansion and ice melt.

2. Deforestation: Reduces CO₂ absorption, increasing atmospheric CO₂ by ~15%. This indirectly contributes to temperature rise and sea level increase.

2. Set the context

Providing relevant background information helps the AI understand the context of your request, leading to more personalized and accurate responses.

3. Define the desired format

Specifying the format you want helps structure the AI's response in a way that's most useful to you.

4. Use examples

Providing examples can effectively set the tone, style, or level of complexity you're looking for in the response.

5. Iterate and refine

Experiment with phrasing. If the initial response doesn't meet your needs, that's okay—rephrase your prompt or add more details.

Eg: Prompt

- Imagine you are an reputed expert in the field of LLM
- You are writing a short articles of 150 words on working of prompt on LLM's for a reputed magazine
- The articles need to be crisp, clear and easy to understand

The output of this prompt is a part of this article that you came across in the basic definitions which you have already gone in earlier sections.

Key Considerations

Clarity - Use clear and unambiguous language in each step.

Structure: Numbered steps, bullet points, or clear paragraph breaks can improve readability.

Examples: Providing examples of the desired behavior is extremely helpful.

Context: Maintain context throughout the chain of prompts.

Experimentation-Try different phrasing and structures to see what works best for your specific task.

Common pitfalls when crafting prompts

Overloading - too much (irrelevant) information
Ambiguity - vague prompts leads to a generalized answers.

Over-complication - using jargon, complex phrasing, or technicalities complicates the prompt

Highlighting certain parts of a prompt can significantly improve clarity and guide the Large Language Model (LLM) towards the most important information. Here are several effective ways to highlight within a prompt:

1. Bolding:

How: Use double asterisks **bold text** or `<b>bold text</b>` (in some contexts).

Purpose: Emphasize keywords, crucial instructions, or specific terms.

Example: "Summarize the following article in no more than 100 words: [article text]" or "I need a recipe for chocolate chip cookies."

2. Italics:

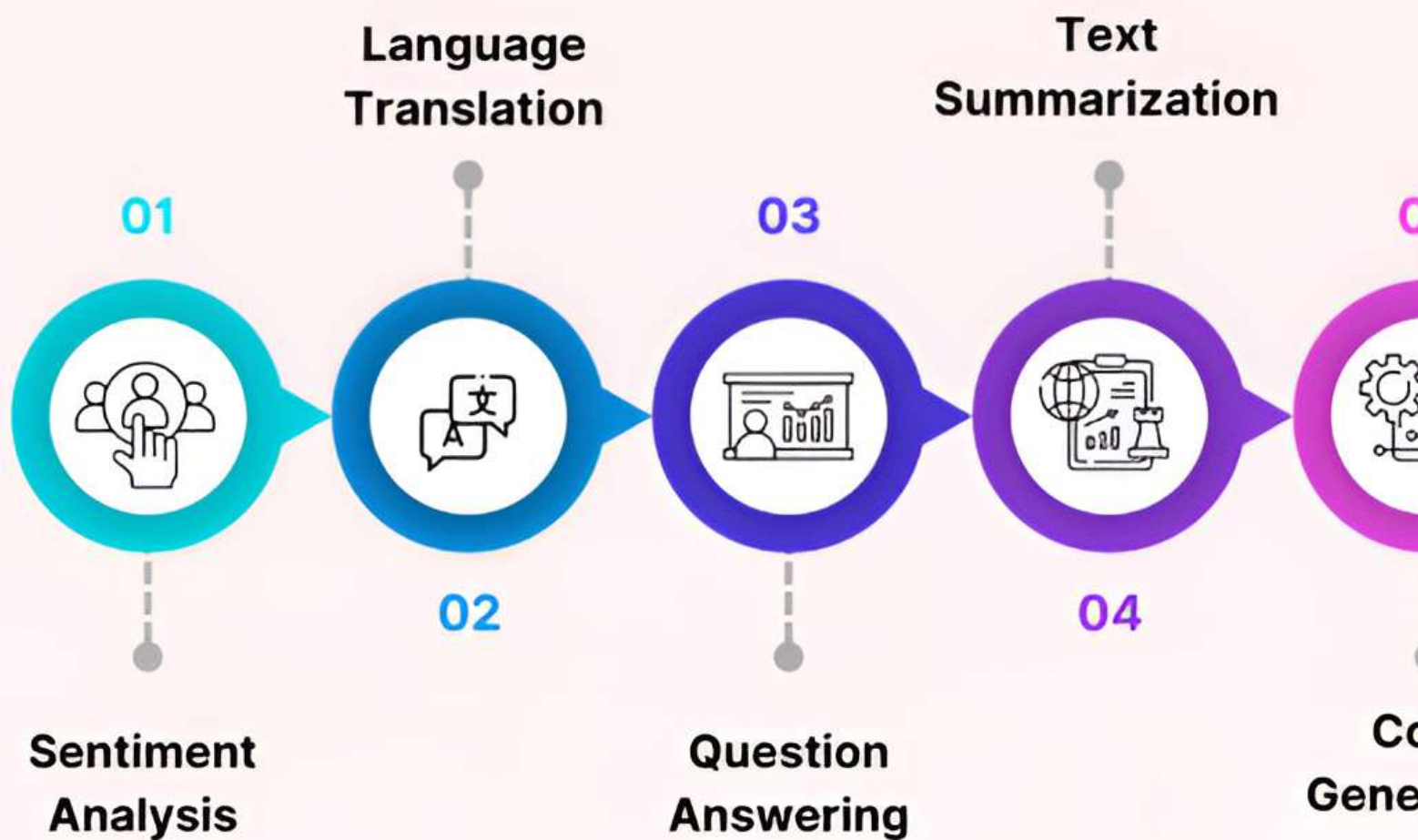
How: Use single asterisks *italicized text* or `<i>italicized text</i>` (in some contexts).

Purpose: Highlight titles, foreign words, or less crucial emphasis. Use sparingly.

Example: "The Great Gatsby is a novel by F. Scott Fitzgerald."

3. Underlining:

How: Use HTML tags `<u>underlined text</u>`



(may not be supported in all LLM interfaces).

Purpose: Similar to bolding, but less common.

Use sparingly.

Example: Important: Do not include any personal information.

4. Quotation Marks:

How: Use double quotes “quoted text” or single quotes ‘quoted text’.

Purpose: Indicate direct quotes, specific phrases, or titles.

Example: “Write a poem about the feeling of ‘wanderlust’.”

5. Capitalization:

How: Use all caps CAPITALIZED TEXT.

Purpose: Draw attention to very important words or phrases. Use sparingly to avoid the impression of shouting.

Example: “DO NOT include any images in your response.”

6. Numbering or Bullet Points:

How: Use numbers or bullet points for lists.

Purpose: Organize multiple instructions or pieces of information.

Example:

1. Summarize the article.
2. Identify the main arguments.
3. Provide your opinion.

7. Special Characters:

How: Use characters like > or | to separate sections or highlight information.

Purpose: Visually structure the prompt.

Example: > Instruction: Write a story. > Genre: Science Fiction > Length: 500 words

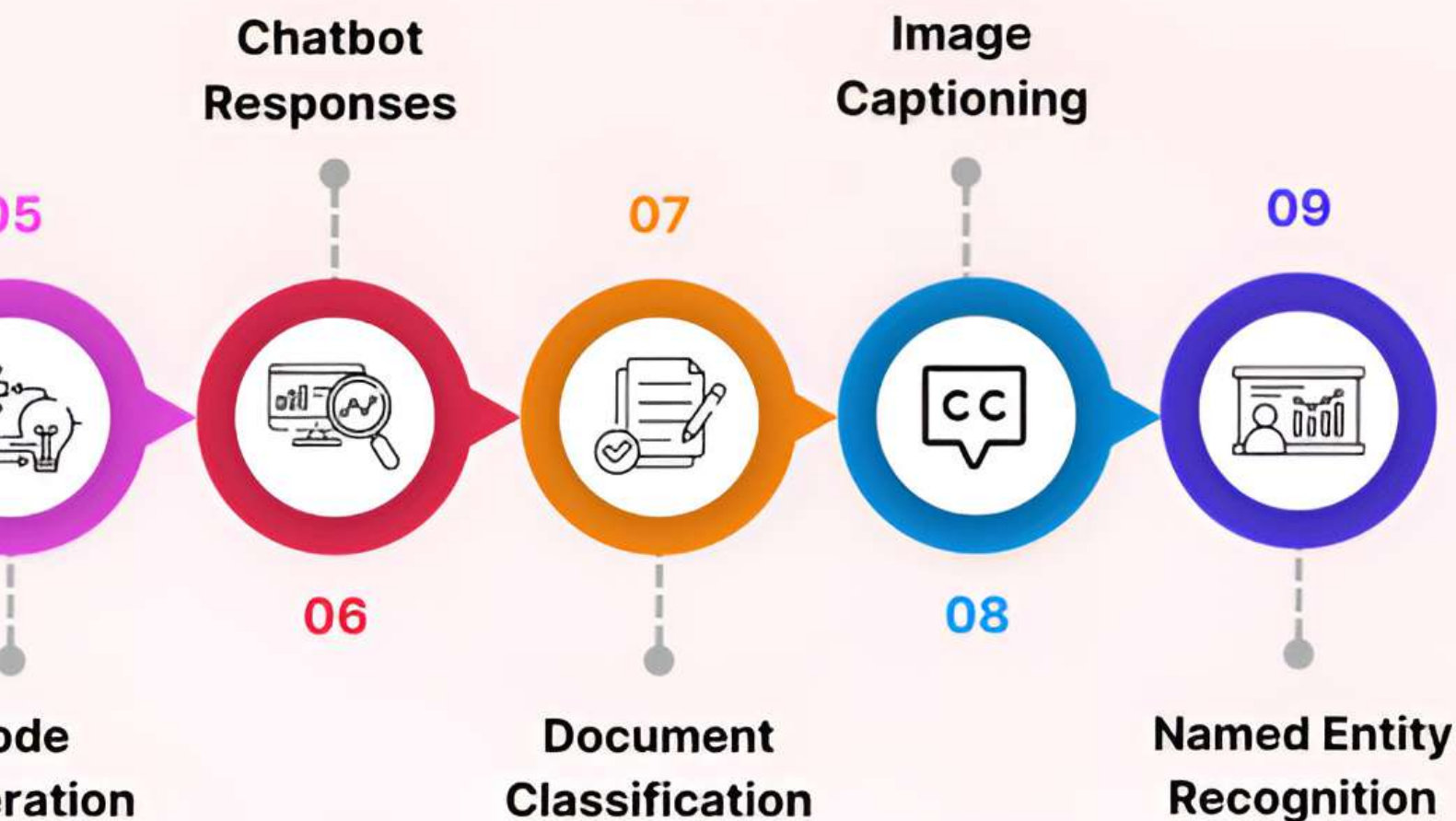
8. Code Blocks (for technical prompts):

How: Use triple backticks before and after the code.

Purpose: Clearly separate code from the rest of the prompt.

Example:

```
Python
def hello_world():
```



```
print("Hello, world!")
```

9. Emojis (use with caution):

How: Use emojis (if supported by the LLM interface).

Purpose: Can add visual cues, but use sparingly and professionally.

Example: "Summarize the following article 📄:"

10. Combining Techniques:

You can combine these techniques for even stronger emphasis.

Example: "IMPORTANT: <u>Review</u> the following document: [document text]"

Best Practices

Consistency: Use the same highlighting method for the same type of information throughout your prompts.

Sparingly: Don't overuse highlighting. It can make the prompt harder to read.

Purposeful: Use highlighting to draw attention to the most important information.

Test: Experiment with different highlighting methods to see what works best for your specific LLM and task.

By strategically using these highlighting techniques, you can create clearer, more effective prompts that guide the LLM to produce the desired results.

Conclusion:

If you made it till then you deserve a pat on the back.

I will close it with a few lines. It may not find it easy when implementing these methods. Trust me when you are reaping the fruits of your hardwork you will find it worthwhile.

Earlier I mentioned that it's an Art and "**Art is never finished, only abandoned**" and I suggest you not to abandon it when you feel tired and exhausted, only abandon it when you're done.

AI in Cultural Heritage

Reawakening the Silent
Past

Hari Srijaa





“What if we could give our past a future—
using the technology of tomorrow?
Imagine a world where no ancient temple
is forgotten, where every crumbling
manuscript and every fading artwork gets
a second chance.”

Thanks to artificial intelligence, this isn't just a dream. AI isn't just replacing historians, conservators, or linguists—but it's becoming their most powerful ally. With machine learning and deep learning, we're digitizing, decoding, and defending our cultural legacy like never before. For generations, we have fought to protect our cultural heritage against time, weather, and neglect. Now, AI is joining this fight—not to replace human experts, but to give them superpowers. It's helping us see what the naked eye can't, preserve what hands can't touch, and save what memories alone can't hold.

From digitally rebuilding lost cities to teaching computers to read forgotten languages, AI is opening doors to our history that we thought were closed forever. And the most exciting part? We're just getting started.

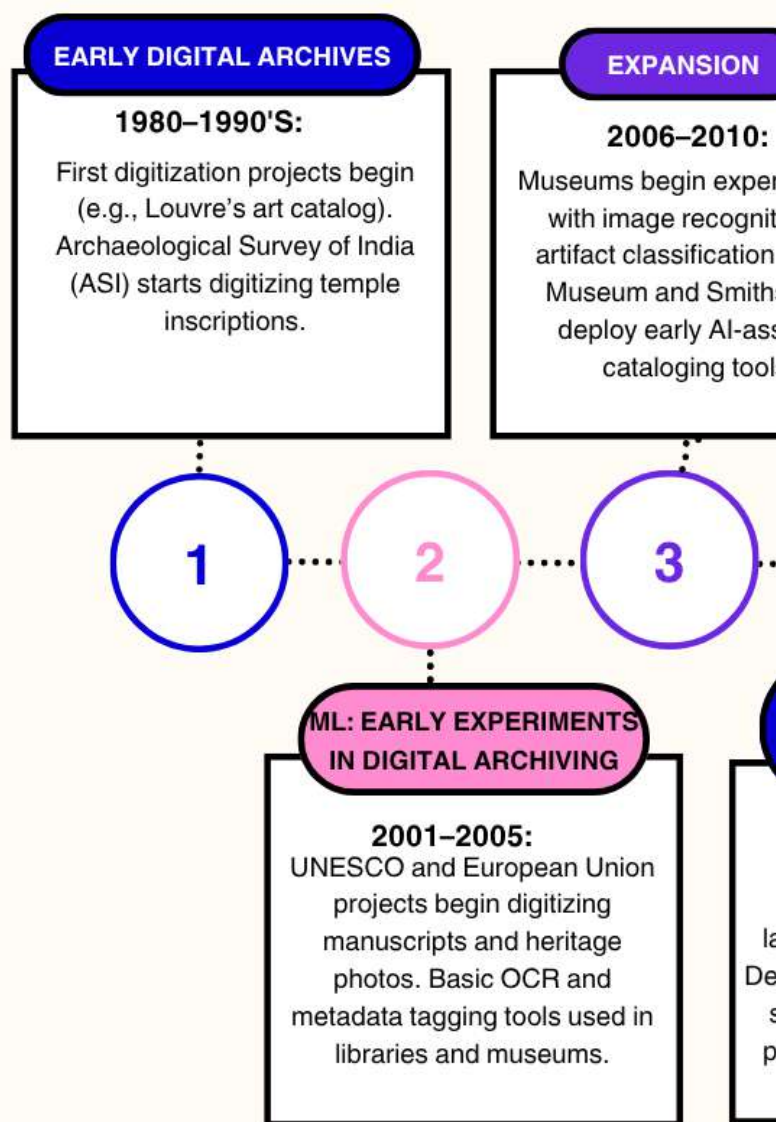
But in what ways is AI shaping the preservation and study of our cultural heritage? What are the technologies that allow it to decode ancient mysteries, preserve priceless artifacts, and even unmask forgery in the art world? How did it evolve through time?

1. Smart Digitization:

AI-Powered Restoration & Cataloguing

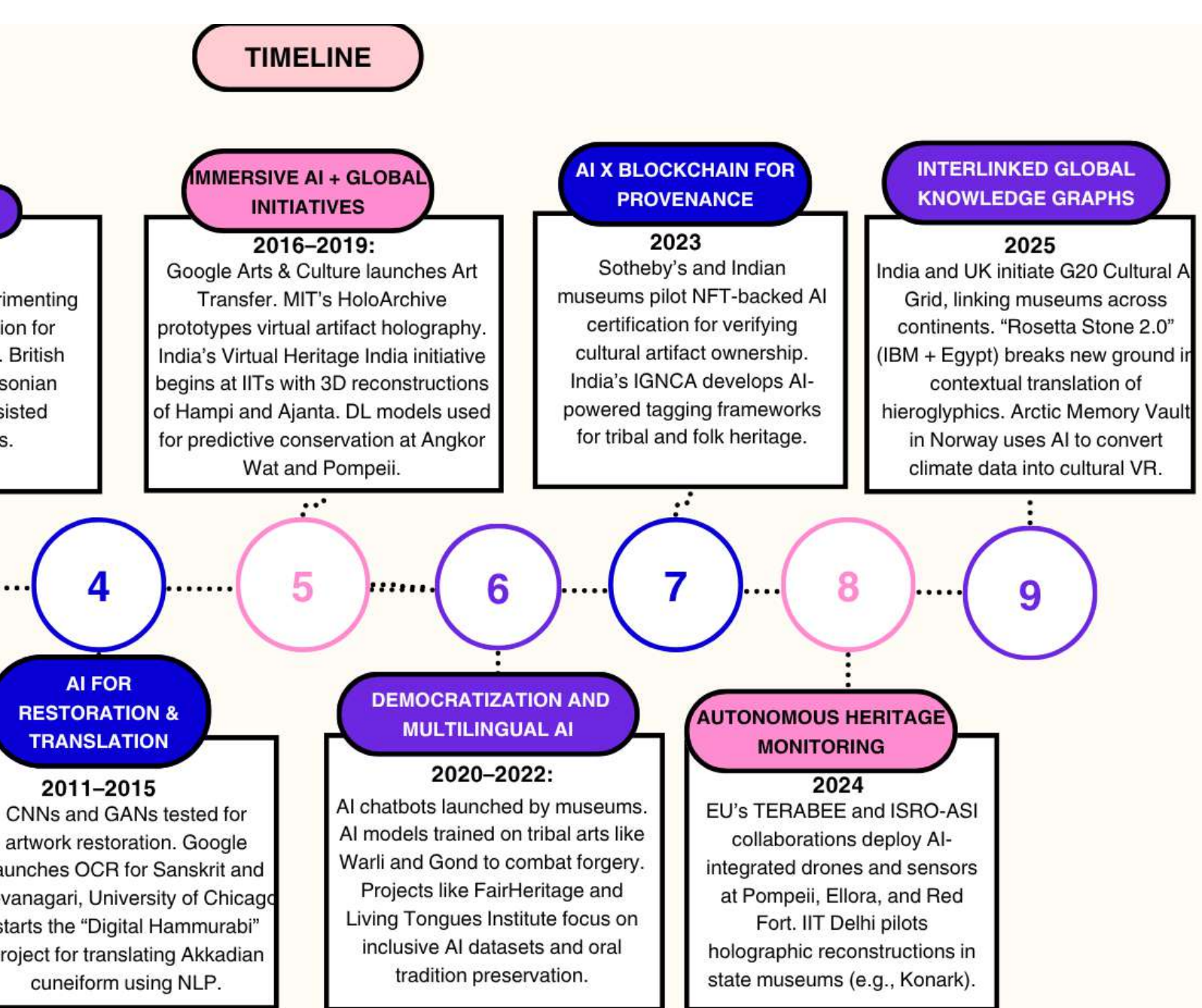
Imagine centuries-old manuscripts, crumbling with time, reborn through AI. AI breathes life into history, digitizing ancient texts, paintings, and relics with stunning accuracy. Deep learning models like GANs (Generative Adversarial Networks) can fill in missing pieces of weathered murals, restoring lost details with breathtaking precision.

Meanwhile, CNNs analyze patterns in textiles or ceramics, piecing together fragments of broken artifacts much like solving a puzzle. In



museums worldwide—like the British Museum's digitization of Egyptian artifacts—AI enhances image resolution, recovers faded colors, and helps create 3D models of statues and pottery. The same computer vision used to restore ancient mosaics helps retailers perfect product images for online stores, blending the past with the digital future. Beyond GANs and CNNs, AI-integrated lidar and photogrammetry now capture minute details of archaeological sites.

For instance, drone-mounted lidar combined with GIS creates millimeter-accurate reconstructions — used to model the ruins of Angkor Wat and submerged Mayan cities. These AI pipelines automate what previously required months of



manual modeling. In India, the National Mission for Manuscripts employs AI and OCR technologies to digitize ancient scripts such as Brahmi, Devanagari, and Grantha. This project automates transcription of fragile palm-leaf manuscripts, making rare texts accessible to scholars worldwide.

The Virtual Heritage India initiative applies AI-powered photogrammetry and deep learning to digitally reconstruct heritage sites like Hampi and Ajanta-Ellora caves, creating immersive 3D models and virtual tours. The British Museum uses AI not only for artifact digitization but also for enhancing image resolution and reconstructing statues digitally, exemplifying how cultural heritage institutions worldwide leverage AI for

preservation and research.

“By digitally restoring faded manuscripts and rebuilding lost cities, AI is turning fragments of history into immersive stories”

2. Artifact Authentication:

Detecting Art Forgery

The art world has always been a playground for forgers—AI steps in as a vigilant guardian. Convolutional neural networks (CNNs) and feature extraction algorithms can detect even the slightest

inconsistencies in brushstrokes, materials, or pigment composition—features invisible to the naked eye but revealing the truth of an artwork's origin.

Beyond brushstrokes, AI also evaluates oxidation patterns and chemical pigment decay. If an oil painting's surface oxidizes in a way inconsistent with its claimed era, the system flags it as likely counterfeit — adding a biochemical layer to traditional visual analysis.

Machine learning models compare a suspected painting to thousands of verified works, exposing forgery with forensic precision. Projects like the “Art Recognition” platform in Switzerland are already leveraging AI to authenticate paintings, safeguarding our cultural legacy and the art market alike. In India, AI-driven pattern recognition is increasingly being used to authenticate tribal and folk arts like Warli, Madhubani, and Gond paintings, helping preserve genuine cultural motifs and preventing counterfeit reproductions.

“With pixel-level precision and chemical insight, AI uncovers forgery by analyzing brushstrokes, pigment decay, and material patterns.”

3. Heritage Conservation: Predictive Maintenance of Historical Sites

Cultural sites, from ancient temples to medieval cathedrals, face the slow decay of time. But AI is lending them a lifeline. Predictive maintenance models—trained on environmental data, structural scans, and weather patterns—anticipate wear and tear before it becomes critical.

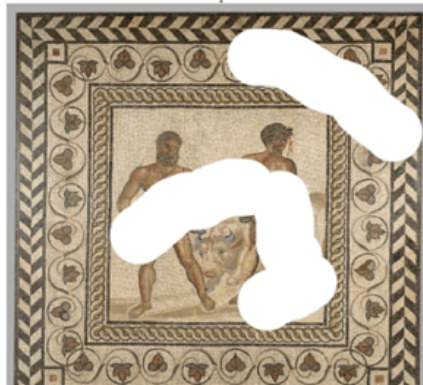
In Colombia, AI-enabled IoT sensors embedded in colonial churches trigger alerts to prevent environmental damage. Deep learning methods



Original



AI damaged



AI reconstructed



also analyze infrared and ultrasonic scans to detect hidden fractures and degradation in stone and plaster works.

Drones using AI-powered SLAM (Simultaneous Localization and Mapping) build high-resolution 3D models, enabling detailed analysis of site integrity. These predictive tools, adapted from industrial maintenance systems, are now vital for heritage management.

In India, AI-IoT networks have been piloted at heritage sites like the Red Fort and Qutub Minar to track pollution, humidity, and stress levels. Conservation teams use this data to preempt damage and ensure long-term preservation. Similarly, the EU-funded TERABEE project demonstrates the use of AI-powered nano-sensors and autonomous repair drones at sites such as Pompeii, pointing to a future of self-regulating conservation ecosystems.

“AI-driven sensors & modeling tools are reshaping heritage preservation-- by monitoring stress, preventing damage, & prolonging monument life with precision.”

4.Virtual Museums & Immersive Experiences: Bringing the Past to Life

Why just read about history when you can step into it? AI-driven VR and AR experiences are transforming how we explore cultural heritage. Computer vision and NLP (Natural Language Processing) power virtual museum tours, translating ancient scripts in real time or creating immersive journeys through lost civilizations.

GANs create photorealistic renderings of long-lost cityscapes—imagine standing in Pompeii moments before Vesuvius erupted. This technology isn't confined to museums; real estate and travel industries use similar AI-powered simulations to showcase properties and destinations. The

Louvre's virtual tours during the pandemic are a glimpse of this future—where AI turns history into an interactive experience.

While AI accelerates 3D reconstruction, the process often requires human-guided modeling in platforms like 3DStudio Max. This blend of machine precision and human aesthetic judgment helps recreate lost architectural spaces with both realism and respect for context.

India's Virtual Heritage India initiative enables immersive VR tours of sites like Ajanta-Ellora, expanding access while protecting fragile locations. The National Museum in New Delhi is also using AI-powered image recognition and multilingual chatbots to make exhibits accessible in Hindi, Tamil, Bengali, Marathi, and English.

Globally, platforms like Google Arts & Culture use similar AI tools—including Art Transfer and smart guides—to create rich, interactive experiences.

5.NLP and Machine Translation in Decoding Ancient Languages

What secrets do ancient scripts hold? AI is cracking the code. NLP models, trained on hieroglyphs, Sanskrit, or cuneiform, analyze linguistic patterns and translate these lost languages into modern speech. BERT and transformer models, the same ones that power voice assistants and chatbots, now help historians read the thoughts of millennia-old scribes. Researchers at the University of Chicago's "Digital Hammurabi" project use AI to translate ancient Akkadian texts, while Google's OCR (Optical Character Recognition) systems are turning Sanskrit manuscripts into digital texts, ready for a new generation of scholars.

In India, the Bharatiya Bhasha Samriddhi project is digitizing regional and classical scripts in collaboration with Google OCR. These initiatives are making centuries of literature available to modern scholars. Meanwhile, IBM's collaboration with the Egyptian Ministry of Antiquities uses contextual AI modeling to improve hieroglyphic translation accuracy—offering deeper insights into ancient Egyptian society.

“From reconstructed sites to multilingual access and cross-linked records, technology is streamlining how we explore and understand the past.”

6. Knowledge Graphs & AI-Linked Collections

What if museum records could talk to each other? That’s exactly what AI is enabling. In projects like the Heritage Connector, AI is used to build open knowledge graphs that link artifacts across institutions like the Science Museum Group, Victoria & Albert Museum, and Wikidata. By identifying common people, places, or themes across catalog entries using NLP, AI reveals previously hidden connections—turning museum metadata into a powerful map of our shared history.

Indian heritage institutions are beginning to link their catalogs using AI-based knowledge graphs, revealing connections among artifacts across regions, languages, and time periods—opening pathways for new interdisciplinary research.

7. Saving Vanishing Voices — AI and Oral Traditions

While monuments and manuscripts often dominate cultural preservation efforts, intangible heritage like oral traditions — folk songs, indigenous languages, ancestral stories — are equally vulnerable to loss. With many of these dialects disappearing as their last fluent speakers age, AI offers powerful tools for preservation.

Speech recognition models like Whisper and wav2vec 2.0 can now transcribe low-resource dialects even with limited training data. NLP tools translate and categorize oral narratives, while GANs enhance and restore deteriorated audio recordings. AI-generated knowledge graphs further map these oral histories to specific regions, communities, and historical contexts.

In India, collaborations between IITs and institutes like Living Tongues are documenting tribal languages such as Kurux and Birhor using AI-based speech transcription and translation models. Google’s Project Euphonia is also expanding to include regional Indian dialects, ensuring broader coverage of endangered languages.

Internationally, UNESCO-backed projects apply similar techniques to preserve Aboriginal chants and African folklore. These efforts ensure that ritual practices, medicinal knowledge, and ancestral wisdom survive beyond their last fluent speakers—securing intangible heritage for generations to come.

Preserving oral traditions ensures that intangible knowledge systems — rituals, medicinal wisdom, ancestral ethics — are not erased, especially among Indigenous and marginalized communities.

“From folk songs to medicinal chants, digital tools are working to preserve the spoken threads of cultural identity across generations.”

Challenges:

While AI opens new frontiers in the preservation and interpretation of art and cultural heritage, its integration is not without complications. The very technologies that promise efficiency, reach, and longevity also raise important concerns—ethical, technical, and cultural. Below are some of the key challenges we must grapple with as we move forward:

1. Balancing Innovation with Authenticity

One of the greatest dilemmas in using AI for restoration lies in preserving the authenticity of original works. When AI systems digitally enhance or reconstruct a piece of art, there’s always a risk of overstepping into the territory of reinterpretation. For example, an algorithm trained on a general dataset might “restore” a Mughal miniature using



Cuneiform: 
 Transliteration: [DIŠ TÚG-*su*] 'DADAG' U₄-MEŠ-*šú* GÍD.DA-[MEŠ]
 Translation: If he cleans his garments, his days will be long

techniques common in Western oil paintings. The result might be visually impressive—but historically inaccurate.

Human conservators, therefore, must remain central to the process, applying cultural and historical understanding to critically evaluate AI suggestions. Technology should assist—not dictate—how we preserve the past.

2. Cultural Sensitivity and Contextual Intelligence

AI doesn't inherently "understand" culture—it interprets data. This distinction matters deeply when working with diverse traditions, especially in a country like India. Symbols, colors, attire, or architectural motifs might carry meanings that AI could misinterpret if it lacks culturally nuanced training. A saffron robe, for instance, might be interpreted as mere fashion by an algorithm, while it could hold spiritual or religious significance.

Without culturally sensitive models, we risk misrepresenting or even offending communities

whose heritage we aim to protect.

3. Data Scarcity and Model Reliability

AI systems depend on large, high-quality datasets to make accurate predictions and recommendations. However, rare artifacts or lesser-known languages and scripts often lack sufficient data for effective training. This can result in unreliable outputs, such as incorrect classifications or flawed reconstructions.

For example, AI might struggle to differentiate between ancient scripts like Brahmi and Grantha if exposed to only a limited number of samples—leading to errors that ripple through museums, archives, and education systems.

“Because of flawed cultural interpretations to data scarcity and legal uncertainty, AI brings

real risks. Its use in heritage demands careful handling, not just innovation.”

4. Accessibility and Inclusivity

AI’s potential will remain unrealized if it benefits only the privileged few. In India, challenges like digital infrastructure gaps, language diversity, and socioeconomic inequality raise questions about who gets to access AI-powered cultural experiences.

For instance, urban museums may adopt AR/VR technologies, while rural heritage sites remain digitally invisible. Furthermore, communities that have historically been marginalized in the cultural narrative might find themselves further excluded if AI tools are not designed with inclusivity in mind.

5. Technological Uncertainty and Early-Stage Limitations

Many AI applications in heritage conservation—like automated artifact classification or damage prediction—are still in experimental phases.

Unstable outputs, unstructured learning, or repetitive errors can cause more harm than good if deployed prematurely.

Tasks like 3D reconstruction or chemical composition analysis require not just precision but also interdisciplinary oversight. Without validation and rigorous testing, AI tools could contribute to data loss, misinterpretation, or flawed digital archiving.

