

3.1 Demystifying the Language of Artificial Intelligence

Sh. Shashank Shekhar Agarwal

Abstract

Artificial Intelligence (AI) has rapidly entered the vocabulary of governance, bringing with it a plethora of rapidly evolving terms, concepts, and acronyms-often making it difficult for administrators and policymakers to keep up with it. This article is an attempt at making the navigation of the complex world of AI, easier. To be precise, it serves four interrelated purposes: first, it aims to familiarise government officers with AI-related jargons in plain and accessible manner. Too many and rapidly evolving technical jargons can delay decision-making and, at times, make it counterproductive. Second, it seeks to differentiate between commonly heard terms often used interchangeably-sometimes correctly, sometimes incorrectly. Using a dialogical approach, this article poses questions, debunks myths, and draws out nuanced distinctions. Third, the article addresses the practical question of how to use AI tools and which one to use for which purpose, acquainting the reader to prompt engineering, the appropriate use of deep thinking and fast response modes; and the specific capabilities of different tools for tasks such as document analysis, research synthesis, and presentation generation. Lastly, the article engages with the limitations of AI and integral topic of responsible AI, specifically putting forth hallucination, and biases as a result of skewed training data and discussions of AI's impact on environment and employment. In an AI-saturated governance landscape, this paper seeks to equip the reader with a conceptually precise vocabulary that enables informed decision-making.

Keywords

Artificial Intelligence (AI); Machine Learning (ML); Deep Learning (DL); Large Language Models (LLMs); Generative AI; Prompt engineering; Hallucination; Retrieval-Augmented Generation (RAG); Agentic AI; Application Programming Interface (API); AI Governance; Public Administration; AI Literacy

Introduction

The language of Artificial Intelligence (AI) is no longer confined to research laboratories or technology companies; it has become the language of everyday use. Management Information Systems reference *AI-driven dashboards*, policy deliberations mention Machine Learning (ML)-based *citizen services*, vendor meetings are populated with terms such as Large Language Model (*LLM*), *multimodal*, *agentic*, *Retrieval-Augmented Generation (RAG)*, and many more. Yet in practice, these terms are frequently used interchangeably and sometimes imprecisely. The objective of this paper is practical and urgent:

- to familiarise with a conceptually precise AI vocabulary,
- separate concepts that are casually considered synonyms-sometimes rightfully so, sometimes incorrectly, and
- understand how to use AI tools intelligently and cautiously, including how to ask better questions and when not to overuse the fancy features.

While doing so, the article will attempt busting commonly held myths along the way, including the very popular *'training AI by chatting with it'*.

The Jumbled Hierarchy

The first question worth pondering over is: Are AI, ML, and Deep Learning (DL) the same thing? Many people, irrespective of proficiency with AI tools, would say yes, or at least use the three terms interchangeably (IBM, 2026a).

A helpful analogy to understand the distinction would be that of a set of nesting dolls. AI can be thought of as the outermost doll-it represents the classification of machines that simulate human intelligence and cognitive processes like learning, problem-solving, and decision-making (IBM, 2026b). Inside it sits ML, a subset of AI aiding optimization and prediction. ML refers to the technology enabling systems to improve by recognising patterns in data rather than following hardcoded rules. A spam filter that gradually learns which emails get marked as junk? That is ML at work. Amazon recommending products based on what has been looked at and bought? ML at work. Inside ML, sits the innermost doll, that is DL-a technique using multi-layered neural networks to detect extraordinarily complex patterns on the unstructured data. This is what powers voice recognition, real-time translation, and, critically, the large language models that