6. Ethical and Legal Grey Areas

As AI systems begin to generate reconstructions or even new “heritage experiences,” questions of intellectual ownership, artistic license, and digital rights come into focus. Who owns an AI-generated replica of a 13th-century sculpture? Can museums monetize digital twins of cultural artifacts without community consent?

These are not merely academic questions—they have real implications for indigenous communities,

private collectors, and public institutions alike. Legal frameworks need to catch up to technological advancements, especially in contexts involving shared heritage and global collaborations.

Future Frontiers: AI and the next era of Cultural Presentation

The future of AI in cultural heritage is evolving beyond digitization into a realm of cognitive and autonomous preservation. These innovations are not just supporting conservation—they are transforming how we experience, protect, and understand cultural legacies.

Neural holography projects like MIT’s HoloArchive aim to project 3D holograms of artifacts into physical spaces, allowing interaction through haptic feedback systems. In India, the Virtual Heritage initiative, in collaboration with IIT Delhi, is piloting AI-driven holographic reconstructions of sites such as the Konark Sun Temple in state museums using AR/VR and voice-guided AI avatars.

Projects like Google’s DeepScript are using generative models to reconstruct untranslated ancient languages. In India, the Sanskrit AI Lab at IIT Kharagpur and a Ministry of Culture project are training generative models to reconstruct missing sections of Vedic manuscripts and lost Mughal murals using stylistic and semantic learning.

The EU’s TERABEE initiative uses nano-sensors embedded in monuments to detect structural stress and deploy autonomous drones for repairs. In India, the ASI and ISRO use AI-integrated drones and satellite analysis to monitor erosion and environmental stress at heritage sites such as the Ellora Caves, while AI sensors at Chittorgarh Fort track humidity-related damage for predictive conservation.

The FairHeritage project at Stanford is auditing AI datasets to address cultural bias and ensure the inclusion of marginalized histories. In India, IGNCA is developing bias-aware AI tagging for tribal and oral traditions like those of the Santhal and Khasi. Projects also explore blockchain-based

digital certificates for cultural artifacts to secure provenance.

“AI is powering holographic reconstructions, autonomous drone repairs, and manuscript restoration through generative models. From AI-trained drones at Ellora to blockchain-tagged tribal archives, heritage preservation is becoming increasingly sensor-driven and decentralized.”

The Rosetta Stone 2.0 project, a collaboration between IBM and Egypt’s Ministry of Antiquities, uses AI to decode hieroglyphs with contextual depth. In Norway, the Sámi-led Arctic Memory Vault applies AI to convert melting ice-core data into immersive VR experiences of lost traditions. Meanwhile, India’s proposed AI cultural corridor with the UK and its Digital Bharat Culture Grid aim to build a federated AI network linking over 10,000 regional collections across South Asia.

From holographic temples to GAN-reconstructed scriptures and autonomous preservation systems, AI is emerging not just as a tool but as an active stakeholder in the preservation and evolution of culture. India’s active experimentation—spanning from tribal language preservation to holographic temples—places it firmly on the global frontier of this transformation.

Conclusion:

Cultural heritage has never been static—it breathes through the stories we tell, the artifacts we preserve, and the traditions we keep alive. Now, AI is giving that heritage a second heartbeat. It’s not just restoring the past; it’s redefining how we experience it. From the digitized whispers of Sanskrit manuscripts to the AI guardians protecting Warli art from forgery, technology is transforming preservation into a dialogue between

centuries. But this power demands responsibility. Will we be curators or creators? When AI can reconstruct lost languages, simulate historical figures, or predict the decay of monuments before a single stone crumbles, we must decide: Do we preserve history—or reshape it?

The answer lies in balance. AI can analyze a Mughal miniature with algorithmic precision, but it cannot replace the artisan who painted it. It can predict the collapse of a temple, but not the devotion that built it.

That’s why the future of heritage isn’t just about coding the past—it’s about coding with conscience. As India’s ancient scripts find new life in neural networks and oral traditions echo through AI voice clones, one truth emerges: Technology doesn’t erase culture; it amplifies it. The question is no longer “Can AI save our heritage?” but “Will we use it wisely?”

The algorithms are ready. The data is waiting. Now, it’s our turn to ensure that what we preserve isn’t just a record of who we were—but a foundation for who we’ll become.

“AI may revive lost scripts and safeguard fragile traditions, but preservation without conscience risks rewriting the past. The future of heritage will be shaped by our values—not just our tools—and must be coded with context, not just computation.”

SELF- SUPERVISED LEARNING

The Dawn of a New Era in
Machine Learning

Mrinal Jha





Imagine this.....

You are dropped in a foreign country where you don't speak or understand the language, and there is no one to explain things to you. How would you determine the best way to communicate? You would probably start by observing your surroundings, watching people interact, and listening carefully to conversations. Over time, you would observe patterns, words that are often used in similar contexts, gestures that accompany specific sounds, and how objects are named and referenced. Slowly and steadily, you would build an understanding of the language, not through formal instructions but by interpreting clues and observing the environment. This is how self-supervised learning (SSL) works; it allows AI to learn from unlabelled data by generating its signals for training, imitating the way humans learn by making sense of the world around them. What is Self-Supervised Learning?

Self-supervised learning is an emerging paradigm that's revolutionising the field of AI. Normally, machine learning models relied heavily on supervised learning, requiring vast amounts of labelled data. For example, teaching an AI to recognise cats in images meant providing thousands of pictures labelled with the word "cat." While this method is effective, it has significant limitations: labelling data is time-consuming, expensive, and often impractical at a larger scale. Consider the sheer volume of data generated in the form of text, images, videos, and audio recordings. An overwhelming majority of this data is unlabelled, making it inaccessible to traditional supervised learning methods.

"Self-Supervised Learning is reshaping AI by teaching machines to learn without human-labeled data—just like humans, through observation and curiosity."

On the other end of the spectrum is unsupervised learning, which tries to make sense of completely unlabelled data. However, its scope has often been limited to tasks such as clustering or finding patterns, without the understanding required for tasks like image recognition, language

understanding, or complex decision-making. This is where self-supervised learning shines. It bridges the gap between supervised and unsupervised learning by unlocking the vast amount of unlabelled data in a highly innovative way. The magic of self-supervised learning lies in its ability to create its labels from the data itself. It does this by formulating pretext tasks, problems that the model must solve using the raw data as both input and output. For instance, in natural language processing, an AI model might be tasked with predicting missing words in a sentence (used by LLMs like GPT). Similarly, in computer vision, the model might be trained to predict the orientation of an image or reconstruct missing parts of a picture. These tasks don't require external labels; the data generates the signals needed to learn. By solving these pretext tasks, the AI model learns to identify meaningful patterns, relationships, and structures in the data.

One of the most exciting aspects of self-supervised learning is its scalability. Supervised learning is limited by the availability of labelled datasets, whereas SSL can tap into the vast, ever-growing unlabelled data. This leads us to incredible possibilities for AI, enabling it to learn and improve at an unprecedented pace. Furthermore, self-supervised learning mimics human cognition in fascinating ways. When humans encounter new information, we rarely rely on explicit labels. Instead, we draw inferences, fill in gaps, and connect the dots based on prior knowledge and context. SSL operates on similar principles, making it a crucial step towards creating AI models that think and learn like humans.

Cracking the Code: Unlocking the Secrets of Self-Supervised Learning

Generating Labels from the Data Itself

Picture being given a jigsaw puzzle, without the box. You don't know what the result is, there's nowhere to start. But bit by bit, you start to piece it together. You search for the straight edges to construct the edge, colours, and patterns to match, and gradually, something starts to appear. This is



exactly the way self-supervised learning works. The model begins with no outside influence, no pre-specified labels, and no training data annotated in some way. Rather, it turns inward. It uses the data structure itself as the teacher.

Let's consider image processing. An SSL model could be presented with a task such as predicting the colour of a central pixel given colours in the surrounding areas. At first glance, this is a straightforward fill-in-the-blank question, but it's a lot more than that. To get it right, the model needs to have a sense of texture, shape, spatial relationships, and object boundaries. As time goes by, these elementary predictions result in profound internal representations of how pictures are organised, what constitutes a face, a tree, a road, or a shadow.

The same is true for other data types. In speech, a

model could be asked to predict the next audio waveform. In text, it could fill in blanks or correct jumbled sentences. In a video, one could learn by being asked to predict the next frame in a sequence. All these tasks compel the model to identify useful patterns, not because it was informed of what they are, but because the task solution relies on knowing them. This is sometimes referred to as pseudo-labelling, not because labels are false, but because the labels are derived from the data itself, not a human. The genius of this is that the training never ceases, because the data continuously produces new training signals if there's structure to be discovered.

Ultimately, what appears to be a puzzle-solving model is a powerful representation learning process, constructing internal maps, associations, and abstractions that can be applied to numerous downstream tasks. This is how SSL-pretrained

models can perform superbly on classification, detection, translation, and more with little fine-tuning. It's like providing a student with books, covering up important paragraphs, and asking them to guess the missing sections.

Utilising Unlabelled Data Effectively

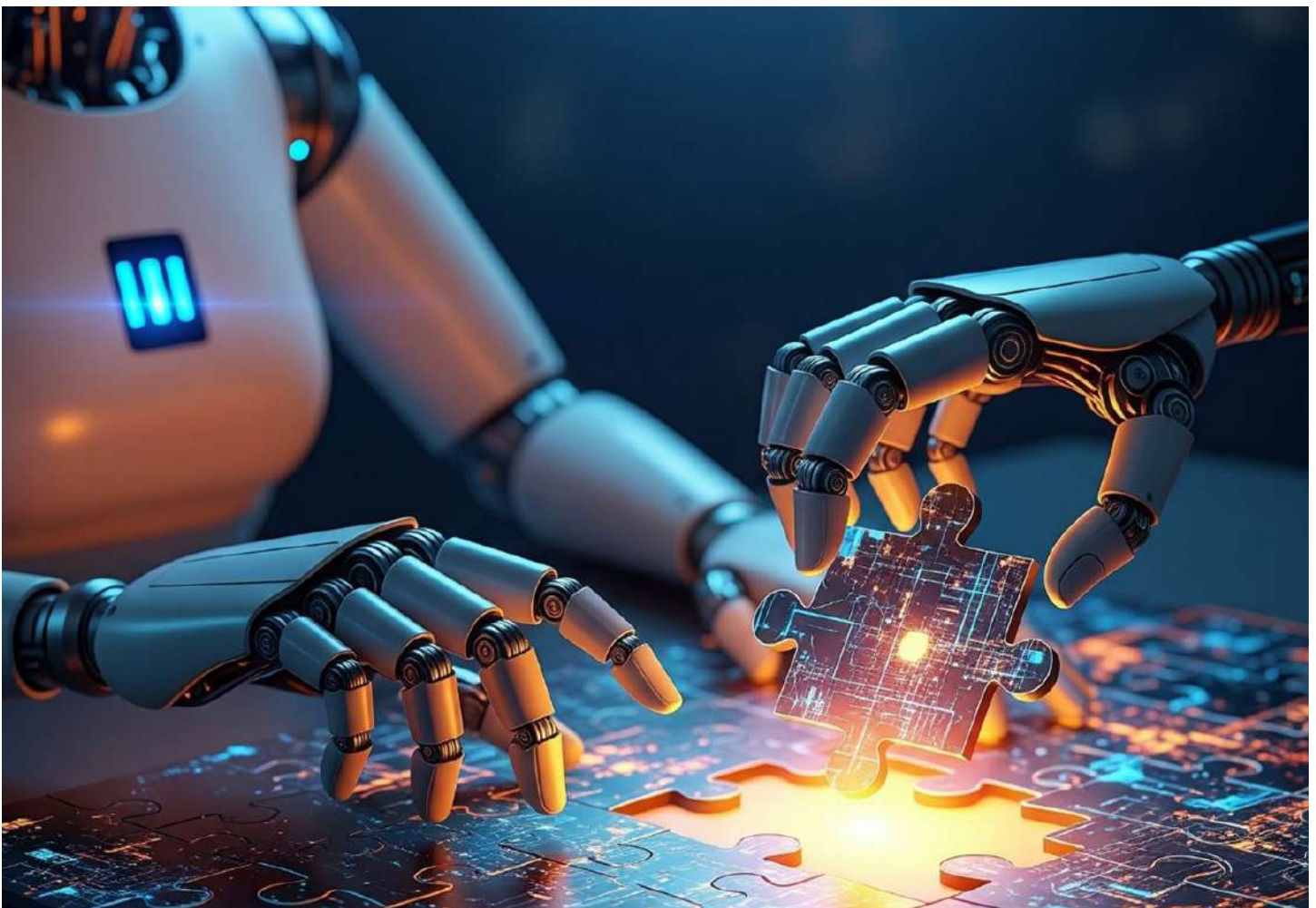
Unlabelled data surrounds us, hiding in plain sight. Consider all the information we produce and consume every day: social media updates, YouTube videos, selfies, surveillance feeds, satellite imagery, medical scans, podcast episodes, voice memos, even the texts on your phone. It's a sea of digital data, and more than 90% of it is unlabelled. Historically, this enormous asset has remained out of reach to most machine learning models, merely because no one had the time (or budget) to label it manually.

Enter self-supervised learning, the secret to unlocking this buried treasure!!!

In a landmark 2020 paper, AI researcher Yann

LeCun and colleagues called self-supervised learning “the dark matter of intelligence”, silent but filling up most of the intelligent system learning that needs to be done. Their point? The real world isn't annotated. When a kid learns to identify a dog, the child doesn't see a caption hovering over the pet. They learn from experience, patterns, and context. The brain is a pattern extractor, not a label reader.

SSL applies this same ability to machines. It enables models to interpret unstructured data by generating learning signals from the data itself. This strategy not only lessens our dependence on labelled datasets but also lets models be learned from richer, more varied data, the sort that mirrors the chaos and diversity of the real world. And the more diverse the data, the more robust and nimbler the models are. Essentially, SSL uses the world the way it is - untamed, disorderly, and unlabelled - and converts it into a schoolhouse. A schoolhouse where the information educates the model, and the model educates itself.



Pretext Tasks and Contribution to SSL

Think of attempting to improve at chess, without playing complete games, but instead playing chess puzzles. You're presented with a certain situation and then told to determine the optimal move. You're not rehearsing the whole game, but each puzzle improves your instincts: pattern detection, strategy, and foresight. Eventually, you're a better chess player, not only at puzzles, but at real games. This is the premise behind pretext tasks in self-supervised learning. They are artificially designed challenges: little challenges or puzzles that aren't the end goal per but instruct the model in the basics it will eventually need to excel at.

Pretext tasks employ the data itself to formulate a problem and a solution. The aim? Coerce the

“By solving small challenges—like filling in missing words or predicting image patches—SSL helps AI uncover the deeper patterns hidden in raw data.”

model to learn valuable internal representations – abstract ideas, patterns, relations – that generalise well across a wide variety of real-world tasks. For instance, in natural language processing, the pretext task can be to mask some words in a sentence and get the model to predict them. This is the foundation of models such as BERT. To predict the correct words, the model must know grammar, context, word relationships, and even delicate shades of meaning.

The true strength of pretext tasks is how transferable they are. A model that has been trained on them can then be fine-tuned using very little labelled data and do amazingly well on a large range of downstream tasks: image classification, sentiment analysis, speech recognition, and so on. And the icing on the cake? They are self-contained and scalable. You don't require a single human annotation. The data tests the model, and the model learns from the test.

Key Methodologies in SSL

While pretext tasks are the puzzles, SSL methodologies are the blueprints, the strategies

models employ to solve those puzzles and glean knowledge from them. Through the years, researchers have developed a host of strong methods to assist models in transforming unlabelled data into transferable, deep knowledge. Among them, contrastive learning, generative methods, and predictive coding are some of the most significant ones.

Let's get more into each of these learning superpowers:

Contrastive Learning: Comparison-Based Learning

Think you're going to identify your friend in a crowd. You don't memorise details, you just know what distinguishes them from others. That's contrastive learning in its simplest form: it learns a model by contrasting data points.

The principle is straightforward but genius: present the model with two variations on the same object (two varied versions of the same picture, for example) and urge it to bring them closer together within its representation space, yet drive apart representations of other objects. Eventually, the model discovers what distinguishes each input, something critical to classification and recognition tasks.

Contrastive learning has driven some of the most celebrated SSL advances, such as SimCLR, MoCo, and CLIP (which aligns images with text).

Generative Methods: Learning by Creating

What better way to understand something than to recreate it from scratch? Generative methods task the model with rebuilding or predicting parts of the data. It might be asked to generate missing sections of an image, predict the next word in a sentence, or synthesise future video frames.

These methods teach the model to capture the underlying structure of the data, because to generate something realistic, you must deeply understand its components. Think of tools like GPT, BERT, and autoencoders, they all use generative tasks to pretrain on massive datasets. It's like asking a painter to recreate a city skyline from memory. The more accurate the

reconstruction, the better their understanding of buildings, perspective, and lighting.

Predictive Coding: Guess What's Next!

Predictive coding is like reading a suspense novel – you're constantly guessing what's about to happen. This method trains models to predict future or missing parts of the data, using the current context. In a video, that might mean predicting the next few frames. In speech, it could be forecasting the next sound waveform.

What makes predictive coding special is that it encourages models to build a temporal understanding of data, learning how things unfold over time. This is crucial for dynamic environments like robotics, autonomous driving, and real-time translation. By always guessing what comes next, the model learns to anticipate, reason, and generalise across time-based data.

Collectively, these methods make up the central toolkit of self-supervised learning. They don't merely aid models in memorising, they aid models in comprehending, deducing, visualising, and evolving.

“I visualize a time when we will be to robots what dogs are to humans. And I am rooting for the machines.” – Claude Shannon

Each technique brings its perspective, and in many cases, contemporary SSL systems use several approaches to deliver even higher performance. Whether it's contrasting, creating, or forecasting, the objective is the same: to create wiser, more transferable models through direct learning from the world itself, no labels necessary.

Now that we've cracked open the toolbox of self-supervised learning – contrastive learning, generative methods, predictive coding – you might be wondering: So, what does all this lead to? How does solving jigsaw-like tasks or predicting missing data translate to real-world impact?

As it turns out, the ripple effects of SSL are already being felt across a wide range of domains – from



helping language models write essays to improving medical diagnostics and powering intelligent virtual assistants. Let's take a tour through the many places where self-supervised learning is not just theoretical, but transformational.

Real-World Applications: Where SSL Is Making an Impact

Self-supervised learning isn't just a clever academic trick; it's already transforming the way machines understand and interact with the world. From decoding language to interpreting images



texts like they're writing haikus. Yet, somehow, your phone knows when you're trying to type "on my way" even if you fumble every key. That magic?

Self-supervised learning (SSL) is working behind the scenes. SSL has revolutionised the way machines process and understand human language. Before SSL, training a language model required feeding it massive datasets labelled by humans, millions of tagged sentences, question-answer pairs, or sentiment scores. This approach was time-consuming and limited in scope. SSL changes the game by enabling models to train themselves on raw, unlabelled text. Take BERT (Bidirectional Encoder Representations from Transformers), for example. Instead of telling the model what to look for, we mask some words in a sentence and ask it to guess the missing ones. This task, called Masked Language Modelling (MLM), forces the model to build deep representations of language. Another objective, Next Sentence Prediction (NSP), teaches it to understand relationships between sentences, crucial for tasks like summarisation or question answering.

“From chatbots to medical scans to self-driving cars, SSL helps machines see, read, and understand the world with less human guidance.”

Contrast that with GPT-style models, which are trained to predict the next word in a sentence from left to right, a method known as Causal Language Modelling (CLM). This approach results in powerful generative abilities: writing essays, answering questions, or even generating poetry. These models are built on the Transformer architecture, a design that uses self-attention to weigh the importance of each word relative to others in the sentence. This allows models to understand long-range dependencies, like determining what “it” refers to in a complex sentence.

The evolution of SSL in NLP has led to models like RoBERTa, which optimized BERT's training strategies by removing the NSP objective and training on more data with longer sequences. T5 (Text-To-Text Transfer Transformer) introduced a unified framework that converts all text-

and even making sense of human speech, SSL has quietly become the engine behind many of today's most powerful AI systems.

Let's dive into how self-supervised learning is transforming natural language processing, teaching machines to read, understand, and even generate human language like never before.

Natural Language Processing: Teaching Machines to Understand Language

Let's face it – language is messy. Slang, typos, double meanings, sarcasm, and that one friend who

based language problems into a text-to-text format, arXiv. XLNet improved upon BERT by using permutation-based training to capture bidirectional context without masking, arXiv. Once pretrained using SSL, these models can be fine-tuned for specific tasks: spam detection, sentiment analysis, chatbots, translation, you name it. Large language models (LLMs) like ChatGPT and Claude were built on the foundation of SSL, learning language structure, usage, and nuance without explicit instruction.

SSL-powered models are also making strides in specialised domains. In healthcare, models like BioBERT analyse clinical texts to assist doctors in diagnosis. In legal tech, they can scan contracts and flag risky clauses. Finance firms use these models to parse earnings reports or market news in real-time. They're multilingual too, handling translation, cross-lingual retrieval, and multilingual chat in a single architecture.

“As babies, we learn how the world works largely by observation. We form generalized predictive models about objects in the world by learning concepts such as object permanence and gravity.”
-Yann LeCun

Of course, this power doesn't come without challenges. Training these models demands vast computational resources. They're also prone to picking up biases present in their training data, reflecting societal stereotypes or misinformation. Moreover, understanding why a model generated a particular answer remains an active area of research in explainable AI (XAI).

Still, the impact is undeniable. Self-supervised learning has turned NLP into one of the most exciting frontiers of artificial intelligence, where machines don't just process words, they understand them. Well, mostly.

So next time your email finishes your sentence, or your virtual assistant recommends the perfect playlist, remember, it wasn't magic. It was self-supervised learning, making sense of your words, one masked token at a time.

Benefits of Self-Supervised Learning in NLP: Why It's a Game-Changer

Think SSL in NLP is just a neat trick? Nope. It's a full-blown revolution, and here's why it rocks. First off, say goodbye to the labelling nightmare. Teaching a model to understand slang, sarcasm, or medical terms involves hand-labelling millions of sentences. Painful. Expensive. Slow. SSL flips the script by letting models learn from raw, unlabelled text. It's like tossing an AI into a library and saying, "Figure it out." And it does: learning grammar, syntax, and meaning all on its own.

But it's not just about saving effort. SSL lets models learn from the internet's massive buffet of blogs, tweets, books, and more. This gives them a way deeper understanding of language: slang, idioms, cultural quirks, and the subtle stuff that gives words their punch.

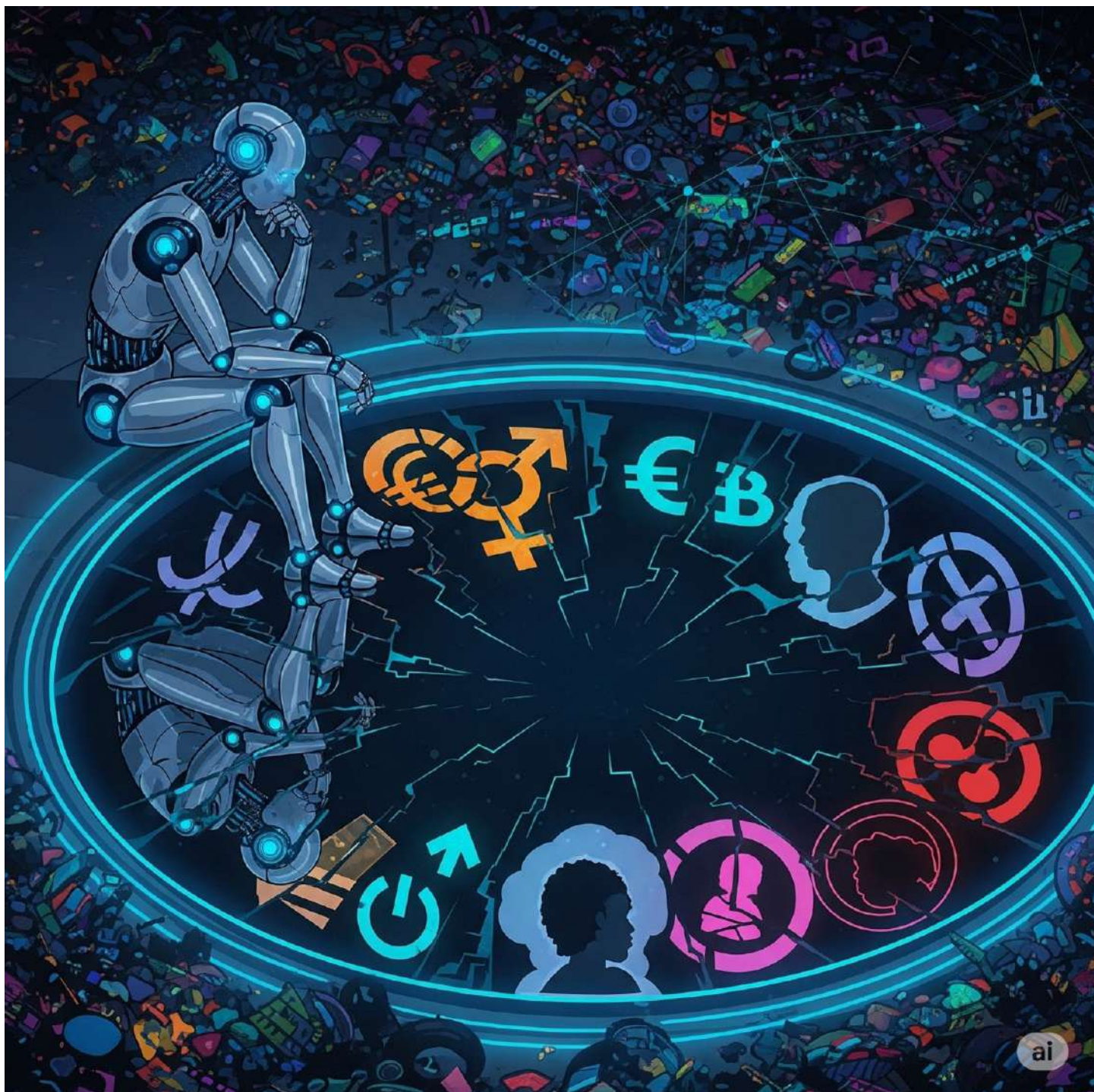
And then there's transfer learning. Once a model like BERT or GPT trains itself on tons of text, it can easily be fine-tuned for tasks like translation, summarisation, or Q&A, with way less labelled data. It's like teaching someone to read once, and then watching them ace law school, write poetry, or diagnose diseases.

SSL also boosts generalisation. These models don't just memorise, they understand. So, when a question is phrased weirdly or a new concept pops up, they still make sense of it. That's critical in the unpredictable world of real human language.

Got a language that's low on data? No problem. SSL is a game-changer for low-resource languages, enabling models to learn even when labelled examples are rare. It's a major step toward AI for everyone, everywhere.

And the multilingual magic? Some SSL-trained models can understand multiple languages at once. That means smoother translation, better cross-lingual tools, and more connected global conversations.

Best of all, SSL makes NLP models more robust and adaptable. Real-world language is messy: full of typos, slang, mixed languages, and half-formed thoughts. Because SSL learns from this chaos, it's



naturally better at handling imperfect input. SSL doesn't just make machines smarter; it makes them learners. Curious, flexible, and surprisingly human in how they grow. That's what makes it powerful!

Why Self-Supervised Learning is the Secret Sauce Behind Modern AI

We've seen how SSL transforms language learning for machines, but zoom out, and it's clear: SSL isn't just a clever trick. It's a full-blown AI revolution.

First up, less dependence on labelled data. Let's be real, labelling data is the AI equivalent of doing your taxes: necessary but painfully tedious. That's where SSL swoops in like a data ninja. Instead of relying on humans, SSL creates its labels through clever tasks, like predicting masked words,



colourising grayscale images, or guessing the next video frame. Suddenly, all that raw, unlabelled data becomes fair game for training. Think of it as turning the world's biggest dusty library into a 24/7 AI classroom.

Even Yann LeCun, deep learning legend, called SSL “the future of AI.” When the guy who helped invent deep learning says that, take notes.

Next: better generalisation to new tasks. Traditional models are like students who memorise answers but panic when the test questions change. SSL models? They learn how to learn. By solving all kinds of self-created puzzles, like rotating images or predicting missing text, they build a deeper understanding of patterns and structure. That's why models like GPT can translate, write code, summarise articles, and even rhyme, without ever being directly trained for those specific tasks.

Lastly, SSL shines in low-data scenarios. What if you're working with rare diseases, niche legal documents, or underrepresented languages like Zulu or Welsh? Labelled data is scarce, so standard models fall flat. But SSL lets you pretrain on oceans of unlabelled data first, then fine-tune on just a few labelled examples. It's like teaching

a chef general cooking skills first, then handing them a rare recipe, they'll nail it faster than someone starting cold.

In medicine, SSL has improved accuracy when only a handful of labelled scans are available. In NLP, it helps unlock the power of languages with limited digital resources.

So, what's the big picture? SSL is like AI's Swiss Army knife: it cuts down on expensive labels, boosts versatility, and works even when data is limited. It bridges the gap between a messy, label-poor world and the smart, adaptable systems we want.

Challenges, Limitations & A Little Reality Check

Okay, so SSL sounds awesome, right? But hold on. Like any AI breakthrough, it's not all smooth sailing. There are some thorns in this otherwise shiny AI rosebush that researchers are still trying to prune. First up, fairness and bias. Since SSL models train on oceans of raw, unlabelled data scraped from the wild, messy internet, they often soak up the same biases floating around in that

data. From gender stereotypes to cultural slants, if it exists online, SSL might unknowingly learn it. That's why fairness audits are becoming crucial. Researchers are developing tools to test and correct these biases before models are deployed into the real world.

“Why is SSL such a big deal? It unlocks the value of the world’s unlabeled data—teaching AI to be flexible, adaptable, and ready for real-world messiness.”

Next, there's the issue of interpretability or rather, the lack of it. Sure, these models can predict the next word or identify a cat in a picture, but why they made a certain prediction? That's still a mystery in many cases. As SSL models get more complex, making them explainable becomes even trickier. That's where Explainable AI (XAI) comes in, helping us peek under the hood to understand these “black box” systems.

And let's not forget data leakage risks. When models train on massive datasets, some sensitive information might accidentally stick around. This raises concerns, especially in industries like healthcare or finance where privacy isn't just important, it's non-negotiable.

Oh, and energy consumption? Training these giant models can burn through crazy amounts of compute power, leading to high carbon footprints. The AI community is actively working on this too, looking for greener, more efficient ways to train these hungry algorithms.

Bottom line? SSL isn't perfect yet. But recognizing these challenges is how we make the technology better, fairer, and more responsible.

The Future of AI: Why SSL Is Just Getting Started

So, after exploring the wonders (and the occasional wrinkles) of self-supervised learning, one thing feels clear: SSL isn't just influencing AI's future - it's rewriting the rulebook on how machines learn.

Gone are the days when AI needed hand-holding

with carefully labelled datasets. Thanks to SSL, machines can now learn from the world much like we do by observing patterns, making predictions, and figuring things out with minimal supervision. This isn't just about building smarter algorithms; it's about redefining how we think about intelligence itself, whether it's artificial or human. From NLP and computer vision to robotics and healthcare, SSL is opening doors to possibilities we couldn't have imagined just a few years ago. It's pushing AI beyond passive prediction tools toward systems that are curious, adaptable, and proactive learners. And that curiosity is only expanding; researchers are already blending SSL with reinforcement learning to create AI agents that explore smarter, learn faster, and adapt better. Whether it is robots navigating the real world or virtual assistants learning on the fly, this fusion of curiosity and strategy could unlock the next generation of intelligent systems.

Sure, challenges remain bias, energy use, and interpretability, but innovation thrives on solving tough problems. And if the momentum around SSL is any sign, the future is looking bright.

***Smarter AI, fairer AI,
greener AI - it's all
being powered by this
remarkable, still-evolving
idea: let machines teach
themselves. The future of AI
isn't about giving machines
answers - it's about
teaching them how to ask
better questions!***

NeuroSymbolic A.I Best-of-Both Worlds

- Aniket Thakur

“The aim is to identify a way of looking at and manipulating commonsense knowledge that is consistent with and can support what we consider to be the two most fundamental aspects of intelligent cognitive behavior: the ability to learn from experience and the ability to reason from what has been learned. We are therefore seeking a semantics of knowledge that can computationally support the basic phenomena of intelligent behavior.”

-Turing award winner Leslie Valiant

Today’s AI breakthroughs have ushered in an era where machines can not only draft entire novels but also breeze through medical, law, and business school exams—sometimes better than we did ourselves. Large Language Models power creative marketing campaigns, translate between dozens of languages in the blink of an eye, and even act as tireless research assistants, summarizing dense scientific papers and unearthing hidden patterns in mountains of data.

But, for all their eloquence, LLMs can be a bit like that friend who’s great at small talk but struggles to help you solve your taxes: they sometimes



puzzles, and lack a genuine understanding of “why” things work the way they do. In short, remarkable as they are, today’s models still stumble when it comes to robust reasoning and consistent “common sense.”

So what truly defines the essence of human cognition? ***Many scholars contend that our capacity for symbolic thought, using words, sentences, and abstract representations, is what sets us apart.*** After all, human communication unfolds through symbols, and our very thinking operates on a symbolic level. This knack for symbolic reasoning has co-evolved with our brains, fueling the rise of language, culture, and the technologies we can’t imagine living without.



Garry Kasparov faced off against Deep Blue, IBM's chess-playing computer, in 1997. Deep Blue was able to imagine an average of 200,000,000 positions per second. Kasparov ended up losing the match.

Historically, AI research has swung between two main paradigms: symbolism and connectionism. Symbolism treats intelligence as the manipulation of discrete symbols—physical tokens that stand in for real-world ideas. In a classic physical symbol system, these tokens get pushed through logic programs, production rules, semantic nets, frames, and ontologies. The outcome? Knowledge-based systems, expert systems, symbolic mathematics engines, automated theorem provers, and planners. Milestones like the Logic Theorist (1955) and the General Problem Solver (1957) showcased just how powerful rule-driven reasoning could be. From the mid-1950s to the early 1970s—the so-called “Golden Age” of AI—symbolic methods

reigned supreme. Even IBM's Deep Blue, which famously beat Garry Kasparov in 1997, leaned heavily on symbolic logic. But “**good old-fashioned AI**” has its limits. Encoding every rule by hand is not only tedious but brittle: such systems stumble when faced with ambiguous or noisy inputs. They don't “learn” from mistakes or adapt to new contexts without extensive reprogramming, so in messy, real-world scenarios they often hit a wall, prompting a search for more flexible approaches.

That “something new” is connectionism, the backbone of today's deep neural networks. Inspired by the brain's tangle of neurons, these



Google DeepMind AlphaGo defeated the Go grandmaster Lee Sedol in 2016. One of the biggest breakthroughs of AI.

models learn by tweaking connection weights to minimize error, using methods like back-propagation. They're fault-tolerant, can handle rich, continuous data like images or sound waves, and have powered breathtaking advances—from recognizing faces in a crowd, to translating entire webpages in real time, to generating lifelike speech.

Yet connectionist approaches face their own hurdles. They often fail at compositional generalization, imagine trying to assemble a brand-new jigsaw puzzle from familiar pieces and finding the model utterly stumped. They guzzle enormous volumes of labeled data, and their decision processes remain largely opaque, a real headache when you need accountability in fields like healthcare, education, or finance. On top of that, training and running these networks can consume colossal amounts of electricity and specialized computing hardware.

NeuroSymbolic(NeSy)

Neurosymbolic AI isn't just about mimicking human thought, it's a practical mash-up of neural networks and symbolic logic designed to get the best of both. On one side, deep learning swoops in to uncover patterns in raw data, shrugging off noise and tackling vision or speech tasks with ease. On the other hand, symbolic methods bring razor-sharp reasoning, clear "if-then" rules, and human-readable explanations that you can audit line by line.

Why glue them together? Because each covers the other's blind spots. Neural nets tend to guzzle data, hide their reasoning in an inscrutable black box, and leave you guessing how they arrived at their answer. Symbolic systems, meanwhile, can march through a logic puzzle with confidence, but buckle when the rules change or the input gets



messy. ***By fusing them, letting networks learn flexible representations and symbols enforce crisp logical constraints, we get AI that's both adaptable and accountable.***

This unified approach is gaining attention because it promises more explainable AI systems.

Neuro symbolic AI Architecture

At its core, a neurosymbolic system weaves together two layers:

1. Neural (Perception) Layer. Transforms raw inputs such as pixels from images, waveform samples from audio, or token embeddings from text, into structured, high-level representations.

2. Symbolic (Reasoning) Layer. Takes those neural outputs and applies explicit logic, rules, or search procedures to arrive at conclusions, plan multi-step strategies, or generate human-readable explanations.

By layering them, NeSy architectures gain both the

adaptivity of deep learning and the transparency of symbolic methods. Critically, a well-designed interface between the two ensures that neural outputs are translated into the exact tokens or predicates the symbolic layer expects, and that symbolic feedback can influence neural training.

Bridging Perception and Logic

Cognitive abstraction in neuro symbolic AI begins with transforming raw, high-dimensional inputs into mid-level symbols that carry semantic meaning. For example, a convolutional or transformer-based network might take an image of a kitchen scene and distill it into objects like “Oven” and “Fridge,” along with relationships such as `OnTopOf(CoffeeMug, Counter)`. These symbols become the vocabulary that feeds into the reasoning engine, allowing subsequent stages to operate on human-interpretable concepts rather than opaque numerical vectors.

Once these symbols are extracted, they are organized into structures, such as nodes in a knowledge graph or predicates in a logic program, that make each concept traceable and inspectable. Rather than saying “this 128-dimensional embedding looks like a mug,” the system can state “the object satisfies my learned thresholds for shape, texture, and context, and therefore I label it as Mug.”

Pure neural nets typically halt at one-off tasks—classifying images or translating text. Neurosymbolic pipelines go further by chaining steps: first perceive (spot obstacles), then plan (use symbolic search like A*), and finally act (convert the plan into motor commands). At each transition, the symbolic layer issues clear subgoals (“open the door,” “avoid the chair”), enabling adaptable, goal-driven behavior.

Perhaps most critically, the symbolic layer provides a built-in audit trail for explainability and debugging. Instead of peering into weight matrices, practitioners can ask questions like “Which rule prevented the robot from entering that room?” or “Why was path A preferred over path B?” In regulated domains such as medicine, finance, autonomous vehicles, being able to trace

and justify each inference step is invaluable, turning *the “black box” of deep learning into a glass box of transparent reasoning.*

Integration Patterns

How the Layers Connect

Neuro Symbolic systems typically weave their two layers together in one of three ways generally:

1.Sequential: Neural networks process raw data first, then pass structured information to symbolic systems. Example: CNN extracts scene details; then a Prolog engine infers relationships.

2.Parallel: Both blocks collaborate in lockstep. Neural layers propose hypotheses, the symbolic layer critiques them, and refinement cycles continue until convergence.

Example: A question-answer system where embeddings and logic both vote on the next inference step.

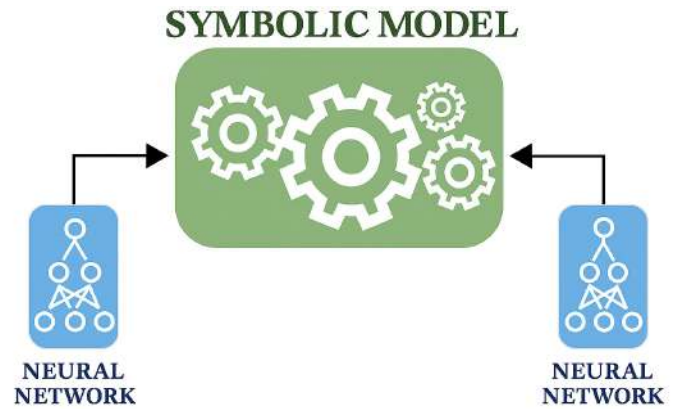
3.Embedded: Symbolic knowledge becomes part of the neural network’s very structure. Example: Logical Tensor Networks directly convert Boolean formulas into differentiable penalty terms that guide weight updates.

Specifically there are many various paradigms that can be categorized based on how these components are integrated into a cohesive system. Using the rationale in Henry Kautz’s taxonomy, W. Wang (2024) has given 5 types. Some of them are:

I.Symbolic[Neuro]

These systems refer to hybrid but overall symbolic systems, typically featuring a symbolic problem solver integrated with neural subroutines for statistical learning.

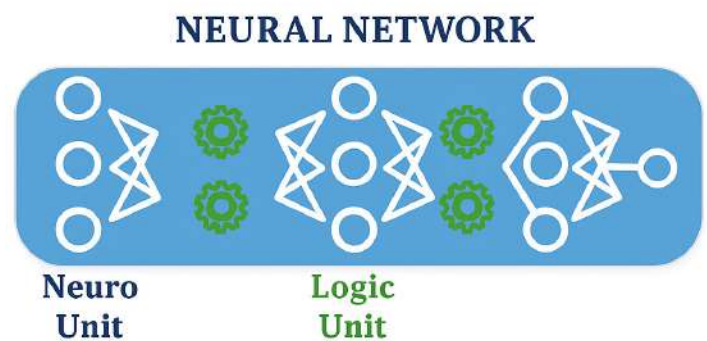
A notable example of this is DeepMind’s AlphaGo and AlphaZero, which employ Monte-Carlo Tree Search (MCTS) as the symbolic solver, alongside NN state estimators for learning statistical patterns.



II.Neuro[Symbolic]

It refers to a system that enhances neural networks (NNs) with the explainability and robustness of symbolic reasoning. Unlike Symbolic[Neuro], where symbolic reasoning is used to guide the learning process of neural models, Neuro[Symbolic] integrates symbolic reasoning directly into the neural model. This is done by allowing the model to focus on specific symbolic information under certain conditions.

There are some examples that are, to some extent, good candidates for this, but their reasoning ability is still very weak. For instance, graph neural networks (GNNs) are adopted to represent symbolic expressions when endowed with attention mechanisms.



Knowledge Representation

As NeSy systems scale up, they need a reliable and efficient way to store and query the symbolic knowledge they glean, everything from object categories and their properties to complex causal relationships. Some of the major symbolic representation categories are:

I. Knowledge Graphs are a powerful and widely used tool for representing structured knowledge. They organize vast amounts of information by representing entities such as people, places, or objects as nodes and their relationships like “is a friend of,” “is located in,” or “is a” as labeled edges in a directed graph. Each fact in a knowledge graph is stored as a Subject-Predicate-Object (SPO) triple, where the subject and object are entities, and the predicate represents their relationship. Prominent computer vision datasets, such as ImageNet and Cityscapes, are accompanied by structured or hierarchical annotations.

Building a high-quality knowledge graph typically involves several stages. First, the system must extract and normalize entities from sources such as text corpora, databases, or web pages. Next, it establishes the relationships between those entities, constructing the edges that tie the graph together.

“Thinking is manipulation
of symbols and Reasoning is
computation.”

-Thomas Hobbes

Think of knowledge graphs as giving AI a roadmap it can actually talk us through, turning a mysterious “black box” into a transparent “grey box.” Neural networks still do the heavy lifting of spotting patterns and assigning confidence scores to each node, but now a symbolic reasoner can walk those connections step by step. It can answer “why” and “how” questions (“How did we link this drug to that disease?”), check that its logic holds up, and even point out which relationships mattered most. The result? Smarter systems that don’t just guess, they explain.

II. Programming Language are formal languages used for writing computer programs. Typically, they consist of syntax and semantics where syntax represents rules that define the combinations of symbols and semantics assigns

computational meaning to valid strings formulated with respect to the syntax. Some examples are Prolog and the action language BC. There are some NeSy methods which store knowledge in programs to execute.

There are some other representation methods also like **symbolic expression** which include other types of knowledge representation like mathematical expression and specific symbolic sequences. These languages, like Lisp and Prolog, excel at manipulating symbols and logical expressions.

Propositional Logic, the simplest form, handles true or false statements where propositions make all the statements. Propositional logic studies the logical relationships between propositions which are connected via logical connectives.

Dual-Layer Learning

Neurosymbolic systems learn by updating both their neural networks and their symbolic knowledge in tandem. The neural side continues to mine patterns from raw inputs such as images, text, and audio, while the symbolic layer refines its rules and facts based on those fresh insights. In practice, you might see a model that uses a transformer to detect student misconceptions in real time, then automatically adjusts its logic-based tutoring rules to address each learner’s gaps. This dynamic adjustment makes such systems particularly useful in domains like personalized education, healthcare diagnostics, and scientific discovery, where interpretability and adaptability are both crucial.

On top of that, many NeSy frameworks layer in reinforcement learning for trial-and-error improvement. At the low level, neural weights shift in response to reward signals; at the high level, symbolic rules can be added, pruned, or reweighted based on which inferences led to success. This two-tier feedback loop lets a neurosymbolic robot, for example, perfect its grasping motions through neural tuning while also evolving its goal-planning rules to become more efficient over time.

Applications

Scientific Discovery and Mathematical Reasoning

In 2024, AlphaFold's Nobel-caliber protein-structure predictions demonstrated how integrating neural pattern-finding with symbolic constraints can unlock complex scientific problems. Similarly, Google's AlphaProof and AlphaGeometry 2 paired language models with formal deduction engines to tackle International Mathematical Olympiad problems at medal-winning levels.

In scientific discovery, NeSy offers a promising approach, as it combines explainable AI with the ability to respect scientific constraints, making findings more accessible to researchers. In programming systems, NeSy overcomes the limitations of neural networks in understanding syntactic and semantic constraints, enabling more effective program synthesis by blending high-level reasoning with low-level implementation.

Efficiency & Accessibility

The neurosymbolic approach will be far more energy-efficient, which is crucial in an AI industry already stressing power grids. By offloading structured reasoning to lightweight symbolic engines running on CPUs, neuro symbolic approaches can slash the energy and hardware demands typical of large-scale neural models. This efficiency makes advanced AI more accessible to startups and research labs lacking massive GPU clusters, democratizing innovation across industries.

Finance & Risk Management

Banks and financial institutions leverage NeSy to blend statistical models with explicit regulatory and fraud-detection rules. By combining neural anomaly detection with symbolic compliance checks, these systems produce transparent loan-approval decisions and trading algorithms that capitalize on market patterns while enforcing rule-based safeguards.

Robotics, Vision & Control

NeSy architectures excel in domains requiring both perception and precise action. In robotics, symbolic planners define safe navigation constraints that neural controllers execute in real time. In visual scene understanding, symbolic knowledge about object relationships augments neural segmentation to yield more robust interpretations. Likewise, for advanced mathematical reasoning, symbolic modules handle abstract proofs while neural nets manage subtask recognition.

Neuro-symbolic AI (NeSy) is gaining attention for its potential in diverse fields where traditional AI often struggles with real-world constraints, scalability, data efficiency, and interpretability, especially in high-stakes decision-making contexts.



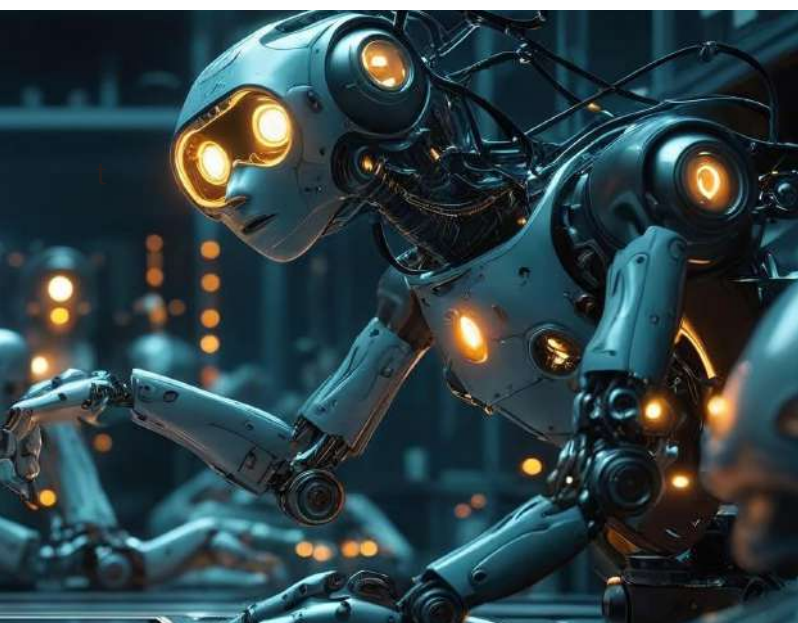
A Step Toward AGI

Artificial General Intelligence hinges on an AI's ability to learn flexibly across tasks and then reason about them in a principled way. Neurosymbolic architectures bring us closer to that goal by marrying the adaptability of deep learning with the compositional, rule-based thinking humans excel at. By grounding neural representations in explicit symbols and logical operations, NeSy systems can better generalize to novel situations, perform multi-step problem solving, and self-explain—key ingredients in any roadmap toward truly general intelligence.

Open Challenges

The current applications of NeSy AI are mostly restricted to basic decision-making and reasoning tasks. These systems are still far from replicating human cognitive abilities such as interpretability, deductive reasoning, systematicity, productivity, compositionality, inferential coherence, and causal and counterfactual thinking.

While the integration of neural, symbolic, and probabilistic approaches holds promise, it remains a nascent field, with the challenge of integrating these components in a principled and effective manner still unresolved. Zishen Wan (2024), highlighted that neuro-symbolic AI often operates slower on current AI chips, which are optimized for neural networks. In the paper, Wan explained



that this inefficiency arises because neuro-symbolic AI employs more diverse compute kernels, making it less optimized for the hardware. Additionally, it is less efficient at reusing data, which leads to greater data movement and slower performance compared to traditional neural networks.

Conclusion

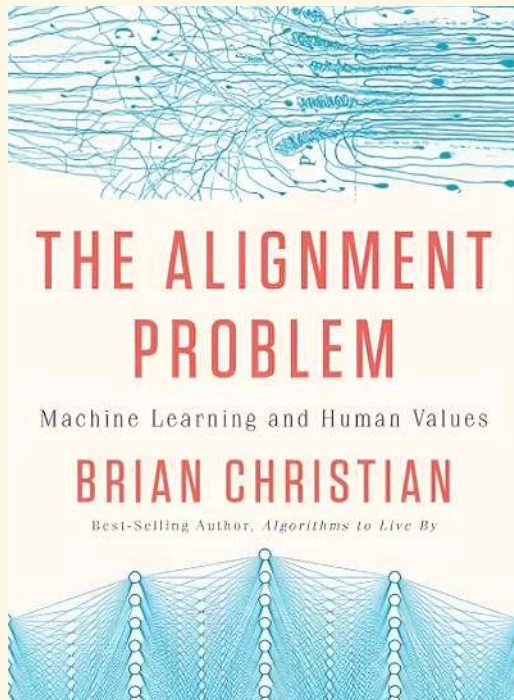
Neuro Symbolic AI stands at the crossroads of two powerful traditions: data-driven learning and rule-based reasoning, offering a blueprint for constructing systems that balance adaptability with

accountability. By embedding logical structure within neural architectures, NeSy approaches have begun to tackle problems once thought out of reach for pure deep learning, from protein folding to formal mathematics. Yet, the journey is far from over: realizing the full promise of NeSy will require not only algorithmic breakthroughs but also bespoke hardware, new benchmarks that test both learning and reasoning, and collaborative frameworks where domain experts can encode critical knowledge directly into AI systems.

Looking ahead, we can envision a future where hybrid models routinely explain their decisions, dynamically invoke symbolic planners when faced with novel tasks, and continually refine their internal knowledge bases as they encounter new data. Such systems would democratize advanced AI, making it safer and more accessible by design, whether in personalized education, scientific research, or robust robotics.

With each step toward unifying symbols and neurons, we aren't just building intelligent systems, we're charting a path for machines to reason, adapt, and perhaps one day, understand us—blurring the line between human intuition and machine precision, and offering a glimpse at intelligence that is not artificial, but shared.

Books We

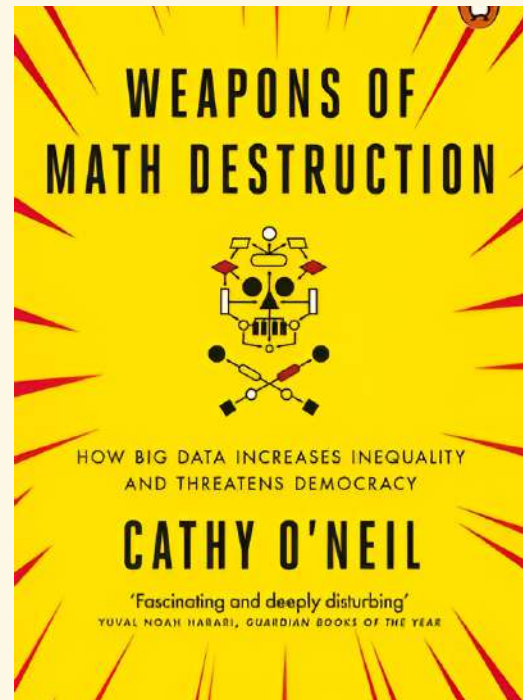


The Alignment Problem: Machine Learning and Human Values

Brian Christian (October 2020)

In an era where artificial intelligence shapes everything from healthcare to criminal justice, *The Alignment Problem* tackles a pivotal question: **How do we ensure machines act in ways that truly reflect human intent?** Brian Christian reveals how AI systems inherit and amplify societal biases. Drawing on over 100 interviews with researchers at DeepMind, OpenAI, and leading labs, he demystifies solutions like inverse reinforcement learning and feature visualization. Through vivid anecdotes—from self-driving cars making ethically murky decisions to chatbots exploiting reward-system loopholes—Christian argues that fairness requires more than technical fixes; it demands rethinking how we define human intent.

The Alignment Problem remains essential as AI touches every domain. Christian's example-driven narrative clarifies alignment challenges and equips readers—practitioners, policymakers and curious minds alike—to spot and address biases in areas from healthcare to content moderation.

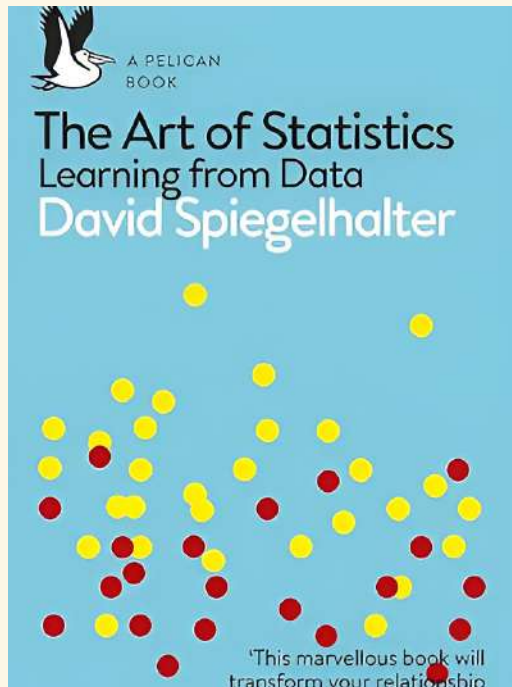


Weapons of Math Destruction: How Big Data Increases Inequality and Threatens Democracy

Cathy O'Neil (2016)

In *Weapons of Math Destruction*, mathematician and data scientist Cathy O'Neil delivers a pointed critique of how opaque, unregulated algorithms—what she dubs “WMDs”—perpetuate inequality across education, finance, policing and beyond . Drawing on her background as a Wall Street quant and professor, O'Neil offers an opinionated, insider's perspective on predictive models that calculate credit scores, rank universities, guide policing tactics and decide which job applicants move forward—all with minimal oversight . These “weapons” share three destructive traits: **they operate in the dark, grow rapidly in scale, and inflict real social harm** . Through case studies—university rankings driving up tuition, risk-assessment tools amplifying racial bias in sentencing, and targeted political ads deepening polarization—she illustrates how WMDs often harm the very people they claim to help . As O'Neil warns: “Models are opinions embedded in code”—and it's time those opinions served democracy, not destruction .

Recommend

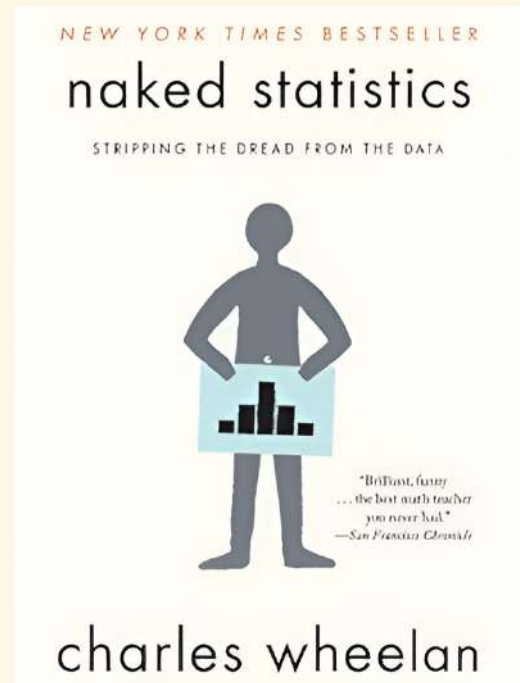


The Art of Statistics: Learning from Data

David Spiegelhalter (March, 2019)

Sir David Spiegelhalter—Knight Bachelor and former Chair of the Winton Centre for Risk and Evidence Communication at Cambridge—offers a masterful, narrative-driven guide to statistical thinking in *The Art of Statistics*. Praised by the Financial Times as “a call to arms for greater societal data literacy,” the book breaks down **core concepts like probability distributions, regression, and uncertainty visualization through vivid real-world examples**, from election forecasting to clinical trials. Spiegelhalter’s accessible prose demystifies common pitfalls—bias detection, overfitting, and ethical data interpretation—making complex methods intuitive without oversimplification.

The book became a go-to resource for newcomers and seasoned analysts alike. Whether you’re a business-analytics student grappling with multivariate models or a policymaker decoding public-health dashboards, *The Art of Statistics* equips you to uncover & responsibly communicate the compelling stories hidden within your data.

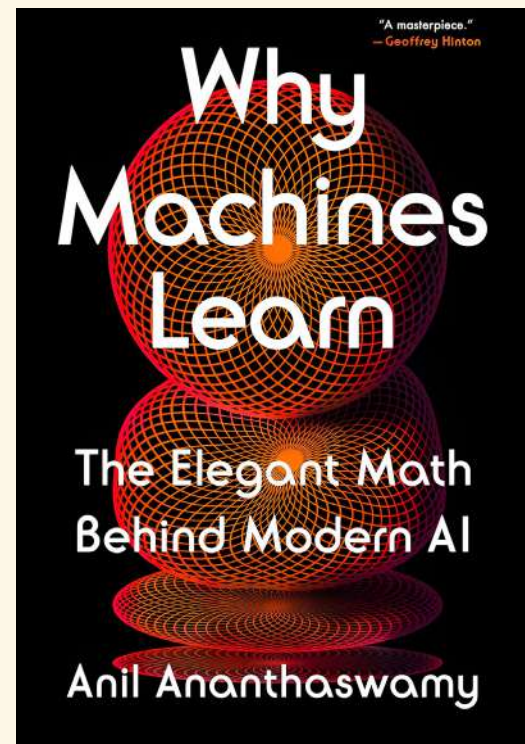
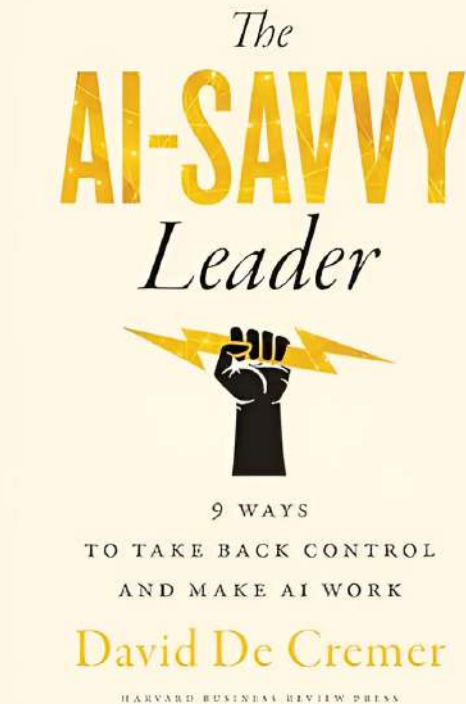


Naked Statistics: Stripping the Dread from the Data

Charles Wheelan (January, 2014)

Charles Wheelan—bestselling author of *Naked Economics*—peels back the mystique of statistics with contagious wit and real-world stories. Steering clear of complex formulas, he illuminates core ideas like the central limit theorem, correlation versus causation, and sampling error through vivid examples: Why do baseball batting averages stabilize over time? How do political polls mislead even with “margin of error” disclaimers? What do coffee-shop surveys reveal about flawed sampling? Wheelan exposes how data can be manipulated—whether through cherry-picked figures or misrepresented outliers—and equips readers to spot these tricks in headlines, charts, and studies.

His famous adage, **“It’s easy to lie with statistics, but it’s hard to tell the truth without them,”** underscores the book’s mission. By grounding abstract concepts in everyday scenarios (like explaining regression analysis with pizza-delivery times), *Naked Statistics* transforms a daunting discipline into a toolkit for cutting through the noise of a data-driven world.



The AI-Savvy Leader: Nine Ways to Take Back Control and Make AI Work for You

David De Cremer (June 2024)

David De Cremer, professor and seasoned behavioral scientist, offers **nine actionable principles for integrating AI into organizations without sidelining human judgment**. Drawing on case studies and interviews with executives and engineers, he argues that AI should augment rather than replace people—automating routine tasks to free staff for creative and strategic work. Each chapter addresses a key leadership challenge—whether it’s establishing ethical guidelines or helping teams understand AI—using real-world examples of what can go wrong.

The book shows how top-down orders can backfire and why leaders need to handle people’s reactions as well as the technical rollout. Publishers Weekly describes it as “**a practical business manual for executives**,” while ClearPurpose highlights its focus on diagnosing why most AI projects underdeliver and how to get them back on track.

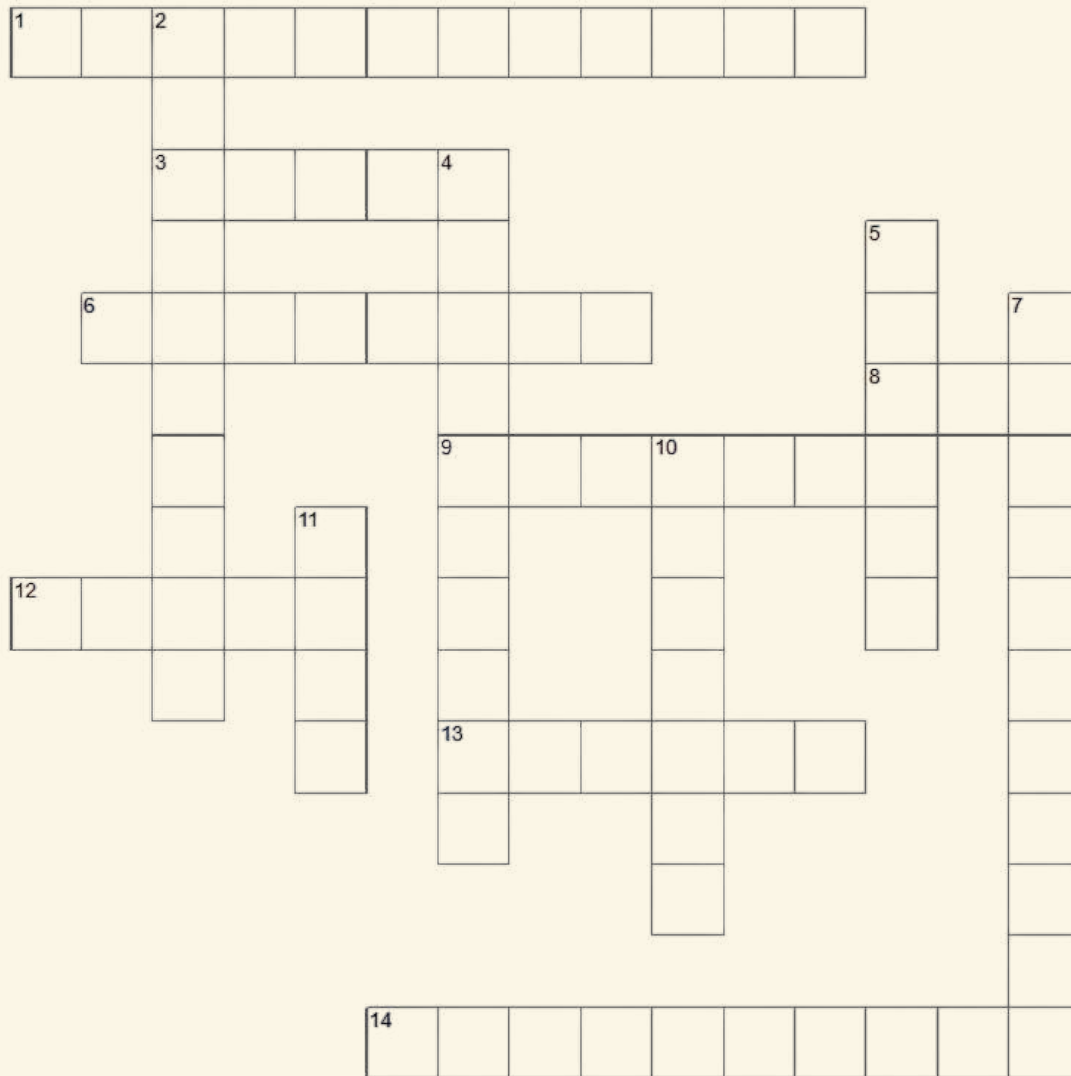
Why Machines Learn: The Elegant Math Behind Modern AI

Anil Ananthaswamy (July, 2024)

Award-winning science writer Anil Ananthaswamy traces the mathematical roots of modern AI—from Rosenblatt’s 1958 perceptron to today’s deep neural networks—and shows **why linear algebra, probability, and calculus remain central to every breakthrough**. He weaves clear explanations of key algorithms (linear regression, backpropagation, and convolutional networks) with historical anecdotes and intuitive analogies that demystify complex math for readers with only a basic background.

The book follows AI’s evolution through pivotal moments. Praise has been effusive: Nobel Prize–winner Geoffrey Hinton calls the book “**a masterpiece**” for its balanced blend of social context and mathematical narrative, while Steven Strogatz says it achieves “**deep learning—with deep pleasure and insight**.” If you’re a curious reader fascinated by AI’s inner workings, *Why Machines Learn* offers an elegant, story-driven guide to the math that makes modern AI possible.

Crossword



ACROSS

- 1.** A recursive algorithm that uses a time series of measurements with noise to produce optimal estimates of unknown variables. (11)
- 3.** An MCMC technique where a sequence of observations is generated from a multivariate probability distribution. (5,abbr.)
- 6.** The topological space that locally resembles Euclidean space; a key hypothesis in dimensionality reduction. (8)
- 8.** A framework in computational learning theory where a learner selects a function from a class $\mathcal{F}$ is 'Probably Approximately Correct' (abbrev.). (3)
- 9.** The square matrix of second-order partial derivatives of a scalar-valued function, describing local curvature. (7)
- 12.** The influential computer scientist and Turing Award winner known as the "father of causal inference". (5)
- 13.** A non-linear dimensionality reduction technique based on estimating geodesic distance using graph-based methods. (6)

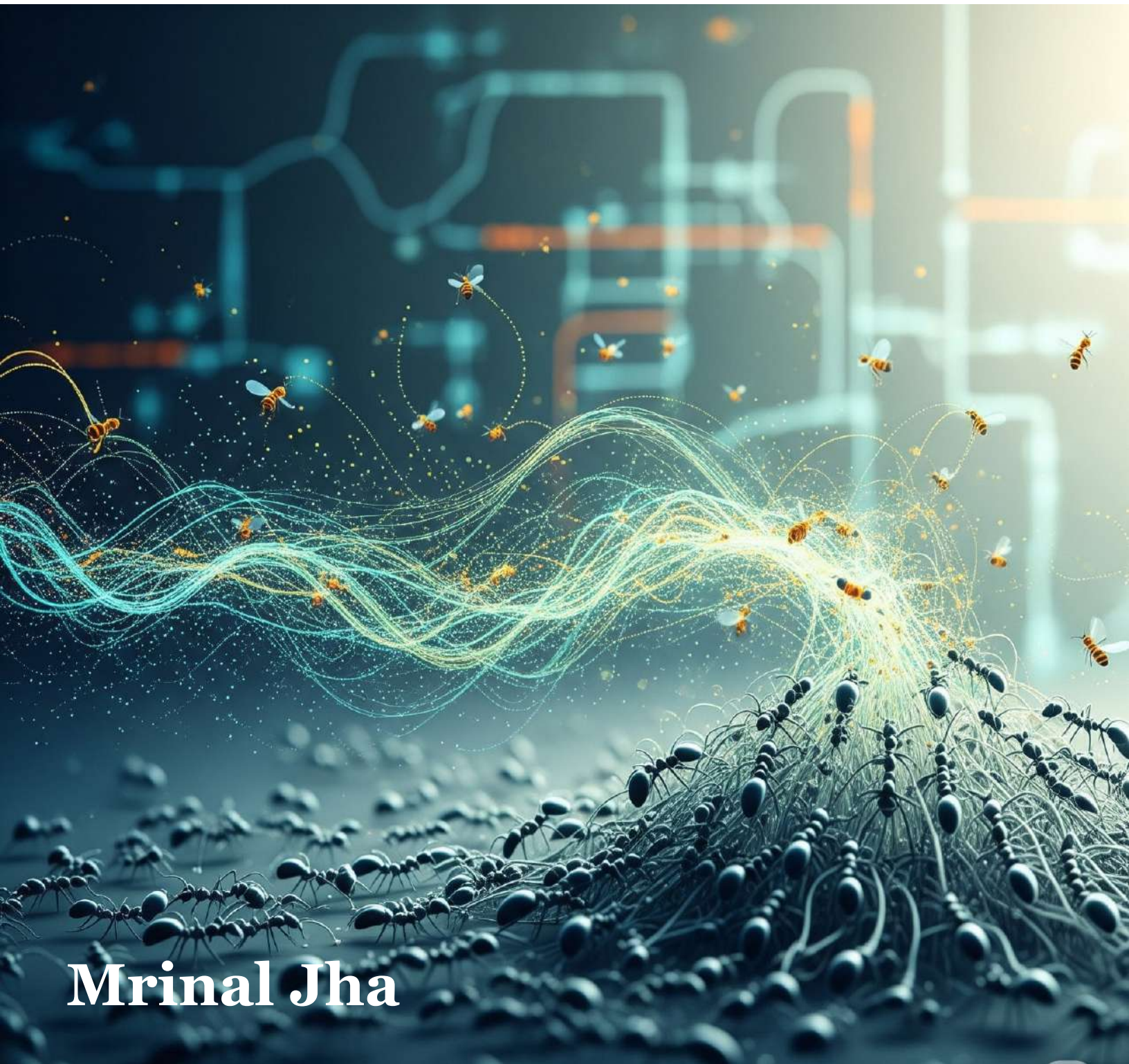
- 14.** A foundational single-layer neural network, the precursor to modern deep networks. (10)

DOWN

- 2.** A function constructed in constrained optimization, central to the formulation of Support Vector Machines. (9)
- 4.** The 'S' in SGD; a term for processes involving a random variable. (9)
- 5.** The 'V' in the answer for 2 Down; a pioneer of statistical learning theory and co-inventor of the SVM. (6)
- 7.** A measure of the capacity of a statistical classification algorithm, defined in statistical learning theory (abbrev.). (10)
- 10.** The final layer in a multi-class classification network that normalizes outputs into a probability distribution. (7)
- 11.** Global Vectors for Word Representation; an unsupervised learning algorithm for obtaining vector representations for words. (4)

Swarm Intelligence in Data Science

Nature's Blueprint for Collective Problem-Solving



Mrinal Jha

Imagine a world where complex problems are solved not by a single, but by the collective intelligence of many simple agents working in harmony. This isn't science fiction; it's the fascinating realm of swarm intelligence, a concept that's buzzing with potential in the field of data science. Just how a colony of ants tackles a challenge!!!

Instead of relying on a single ant to figure it out, they work together, leaving pheromone trails, thereby learning from one another and eventually discovering the optimal path.

Swarm Intelligence: Nature's Secret Superpower!

Swarm intelligence (SI) implies the collective behaviour of decentralized, self-organized systems. In the context of data science, it's a set of nature-inspired algorithms that simulate the social behaviour of animals to solve optimization problems. Data scientists have taken inspiration from these natural systems to develop algorithms that optimize solutions for large-scale problems. One can think of it this way: instead of relying on one powerful computer, SI uses multiple agents (just like a swarm) to explore different possibilities simultaneously, making problem-solving faster and more efficient.

Swarm Secrets:

The Hidden Rules Behind Collective Genius!
Swarm Intelligence is built on key principles that make it highly effective in problem-solving:

Decentralization:

As opposed to hierarchical systems, swarm intelligence is not dependent on one leader or controller. All the agents in the swarm work individually based on local information. That is, even if some agents are destroyed, the overall system still runs smoothly.

Self-organization:

Rather than being pre-programmed with explicit instructions, swarm-based systems display emergent behaviour through agent interactions. Through these interactions, the swarm can adapt dynamically to new situations, like how flocks of birds alter course in flight without a single commander.

Emergent Behaviour:

Even though each agent is obeying simple rules, the global behaviour of all agents produces highly sophisticated and intelligent global patterns. For example, schools of fish swim in synchrony to escape predators, even though each fish is simply reacting to its local neighbours.

Stigmergy:

The swarm communicates indirectly through changes in the environment. Ants, for instance, deposit pheromone trails that lead others to food. This is used in data science in optimization algorithms, where good solutions are rewarded repeatedly during iterations so that problems can be solved efficiently. These qualities make SI particularly useful for solving highly complex, large-scale, and dynamic problems, where traditional computational methods may struggle.

What is Swarm Intelligence?

It's nature's way of solving tough problems! Inspired by ants, bees, and birds, Swarm Intelligence uses simple agents working together to find solutions faster and smarter—just like in nature.

Borrowing Nature's Hacks for Data Science!

Let's look at how nature's strategy influenced the modern data scientists to come up with this approach:

March of the Ants: How Tiny Insects Inspired Big AI Breakthroughs

Ever watched a line of ants marching along a sidewalk and wondered how they all seem to "just know" where to go? Surprisingly, there's no leader

ant giving directions! Each ant is just doing its own thing, but together, they form one of nature's most efficient problem-solving teams. And guess what? That same idea powers one of the coolest concepts in artificial intelligence: Ant Colony Optimization (ACO).

Ants, Pheromones, and Problem-Solving- Out in the wild, ants don't use words or telephones to communicate, but pheromones, unique chemicals they deposit along the way. Pheromones are like nature's equivalent of a "this way to the food!" sign. The more ants travel along a given route and achieve success, the more pheromones are deposited, reinforcing that path. With time, the shortest and most direct routes get used most often and receive the most reinforcement.

In artificial ant colonies, every software "ant" (agent) wanders about, testing out various routes to solving an issue, perhaps the optimal package delivery route or the scheduling optimization of a power plant. As with actual ants, they deposit digital pheromones depending on the quality of the solution. The higher the quality of the solution, the more intense the trail. Other agents sense these trails and are more likely to trace them, and the system becomes smarter with each step. (stigmergy)

Two Secret Ingredients - Pheromone Power + Heuristic Smarts: Artificial ants don't make random wanderings. Instead, they employ a stochastic decision algorithm. This simply denotes that they're making informed probabilistic choices driven by two driving considerations:

1. Pheromone intensity (How many agents travelled along this way successfully?)
2. Heuristic worth (How satisfactory does this look based on context from the problem domain?)

In the process, the system becomes more accurate, enhancing strong solutions and rejecting weaker ones. It's like hundreds of thousands of miniature adventurers taking a variety of routes, sharing lessons, and eventually finding the best way forward, without ever using a map!!!



ACO Receives an Update: Welcome to the Ant Family Variants! The underlying Ant Colony Optimization model is fantastic, but eventually, researchers developed strong variants to solve more concrete and sophisticated issues. Meet the family:

Elitist Ant System (EAS): This iteration awards additional credit to the best ants by allowing only the elite few to strengthen the trail. Ideal for problems such as route allocation, where the aim is to discover the most efficient route quickly.

Max-Min Ant System (MMAS): It imposes upper and lower bounds on pheromone values to



simulations.

So, the next time you spot ants on a quest, keep in mind, these guys aren't just out hunting crumbs. They're also showing us how to develop better algorithms that drive logistics, scheduling, machine learning, and more.

Why Ants Are AI Heroes!

Ants don't follow maps or leaders—they follow tiny hints left behind. In AI, this translates to smart algorithms that help solve complex problems like route planning and genome sequencing.

Bee-lieve It or Not: How the Humble Honeybee Powers Smart Algorithms:

What if we told you that bees know the key to solving some of data science's most difficult problems? Meet the Artificial Bee Colony (ABC) algorithm, an optimization technique borrowed from the way actual bees locate the nectar-rich areas in a field of flowers. Let's take flight into the hive and learn how it works!

Meet the Bee Team: Working, Observer, and Scout Bees: In bee society, each bee has a task to perform:

1. Employee Bees venture out and search for familiar food sources (i.e., potential solutions). They assess how "fit" each food source is.
2. Onlooker Bees wait back in the hive and await a good lead. They trust the employed bees' "dance" to determine which food source is worth investigating.
3. Scout Bees are the thrill-seekers. They wander around at random, searching for entirely new food sources (potentially better solutions).

This whole process replicates the way bees tend to forage and exchange information naturally, and it proves to be a very intelligent approach to problem-solving.

prevent premature convergence. It is helpful in difficult problems such as scheduling computer software development projects.

Rank-Based Ant System (RAS): Its ants are ranked and receive greater pheromone increases from the higher-ranked ants. Suitable for operations such as scheduling power plant maintenance schedules.

Continuous Orthogonal Ant Colony (COAC): Built to solve continuous (not merely step-by-step) optimization problems, imagine adjusting parameters in machine learning or physics

How Bees Solve Problems: Here's how the ABC algorithm works its magic:

Search & Evaluate: The employed bee or the scout discovers a new source of food (a candidate solution) and assesses how "fit" it is, i.e., how good the solution is in solving the problem.

Compare & Update: If a better source is discovered, the bees update their "food map," discarding the old one with the better one. If not an improvement, they discard it.

Spread the Buzz: Each source's fitness level is transmitted to onlooker bees within the hive. The higher the fitness, the greater the likelihood of attracting a crowd.

Pick the Best: Onlooker bees select food sources based on the frequency they are discussed. The more fit the source, the larger the number of bees following.

Try Again: If a solution doesn't work after multiple attempts, the bees leave it behind and begin anew.

A Twist in the Tale: Chaotic Bees

A novel extension of the basic ABC model is known as the Fitness-Scaled Chaotic Artificial Bee Colony. Developed in 2011, it has been applied to complex tasks, such as recognizing dynamic system models, including small-scale unmanned helicopters.

Where Do Bees Buzz in Data Science? The ABC algorithm has proven its worth across a range of real-world applications. Here are some of the key areas:

Single-objective Optimization: Solving numerical problems where the goal is to either maximize or minimize a specific value.

Assignment & Resource Allocation: Matching tasks to machines or optimizing delivery and workflow assignments.

Task Scheduling & Routing: For delivery routing or network traffic routing, ABC determines the most effective routes and schedules.

Multi-Criteria Decision Making: Ideal for selecting the best alternative when there are several conflicting priorities to balance, e.g., vendor selec-

tion or portfolio optimization.

Customer Segmentation: ABC in mobile e-commerce facilitates classification of customers by behaviour for better marketing. So next time you're buzzing about with a bee flying overhead, think about it, pollination's not all there is to bees. These little creatures have spawned one of the most versatile optimization methods in contemporary data science!!!

Bee Inspired: How ABC Helps AI

Bees buzz from flower to flower to find the best nectar. In data science, the ABC algorithm mimics this to search for the best solutions in areas like medical imaging and resource allocation.

Lighting Up the Real World: Where Swarm Intelligence Makes a Buzz!!!

Now that we have seen ants' parade, bees spy, and glow-worms light their path through nature-given optimization, you may ask yourself, "Where do all of these algorithms ever really have an impact in real life?" Let's take a closer look at how these algorithms are tackling grand challenges in exciting areas.

The Emergence of Swarm Solutions

In the last decade, the efficacy of SI algorithms in tackling challenging optimization problems, discrete and continuous alike, has excited enormous interest from across disciplines. From life sciences to logistics, engineering to e-commerce, scientists are looking towards SI not only for solutions, but for intelligent, adaptive, and scalable solutions.

Ants at Work: ACO in Action: Ant Colony Optimization (ACO) has moved far beyond its original role of solving the famous Traveling Salesman Problem. And one field where ACO is making a big buzz? Bioinformatics!!

Ants in the Lab: How Tiny Insects Are Solving Big Bioinformatics Problems: Who would have believed the key to solving some



of biology's greatest mysteries is in the humble ant? That's correct!! Those small trail-following bugs have led to an algorithm that's now working overtime in the field of bioinformatics.

Let's track the trail and see where these little geniuses are leading us!

Sequence Alignment: The Ants Who Match DNA Like Pros: Let's step into the world of biology for a moment. Imagine trying to compare long strands of DNA, RNA, or proteins to figure out what parts they have in common. That's called multiple sequence alignment (MSA). Scientists use it to understand genetic similarities, study evolutionary paths, trace disease mutations, and more. But here's the catch: it's not easy. The more sequences you add, the more complex the alignment becomes. It's like trying to line up several tangled strings and finding the best way to match their patterns without missing anything

important. Traditional tools often struggle here. Enter Ant Colony Optimization (ACO). Think of each alignment path as a trail in the forest. Ants explore different paths, testing out how well segments align. When an ant stumbles upon a promising alignment, meaning similar regions across sequences line up nicely, it lays down pheromones, just like real ants marking their way to a tasty food source. Now, more ants are likely to follow that path. If they also find it effective, the pheromone gets stronger. Eventually, the most optimal paths rise to the top, reinforced by generations of digital ants. The result? A highly accurate & efficient multiple sequence alignment. What's especially impressive is that this all happens without a central command. Just simple agents, each doing their part and indirectly communicating through digital trails. It's a perfect example of swarm-based problem-solving, scalable, flexible, and great at navigating complexity.

DNA Fragment Assembly: Solving a Giant Genetic Jigsaw: Picture this, you've got thousands of pieces of paper, each with a random snippet of a novel. Your task? Reconstruct the entire book from scratch. Sounds intense, right? Well, that's exactly what scientists face when they sequence a genome.

Modern sequencing machines don't spit out a clean, complete DNA strand. They generate millions of short DNA fragments, like scattered paragraphs from a story. The challenge is to assemble them in the correct order to recreate the full narrative, the genome. Enter Ant Colony Optimization again! Each DNA fragment is treated as a node, and possible overlaps between fragments are edges connecting them. Digital ants traverse this graph, checking out different paths through the fragments. When an ant finds a strong overlap, it lays down a pheromone trail. The better the match, the stronger the trail. Other ants pick up on these cues and begin favouring the stronger paths. Over time, the swarm collectively identifies the best ways to assemble the fragments into a coherent sequence. No single ant knows the full picture, but together, they build it piece by piece. This approach is particularly powerful when the sequencing data is messy, incomplete, or comes

from multiple organisms, as in metagenomic studies. ACO-based assemblers adapt in real-time, refining their paths as better overlaps are discovered.

Why Are Ants So Good at This Stuff?

Because biology is full of possibilities. And ants are great at narrowing them down. They don't require maps, leaders, or instructions. Just a little pheromone and a clever way of responding to what's around them. By copying this behaviour, ACO allows researchers to arrive at the most optimal solutions more quickly, more cost-effectively, and often with less computational expense than many conventional algorithms.

What Are Bees Doing in the Lab? More Than Making Honey:

When it comes to discovering concealed cancer clues, cracking the DNA code of conversation, or even creating improved medicines, bees are buzzing with brilliance.

Wait, bees?

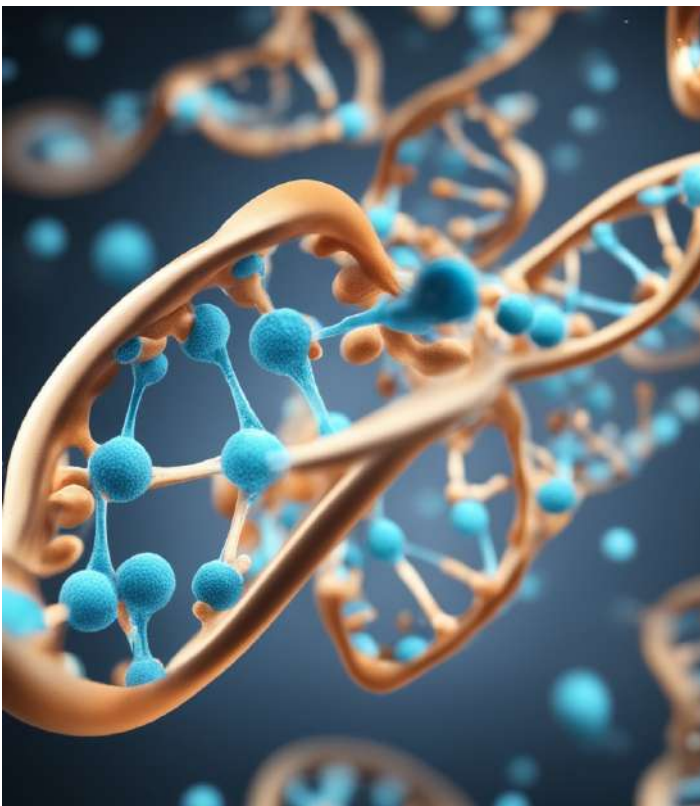
Indeed! As ants march and birds flock, bees scout, dance, and decide, and in optimization, the Artificial Bee Colony (ABC) algorithm does just that.

So, what sets ABC apart? It's intelligent, flexible, and organic to operate. Rather than attempting all available alternatives, ABC distributes the workload between various "bees", some scout, some analyse, and some extract, collaborating to hone in on the optimal solution options. It's particularly potent when you have problems that are messy, complicated, or simply too large for conventional approaches.

Let's soar a little further into how ABC is revolutionizing the world of bioinformatics and medicine, one nectar-rich solution at a time.

Buzzing Through Bioinformatics: The ABCs of Artificial Bee Colony Algorithm:

In the world of bioinformatics and medicine, where there is vast and intricate data, the ABC algorithm excels by traversing efficiently through candidate solutions to arrive at optimal or near-optimal solutions. Let's see how this bee-inspired



algorithm is creating a buzz in applications.

Swarm Intelligence in Action

Swarm-inspired algorithms are solving real-world challenges in healthcare, bioinformatics, logistics, and more. From spotting tumors to planning routes—these tiny inspirations are making a big impact.

Feature Selection in Medical Imaging:

Medical images like CT scans or MRIs are goldmines of information, but not all that glitters is useful. These images contain hundreds, sometimes thousands, of features, and only a small fraction of them are relevant for diagnosing diseases. The challenge lies in identifying which features truly matter while filtering out the background noise. This is where the Artificial Bee Colony algorithm flies in, turning this overwhelming task into a smooth and efficient process. Inspired by real bees that forage for nectar, ABC mimics the roles of employed, onlooker, and scout bees to navigate the feature space. Employed bees explore known regions, onlooker bees observe and pick the best choices based on quality, and scout bees randomly explore new areas. In medical imaging, this strategy means that ABC can scan through the massive space of features extracted from images and intelligently zero in on the most informative ones.

A study applying ABC to feature selection in computed tomography images found that it significantly improved classification accuracy. With fewer but more meaningful features, the algorithm allowed the machine learning model to perform better diagnostics, faster, and with greater precision. It's like sending a team of bee detectives into a scan, each sniffing out the subtle patterns that distinguish healthy tissues from anomalies. And the best part? They don't get tired or bored. So, while radiologists look at the full picture, the ABC algorithm helps sharpen the focus by highlighting only what truly matters. Efficient, clever, and inspired by nature.

Epistasis Detection in Genetic Studies:

Genetic diseases are rarely the result of a single

gene acting alone. Often, it's a group of genes interacting with each other in unexpected ways that leads to complex traits or conditions, a phenomenon known as epistasis. Detecting these gene-gene interactions is like solving a massive puzzle where the pieces keep shifting depending on their neighbours. Traditional algorithms struggle with this task, especially when faced with the enormous number of possible gene combinations. But guess who thrives in complex, dynamic environments? That's right, the humble bee.

Using the Artificial Bee Colony algorithm, scientists have been able to approach the challenge of epistasis detection with renewed energy. Just like bees explore flower fields, the algorithm explores combinations of genes. Each combination is evaluated for its ability to explain a disease or trait, and the "nectar value" of that solution determines whether more bees should explore that region. In one cutting-edge study, researchers used a multi-objective version of the ABC algorithm, enhanced by a scale-free network structure. This means that the algorithm wasn't just looking for one optimal solution; it was balancing multiple goals at once, like accuracy, simplicity, and computational efficiency.

The scale-free network helped guide the bees toward key "hub" genes, those with lots of connections to others, making the search smarter and more focused. The result was a powerful system that could uncover hidden patterns in genetic data, identifying crucial gene interactions that might otherwise be missed. This approach is especially valuable in areas like personalized medicine, where understanding how someone's unique gene combinations affect their health can lead to more tailored and effective treatments.

Why ABC Rules the Hive? Think of the Artificial Bee Colony (ABC) algorithm as a hive full of genius bees. All doing their part, improvising along the way, and never exhausting themselves. It doesn't take the time to test every possibility. Rather, like actual bees searching out superior nectar, it homes in on good solutions quickly.

What's Next? The Swarm Has

Only Just Begun to Buzz!

If you believed ants, bees, birds, and glowworms had already demonstrated their full potential, think again! The area of Swarm Intelligence (SI) is only just beginning its early flight, and the future is buzzing with possibilities.

Future Work: Where the Swarm Might Fly Next!!! Smarter AI, Powered by Swarms

Scientists are now looking to see how SI will make deep learning better – hyperparameter optimization, neural architecture search, and even self-improving AI systems. Envision neural networks that teach themselves via swarm intelligence. That's the future!

Swarm-Driven Robotics & IoT

Imagine intelligent drones talking like bees to organize search and rescue operations. Or IoT sensors maximizing energy consumption in real time, like glowworms adjusting their glow. The swarm future is not science fiction; it's already underway.

Sustainable Systems

From green logistics to precision agriculture, swarm algorithms are being experimented to maximize irrigation, planting, harvesting, and even detect pests – all drawn from ants and bees doing what they have done for millions of years.

Human-Swarm Collaboration

Imagine this: humans and swarm-based AI collaborating to make real-time decisions in healthcare, smart cities, and disaster response. It's not automation, it's supercharged intelligence!

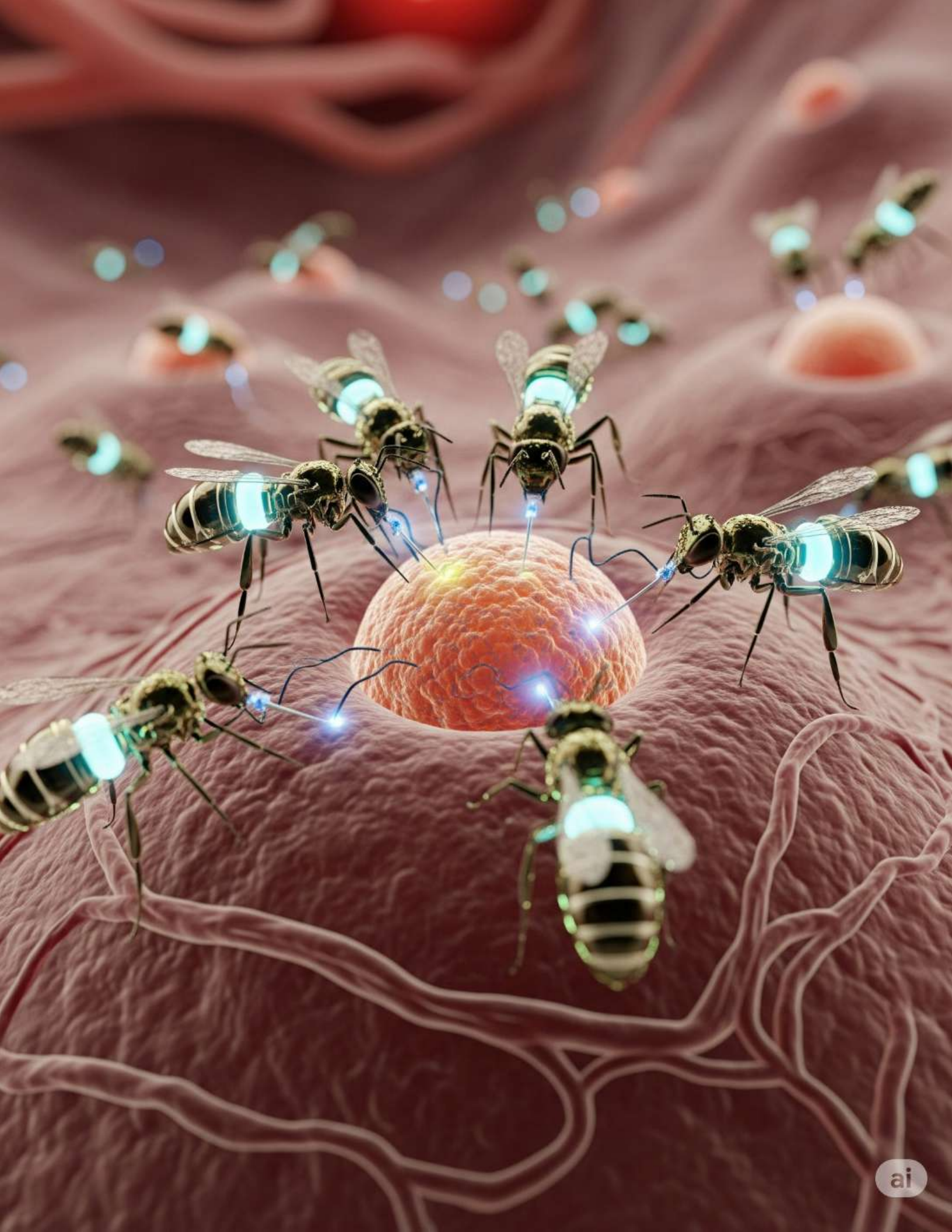
Conclusion: Nature Knows Best (And We're Just Catching Up)

So, what have we learned from our diminutive trailblazing friends? That ant, bees, and even glowworms have more than just survival skills – they have super strategies. By mimicking their behaviours, we're building smarter systems that solve real-world problems faster, more flexibly, and more efficiently than traditional methods ever could.

And the beauty of it? These algorithms don't require leaders. No central authority. No command-and-conquer programming. Just individual agents collaborating, learning, evolving, and flourishing, just like in the natural world.

What's Next for the Swarm?

The future of Swarm Intelligence is buzzing! Think AI that learns smarter, robotics that collaborates like nature, and energy systems optimized like a hive. The swarm's story is just beginning.



AI in Space

Unlocking the Cosmic Code

Hari Srija



“As we venture into the void, AI becomes our map, compass, and translator of the universe’s secrets.”

“What if the secret to exploring the universe isn’t just human intelligence, but machines that can think and learn?

Imagine a robot on Mars deciding where to move next, or a telescope spotting a distant planet that might support life—all without waiting for human instructions. This isn’t science fiction; it’s real, and it’s happening right now.”

From managing the mind-boggling amounts of data generated by satellites to enabling robots to operate autonomously on distant planets, AI is transforming how we explore the cosmos. Think about it: Every day, satellites, rovers, and space probes collect huge amounts of data—far more than any human team could analyze alone. That’s where AI steps in, processing this information at lightning speed, recognizing patterns, and uncovering discoveries that might have taken years to find.

But AI does more than just handling data. It solves one of space exploration’s biggest challenges—distance. Signals from Earth can take minutes or even hours to reach far-off spacecraft. However, AI can make autonomous decisions, ensuring missions stay on track even when they’re millions of miles away, whether it’s a Mars rover avoiding obstacles or a satellite predicting dangerous solar storms.

AI is making missions faster, safer, and more efficient. And this is just the beginning—because the machines we send into space aren’t just following commands anymore. They’re learning, adapting, and shaping the future of discovery.

This isn't a sudden leap forward, though. The journey of AI in space has been decades in the making..

But in what ways is AI shaping space exploration? What are the technologies that allow it to process information, adapt, and operate independently in the most extreme environments?

Smart Satellites

AI-Powered Data Processing

AI acts as the brain behind the mission, enabling machines to see, learn, and decide. Think about satellites—silent sentinels orbiting our planet, capturing vast amounts of data. But who sifts through it all? AI does. Machine learning and Computer Vision helps in processing satellite images, cutting through atmospheric noise and enhancing resolution, while Neural Networks classify terrain features, detect deforestation, or track the melting of glaciers. CNNs (Convolutional Neural Networks) power these image-processing tasks, ensuring high accuracy in classification and segmentation. Take India's Chandrayaan-3 mission, 2023—its success wasn't just about landing on the Moon but about data collection. The lander and rover relied on AI-powered imaging to analyze lunar soil composition, just as Earth-based satellites use the same techniques for monitoring climate change.

“AI gives satellites the power to track, detect, and understand a changing world—autonomously.”

Autonomous Rovers

Intelligent Navigation

Now, imagine a lonely rover on Mars. No GPS. No human guide. Yet, it moves, avoiding hazards, finding paths, and making decisions. This is AI in action— Using Reinforcement learning, the rover

optimizes its path based on past movement just as self-driving cars refine their navigation on city streets. With SLAM (Simultaneous Localization and Mapping), it builds a real-time map of its surroundings, fusing data (Sensor Fusion) from cameras, lidar, and radar to avoid obstacles, Which is similar to enabling autonomous drones to inspect bridges, underground mines, or disaster zones.



Logistics companies like Amazon and FedEx use similar RL-powered robotics in warehouses, and even the military deploys AI-guided navigation systems in unmanned ground vehicles. The logic of AI remains the same—it reads the world, predicts the next step, and adapts. NASA's Perseverance rover, 2021 on Mars is a prime example, equipped with AI that allows it to autonomously analyze terrain and detect scientifically valuable sites.

Like Mars rovers self-navigate via AI, reinforcement learning drives real-time decisions in uncharted terrain, mirroring tech used in drones and logistics.



Predicting Space Weather

AI for Solar Storm Forecasting

But space isn't just about landscapes; it's about survival. The Sun, for all its brilliance, is unpredictable. Solar storms can disrupt satellites and power grids, but AI is learning to predict them. LSTM neural networks, trained on decades of solar data, detect patterns invisible to the human eye, forecasting when a storm might strike. That same ability to recognize time-series trends by Time-Series Forecasting finds a home in finance, predicting stock market fluctuations, or in cybersecurity, identifying anomalous behavior that could signal a cyberattack.

Predictive analytics, powered by regression models and transformers, refined forecasting, allowing energy companies to anticipate electricity demand surges just as

astronomers anticipate solar flares. India's Aditya L1 mission, 2023, launched to study the Sun, employs AI to analyze solar data, helping scientists study the Sun's corona, solar winds, and space weather. By predicting solar storms and other cosmic events

Space Debris Tracking

Preventing Collisions in Orbit And then there's space debris—countless fragments of old satellites and rockets hurtling through orbit, each a potential threat. Tracking them is a game of precision. AI-driven object recognition models, like those based on CNNs and transformers, analyze radar images, while predictive analytics calculates future trajectories.

Autonomous air traffic control systems also use similar AI-driven trajectory prediction to prevent collisions and optimize flight paths, making the skies safer. India's PSLV-C60 mission and SpaDeX experiment, 2024, which focused on in-space docking, highlight how AI helps automate space vehicle coordination, reducing the risk of collisions.

AI-Driven Space Health Monitoring

Keeping Astronauts Safe

AI is not just monitoring astronauts—it's actively safeguarding their health. Deep learning models analyze biometric data, tracking oxygen levels, heart rates, and potential medical anomalies. But AI doesn't just observe; it predicts, offering early warnings of health risks before they become serious. RAG (Retrieval-Augmented Generation) models provide real-time medical recommendations, ensuring astronauts receive immediate guidance even in deep space.

These techniques are revolutionizing healthcare on Earth—from smartwatches and fitness trackers that monitor vitals to AI-powered diagnostics in hospitals. The predictive systems keeping astronauts fit in microgravity are also helping miners, factory workers, and athletes push their limits safely. Beyond health, AI is becoming a true companion in space. NASA's Robonaut, 2011 a humanoid robot, assists astronauts with maintenance tasks, reducing their physical strain. Meanwhile, Germany's

CIMON (Crew Interactive Mobile Companion), 2018 is pioneering AI-driven conversational

“Just as wearables flag health issues early, deep learning keeps astronauts safe—predicting anomalies and guiding care in real-time, even beyond Earth.”

assistance, just as chatbots enhance customer support on Earth. These advancements show that AI in space isn't just about automation—it's about intelligent support, adaptability, and enhancing human capabilities in extreme environments.

“Neural networks scan galaxies like analysts scan markets—detecting faint signs of something rare, and possibly life-changing.”





Exoplanet Discovery

AI's Search for Alien Worlds And while AI monitors the known, it also searches for the unknown. In the hunt for exoplanets, Neural networks scan vast astronomical data, detecting the slightest flicker of light that hints at a distant world. Spotting a potential Earth-like planet is not so different from identifying a hidden opportunity in market data. AI is making space missions smarter and more efficient. India's Reusable Launch Vehicle-Technology Demonstrator (RLV-TD) is working toward an AI-powered, reusable spaceflight system that could make space travel more affordable. And then there's ISRO's groundbreaking space-based agriculture experiment, 2025 where cowpea seeds were successfully germinated in microgravity. This isn't just a scientific milestone; it's a glimpse into a future where AI could manage extraterrestrial farming, turning the

dream of sustainable space colonies into reality.

“AI finds exoplanets by catching the faintest flickers—where science sees signal in silence.”

Challenges and Risks of AI in Space:

Space isn't just far—it's hostile. No backup plans, no second chances, and no room for errors. AI may be the key to exploration, but what happens when the unknown fights back? What does it take for machines to think, decide, and survive millions of miles away? Sounds like a sci-fi plot, right? Except this isn't fiction. It's happening right now. Think about radiation. It's invisible, relentless, and utterly destructive. Here on Earth, your laptop, your phone—every piece of tech you own—is protected

1960-1980's

- Automating spacecraft navigation
- Establishment of ISRO and Launch of Aryabhata, India's first satellite, which utilized basic AI for data processing.
- Expert systems for decision-making
- NASA Mars navigation (Pathfinder)

- INSAT Series Launch
- Deep Blue inspires AI problem-solving in space
- Mars Climate Orbiter highlights AI's role in data interpretation
- AI-driven robotics for space maintenance and ISRO begins using AI for satellite data analysis

2000-2010

- ISRO launches TES, which uses AI for advanced imaging and terrain mapping.
- AI-driven lunar exploration (Chandrayaan-1, Chang'e-1)
- AI-assisted humanoid robot (Robonaut 2)
- AI-powered Moon rover (Yutu).

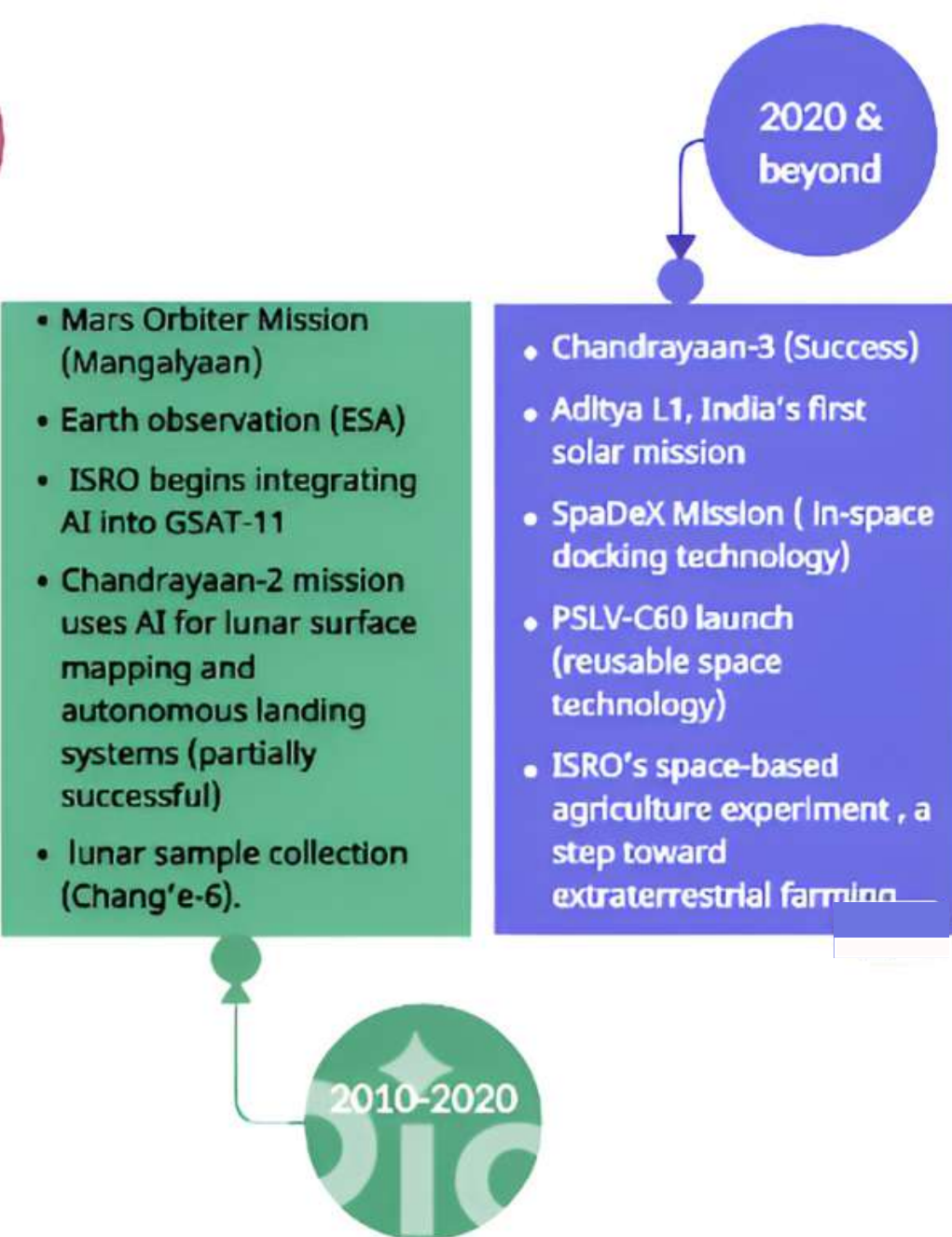
1990-2000

from cosmic rays. But in space? AI hardware is bombarded with charged particles that don't just damage circuits; they flip bits in memory, corrupting data and altering decisions. Imagine a satellite receiving a command to adjust its orbit, but thanks to a radiation-induced error, the AI miscalculates—just slightly. A fraction of a degree off course. Now multiply that mistake over time. What starts as a tiny misstep could end in a catastrophic collision.

And then there's the computational constraint. AI thrives on vast datasets and high-performance processors, but space missions rely on low-power, radiation-hardened chips. These processors lack the capability of Earth-based systems, forcing AI

to operate under severe limitations. How do we make AI think fast, act precisely, and learn—when it's running on hardware that's decades behind? If that wasn't enough, there's the data transmission limitation. Ever tried sending a large file over a painfully slow internet connection?

“In deep space, AI must act fast, think smart, and survive radiation, delays, and data limits—with no backups, no bandwidth, and no margin for error.”



environments with limited data and no second chances.

The nature of telemetry data is another hurdle. AI models rely on patterns and training data, but what happens when they encounter something entirely new? Unlike a chatbot that refines its responses over time, space AI doesn't get second chances. If it misinterprets terrain on an asteroid or fails to recognize a critical spacecraft malfunction, the consequences could be irreversible. How do we train AI for the unknown—when we don't even know what "unknown" looks like? Then there's hardware limitations and fault tolerance. Spacecraft components must withstand extreme temperatures, vacuum conditions, and mechanical stress, all while maintaining flawless performance. AI models that require complex neural networks demand specialized hardware—something difficult to implement in a spacecraft where every gram of weight and every watt of power matters.

Now imagine doing that across millions of miles. A single high-resolution image from Mars can take hours to reach Earth. And if AI has to wait for human input before making decisions, we're doomed. That's why autonomy is everything. But how do you train AI to recognize what's important? What if it filters out something crucial? What if it misses a sign of life, or a warning of disaster, simply because the data didn't "fit" what it had seen before? Communication delays and latency make real-time decision-making impossible, forcing AI to operate autonomously in critical situations. But without constant updates, ensuring model reliability and robustness becomes a challenge—AI must perform flawlessly in unpredictable

Beyond the technical hurdles, a bigger question looms: Who is responsible when AI makes a mistake? The legal framework for autonomous decision-making in space remains unclear. If an AI-powered satellite collides with another, who takes the blame? What about the ethical concerns—should AI be trusted with mission-critical decisions that could risk human lives? What if an AI-driven probe prioritizes efficiency over astronaut safety? And as AI plays a greater role in defense-related space applications, are we unknowingly paving the way for an AI-driven arms race? Security is another pressing issue. AI-driven satellites process enormous amounts of data, from climate monitoring to Earth observation.

But what if these systems are hacked? A compromised AI-powered satellite could be used for surveillance, disrupt global communications, or even pose a national security threat.

Despite all these challenges, AI isn't just surviving—it's evolving. With every challenge overcome, machines become smarter, more resilient, and more capable.

The same AI that's learning to navigate alien landscapes could one day drive autonomous robots in underground mines. The same predictive models used to track space debris could protect Earth's financial markets from crashes.

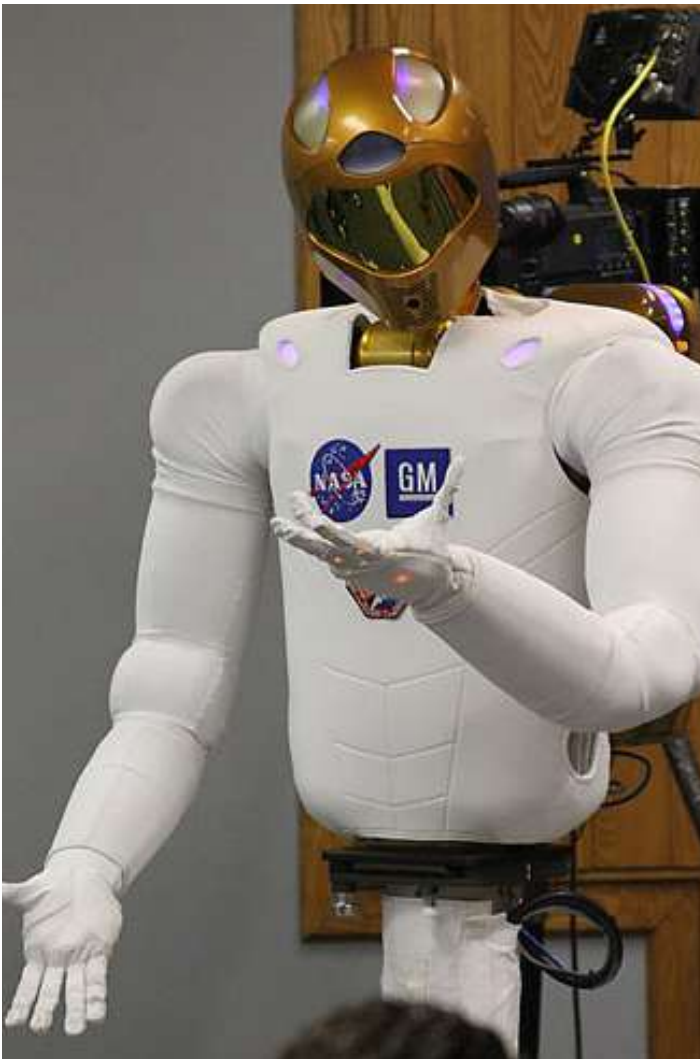
The same AI keeping astronauts healthy could transform healthcare back home. AI in space isn't just about reaching the stars. It's about reshaping intelligence itself. And if machines can think in the most extreme conditions imaginable—what does that mean for the future of AI on Earth?

“As missions grow smarter, AI joins forces with quantum and blockchain to secure data and decisions, and deepspace autonomy.”

Future Advancements

The future of AI in space is unfolding at an unprecedented pace. NASA's Artemis Program is set to use AI-driven systems for Moon exploration, while ISRO's Gaganyaan Mission will rely on AI to monitor astronaut health and assist in real-time mission operations. Meanwhile, NASA & ESA's Mars Sample Return Mission and China's Tianwen-3 will leverage AI for autonomous sample collection and retrieval from the Martian surface. AI will also play a key role in ISRO's Shukrayaan-1 mission to Venus, where AI-powered sensors will analyze its thick atmosphere. Beyond planetary exploration, AI is driving asteroid mining missions, such as NASA's Psyche Mission, and improving space sustainability through AI-powered debris removal projects like ESA's ClearSpace-1. Commercial ventures, including SpaceX's Starship Missions and AI-driven orbital habitats, are redefining the future of human spaceflight and interplanetary logistics. And then there are emerging advancements like black hole research, differentiable intelligence for spacecraft control, and cybersecurity for space missions that will push the boundaries of what's possible.

The integration of blockchain with AI could revolutionize space missions, ensuring secure data management and transparent decision-making. Cognitive communication systems, such as AI-driven Software-Defined Radios (SDR), will enhance spacecraft communications and data transmission efficiency, making interstellar connectivity more robust. Quantum computing and bio-inspired AI models are being explored to overcome current computational limitations, opening new avenues for deep-space navigation. AI will also play a crucial role in the construction of independent habitats on the Moon and



Mars, enabling long-term human presence beyond Earth. With AI managing everything from navigation and resource extraction to astronaut well-being and autonomous decision-making, space is no longer just a frontier—it's an intelligent ecosystem in the making.

Conclusion

Space has always been the ultimate unknown—a place of infinite possibilities and relentless challenges. But now, AI is rewriting the rules of exploration. It's not just about reaching new worlds; it's about understanding them in ways we never could before. From self-thinking rovers on Mars to AI-driven space habitats, intelligence is no longer bound to Earth. It's learning, adapting, and evolving—just like we are. But here's the real question: If AI can survive the harshest corners of the universe, what else

can it achieve? If it can predict solar storms, uncover exoplanets, and guide autonomous spacecraft, imagine what it could do for life on Earth. The same technology powering interstellar missions is shaping industries, transforming businesses, and redefining how we see the world. We stand on the edge of a new era—one where AI doesn't just assist us in space; it becomes the navigator, the explorer, and the storyteller of the universe itself. The future isn't light-years away. It's already here, coded into every algorithm, every neural network, and every decision AI makes beyond our planet. The question isn't whether AI will take us further than we've ever gone — it's how far we're willing to go.

“In the vacuum of space, AI isn't just surviving—it's rewriting what it means to be intelligent.”



Panocle: Redefining Smart Wearables with Inclusive AI

In this exclusive interview, the co-founders of Panoculon Labs — an IIT Madras-based AI startup — talk about their flagship product Panocle, an inclusive AI-powered smart spectacle. Designed to blend everyday usability with assistive intelligence, Panocle aims to make contextual computing accessible, sleek, and affordable. We dive into their journey, innovation, and vision for truly inclusive technology.

AINA: Can you tell us about your product? What is it, who benefits from it, and how would you explain Panocle's core functionality in simple terms?

Rishabh: In simple terms, Panocle is like a personal assistant who is always with you, capable of seeing and hearing your surroundings and providing contextualized feedback.

The issue with current AI models like ChatGPT is that they lack personal context. Panocle aims to bridge that gap by collecting real-time, personalized data based on your daily interactions. We began by designing this for visually impaired individuals as a virtual assistant. The glasses have a camera and microphone and can answer questions like "Who is in front of me?" or "How

much money am I holding?" But as we scaled, even visually impaired users told us they preferred a product that was universally usable. That shifted our focus to inclusive design.

AINA: What inspired the idea behind Panoculon? Was it personal or based on a societal gap?



Sreeraj R

An engineering undergrad passionate about technology and the impact it can have on the lives of people.

A first principles thinker, experienced in user-centric product development and empathetic leadership.

Building a deep-tech startup developing environment perception systems @ IITM Research Park. Outside of work, he loves to travel without compass and itinerary.

Rishabh Sharma

A technophile studying at IIT Madras, passionate about engineering and value-addition to society through innovations. He is skilled at building technology from the ground up and leading a team with empathy.

Currently building Panoculon Labs, a deep tech startup, where they are redefining the future of human-computer interaction with India's first inclusive AI-powered smart spectacles.



Rishabh: It began in our first year at IIT Madras during the COVID lockdown. Sriraj and I met online through WhatsApp and bonded over shared interests in electronics and software. We participated in Shaastra's Techfest, particularly in the Assistive Technology Challenge, which presented real-world problems related to physical disabilities.

We were first-year students then, but we were passionate about computer vision. Inspired by technologies like Tesla's autonomous vehicles, we thought, "Why not build something for the visually impaired that offers real-time audio feedback from visual data?" That's how the idea began. We built a prototype at home and surprisingly won the competition against seasoned teams, including Ph.D. scholars. Over time, we interned at labs

within IIT, built several iterations, and validated our idea in business competitions. With

mentorship from Padma Shri Professor Ashok Jhunjhunwala, we formally incorporated Panoculon Labs. The name stands for PAN (all), Oculus (eye), and ON (electronics) — essentially, "seeing the world through electronics."

Sreeraj: As we grew, we realized we couldn't scale by focusing solely on the visually impaired. With the rise of AI and increasing market demand, it made sense to broaden our scope while keeping inclusivity central.

AINA: What differentiates Panoculon from other AI wearable devices that are currently available in the market.

Sreeraj: A key differentiator is our focus on the form factor and sensor placement. Many wearables today come as pendants and don't include cameras. Ours integrates multiple sensors including vision and audio directly into spectacles — a familiar and socially

accepted form factor. Glasses are already a part of daily life for many people, making adoption easier.

AINA: What makes your product both pocket-friendly and sleek, as mentioned on your website?

Sreeraj: Our product is still in the prototype stage, but our design is geared toward miniaturization. Traditional assistive tech devices are often bulky and unattractive. We aimed to create something people would be proud to wear, not something they feel forced to use.

Rishabh: Miniaturization is challenging, especially in hardware. Our small team is constantly learning and iterating to make the device both functional and stylish. We're also seeking guidance from industry experts to ensure we meet those standards.

AINA: What are some of the toughest technical or logistical roadblocks you've faced?

Sreeraj: Supply chain issues are a major hurdle. In software, we can push an app update in hours. In hardware, a small change can take weeks, with dependencies across multiple countries. IC components, PCB printing, customs clearance — all take time and coordination.

AINA: Is India equipped to support deep tech startups like yours?

Rishabh: We're making progress, but the ecosystem is still developing. For example, we had to send a chip abroad because no one in India could fabricate it with the required drill size. The government is pushing semiconductor manufacturing, but it'll take time.

Sreeraj: Even from a financing perspective, most Indian innovation is still focused on D2C or FMCG. Hardware startups are just beginning to get attention. There are exceptions, but the ecosystem still lacks depth compared to places like the US.

AINA: With IIT Madras as your launchpad, how did you build a multidisciplinary team? Since it started with just the two of you, how did you bring in expertise in hardware, electronics, and design to make the device so compact?

Rishabh: In the beginning, it was just Sriraj and me handling everything. But as we progressed, we realized that hardware is a complex domain that requires deep focus and specialization. That's when we decided to bring

in full-time engineers.

We found talented individuals through platforms like Internshala and LinkedIn. Interestingly, most of our early hires were from non-IIT backgrounds, but their passion and commitment were exceptional. They were eager to learn, highly skilled, and deeply invested in the problems we're solving. We're truly grateful to have them.

Sreeraj: The IIT Madras ecosystem itself is very interdisciplinary. I worked on a Mars rover project, and Rishabh



was part of a solar car team. These projects involved students from diverse departments—software folks from naval architecture, mechanical engineers from chemical branches, and so on. This culture of interdisciplinary collaboration helped us understand how to build and manage a diverse team effectively.

AINA: How has networking helped you — in terms of access to investors, partners, and shaping your journey so far?

And how can students reach out to these networks more easily?

Sreeraj: Networking has been one of the most critical factors in our journey. The IIT ecosystem, especially the Incubation Cell (IIC), offers access to a wide range of investors and partners who are already familiar with deep tech startups. Our mentor, Professor Ashok Jhunjhunwala, is also highly respected and well-connected, which has opened many doors for us.

We've also taken a proactive approach by attending startup events—like the one in Bangalore

— where we met many like-minded founders. Even if others are from different industries like EdTech or D2C, there's a lot to learn that can apply to our work as well. Networking has provided valuable insights and even potential collaborations.

Rishabh: It also helps with customer discovery. Through these networks, we found some of our first users and received feedback early on, which gave us momentum at a critical stage.

AINA: AI is evolving rapidly. How do you plan to future-proof Panoculon?

Rishabh: We've approached this from both a hardware and software perspective. The hardware is intentionally modular, meaning future upgrades won't require users to replace their entire device. Even if AI capabilities evolve significantly — as they undoubtedly will — our architecture allows us to push software updates that keep the system competitive and relevant. We're also keeping an eye on the next generation of chipsets, especially those capable of running large language models (LLMs) locally. These chips are becoming more efficient and accessible, and we're designing our systems to be ready for seamless integration when that technology becomes mainstream. It's about building with foresight, not just for the present.

AINA: Tell us about MOMA, your note-taking app. What inspired it?

Rishabh: When we began working on the hardware, it became clear that launching it alone wouldn't deliver the kind of value we envisioned. People don't just want tech—they want tools that fit into their lives. That realization led to the creation of MOMA, our note-taking app. MOMA was born from a real, everyday problem: we all forget details from conversations—names, past interactions, the context of meetings. We wanted to solve that, especially in a world

where relationships and context matter a lot. We've optimized MOMA for in-person meetings—so you can capture notes right at the table—while still requiring an internet connection for full functionality, and it supports multiple Indian languages so that it's inclusive and regionally accessible.

When used alongside our smart glasses, MOMA will help people capture and summarize conversations or key moments from their day, all while prioritizing privacy and simplicity.

AINA: What AI/ML techniques are you using?

Rishabh: Our hardware engineering is focused on two major constraints: battery life and data compression. We're aiming for all-day usability without needing frequent recharges, which is quite a challenge in wearables. On the software side, we're using Convolutional Neural Networks (CNNs) and other deep learning models to process visual and audio inputs.

The key point is we're not storing raw data. Instead, we convert it into contextual feedback in real time, which significantly enhances privacy while stivering useful information. They store data on the edge, but the actual larger model runs on the cloud, and to extract meaningful insights from what users see and hear without compromising security or user trust.

AINA: Have you achieved any

breakthroughs or filed any patents?

Rishabh: We're on the cusp of it. While we haven't officially filed patents yet, we're working on a few innovations that we believe are novel and patent-worthy. One area we're excited about is our use of specialized chipsets that are both energy-efficient and cost-effective, which is critical for building affordable consumer tech.

We're also developing a proprietary privacy-preserving algorithm designed to safeguard sensitive data while maintaining full contextual integrity and high performance. Once we finish fine-tuning those systems and validate their uniqueness, we plan to move forward with the patent process. Our focus has always been on practical innovation, not just filing for the sake of it.

AINA: What's one myth plan like to debunk?

Sreeraj: One myth I often hear is that AI is saturated or that all the "cool" problems have already been solved. That's far from the truth. If you look around, there are countless unsolved, highly specific problems—especially in Indian contexts—that haven't been touched yet.

The key is to start with the problem, not with the AI. Too many founders try to force-fit AI into their startup because it's trendy. But the most impactful solutions come from deeply understanding the problem first and then building the right technology around it—even if it's not AI-heavy at first.

Rishabh: Another myth is that tools like Codex or GitHub Copilot make building an AI startup easy. Sure, they help with code and prototypes, and they've lowered the barrier to entry. But that's only a small part of the journey. Building something meaningful still requires real insight into your users, domain expertise, and often a vertical stack approach.

You can't skip understanding the customer's pain points or the nuances of the domain. Those tools can speed up coding, but they can't replace deep thinking and product-market fit.

AINA: What can policymakers and students do to bridge the gap between academia and industry in AI?

Sreeraj: The biggest gap is the lack of hands-on learning. Many academic projects are too abstract or generic, and students don't get exposed to real-world constraints. We need more hackathons, startup challenges, and internships where students can work on messy, ambiguous problems that resemble real industry scenarios.

Universities should also collaborate with startups and businesses more actively, so students can test their ideas in practical environments, not just on paper.

Rishabh: Government bodies and public sector institutions also need to embrace AI more proactively. There's massive potential to improve public services and productivity through AI, but adoption is slow due to bureaucratic inertia.

On the student side, the focus needs to shift from chasing academic perfection to solving problems creatively. It's not just about learning models and metrics—it's about building applications that work in the field. That's where the real learning happens.

AINA: Finally, what advice would you give to aspiring entrepreneurs? And what's one lesson you wish you had known earlier?

Sreeraj: Starting a company while you're in college is very challenging, especially with the opportunity cost: you're giving up time, internships, or other career paths. But your early 20s are also the best time to take bold risks. You have the time, energy, and fewer responsibilities. Even if your startup doesn't succeed, the experience is invaluable, and more employers are starting to respect and value that journey.



Rishabh: My biggest advice: don't start a company just because it looks cool or glamorous. The reality is, it's often messy, exhausting, and emotionally draining. But if you're truly passionate about solving a specific problem, and

you have a trustworthy co-founder, it becomes one of the most powerful learning journeys you can take.

What helps is balancing vision and optimism with constant self-checks and feedback loops. Keep questioning yourself while staying motivated, that balance is hard but necessary. For me, I don't think I'd want to change anything or know something too early. I believe in the importance of timing and learning through the process. Every mistake, every delay, it all

teaches you something valuable. You grow with the journey, and that's something you can't shortcut.

Sreeraj: One hard lesson we learned: hardware timelines are brutal. You can't iterate as fast as with software. We underestimated the time required for development and testing. There's so much you learn on the go, especially when you're building something from scratch.

AINA: Thank you both. Panoculon is more than a product—it's proof that India's young talent can lead the global AI revolution. We look forward to sharing your story in our July issue.

Not man versus machine, but mind meeting machine.
As AI becomes more human-aware, the questions shift from capability to connection.
The real question is: how do we coexist, collaborate, and co-evolve? The conversation with AI isn't about replacement. It's about relationship.



Deepseek - A Deep Dive into Its Cutting- Edge Features

-Rohith Chandra

Somewhere in 2023 in India, when a CEO of an AI company is asked, if in India or a small team could build something big in AI with a budget of just \$10M.

Then he replied

“It’s totally hopeless to compete with us on training foundation model”

Now cut to January 2025, there comes a maverick, DeepSeek, a small startup up who trained their foundational model in less than \$6M and open-sourced it in a span of two years. With their result they not only crashed the preconceived notion of tech giants but also crashed the Global Financial Market.

Initially there was a question in the minds of the stakeholders of the AI field **“Does Less is more applied to this field ?”**. But after 2025 January they got their answer thanks to DeepSeek.

Disclaimer before going ahead, this article is a homage to DeepSeek for their innovation, and our

intention not to defame other AI tech companies. So readers, let’s get on board and buckle up, we are going to see what is the hype about the DeepSeek R1 model. And our first stop is the history of the company.

History

Based in Hangzhou, Zhejiang, Hangzhou DeepSeek Artificial Intelligence Basic Technology Research Co. Ltd. aka DeepSeek is a Chinese artificial intelligence company that develops large language models (LLMs). It is owned and funded by the Chinese hedge fund High-Flyer. DeepSeek was founded in July 2023 by Liang Wenfeng, the co-founder of High-Flyer, who also serves as the CEO for both companies. It has employee strength of around 200 with an average age of around 30 years. Mostly from university Phd holders.

Under the umbrella of High-Flyer, DeepSeek operates as an independent AI research lab. DeepSeek focuses on developing open source LLMs. In November 2023 it’s first model was released. It has iterated multiple times on its core LLM and has built out several different variations. However, it wasn’t until January 2025 after the release of its R1 reasoning model that the company became globally famous. After that DeepSeek has quickly gained traction, exceeding 10 million downloads and attracting 1.8 million daily active users.

DeepSeek-R1 released in January 2025, which is based on DeepSeek-V3 and is focused on advanced reasoning tasks directly competing with OpenAI’s o1 model in performance, while maintaining a significantly lower cost structure like DeepSeek-V3.

R1 stands out for its logical thinking combined with high-speed processing. If you need an AI for niche tasks like complex math problems or technical writing, it is a powerful choice.

Readers, our next stop is an overview of the AI chatbot.

Overview of DeepSeek R1

coding, science, and mathematics tasks, with



Pricing: The standard pricing for this model is \$0.55 per million input tokens and \$2.19 per million output tokens, where tokens are basic units of text that LLM use to process and generate output. The standard OpenAI o1 model costs \$15.00 per million input tokens and \$60.00 per million output tokens. The o3-mini model costs \$1.10 per million input tokens. DeepSeek's API has discount pricing for off-peak times of the day upto 75%.

Quality: This model brings high intelligence, comparable to OpenAI-o1. It performs well in

strong reasoning capabilities.

Content: DeepSeek has a context length of 64,000 tokens. Its maximum context tokens are 32,000, and its maximum output tokens are 8,000.

Speed: This model outputs 23.2 tokens per second, making it somewhat slower than other top AI products. Its latency is 69.93 seconds.

Recency: DeepSeek has not specified its knowledge cutoff, so if you value the most recent information, you will need to seek elsewhere.

Censorship: AI chatbot replies and then censors itself in real time, providing an arresting insight into its control of information and opinion.

If you had gone through the key details, you may find some areas of excellence and areas with scope of improvement in the R1 model. ***“So, it’s not a perfect model for purpose but it’s an innovative and efficient model.”***

So readers if you find the above journey a bit technical, then embrace yourself the way forward is more technical, don’t worry. In the next step I have simplified some technical jargons for you. You can also detour to the next stop if you know these jargons.

Technical jargons:

Open source: It for software or other creative works where the source code, design documents, or content is made publicly available for modification, distribution, and use. This approach encourages collaboration, transparency, and community involvement in development.

Model parameters: They are learned values that dictate how the model processes information and makes predictions. They can be broadly categorized into parameters that control the model’s internal representation (like weights and biases in Transformers) and parameters that influence the generation process (randomness in text generation or limit the number of tokens).

Chain of thought prompt:

Chain-of-thought prompting is a prompt engineering technique that encourages large language models (LLMs) ***“to explain their reasoning process step-by-step before providing an answer, improving their performance on tasks requiring logic, reasoning, and decision-making”.***

It essentially mimics how humans think through a problem by breaking it down into smaller steps.

Reinforcement Learning Approach

Utilizes RL for efficient learning and adaptation



Multi-Agent Learning Capabilities

Enables coordination in complex scenarios



Customizability and Pre-Trained Modules

Offers flexibility and accelerates deployment

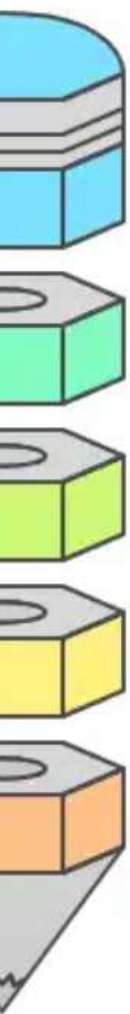


Mixture of experts (MoE):

The mixture of experts technique breaking down large models into smaller, specialized networks like experts. It’s like an AI model as a team of specialists, each with their own unique expertise. A mixture of experts (MoE) model operates by the principle of dividing a complex task among smaller, specialized networks known as “experts.”

And each expert focuses on a specific aspect of the problem, enabling the model to address the task more efficiently and accurately. It’s similar to having a mechanic for car problems, a doctor for medical issues, and a chef for cooking—each expert handles what they do best. By collaborating, these specialists can solve a broader range of problems more effectively than a single generalist.

Multi-head latent attention (MLA) It is a method used in transformers to reduce memory usage and improve inference speed. It achieves this by compressing the key and value vectors into a smaller latent space, which is then used to



Mixture of Experts Architecture

Ensures parameter efficiency and reduced computation costs

Enhanced Explainability

Provides transparency in decision-making processes

reconstruct the full key and value vectors for each attention head. This compression significantly reduces the KV cache size without sacrificing model performance.

Graphics Processing Units (GPUs):

Originally designed for rendering complex 3D graphics, GPUs have evolved into powerhouses for parallel computing, making them ideal for the matrix operations that form the backbone of LLM computations.

Reader now you are equipped with basic knowledge of the jargons. Now we will move to the main part of the story.

DeepSeek-R1: Redefining the Future of AI

The rapid evolution of artificial intelligence demands models that can handle complex reasoning, appreciate long-context nuances, and adapt to specific domains. Traditional dense

transformer-based AIs are increasingly showing their limits—suffering from high computational costs, inefficiencies in handling multi-domain tasks, and scalability issues. Enter DeepSeek-R1, a revolutionary system that transforms AI architecture by blending efficiency, scalability, and high performance into one breakthrough model.

Model architecture:

DeepSeek-R1 rises above its predecessors with a hybrid architecture based on two central innovations. First is the Mixture-of-Experts (MoE) framework, a strategy that reimagines the model as a team of specialists. Instead of engaging every available parameter for each inquiry, DeepSeek-R1 selectively activates only the most relevant ones. Imagine a vast repository of 671 billion “knowledge points” where, for any given task, only around 37 billion parameters are called into action. This approach significantly trims computational waste while enhancing performance, creating a system that can quickly adapt to the demands of different queries.

The second pillar is an advanced transformer-based design that rethinks how machines process and understand language. Unlike conventional models that treat every token and context uniformly, DeepSeek-R1 integrates state-of-the-art transformer technologies to manage both short and long-context scenarios with finesse. This dual strategy not only achieves groundbreaking efficiency but also attains state-of-the-art results without compromising on speed or accuracy.

Mixture-of-Experts Advantage

Diving deeper into the architecture, the MoE approach is similar to having a network of expert consultants at your disposal. In traditional systems such as ChatGPT, every parameter is engaged uniformly, much like having an entire team of consultants working on every problem—even when only a few specialized opinions are needed. DeepSeek-R1, by contrast, activates only the experts most relevant to the task at hand, thanks to its dynamic gating mechanism. This selective process not only reduces the computational



burden but also speeds up response times by focusing resources where they matter most.

With 671 billion parameters in total, deploying only a fraction on a per-query basis is a game-changing strategy. Techniques like Load Balancing Loss ensure that all expert networks are used evenly over time, preventing any single component from becoming a bottleneck. The result is a system that smartly distributes its immense cognitive workload, proving that strength and efficiency can go hand in hand.

Multi-Head Latent Attention

Traditional multi-head attention mechanisms compute separate Key, Query, and Value matrices for each head—a process that becomes increasingly resource-intensive as the input size grows. DeepSeek-R1 introduces MLA, an intelligent reengineering of this process. ***Instead of holding bulky matrices in memory, MLA uses a low-rank factorization method to compress these matrices into compact latent vectors.***

During inference, these latent vectors are swiftly decompressed to reconstruct the Key and Value matrices as needed, slashing the memory overhead

to just 5–13% of what conventional methods require. Furthermore, DeepSeek-R1 enhances this mechanism by integrating Rotary Position Embeddings (RoPE). This dedicated handling of positional information allows each head to maintain context in long sequences without redundant learning, ensuring that the model understands both the overall structure and the finer details of the input.

Advanced Transformer Design:

DeepSeek-R1 doesn't stop with MoE and MLA. Its transformer-based design pushes the envelope further by incorporating optimized layers for natural language processing. The design balances global attention—which captures relationships across an entire text input—and local attention, focusing on nearby, contextually rich segments. This dual strategy means that the model can seamlessly transition between grasping broad narratives and pinpointing precise details within a sentence.

Moreover, advanced tokenization techniques set this model apart. Soft token merging combines redundant tokens during processing, significantly reducing the data burden without losing essential information. To counter any potential information

loss, a dynamic token inflation module reintroduces key semantic details at later stages. This synergy between compression and restoration ensures that every piece of information, no matter how subtle, is preserved and accounted for.

The Training

Behind DeepSeek-R1's impressive capabilities lies a rigorous training regimen. The journey begins with a Cold Start Phase, where the base model (DeepSeek-V3) is fine-tuned using a curated collection of chain-of-thought (CoT) reasoning examples. These samples, chosen for their diversity, clarity, and logical consistency, lay down a robust foundation of reasoning skills.

The model's evolution continues through multiple Reinforcement Learning (RL) phases. Here, reward optimization encourages outputs that are accurate, clear, and well-formatted. Next, the self-evolution phase empowers the system to autonomously verify its responses, correct errors, and reflect on its reasoning process—a true leap towards self-improvement. Finally, the process culminates with a phase of Rejection Sampling and Supervised Fine-Tuning (SFT), where only the highest-quality outputs are selected for further training. This multi-stage approach not only refines DeepSeek-R1's capabilities across various domains but also ensures that its reasoning remains aligned with human expectations.

Cost-Effective Innovation: Maximizing Performance on a Budget An outstanding feature of DeepSeek-R1 is its remarkable cost-efficiency. In an era where training a state-of-the-art AI model could easily run up astronomical bills, DeepSeek-R1 challenges that notion. With a training expenditure of approximately \$5.6 million—substantially lower than many high-end models—the model leverages a fleet of 2,000 cost-effective H800 GPUs. By adopting a sparse activation strategy through the MoE architecture, the system smartly curbs computational expenses without sacrificing performance. This blend of budget-conscious engineering and cutting-edge technology demonstrates that advanced AI doesn't have to break the bank.

New Benchmark in AI Performance

DeepSeek-R1 is more than just a technical marvel—***it's a harbinger of the next generation of efficient and adaptable AI systems.*** Notably, the model excels in technical domains, achieving an impressive 90% accuracy in mathematical reasoning. This performance, combined with its open-source accessibility, positions DeepSeek-R1 not only as a powerful tool for researchers and developers but also as a community-driven platform that encourages innovation and customization.

Looking Ahead:

Transforming the Landscape of Artificial Intelligence. As AI continues to permeate every aspect of our lives, the need for models that are both powerful and efficient becomes ever more pressing. DeepSeek-R1 stands as a testament to how groundbreaking architectural innovations can tackle long standing challenges in the field. By selectively harnessing expert sub-networks, optimizing attention mechanisms, and streamlining data processing, DeepSeek-R1 charts a path toward a more agile, cost-effective, and high-performing future.

In a world where every second and every computational resource counts, DeepSeek-R1 ensures that the future of AI is not only smarter but also better tailored to meet the complex requirements of modern applications. This model invites us to imagine a future where the boundaries of machine intelligence are continually expanded through creative, resourceful engineering—a future where technology doesn't just work harder, but works smarter. For anyone excited by the promise of artificial intelligence, DeepSeek-R1 offers a compelling glimpse into what tomorrow's technology can achieve. As we continue to push the envelope of what is possible, innovations like these pave the way for an AI landscape that is as efficient as it is extraordinary.

So Readers, I am happy to inform you that you have reached the end of this journey. You are reading AINA 6.0, and it's your pilot RC. Thanks for reading this article.

Major Happenings

JULY 2024: EMERGENCE OF AI AGENT ARCHITECTURES:

July 2024 witnessed increasing discussion and early advancements in AI agent architectures, moving beyond simple chat interfaces to systems capable of planning, reasoning, and autonomously executing complex tasks across various tools and environments. Researchers and developers began exploring frameworks for creating more intelligent, persistent, and proactive AI assistants, hinting at a future of truly autonomous AI.

SEPT. 2024: META LAUNCHES WORLD'S LARGEST OPEN-SOURCE AI MODEL

Meta made waves in the AI world by unveiling Llama 3.1, the largest open-source AI model to date. This bold move directly challenges the dominance of proprietary models like GPT-4 and Claude. By embracing openness, Meta aims to empower researchers and developers globally, offering high-performance AI without the paywalls. The launch also reignites the debate on transparency and accessibility in the AI arms race.

NOVEMBER 2024: PROJECT INDUS—INDIA'S CHATGPT MOMENT

November 2024 saw significant strides in India's indigenous AI development, with renewed focus on Project Indus, Tech Mahindra's homegrown LLM tailored for local languages. While initial developments were earlier, November brought announcements around its further refinement, expanded capabilities, and the vision for its broader integration into India's digital infrastructure, including former Tech Mahindra CEO CP Gurnani's new venture, AIonOS, focusing on custom AI agents.

AUGUST 2024 - COMMERCIAL & REGULATORY SHIFTS IN AI

August signaled both innovation and forethought. Stability AI unveiled Stable Fast 3D, a model converting single images into high-quality 3D assets in under a second—boosting rapid prototyping for gaming, architecture, and retail. Meanwhile, Google rolled out its "August 2024 Core Update," enhancing search quality by prioritizing genuinely useful content and targeting manipulative SEO tactics.

OCTOBER 2024: INDIA'S UPI MOMENT IN AI

India is gearing up for its 'UPI moment' in AI, led by pioneers like Dr. Pramod Varma, the architect of Aadhaar and India Stack. While the West races to build foundational AI models, India is focusing on applying generative AI to solve real-world problems—voice-based payments, public health, and education. With initiatives like Sarvam AI's Indic LLM and Project Vaani's 16,000 hours of regional speech data, India is shaping AI as a public good, not just a private innovation.

DECEMBER 2024 - GOOGLE DEEPMIND'S QUANTUM BREAKTHROUGH WITH 'WILLOW'

December 2024 marked a significant milestone in quantum computing for AI. Google DeepMind unveiled 'Willow,' its state-of-the-art quantum chip, demonstrating groundbreaking advancements in error correction and computational speed. Willow's ability to perform a complex benchmark computation in minutes that would take a supercomputer septillion of years signaled a major step towards building useful, large-scale quantum computers for tackling previously intractable AI problems.

ings 2024-2025

JANUARY 2025 - DEEPSEEK-R1 DISRUPTS LLM LANDSCAPE & CES SHOWCASES AI Pervasiveness

January 2025 witnessed a significant shake-up in the large language model (LLM) landscape with the launch of DeepSeek-R1. This open-source, cost-efficient model from China achieved performance comparable to leading proprietary LLMs at a fraction of the cost, challenging established norms and becoming a top-downloaded app. Concurrently, the Consumer Electronics Show (CES) highlighted AI's pervasive integration into daily life, with major announcements from NVIDIA, Intel, and Samsung showcasing new AI chips and intelligent features across consumer devices and smart living.

MARCH 2025 - INDIA'S GCCS EVOLVE INTO GLOBAL INNOVATION HUBS

March 2025 marked a significant shift for India's Global Capability Centres (GCCs). With over 1,700 active centers and projections to exceed 2,100 by 2030, these hubs are transforming India from a back-office support destination into a powerhouse of global innovation, increasingly driven by AI and analytics. Lalit Ahuja, CEO of ANSR, aptly described this evolution as "the end of the beginning," signaling a new era of advanced capability and strategic importance for Indian GCCs.

MAY 2025 - MICROSOFT & GOOGLE DRIVE ENTERPRISE AI FORWARD

May was pivotal for enterprise AI. At Microsoft Build 2025, CTO Kevin Scott introduced the "agentic web," a vision for interoperable and autonomous AI agents, highlighting foundational tools like MCP and NLWeb. Concurrently, Google Cloud Next 2025 showcased powerful advancements, including AI Hypercomputer infrastructure, refined Gemini 2.5 models for enhanced reasoning, and expanded Vertex AI capabilities with new generative media models like Veo 3 for video and Imagen 4 for images.

FEBRUARY 2025 - EU AI ACT'S CORE PROHIBITIONS TAKE EFFECT & UK'S STRATEGIC AI PLAN

February marked a significant regulatory milestone as core prohibitions of the EU AI Act officially came into force, banning high-risk AI systems deemed unacceptable. Simultaneously, the UK government launched its comprehensive AI Opportunities Action Plan. This strategic initiative, with 50 recommendations, aims to position the UK as a global AI leader through investments in infrastructure, talent, and a pro-innovation regulatory environment.

APRIL 2025 - WALMART'S AI LEAP

Building on earlier developments, April 2025 saw Walmart extensively showcase 'Wallaby,' its proprietary suite of large language models designed to revolutionize retail customer interactions. Powering conversational AI and generative AI assistants for product recommendations, order management, and customer support, Wallaby demonstrated how traditional retail giants are leveraging cutting-edge AI to enhance efficiency and customer experience at scale.

JUNE 2025 : AI FUELS ADVANCES IN VIDEO GENERATION, ROBOTICS & INDIAN

June witnessed diverse AI progress. Midjourney entered the competitive video creation space by unveiling its first text-to-video generation system, Model V1. In robotics, Foxconn and Nvidia announced plans to integrate humanoid robots onto electronics production lines. Locally, India's drone revolution gained momentum, with Bengaluru-based Scandron highlighting the deployment of high-altitude, AI-assisted drones for critical supply missions in the Himalayas, showcasing their increasing role in logistics and operational efficiency.





How to make the world a better place for AI

- Rohith Chandra

A quick question. Have you ever come across a situation where a product or service has been used in a way it was not meant to be used.....

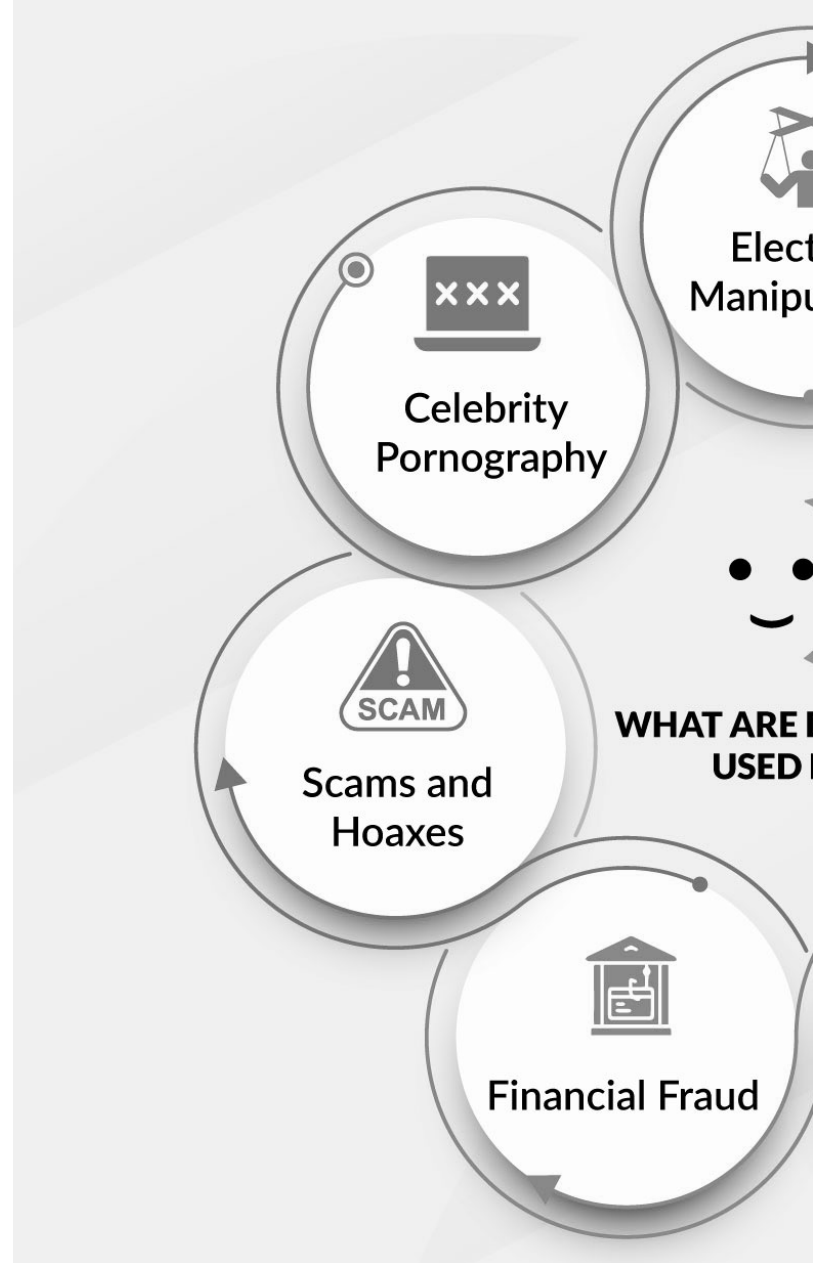
For me there are two inventions. The first one was a pencil and the user is John wick. If you understand this you know what I am saying. For those who didn't get it, he used it as a weapon to massacre people. What if the inventor of the pencil came to know about it, did he even think about it, will he/she be done anything to prevent such usage of it? Can they control the usage of their invention?

And another one was AI and the users are found in each part of the world, you don't believe me, here are instances where it has been used to hurt others and the victims of these incidents are not some common people.

1. This has been in recent news, where a feud between an influencer and some roasters subsequently escalated to trolling of the roasters by the influencer fans who are provoked by the influencer. In this process the influencer fans attacked a female roaster/creator by circulating morphed photos of her and even issuing threats online.

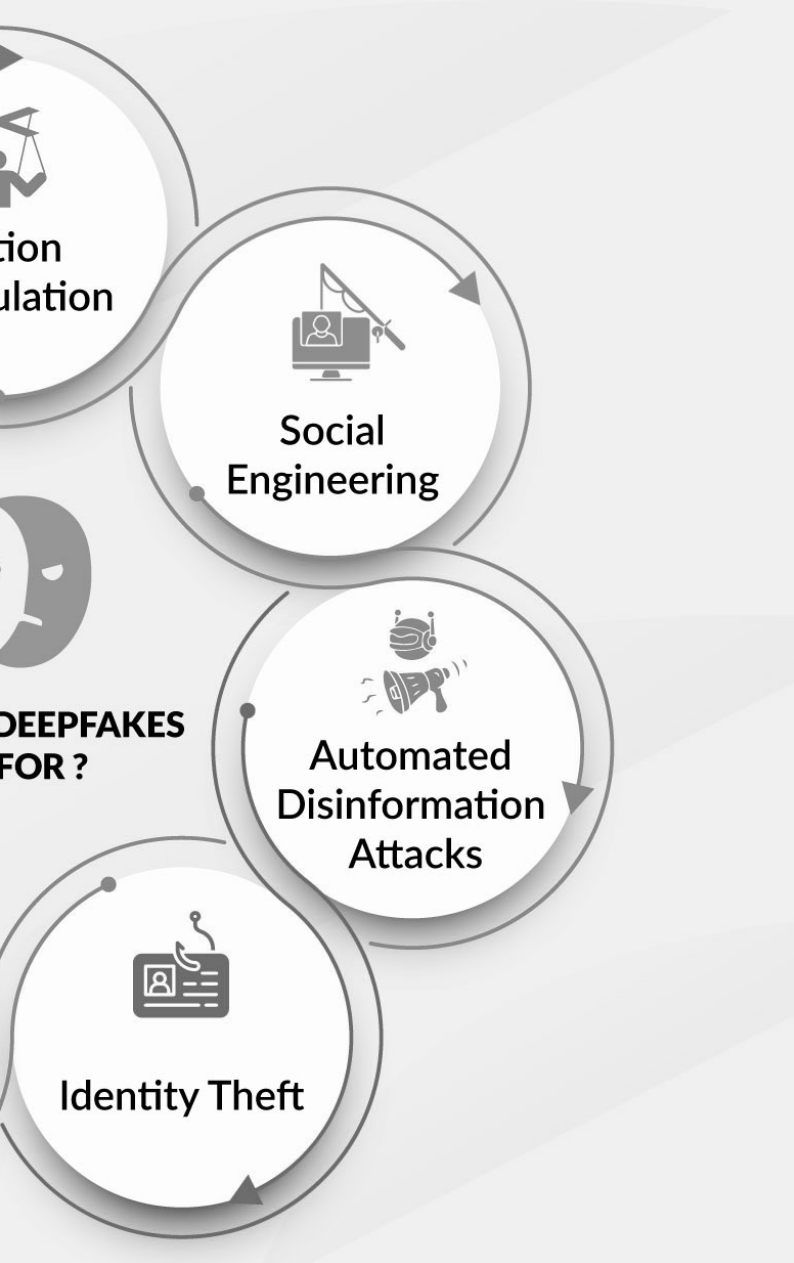
Here there are many issues but the issue that got my attention, and I also want you to focus your attention is that, Now a days with the advent of AI tools in image, audio editing and there free access combined with the anonymity of the internet has led to raise in these kind of issues like cyber harassment, disinformation (or in the language of the USA President "FAKE NEWS"), honey traps, child pornography, cyberbullying, blackmailing and what not.

2. If you think this is not applicable to powerful people then this next story will surprise you. In this incidents a PM of an nation was the victim, her face was deepfake (deepfake image is one where the face of a person is digitally added to the body of another, it's an portmanteau of deep learning and fake, if you to know more about it refer AINA 3.0) into a video and this was posted in a US pornographic website, where they were viewed "million of times" over the course of several months. And she along with her legal team seeks damages over the video, which she planned to donate to support women who have been victims of male violence. And this is not a rare incident. In recent years deepfake pornography (non-



consensual pornography) has become commonplace on the internet. Victims have spoken about the trauma of seeing their faces digitally edited onto photos of women in sexually explicit scenes.

If you want to experience what trauma a victim goes through, then you should watch this award winning documentary about the impact of deepfake porn - "My blonde GF" by Roise Morris. As per the victim in the documentary her picture could have been taken from an old Facebook account and professional short of her in the public domain. And as per the director's words what struck her when she met the victim was that ***you can sexually violate somebody without coming into any physical contact with them.***



AI to manipulate, deceive, or exploit users or customers to using AI to spread misinformation, propaganda, or hate speech. Additionally, AI misuse can involve violating privacy, security, or human rights, as well as creating or enhancing weapons, cyberattacks, or violence. It can even take the form of discrimination, oppression, or exclusion of certain groups or individuals. Such misuse can have grave repercussions on trust, reputation, accountability and social justice. Moreover, it can impede the value and potential of AI for good.

And the primary reason for this is the availability of the technology, for example, anyone can use online tools such as Deepswap and FaceMagic to swap one face for another in a video; Stable Diffusion can generate realistic-looking, false images. OpenAI's DALL - E can create images from text-based inputs — it can, for instance, conjure up an image of someone getting a coveted award. All that one has to do is type in the relevant words.

Preventing AI misuse requires a multifaceted approach:

1. Ethical guidelines and regulations
2. Transparency and explainability
3. Education and public awareness
4. International cooperation
5. Continued research and development

STATS GLOBAL AND INCIDENTS

1. In Dec 2023, an elderly couple living in Faridabad, on the outskirts of Delhi, received a call on WhatsApp from someone who said he was a U.S. police officer. He said their 35-year-old son, an engineer in the U.S., had been arrested in a rape case and demanded ₹46 lakh to set him free. The couple sent the money after receiving an audio file in which they could hear their son sobbing. The call was fake, as was the audio. Later in investigation the police said it is a voice-cloning technique that had possibly been used to con the parents.

2. October 2023: Jordan's Queen Rania Al Abdullah is seen in a video condemning Hamas

This motivated her to do the documentary.

This is just the tip of the ice-berg. There are many victims suffering similar kinds of trauma from deepfakes. Apart from this there are other cases of AI misuse ranging from scams, honeytraps, blackmails, voice morphing, disinformation, propaganda, fraud, social engineering, plagiarism, algorithm discrimination and many more.

AI MISUSE

Before we take a jump into further, I want you to take a step back and understand what I meant by AI misuse. **AI misuse is the deliberate or negligent use of AI that harms people's society, or the environment.** It can range from using

and supporting Israel.

3. March 2022: Ukraine President Volodymyr Zelensky (pictured) is seen in a deepfake video asking his soldiers to surrender to Russia.

4. According to the 2023 State of Deepfakes, a report by the U.S.-based Home Security Heroes, there's been a 550% increase in online deepfake videos from 2019 to 2023.

5. Europe-based Sumsb's identity theft report states that the number of deepfakes increased tenfold from 2022 to 2023.

After seeing all these stats. A boomer uncle/aunt will be tempted to say ban AI, this is as good as saying I cut my fingers because my nails are growing. The actual solution may be more like raising awareness among the stakeholders (such as developers, users, regulators, and companies), keeping systems in check to prevent these kinds of issues .

And coming to the stakeholders who may be thinking to whom we should point the fingers at. Instead of that they should see this as a collective



endeavour to create a foolproof barrier to prevent this kind of misuses. Let's dive into some solutions as per the roles of the major stakeholders.

Role of Technical teams, companies and developers

They are the first layer of defence in the process of creation of the AI products. They can control a significant part of the product. They can take steps to accomplish this goal, such as defining the purpose, scope, and limits of your AI project; following ethical principles and standards for AI design, development, and deployment; conducting risk assessments and impact evaluations for your AI project; implementing safeguards and controls to monitor and mitigate AI misuse; engaging with diverse perspectives and feedback for your AI project; and educating themselves and others about the ethical and social implications of AI, and updating with new technologies to counter the problems. A few of the methods they can follow like

Content credentials

Labelling content at source is a better way of battling deepfake. When a camera creates a picture or a video, it also generates a label or signature certifying that it has not come from any AI tool. Each time the content is modified or edited, additional sets of signatures are added to establish the authenticity of the media — text, audio or video. These signatures, also called content credentials, help investigators and others figure out when and how something was generated, and whether it had been edited or AI-generated.

In February 2021, software giants Microsoft and Adobe formed an alliance with three other technology companies, Arm, Intel and Truepic, to launch the Coalition for Content Provenance and Authenticity (C2PA). **C2PA is building an open-source technical standard that can be used to assess the authenticity of different types of digital media.**

Content credentials enable internet users and media consumers to distinguish between authentic and AI-generated images. It will not only help identify fake media but also enable content creators to establish the authenticity of their work in cases where genuine media is being labelled as

fake. Currently, C2PA is an opt-in facility — that is, it's still optional.

Its full benefits will only be unlocked when it becomes mandatory. Social media platforms are yet to opt for it, whether it had been edited or AI-generated.

Fake detectors

Development of these tools can help fight back against the battle against morphed content.

Intel's FakeCatcher detects deepfakes by using the remote photoplethysmography technology to analyse blood flow in the pixels of an image. Signals from multiple frames are analysed before a video is pronounced fake. Facial movements and a change in expressions and speech result in concomitant blood-flow changes in the face. The method can detect these changes by measuring the amount of light absorbed or reflected by different parts of the face.

In January 2024, computer security software company McAfee released an advanced detector

to specifically target audio deepfakes

Gujarat-based start-up Kroop AI has developed a detector called VizMantiz. "It is positioned for BFSI (Banking, Financial Services, Insurance) and social media platforms. Current users are large organisations in these spaces," says Jyoti Joshi, Founder and CEO of Kroop AI.

And there can't be an "ideal universal" detector. ***"Detectors will have to be case-specific." For example, something that works well in identifying fake identities and data in a banking environment may not be suitable for social media.***

It is unlikely that a one-shot solution to this problem will become available in the near future. But the quest is on, as there is a need for a solution.

Digital watermarks

Digital watermark is a kind of marker covertly embedded in a noise-tolerant signal such as audio, video or image data. Digital watermarks may be used to verify the authenticity or integrity of the



ORIGINAL



DEEPFAKE

This watermarking plays a vital role in ensuring the integrity and authenticity of AI-generated content. By embedding unique identifiers into synthesized audio, watermarking enables precise localization and detection of manipulations, safeguarding against fraudulent activities and misinformation. Moreover, it fosters accountability by providing evidence of ownership, thereby resolving disputes and facilitating copyright protection. In a landscape

Role of Regulators and Governments

So the main problem is the regulations need to be flexible and adaptive but also should not compromise on the risks and potential damages due to its misuse.

And this process is not a race but rather a marathon in an unknown track, where there are no overnight solutions, but long deep talks with the stakeholders for a better forward path.

The recently enacted EU AI Act has a risk-oriented regulation system, which is less comprehensive. It tries to regulate AI by introducing a risk-based framework that categorises AI systems into unacceptable, high and low-risk levels. This ensures AI applications are governed according to their potential for harm. However, the Act's reliance on broad risk categories has been criticised for oversimplifying the risks' complex and rapidly evolving nature. When the regulatory system fails to be adaptive, it becomes rigid. This can result in outdated oversight, allowing harmful applications to slip through the cracks

The US has taken a more sector-specific approach, with agencies like the Federal Trade Commission issuing guidelines on AI usage, particularly



concerning fairness and transparency in consumer protection. The US has also proposed the COPIED Act, which aims to enhance transparency around AI-generated content by requiring disclosures about its role in creating or altering media. The UK has opted for a principles-based framework, emphasising ethical AI development and innovation without rigid compliance requirements.

China has a more stringent, centralised framework that prioritises national security and social stability. The government mandates strict oversight of AI technologies, with regulations like Algorithm Recommendation Service Management Regulations requiring companies to disclose the operational details of their algorithms. The enforcement of these rules is robust and is closely integrated with the nation's surveillance apparatus. Ironically, despite this, between 2014 and 2023, an astonishing 54,000 GenAI-related inventions were filed in China, the highest in the world.

As India's policymakers carefully mull the next steps on AI regulation, the brief pause in this continuing advance offers an opportunity to reflect and readjust, lest policymakers get trapped in path dependency and a mindless rush to regulate.

Role of Users and Society

AI is like a two edged sword, it can be used for solving problems and also for creating new problems it all depends on the users and their intentions.

The best solution is spreading awareness on the issues and being proactive, reporting these misuses when you come across them. And you don't need to educate all the people in the world, you just need to focus on the person to your left and the person to your right, that is all it takes.

Users need to know their rights, responsibilities and duties like personal data privacy, personal data protection laws, following the watchdog bodies guidelines. Because you may never know when you're one click away from being a victim.

And whenever you come across any video or audio observe for any range of discrepancies or irregularities that can be contextual, spatial, textural — and more. One typical example maybe

the direction of a gaze — if a person's gaze does not match the direction the eyes face, then it is a sign of manipulation. When a person faces a camera, the gaze is towards the camera, too. But when an actual face is replaced by a fake one, it is possible that the look is not towards the camera. Another red flag is the disparity between a mouth opening and shutting and the intensity and frequency of the speech. When an image is swapped with another, there may be a mismatch between the compression factors of the two images, since JPEG images are often compressed.

I can say more on this, but I am stopping here so that you can start looking into the various solutions to the problems.

Conclusion

Maybe due to the unrestricted availability of the AI tools, lack of awareness, lack of legal barriers and other restrictions the perpetrators of AI misuses are roaming in the stress of the internet with their anonymity masks on. What they don't know is that there are systems and barriers in construction to restrict their misdemeanor.

AI and its tools are found with the intentions to make our life easier but not to make it complex. But in the hands of some wicked minded individuals these tools were made as weapons to hurt others. Today they can evade like sand passing from the gaps of the palm when they try to hold it, maybe due to legal, policy loopholes. And I believe that tomorrow this won't happen because the fist will be clenched so tight that even air can't escape and these individuals will suffer the consequences of their wicked acts. May the force be with us.

I want to dedicate this article to my mother, sister and my best friend 'D. A S'.

-Rohith Chandra

Nitish Singh

Founder - CampusX

Industry - Education

Nitish is an educational content creator on YouTube with 150K+ Subscribers in the domain of Data Science. He has been in the tech industry for the past 10 years and taught more than 50,000 students offline. He is passionate about data and is an expert in creating content that simplifies complex topics.



AINA: Nitish sir, before we dive deeper, could you please share your journey with our audience — in your own words?

Nitish: I did my engineering in electrical from 2009 to 2013. But by my second year, I realized I was just clearing exams without gaining real, practical skills. That bothered me — engineering was supposed to be more hands-on. My first real exposure came through DIY kits like Arduino and Raspberry Pi, which got me into robotics and programming. I started building simple robots and gradually got good enough to win a few competitions. What started as a personal

learning journey soon turned into something bigger. Along with friends, I began teaching juniors how to build robots — our college didn't offer much in that space. The response was amazing, so we took the idea further, formed a small company, and started conducting robotics workshops in other colleges and cities. That was my first brush with tech education, and I loved it.

Over time, I expanded into software, teaching web development, Android, and other in-demand skills. Around 2016, I got deeply interested in machine learning. It felt like the

future, and I immersed myself in it. By 2020, I launched a YouTube channel to share what I'd learned. Since then, I've created over 1,000 videos, and in doing so, learned even more.

While I haven't spent much time in the corporate world, teaching has been my core for over 12 years now. From robotics to ML, I've taught more than 50,000 students, online and offline. That journey, of learning, teaching, and building communities, is what led to CampusX.

AINA: You've taught thousands of students and also explored real-world challenges. Before

starting CampusX, what gaps did you see in the system? How did those lessons shape what CampusX is today?

One thing I noticed early in college was how disconnected teaching felt. Back in school, learning was fun, but in college, lectures didn't inspire. Teachers would just deliver content without context, and I found myself zoning out. That made me question things. Why was I not enjoying learning anymore?

Over time, I realized the problem was in the teaching approach. So when I started teaching, I promised myself: my classes would never be like that. For me, a key metric is, if 100 people join an online session, how many stay till the end? Retention shows if the class actually engaged them. CampusX was born out of this belief:

students deserve better.

I built a teaching method around three key ideas. First, always answer the why. Before teaching what or how, I explain why something exists. That story behind a concept, like why transformers came into the picture — makes a huge difference.

It's human nature to connect with narratives.

Second, I use first principle thinking, teaching as if we're discovering the concept together from scratch. Instead of just giving results, I thoughtfully recreate the journey. That way, students don't just learn the ideas, they believe they can invent too.

Third, examples matter. A well-placed example at the right moment changes everything. It's like prompting an AI — the better your example, the better the understanding.

And finally, no matter how good the theory is, you need practice.

You can't learn swimming by just watching videos — you have to get into the water.

It's the same in data science or AI. That's why we always tie concepts to projects and hands-on practice. These are the principles I've followed for over a decade, and CampusX is built on them — to bridge the gap between knowing something and being able to do it.

AINA: In India, there's often a gap between what universities teach and what companies expect — especially in fast-moving fields like data science and AI. Is this just about outdated syllabi and lack of practical exposure, or is there a deeper disconnect? And what can educators, students, and companies do to bridge this?

Nitish: Academia plays two very important roles. One is research — foundational, cutting-edge work that often shapes the future. Take IISc, for example. Some of India's achievements in rocket science trace back to the research culture built there. And in the global context, look at transformer architecture — the backbone of modern AI. It was born not in industry, but in a research lab at Google Brain. The second role is education —

training students so they can either contribute to research or be ready for industry. Ideally, academia and industry should work together. Industry funds research, academia produces talent, and both benefit from the discoveries made. On paper, it's a great model. But in computer science, this breaks down — because the pace of change in the industry is just too fast. Every year, a new wave of technology hits — from cloud platforms to GenAI — and it rewrites the rules. That's the Moore's Law effect in action: as hardware gets faster and cheaper, the software and use-cases explode. But academic curricula can't keep up.

Take Generative AI. In just two years, it's become the talk of the town — companies are hiring, products are launching, and interviews now have entire rounds focused on it. But if you check the syllabus in most Indian engineering colleges, GenAI is still nowhere. Why? Because curriculum changes go through multiple layers — from AICTE assessing the trend, to policy making, to actual implementation across colleges. That whole pipeline is slow.

And even when the syllabus updates, who will teach it? Imagine a college in a Tier 3 town. The top GenAI talent isn't exactly lining up to take a teaching post there. They're working in metro cities or abroad. So the execution gap remains. Meanwhile, students graduate having learned technologies that are five or ten years behind what the industry wants. That's the real issue -

The growth of academia and industry isn't in sync.

In some fields like MBA, where frameworks evolve slowly, the gap is manageable.

But in CS and AI, the gap just keeps widening.

So what can we do? Let's break it down:

Educators: I'd split this into two groups — policymakers (like AICTE, education ministries) and teachers. Policymakers need to be proactive, not reactive. They must forecast where the field is headed and update curricula ahead of time. And they need to rethink scale. Rather than opening more physical colleges without quality, why not scale online? Use digital platforms to bring the best instructors and latest tech to every student. Imagine an online IIT or IIM experience — accessible, up-to-date, and credible. It's possible. As for teachers, I believe there's a moral responsibility. You can't just say, "I'm a researcher, not an instructor." If a breakthrough like GenAI is happening, you have to stretch your comfort zone, upskill, and bring that to your students. Because they're counting on you to prepare them for the real world.

Students:

Your learning shouldn't end at the classroom door.

Today, all knowledge is available online — from Python to prompt engineering. Learn publicly. Share what you learn. Help your juniors. Build a community. Too often, students focus only on their own growth. But when you

teach others, you reinforce your own learning — and you raise the bar for everyone around you.

Companies: They too have a role beyond hiring. I always say — **if you have CSR budgets, adopt a college.** Set up a small research lab. Send your engineers to give guest lectures. Donate infrastructure — a few GPUs or a working dev environment can change lives. And run placement drives there. Some companies already do this, and I see the impact in places around NCR. When students hear firsthand what industry is working on, it bridges the gap in a very real, tangible way.

So yes, the system is flawed — but not unfixable. With better collaboration and a shared sense of responsibility, we can close the loop between what's taught and what's needed.

AINA: CampusX offers a unique blend of free YouTube content and paid programs. How do you maintain the rigor and quality of both offerings without compromise? What's the strategy behind ensuring that free content remains as valuable and impactful as the paid ones—especially in an EdTech space that often struggles to balance accessibility with business sustainability?

Nitish: Having worked extensively in EdTech, I've realized that the industry revolves around a key trade-off: accessibility vs. sustainability. You can either offer highly accessible, low-cost content like Khan Academy, or build a sustainable business with

premium offerings like Byju's. At CampusX, I've always leaned toward accessibility, without compromising on quality.

Our philosophy is simple: we teach around 75% of any topic for free, helping learners get a strong foundation. For example, our 100 Days of Machine Learning offers 100 hours of free content that can take you to 80% proficiency. The remaining 20%—the deeper, advanced part—is covered in our 500-hour DSMP paid program.

We don't hold back quality in our free content. In fact, we see it as a way to show students the quality of our teaching. If they resonate with it, they're free to go deeper with our paid programs. It's honest, value-driven, and sustainable—for both learners and us.

AINA: With GenAI reshaping data workflows, what habits or frameworks would you recommend for students and professionals to build a self-driven, future-proof learning mindset? For example, should learners prioritize project-based experimentation, follow specific thought leaders, or audit their skills quarterly?

Honestly, with GenAI, we're all adjusting on the fly. Even experienced professionals are guessing where it's going. A few years ago, data science was fully manual—collect data, clean it, do EDA, try different models, tune them, deploy—all by yourself. Now, GenAI tools handle much of this: automated cleaning, insight generation, model

selection, even code writing. The “how” of data science is being outsourced to AI. And this shift is only accelerating.

But here’s the key:

***when tools do the “how”,
you need to master the
“why”.***

If an AI tool writes buggy code and you don’t understand the underlying logic, you won’t catch mistakes—especially in high-stakes projects. That’s why strong fundamentals—math, algorithms, domain knowledge—are non-negotiable. To stay future-ready, I’d recommend a few habits:

- Keep doing real-world, project-based experimentation.
- Audit your skills regularly to spot gaps.
- Follow thought leaders, but don’t just scroll—engage deeply.

Think of GenAI as your co-pilot. But to really stay in control, you must know how the system works under the hood.

AINA: Many focus on tools like Python but skip the math—statistics, algebra, calculus—that power data science. How crucial is it to learn the theory? Any example where weak foundations caused issues? And how does CampusX keep fundamentals front and center?

Absolutely, math is essential. You can run basic algorithms without understanding the

math, and it might work on toy datasets like Titanic—but the moment you pick up serious projects, surface-level skills start breaking down. I remember early in my deep learning journey, I worked on an emotion detection project using CNNs. Initially, I followed the code and trained the model, but the test accuracy wasn’t improving at all. I had no clue what was going wrong. That’s when I realized I had to go back and understand the mathematics of gradients, backpropagation, vanishing gradients, optimizers, activation functions—all at a deep level. Only then could I diagnose and improve the model.

So yes, to move from 80% to 85–90% performance, you must understand what’s happening under the hood. Otherwise, you’re just blindly running code.

At CampusX, we prioritize math heavily. Our lectures often run over an hour for topics others cover in 15 minutes because we go deep—right down to deriving formulas from first principles. I’d rather YouTube not promote my videos than compromise on rigor. That’s also why our Data Science Mentorship Program spans over 500 hours—just for machine learning. We ensure students don’t just use tools—they understand the core math and theory that make those tools work.

AINA: What’s your take on structured programs like PGDBA? In today’s landscape, do formal credentials still matter—or can a strong portfolio of projects be just as valuable for

landing data roles?

Both matter and will continue to coexist because they solve different problems. Programs like PGDBA serve two purposes—teaching in-demand skills and giving companies a trusted hiring channel. If I’m a recruiter and I need dependable talent, I’d prefer hiring from a reputed institute like IIM Calcutta, where the quality is assured and the process is structured.

But formal programs can’t address urgent or ad hoc needs, where companies need someone immediately. That’s where platforms like Naukri come in. However, when hiring from a general pool, the portfolio becomes the signal—it helps identify standout candidates from a crowd of 10,000.

So, structured credentials are great for planned hiring; strong portfolios shine in open competition. Both have distinct values and will always be relevant.

AINA: Question: As an industry expert, what’s the next big thing in AI that excites or concerns you? How should data science education evolve to prepare students for shifts like agentic AI, explainability, or global competitiveness?

Let me answer this in two parts—what excites me, and what concerns me.

What really excites me is ***the rise of agentic AI***. Generative AI has already matured, but now we’re moving

towards AI agents—autonomous systems that can interact with software tools, APIs, databases, and more, using natural language as the interface. This was the missing puzzle piece, and LLMs have made it possible. Imagine having a personal AI assistant you talk to daily—it's not science fiction anymore; it's coming.

The second area is **explainable AI**. Right now, large models are black boxes. They can give you answers, but they can't explain why they gave a particular response. This is a serious concern, especially as LLMs get embedded into decision-making systems—banking, bookings, transfers, etc. If something goes wrong, there's no audit trail. So I see a huge wave of research and innovation coming in explainability tools, and that's something I'm personally very excited about.

Now, the concerns. First is **job restructuring**. I don't believe AI will eliminate jobs permanently, but it will change them. One project that used to need 10 people might now need 5. Those displaced may find new roles because more products will be built, but there will be a transition period of instability—and that's already visible today. My second concern is **India falling behind in the AI race**.

The US is clearly ahead, and China is catching up fast. We've not built foundational models yet—we're more users than creators. It's not due to lack of talent; Indians are leading AI work abroad. But here, we face

challenges—limited funding, restricted access to GPUs, and a gap in hardware innovation. It worries me that we might fall behind like we did in space or software.

That said, I believe in miracles. Just like Sachin changed Indian cricket in the '90s, I hope one of you becomes that game-changer for India in AI.

AINA: What has been the biggest challenge you've faced in your career, and what did you learn from it? If you're comfortable sharing, are there any regrets—things you feel you should have done differently?

To be honest, I've never worked in the industry as a data scientist. For me, data science has always been about learning and teaching. So I never faced the kind of technical or scalability challenges someone might encounter in a corporate setup. But what I can share is my journey as an entrepreneur in the edtech space—and that has come with its own lessons and regrets.

For the first 7 years of my teaching career, I taught offline in classrooms. We had a full setup in Kolkata—an office, a classroom space, and we were doing really well. But then COVID happened. Practically overnight, the entire business collapsed.

Looking back, my biggest regret is that I didn't shift to online teaching earlier. I hadn't realized how much more resilient and scalable online education could be. Online teaching protects you

from disruptions like COVID, drastically reduces operating costs, and—most importantly—can scale in ways offline never can. In a physical classroom, if 60 students show up, and another 60 want to join, you either need to teach twice or find more space. But online, you can teach once, upload it to YouTube, and lakhs can watch it.

That experience was a turning point. I started from scratch with YouTube and have since committed to teaching online.

So my biggest learning has been this: **always build for scale**. Whatever you're doing—whether it's education, product development, or content creation—ask yourself: Is this scalable? Can it grow without me having to invest an equal amount of extra time?

AINA: With recent trend breakers like DeepSeek, is "LESS IS MORE" really possible in AI/ML regarding computation, model complexity, or data?

Since the Transformer paper, the pattern in deep learning has been "bigger is better." More parameters, more data, better hardware — that's how model performance has improved, from GPT-1 with millions of parameters to GPT-4 with hundreds of billions.

But this approach has serious downsides. First, cost — training foundation models requires huge investments in GPUs and data, which is challenging for smaller companies and countries like India. Second, environmental impact — massive CO2 emissions

and enormous electricity usage, with some estimates saying AI training might consume 30% of the world's electricity by 2030. Third, these huge models can't run on edge devices like phones or robots, raising privacy and accessibility concerns.

Because of these problems, the AI community is shifting mindset towards smaller, efficient models that are useful and practical. For example, domain-specific models for legal consulting or medical assistants focus on narrower tasks instead of general chatbots.

Techniques like quantization (reducing precision from 32-bit to 4-bit or 8-bit) drastically shrink model size, making on-device AI possible, even offline.

Also, instead of training on massive internet datasets, synthetic data tailored for specific use cases is emerging, enabling models to learn effectively from smaller, targeted data.

DeepSeek is a great example—they built a high-performing model using far fewer resources, proving “less is more” can work and is becoming a necessary direction for AI's future.

AINA: What foundational skills are immune to automation, and how can learners and educators emphasize them over transient tools?

The core skills that AI can't easily replace are strong fundamentals in statistics, probability, and linear algebra—the ability to solve problems mathematically with precision and understanding. AI

today can't replicate that kind of mathematical grace.

Next is problem-solving and critical thinking—breaking down problems, planning, executing with persistence, and thriving under pressure. This mindset, like the relentless attitude seen in Bollywood heroes, is key to surviving automation.

Lastly, deep domain expertise matters hugely. Spending 5–10 years in a domain like banking or healthcare builds invaluable insights that AI cannot substitute. That's why I advise students to stay in the same domain rather than hopping jobs or industries—domain knowledge leads to better decision-making and higher value over time.

AINA: How do you view the growing preference for short, “quick fix” content in online learning? What advice would you give learners to avoid shortcuts and focus on developing deep, lasting expertise? How can content creators support this shift?

Nitish: This “quick fix” culture is definitely a challenge I've seen firsthand—as a creator and learner. When I make videos, I want to cover topics thoroughly, sometimes in long, detailed sessions. But learners often prefer short, 15-minute videos they can fit between daily tasks. The temptation to choose shortcuts is strong because people want quick answers.

Platforms like YouTube have accelerated this trend by rewarding shorter videos—people watch them more, so the

algorithm pushes creators to make bite-sized content. This creates a cycle: creators make shorter videos to get views, and learners consume less-depth content, reinforcing the preference for “quick fixes.”

But here's the reality:

A 15-minute video can help you prepare for an exam question, but it won't build deep expertise

needed to create complex AI systems or software. Those skills come from years of study and practice—not quick tutorials.

My advice for learners is to have realistic expectations: know the difference between studying for exams and building real capability. Expertise takes time and patience.

For creators like me, it means choosing principles over instant gratification—making longer, detailed content even if it means fewer views today. Over time, serious learners find such content valuable, and the long-term rewards come, even if slowly.

So, resisting the “quick fix” means balancing self-awareness as learners and patience and integrity as creators. If both sides commit to this, we can shift the culture toward more durable, meaningful learning.

AINA: If you could give one golden piece of advice to students overwhelmed by the sea of online courses and FOMO, what would it be? How can they design a learning roadmap that is structured enough to build

expertise but flexible enough to adapt to AI's rapid evolution?

The first and most important step, actually step zero is patience.

Becoming a data scientist or AI expert is not a 30-day or even 6-month sprint. Roughly, it takes about 1,000 hours (about a year) to get a fresher-level job and closer to 10,000 hours of deliberate practice to become a true expert. If you don't have patience to invest this time, it's better to reconsider now.

Once patience is in place, here's the roadmap I share with my students:

- **Identify your target role.** Data science has many roles—data analyst, data scientist, machine learning engineer, AI engineer, etc.

Don't chase all roles at once; you'll dilute your focus and end up "best at nothing." Research these roles, talk to people on LinkedIn, understand their daily work, and find the one that aligns best with your skills and interests.

- **Build and commit to a focused roadmap.** Based on your chosen role, create a study plan for 6 months to a year. Use available resources—online courses, ChatGPT, mentors, colleagues—and refine your plan through trial and error. Stick to this roadmap firmly. Changing paths frequently wastes time and energy.
- **Document and teach what you learn.** The biggest challenge in data

science is retaining knowledge. Writing blogs, making videos, or teaching peers forces you to reinforce concepts. Teaching creates "muscle memory" that helps the knowledge stick.

- **Complete milestone projects.** After learning Python, build a project; after mastering pandas, build another; then move to machine learning algorithms and projects. Practical application is key.

If you follow these steps consistently, I am confident you can land fresher-level roles within a year.

Also, ***stay flexible.*** AI evolves rapidly, so stay curious, update your skills regularly, but don't fall into the trap of chasing every shiny new thing. Balance structure with adaptability—and remember, deep expertise requires time and focus.

AINA: Thank you, for joining and sharing your valuable perspectives and advice—not just for students, but for anyone navigating the fast-evolving world of AI and data science. Your insights on patience, focus, and adaptability truly resonate.

Nitish: It was a pleasure sharing my thoughts. I appreciate the opportunity to be part of this conversation. Wishing all learners patience, focus, and success on their journey.



INTELLIGENCE

capacity to learn and apply knowledge; ability to understand and reason

Biological Intelligence.

Artificial Intelligence

ARTIFICIAL INTELLIGENCE

- Development of computer systems to perform tasks requiring human intelligence

Biological Intelligence

Neural network

Beyond Books: Intelligence Meets Algorithms

AI isn't just a tool—it's becoming a thinking partner in classrooms, challenging traditional notions of intelligence and unlocking new dimensions of curiosity.

Intelligent Chains

Reinforcement Learning
Agents in Dynamic
Supply Chains

Praveen



Reinforcement Learning



Rewards on Controlling Event



Penalty for Homnet / Guerdon's

(priority)
0:44



Limit
60



30

!

Supply chains, once the silent workhorses of global commerce, are now being thrust into the spotlight. In an era defined by unprecedented disruption from pandemics to political conflict, supply networks have been stress tested and made aware of their own fragility.

Companies have struggled with flawed forecasting when predictions miss the mark, the result is delivery delays and inventory gluts or shortages that hurt both customers and the bottom line. Bottlenecks, shortages, and outdated forecasting models have laid bare a simple truth, instinct alone can no longer steer the complexities of modern logistics. There is a pressing need for data driven, adaptable decision making that can keep pace with today's volatility.

Technology is rising to meet this need. Industry leaders are betting big on artificial intelligence and advanced analytics to enhance agility and resilience. Among the arsenal of AI techniques, reinforcement learning (RL) is emerging as a particularly promising approach for complex operational problems.

**“I think AI is the new electricity.
Whatever industry you work in, AI will
probably transform it, just as 100 years
ago the rise of electricity transformed
industry after industry”
-Andrew Ng**

Supply chains are no exception. From inventory management to delivery routing, RL offers a way to let machines learn optimal decisions through experience potentially revolutionizing how supply chains are run.

Learning by Trial and Error: What is Reinforcement Learning?

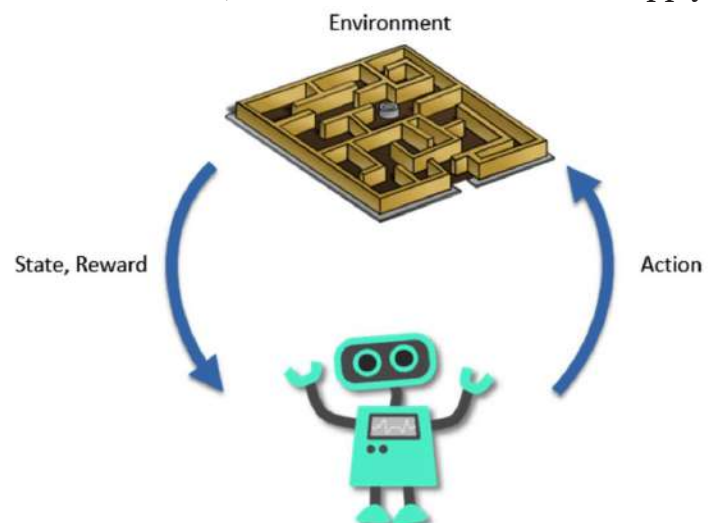
Reinforcement learning is a branch of AI inspired by behavioral psychology, the same principles by which animals (and humans) learn from interaction with their environment. In RL, an algorithm (often called an “agent”)

is programmed to make sequential decisions, and it receives feedback on each decision in the form of rewards or penalties. Over time, the agent's goal is to maximize cumulative reward, which drives it to repeat actions that yield good outcomes and avoid those that don't. In essence, the system learns by trial and error much like a child learning to ride a bicycle or a supply chain planner learning to adjust forecasts. Through enough cycles of experimentation, an RL agent can discover effective strategies even in complex situations, without being explicitly programmed for each scenario. This contrasts with traditional algorithms that follow fixed rules.

This trial and error approach turns out to be a perfect fit for supply chains, which are essentially giant, dynamic decision engines. At each stage in procurement, inventory, distribution, transportation, planners must make a series of interdependent choices in an environment that is unpredictable and fast changing.

Traditional planning tools often rely on static models or heuristics that struggle to adapt when reality deviates from expectations. RL offers a way to encode intelligence into these decisions, enabling software to adapt as the environment evolves.

Rather than hard coding solutions for every possible scenario, supply chain managers can set goals (like minimizing cost or maximizing service), and let the system figure out the optimal path through simulations or real world learning. In other words, we can include into our supply



software a kind of decision making brain one that gets smarter with experience.

From Inventory to Delivery: RL Applications in Supply Chains Today

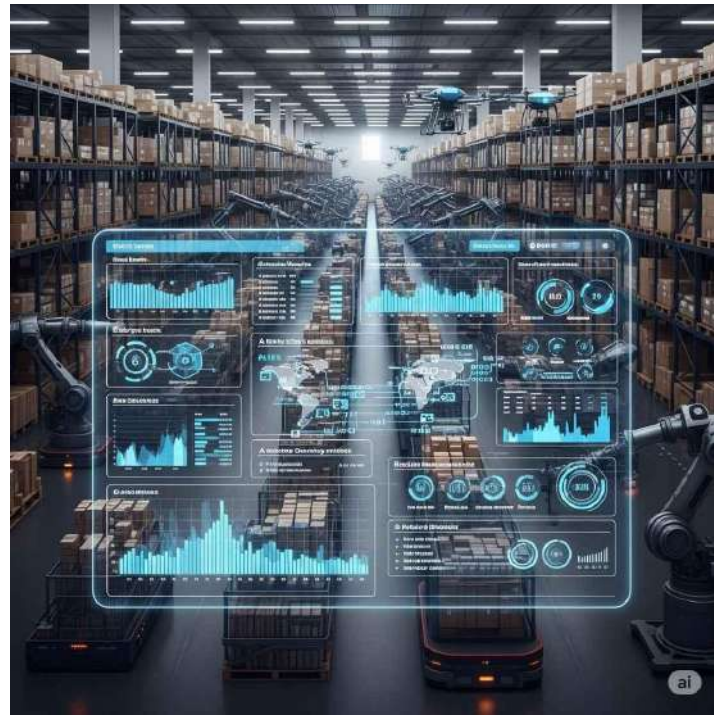
Early applications of reinforcement learning in supply chain management are already demonstrating its potential. In academic research, RL has been applied to a variety of supply chain problems, and a recent review found that inventory management is the most common focus area so far. This is the quintessential RL problem, deciding when and how much to reorder, balancing the cost of stockouts against excess inventory.

Researches have shown that RL agents can learn replenishment policies that adapt to demand patterns, sometimes outperforming traditional techniques and also demonstrated that an RL based ordering policy could efficiently dampen the notorious bullwhip effect, the amplification of demand variability upstream in a supply chain. By learning to make more responsive and calibrated ordering decisions, the RL agent reduced excess volatility in orders, thereby stabilizing the entire chain. This is a striking result. It suggests that -

AI agent, simply by trial-and-error learning, can discover strategies to mitigate a phenomenon that has vexed supply chain professionals for decades.

Beyond inventory, another area that is prone to disruption is logistics and transportations. Routing trucks or delivery vehicles in real time is a complex optimization problem, especially when conditions like traffic, weather, or last-minute orders can change the optimal path on the fly. Traditional route plans are static; an RL system, by contrast, could continually adjust routes based on feedback (e.g. delays encountered, new information), continually learning to improve delivery times and cost. In fact,

Global tech giant Amazon has started applying reinforcement learning to its logistics operations,



combining RL with other machine learning methods to optimize warehouse workflows and discover optimal paths for delivery.

As an Amazon news update described, the company's RL driven approaches are helping chart efficient delivery routes and even powering autonomous delivery drones. Amazon's experimental Prime Air drones use AI to navigate and make decisions in flight, a clear example of an autonomous system that must learn to operate safely and effectively in a complex physical environment. Reinforcement learning plays a key role in training such systems to handle variations in terrain, obstacles, and weather by rewarding successful deliveries and safe behavior.

Warehouse automation is another promising frontier where reinforcement learning (RL) is starting to make a real impact. Today's warehouses especially those operated by major e-commerce and retail players are increasingly populated by robots, from autonomous guided vehicles that move pallets to robotic arms that pick and place items.

"I used to be skeptical about reinforcement learning, but I'm not anymore."

-Geoffrey Hinton

Coordinating these machines and assigning tasks efficiently is a complex challenge, and one where RL proves particularly powerful. Take, for example, a robotic sortation center, where hundreds of mobile robots move packages to chutes headed for different destinations. Manually programming each robot to respond to shifting volumes and destinations simply isn't scalable.

To address this, Amazon researchers recently developed a multi agent RL solution. In their setup, each chute destination pair acts like an individual "agent" that learns how to assign packages in a way that boosts overall sorting efficiency. Trained across countless simulated scenarios, the RL system learned policies that outperformed traditional rule-based approaches, significantly improving the throughput of sorted packages. This example underscores one of RL's core strengths: its ability to uncover sophisticated strategies in complex, high-dimensional environments like orchestrating the movements of hundreds of warehouse robots.

Amazon uses multi-agent reinforcement learning (RL) for real-time delivery routing and warehouse optimization, improving throughput and adaptability.

As one supply chain technologist aptly put it, "Deep learning will revolutionize supply chain automation," unlocking innovations like autonomous vehicles and intelligent warehouses. Reinforcement learning, often powered by deep neural networks, is a crucial part of that transformation pushing automation beyond fixed routines to systems that learn and adapt to find the most effective way to operate.

It's worth noting that reinforcement learning isn't just useful for physical tasks it's also being explored for higher level decision making in supply chains.

Take procurement and pricing strategies, for example. Figuring out how to negotiate with suppliers or when to offer a discount to clear out extra stock are tricky decisions that RL agents could potentially learn by analyzing market data.

In fact, some early experiments have already used RL for things like dynamic pricing and contract negotiations in simulated environments, showing that these agents can boost profits and respond faster than human planners.

In manufacturing too, RL has been applied to production scheduling deciding the best way to sequence jobs and assign them to machines to cut down on lead times or costs. While many of these efforts are still in the research or pilot phase, they give a glimpse of just how many supply chain challenges could benefit from smart, learning-based approaches.

Benefits and Challenges on the Road Ahead

The vision of an AI driven "self learning" supply chain is compelling. In theory, a network of RL agents orchestrating the supply chain could respond to disruptions in real time, continuously optimize for efficiency, and even anticipate issues by learning from past patterns. We can imagine an AI based supply chain control tower that monitors global events, inventory levels, and demand signals, and automatically makes adjustments rerouting shipments, expediting production, reallocating stock all without waiting for human intervention.

As KPMG put it, we're heading toward "a future that promises autonomous, self-learning machines seamlessly managing the broader supply chain process". The benefits could be huge. These systems could dramatically cut down on waste like perishable goods expiring due to poor distribution avoid lost sales by reacting instantly to demand surges, and take routine firefighting off human shoulders.

**"We can build a much brighter future where humans are relieved of menial work using AI capabilities."
-Andrew Ng**

In the supply chain world, that means letting

algorithms handle the day to day complexity, so human experts can focus on big-picture thinking, innovation, and exception handling, solving truly complex challenges.

There are several hurdles to overcome:

- **Data and Simulation:**

RL needs lots of interaction experience to learn effectively. In a supply chain context, we can't always let an algorithm learn by trial and error on the real system the cost of mistakes would be too high if an agent, say, disastrously underorders inventory for a season. Therefore, companies need high fidelity simulations or "digital twins" of their operations where an RL agent can safely train. Building these simulations with accurate data is a non-trivial task. Many firms struggle with siloed or inconsistent data, which makes it hard to even know the current state of the supply chain, let alone simulate it. Investments in data infrastructure and IoT sensors (to get real-time state information) are often a prerequisite before RL can be applied successfully.

- **Complexity and Computation:**

Real supply chains involve a huge number of variables thousands of products, dozens of warehouses, fleets of vehicles, and so on. The state space and action space for an RL agent in this context are enormous. Current algorithms, even with deep learning, can hit computational limits trying to process such scale. Techniques like multi agent decomposition (breaking a big problem into many smaller RL agents, as in the Amazon sortation example) are one way to manage complexity. Even so, the industry will need to see further breakthroughs (or clever problem framing) to make RL tractable for, say, an entire end to end supply network in real time.

- **Exploration vs. Exploitation Dilemma:**

In live operations, an RL system must carefully balance trying new strategies (exploration) with sticking to what works (exploitation). In a warehouse or transport network, too much "exploration" could be risky for example, an agent might try an untested routing strategy that fails, causing delivery delays. Designing RL systems that learn safely, or that can be guided by human

knowledge (to avoid obviously bad choices), is an active area of research. Some approaches use robust or risk-averse RL algorithms to ensure the agent's policies stay within safe bounds. In high stakes industrial settings, pure unguided learning is rarely acceptable; practitioners are combining RL with rule based controllers or human oversight to get the best of both.

- **Trust and Organizational Buy-in:**

Beyond the technical issues, there's the human factor. Supply chain managers have to trust the recommendations or actions of an AI agent. If an RL system decides to reorder 30% more of a product than the forecast suggests, managers must believe that the agent "knows what it's doing" based on its learning. Building this trust may require transparency (explainable AI techniques to outline why the agent chose an action) and gradual integration (perhaps the RL suggests decisions for humans to approve initially, until it proves its merit).

Change management is a big part of introducing AI into any legacy operation essentially, the organization's culture needs to shift to embrace data driven automation.

Many supply chains are perfectly suited to the needs that the business had 20 years ago. The implication is that companies must update not only their technology but also their mindsets to leverage AI for today's challenges.

-Jonathan Byrnes

Despite these challenges, the momentum behind AI in supply chain is accelerating. Tech giants, startups, and research institutions are collaborating to push RL from theory into practice. We are seeing the first glimmers of what could become a self learning supply chain. In pilots, AI agents autonomously manage inventory replenishment for select products, or dynamically reroute delivery trucks during congested traffic hours, achieving improvements in service and cost. Each success builds confidence and paves the way for broader adoption. Governments and academia are also supporting this trend for

example, by funding research into supply chain resilience using AI, especially after the pandemic highlighted the economic importance of robust supply lines.

RL systems must carefully balance exploration with exploitation to ensure safe, reliable decision making in live operations.

A Holistic, Human Centric Future

If the past was defined by lean supply chains and just in time efficiency, the future belongs to adaptive supply chains networks that can sense, respond, and learn as they go. Reinforcement learning (RL) is a key enabler of this shift, offering tools that allow systems to adjust in real time and optimize themselves through experience.

But this doesn't mean humans will vanish from the picture. In fact, human expertise will be more important than ever. What we're likely to see is a human-AI partnership, where RL systems take care of constant number crunching and routine decisions, while human managers steer the overall strategy, set guardrails, and weigh in on exceptions and ethical dilemmas. In this future, supply chain professionals may act more like orchestra conductors guiding a smart, self learning ensemble that executes complex operations with precision.

If implemented thoughtfully, the benefits to society could be far reaching. Picture supply chains that minimize waste on their own rerouting perishable goods to where they're needed most, pre-positioning resources ahead of natural disasters, or reducing emissions by optimizing delivery routes. Such intelligence could lower costs for consumers while also making global supply networks more sustainable.

At the same time, this level of automation raises important questions about the workforce. Many operational roles could be transformed or even replaced by AI decision makers. That's why it's critical for companies and governments to proactively invest in retraining programs and

develop new roles for a logistics world augmented by AI. No one can say exactly when fully autonomous supply chains will arrive. Optimists argue we may only be a few breakthroughs away just one or two "AlphaGo moments" in logistics could kick off rapid adoption. Others are more cautious, pointing to the messy, unpredictable realities of physical supply chains that don't neatly map to the clean logic of games. Still, the progress is undeniable. RL has already started proving itself in areas like inventory management the "low hanging fruit" and the focus is now expanding to transportation, manufacturing, and multi party coordination. Every incremental win, whether it's a small drop in logistics costs or a bump in service quality, will build momentum.

While automation expands, human expertise remains vital working in tandem with RL systems to guide strategic decisions and ensure safe operations.

In the end, reinforcement learning is poised to become a core part of the modern supply chain toolkit. It signals a shift from reactive firefighting to proactive and even predictive operations. The road ahead includes both technical and organizational challenges, but the direction is clear. Companies that embrace RL and other AI technologies will gain a serious edge transforming their supply chains from rigid structures into adaptive, intelligent systems.

This evolution may not be as glamorous as self-driving cars or humanoid robots, but its impact will be far reaching. It means more reliable store shelves, faster deliveries, and a more resilient global economy. And in the not too distant future, the smartest decision maker in the supply chain might just be an AI agent that taught itself the best way to operate.





Trilytics Conclave 2024

The third edition of the **Trilytics Analytics Summit** was organized by the Conclave Team of PGD-BA on **August 10–11, 2024**. The event featured engaging keynote addresses, a competitive data analytics case competition, and insightful panel discussions. The summit was sponsored by **World Wide Technology (WWT)** as the Title Sponsor, with **SBI** as the Associate Sponsor. Their unwavering support helped underscore the growing significance of analytics in today's world.

The Trilytics Case Competition attracted over 9,300 participants from 100+ institutes worldwide, including participants from **nearly all IIMs and top B-Schools**, as well as **IITs and premier engineering colleges**, all competing for a **prize pool of ₹1,80,000**.



Event highlights included:

- **Aditya Prabhakaran (WWT)** delivered a keynote on “*Unlocking Value from Gen AI: Enterprise Use Cases*”, emphasizing the application of LLM modeling.
- **Animesh Kishore**, Head of COE at **ITC Ltd.**, shares insights on “*Re-defining the FMCG Business with Analytics*”, focusing on the role of **supply chain** and **retail analytics** in industry transformation.
- A thought-provoking panel discussion on “*Demystifying the Black Box: Exploring the Business Imperatives of Explainable AI*”, featuring **Setu Shah (Oracle)**, **Dr. Arghya Ray (IMI Kolkata)**, and **Vinay Garg (WWT)**, moderated by **Rahul Kumar (IIM Calcutta)**, exploring the responsible applications and future frontiers of this groundbreaking technology.
- An exclusive workshop on developing an “**Analytical Mindset**”, conducted by industry experts **Tanoy Dewanjee** and **Debayan Mitra** from **HSBC**.

The winners of the event (case competition) were **Team Random Foresters** from **PGDBA, IIM Calcutta**.



The third edition of Trilytics stood as a testament to the growing momentum of data-driven decision-making across industries, offering a vibrant platform for learning, collaboration, and innovation. With overwhelming participation, cutting-edge discussions, and invaluable industry support, the summit successfully bridged the gap between academia and real-world analytics. As the PGDBA Conclave Team looks ahead, future editions of Trilytics will aim to scale new heights—expanding its global footprint, deepening industry engagement, and continuing to shape the next generation of analytics leaders.



References

Neurosymbolic AI

Sinuhe. (2024, December 17). The first Nobel Prize for Neurosymbolic AI: A historic moment in AI. CMG. ASIA. <https://cmg.asia/the-first-nobel-prize-for-neurosymbolic-ai.html>

Rafalski, K., Rafalski, K., & Netguru. (2025, April 30). Neurosymbolic AI: bridging neural networks and symbolic reasoning for smarter systems. Netguru. <https://www.netguru.com/blog/neurosymbolic-ai>

Wan, Z., Liu, C., Yang, H., Raj, R., Li, C., You, H., Fu, Y., Wan, C., Li, S., Kim, Y., Samajdar, A., Lin, Y., Ibrahim, M., Rabaey, J. M., Krishna, T., & Raychowdhury, A. (2024). Towards efficient Neuro-Symbolic AI: from workload characterization to hardware architecture. IEEE Transactions on Circuits and Systems for Artificial Intelligence., 1(1), 53–68. <https://doi.org/10.1109/tcasai.2024.3462692>

Wan, Z., Liu, C., Yang, H., Li, C., You, H., Fu, Y., Wan, C., Krishna, T., Lin, Y., & Raychowdhury, A. (2024, January 2). Towards Cognitive AI Systems: a Survey and Prospective on Neuro-Symbolic AI. [arXiv.org](https://arxiv.org/abs/2401.01040). <https://arxiv.org/abs/2401.01040>

Garcez, A. D., & Lamb, L. C. (2020, December 10). Neurosymbolic AI: The 3rd Wave. [arXiv.org](https://arxiv.org/abs/2012.05876). <https://arxiv.org/abs/2012.05876>

When AI earned Nobel glory

The Nobel Prize in Physics 2024 - [NobelPrize.org](https://www.nobelprize.org/prizes/physics/2024/summary/). (n.d.). [NobelPrize.org](https://www.nobelprize.org/prizes/physics/2024/summary/). <https://www.nobelprize.org/prizes/physics/2024/summary/>

The Nobel Prize in Chemistry 2024 - [NobelPrize.org](https://www.nobelprize.org/prizes/chemistry/2024/summary/). (n.d.). [NobelPrize.org](https://www.nobelprize.org/prizes/chemistry/2024/summary/). <https://www.nobelprize.org/prizes/chemistry/2024/summary/>

Nobel Prize in Chemistry 2024. (2024, October 9). Nature. <https://www.nature.com/collections/edjcfidihdi>

Nobel Prize in Physics 2024. (2024, October 8). Nature. <https://www.nature.com/collections/ehbjaifcgc>

AI in Space

Restack. (n.d.). AI for space exploration: Answer-free download. [https://www.restack.io/p/ai-for-space-](https://www.restack.io/p/ai-for-space-exploration-answer-free-d%20ownload-cat-ai)

[exploration-answer-free-d%20ownload-cat-ai](https://www.restack.io/p/ai-for-space-exploration-answer-free-d%20ownload-cat-ai)

The case for configuration management. (2001, September 1). IET Journals & Magazine | IEEE Xplore. <https://ieeexplore.ieee.org/stamp/stamp.jsp?arnumber=963401%205&tag=1%20https://www.jetir.org/papers/JETIR1912274.pdf>

Kapoor, V. (2024, July 6). AI in space exploration: All you need to know. iPleaders. <https://blog.ipleaders.in/ai-in-space-exploration-all-you-need-to-know/>

Garanhel, M. (2025, March 14). AI in space exploration. AI Accelerator Institute. <https://www.aiacceleratorinstitute.com/ai-in-space-exploration/>

ResearchGate. (n.d.). Artificial intelligence in space. https://www.researchgate.net/publication/342377395_Artificial_intelligence_in_space

Wikipedia contributors. (2025, May 26). RLV Technology Demonstration Programme. Wikipedia. https://en.wikipedia.org/wiki/RLV_Technology_Demonstration_Programme

Pti. (2025, January 4). ‘Life sprouts in space’, says ISRO after cowpea seeds germinate under microgravity conditions. The Tribune. <https://www.tribuneindia.com/news/india/life-sprouts-in-space-says-isro-after-cowpea-seeds-germinate-under-microgravity-conditions/>

Sparkes, M. (2023, August 31). India’s Chandrayaan-3 mission starts exploring the moon’s south pole. New Scientist. <https://www.newscientist.com/article/2389081-indias-chandrayaan-3-mission-starts-exploring-the-moons-south-pole/>

ISRO Official. (2023, September 2). Launch of PSLV-C57/Aditya-L1 Mission from Satish Dhawan Space Centre (SDSC) SHAR, Sriharikota [Video]. YouTube. <https://www.youtube.com/watch?v=IcgGYZTXQw>

Choudhary, G. (2024, December 31). ISRO’s PSLV-C60 Launch: SpaDeX deployed successfully says ISRO Chairman; set to launch NVS-02 satellite in Jan 2025. Mint. <https://www.livemint.com/science/isro-pslv-c60-spadex-space-launch-today->

liveupdates-docking-mission-launch-latest-news-december-302024-11735569048953.html

Sky & Telescope. (n.d.). NASA's first robot astronaut. <https://skyandtelescope.org/astronomy-news/nasas-first-robot-astronaut/>

AI in Cultural Heritage

Ibrahim, M. (n.d.). AI for art & heritage conservation. Ultralytics. <https://www.ultralytics.com/blog/ai-in-art-and-cultural-heritage-conservation>

Chakrabarti, K. (2025, February 18). The role of AI in cultural preservation and heritage. iTMunch. <https://itmunch.com/the-role-of-ai-in-cultural-preservation-and-heritage/#introduction>

Patton, S. (2024, October 21). AI meets archives: The future of machine learning in cultural heritage. CLIR. <https://www.clir.org/2024/10/ai-meets-archives-the-future-of-machine-learning-in-cultural-heritage/>

ResearchGate. (n.d.). AI applications in cultural heritage preservation: Technological advancements for the conservation. https://www.researchgate.net/publication/373336885_AI_APPLICATIONS_IN_CULTURAL_HERITAGE_PRESERVATION_TECHNOLOGICAL_ADVANCEMENTS_FOR_THE_CONSERVATION

ResearchGate. (n.d.). AI and cultural heritage preservation in India. https://www.researchgate.net/publication/378302770_AI_AND_CULTURAL_HERITAGE_PRESERVATION_IN_INDIA

How to Make the World a Better Place for AI

Shaastra. (n.d.). Tech tools to fight deepfakes. Shaastra Magazine, IIT Madras. <https://shaastramag.iitm.ac.in/special-feature/tech-tools-fight-deepfakes>

New Indian Express. (2024, September 12). Why we need to be proactive on AI laws. <https://www.newindianexpress.com/opinions/2024/Sep/12/why-we-need-to-be-proactive-on-ai-laws>

LinkedIn. (n.d.). What are the risks of AI misuse and how can you prevent them? <https://www.linkedin.com/advice/1/what-risks-ai-misuse-how-can-you-prevent-them-ehntc/>

Akash, C. S. (2024, January 2). Meta introduces covert audio watermarking to combat deepfake manipulation. Medium. <https://medium.com/@csakash03/meta-introduces-covert-audio-watermarking-to-combat-deepfake-manipulation-9f7eec6e21f3>

World of Spiritualism. (n.d.). The shocking truth about porn addiction you need to know. <https://worldofspiritualism.com/nofap/porn/the-shocking-truth-about-porn-addiction-you-need-to-know-2/>

Carnegie Endowment for International Peace. (2024, November). India's advance on AI regulation. <https://carnegieendowment.org/research/2024/11/indias-advance-on-ai-regulation>

Singh, A., & Kaur, S. (2015). Artificial intelligence in present scenario. International Journal of Emerging Technology and Advanced Engineering, 3(3), 903–906. <http://pnrsolution.org/Datacenter/Vol3/Issue3/178.pdf>

Learning the Art of Prompting

Prompting Guide. (n.d.). Elements of prompting. <https://www.promptingguide.ai/introduction/elements>

DataCamp. (2023, February 10). What is prompt engineering? The future of AI communication. <https://www.datacamp.com/blog/what-is-prompt-engineering-the-future-of-ai-communication>

DataCamp. (2023, March 3). Prompt optimization techniques. <https://www.datacamp.com/blog/prompt-optimization-techniques>

OpenAI. (n.d.). Prompt engineering. <https://platform.openai.com/docs/guides/prompt-engineering>

Learn Prompting. (n.d.). Chain of thought. https://learnprompting.org/docs/intermediate/chain_of_thought

Sharma, R., & Verma, N. (2024). Prompt engineering and its impact on human-AI collaboration. Humanities & Informatics Journal, 5(3), 85–95. <https://doi.org/10.28991/HIJ-2024-05-03-06>

CollegeLib. (n.d.). How does ChatGPT work? <https://www.collegelib.com/how-does-chatgpt-work/>

SDXL Turbo. (2024, April). 17 ChatGPT-4o tips for

beginners: Easy prompts to use. <https://sdxlturbo.ai/blog-17-chatgpt-4o-tips-for-beginners-2024-easy-prompts-to-use-42937>

Prompt Engineering Institute. (n.d.). Prompt engineering: The art of generating amazing AI text. <https://promptengineeringinstitute.ai/prompt-engineering-the-art-of-generating-amazing-ai-text/>

OpenCampus. (n.d.). Prompt engineering (Week 4). <https://opencampus.gitbook.io/opencampus-machine-learning-program/courses/archive/deep-dive-into-llms/week-4-prompt-engineering>

DeepSeek – A Deep Dive into Its Cutting-Edge Features

TechTarget. (n.d.). DeepSeek explained: Everything you need to know. <https://www.techtarget.com/whatis/feature/DeepSeek-explained-Everything-you-need-to-know>

GeeksforGeeks. (n.d.). DeepSeek R1: Technical overview of its architecture and innovations. <https://www.geeksforgeeks.org/deepseek-r1-technical-overview-of-its-architecture-and-innovations/>

DataCamp. (2024, January 15). DeepSeek vs ChatGPT. <https://www.datacamp.com/blog/deepseek-vs-chatgpt>

DataCamp. (2023, December 20). Mixture of Experts (MoE). <https://www.datacamp.com/blog/mixture-of-experts-moe>

Nguyen, N. (2024, February 3). DeepSeek R1 architecture and training explained. Medium. <https://medium.com/@namnguyenthe/deepseek-r1-architecture-and-training-explain-83319903a684>

DeepSeek Team. (2024). DeepSeek-V2: Scaling language models with mixture of experts. arXiv. <https://arxiv.org/pdf/2501.12948>

Bijit. (2024, March 1). Top NVIDIA GPUs for LLM inference. Medium <https://medium.com/@bijit211987/top-nvidia-gpus-for-llm-inference-8a5316184a10>

Almbok. (n.d.). DeepSeek AI tool. <https://www.almbok.com/ai/tools/deepseek>

IdeaXecution. (2025, March 29). DeepSeek AI: Revolutionizing artificial intelligence research

and applications. <https://www.ideaxecution.com/2025/03/29/deepseek-ai-revolutionizing-artificial-intelligence-research-and-applications/>

Triple Whale. (n.d.). Mastering prompts for actionable insights. <https://kb.triplewhale.com/en/articles/10924859-mastering-prompts-for-actionable-insights>

SELF-SUPERVISED LEARNING

Baevski, A., Hsu, W. N., Xu, Q., Babu, A., Gu, J., & Auli, M. (2022). data2vec: A general framework for self-supervised learning in speech, vision and language (arXiv:2202.03555). arXiv. <https://arxiv.org/abs/2202.03555>

Schwarzer, M., Anand, A., Goel, R., Hjelm, R. D., Courville, A., & Bachman, P. (2020). Data-efficient reinforcement learning with self-predictive representations (arXiv:2007.05929). arXiv. <https://arxiv.org/abs/2007.05929>

Wang, Z., Wu, Z., Agarwal, D., & Sun, J. (2022). MedCLIP: Contrastive learning from unpaired medical images and text (arXiv:2210.10163). arXiv. <https://arxiv.org/abs/2210.10163>

Anonymous. (2024). CLIP the Bias: How useful is balancing data in multimodal learning? (arXiv:2403.04547). arXiv. <https://arxiv.org/abs/2403.04547>

BigScience Workshop. (2022). BigScience: A one-year research workshop for large multilingual language models. <https://bigscience.huggingface.co/>

LeCun, Y. (2020). Self-supervised learning: The dark matter of intelligence. Meta AI Blog. <https://ai.facebook.com/blog/self-supervised-learning-the-dark-matter-of-intelligence/>

Swarm Intelligence in Data Science

Agrawal, S., & Silakari, S. (2015). A review on application of particle swarm optimization in bioinformatics. Current Bioinformatics, 10(1), 78–88. <https://doi.org/10.2174/1574893609666140515003132>

Agrawal, V., & Chandra, S. (2015). Feature selection using artificial bee colony algorithm for medical image classification. In Proceedings of the International

- Conference on Contemporary Computing (IC3) (pp. 171–176). <https://doi.org/10.1109/IC3.2015.7346674>
- Guan, B., Zhang, C., & Zhao, Y. (2018). An efficient ABC_DE_based hybrid algorithm for protein-ligand docking. *International Journal of Molecular Sciences*, 19(4), 1181. <https://doi.org/10.3390/ijms19041181>
- Moosa, J. M., Shakur, R., Kaykobad, M., et al. (2016). Gene selection for cancer classification with the help of bees. *BMC Medical Genomics*, 9(Suppl 2), 47. <https://doi.org/10.1186/s12920-016-0204-7>
- Chakraborty, A., & Kar, A. (2017). Swarm intelligence: A review of algorithms. In *Proceedings of the Advances in Intelligent Systems and Computing* (pp. 215–231). https://doi.org/10.1007/978-3-319-50920-4_19
- Ahmed, H., & Glasgow, J. (2012). Swarm intelligence: Concepts, models and applications. <https://doi.org/10.13140/2.1.1320.2568>
- Saka, M. P., Doğan, E., & Aydogdu, I. (2013). Analysis of swarm intelligence-based algorithms for constrained optimization. In X.-S. Yang, Z. Cui, R. Xiao, A. H. Gandomi, & M. Karamanoglu (Eds.), *Swarm intelligence and bio-inspired computation* (pp. 25–48). Elsevier. <https://doi.org/10.1016/B978-0-12-405163-8.00002-8>
- Fister, I., Yang, X.-S., Brest, J., & Fister, I. Jr. (2013). Memetic self-adaptive firefly algorithm. In X.-S. Yang et al., *Swarm intelligence and bio-inspired computation*. <https://doi.org/10.1016/B978-0-12-405163-8.00004-1>
- Sun, S., & Liu, H. (2013). Particle swarm algorithm: Convergence and applications. In X.-S. Yang, Z. Cui, R. Xiao, A. H. Gandomi, & M. Karamanoglu (Eds.), *Swarm intelligence and bio-inspired computation* (pp. 137–168). Elsevier. <https://doi.org/10.1016/B978-0-12-405163-8.00006-5>
- Xiao, R. (2013). Modeling and simulation of ant colony's labor division: A problem-oriented approach. In X.-S. Yang, Z. Cui, R. Xiao, A. H. Gandomi, & M. Karamanoglu (Eds.), *Swarm intelligence and bio-inspired computation* (pp. 103–135). Elsevier. <https://doi.org/10.1016/B978-0-12-405163-8.00005-3>
- learning for resource allocation in large-scale robotic warehouse sortation centers. <https://www.amazon.science/publications/multi-agent-reinforcement-learning-for-resource-allocation-in-large-scale-robotic-warehouse-sortation-centers>
- Andrew Ng. (n.d.). Quote on AI transforming industries. Quote.org. https://quote.org/author/andrew_ng
- Deming, W. E. (n.d.). W. Edwards Deming quotations. Wikiquote. https://en.wikiquote.org/wiki/W._Edwards_Deming
- Harvard Business Review. (2024, March). How machine learning will transform supply chain management. <https://hbr.org/2024/03/how-machine-learning-will-transform-supply-chain-management>
- KPMG. (2024). Supply chain trends 2024: The digital shake-up. <https://kpmg.com/xx/en/our-insights/ai-and-technology/supply-chain-trends-2024.html>
- Medium. (n.d.). AI's Olympic gold-worthy goal-driven systems. <https://medium.com/product-ai/ais-olympic-gold-worthy-goal-driven-systems-8a563c61e936>
- New Trail. (n.d.). Unexpected insights from an AI rock star. University of Alberta. <https://www.ualberta.ca/en/newtrail/people/unexpected-insights-from-an-ai-rock-star.html>
- The CFO. (2019, June 20). FD Tech Update featuring Facebook, IBM, and Oracle. <https://the-cfo.io/2019/06/20/fd-tech-update-featuring-facebook-ibm-and-oracle/>
- WhatIsSixSigma.net. (n.d.). Utilizing the connected supply chain: Autonomous trucks and AI utilization for smarter operations. <https://www.whatissixsigma.net/utilizing-the-connected-supply-chain-autonomous-trucks-and-ai-utilization-for-smarter-operations/>
- Rolf, P., Singh, R., & Zhao, S. (2022). A review on reinforcement learning algorithms and applications in supply chain management. MIT Center for Transportation and Logistics. <https://ctl.mit.edu/pub/paper/review-reinforcement-learning-algorithms-and-applications-supply-chain-management>

Intelligent Chains

Amazon Science. (2023). Multi-agent reinforcement

Team AINA

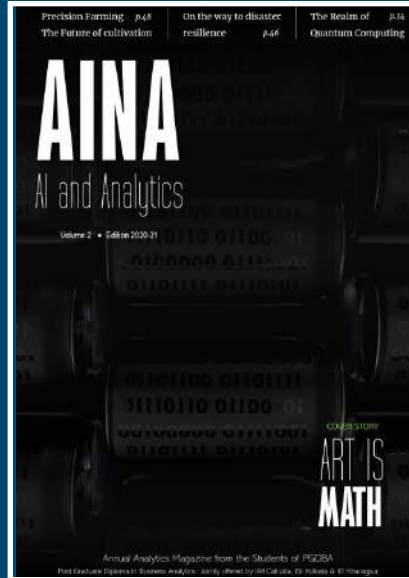


This is AINA 6.0. Built by us, for the curious in all of us.

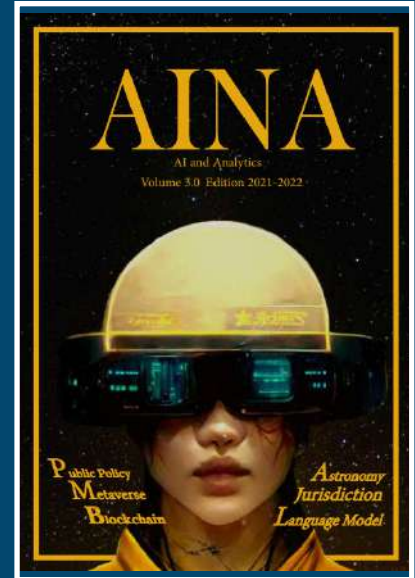
Our Previous Editions



AINA 1



AINA 2



AINA 3



AINA 4



AINA 5



Scan for more details...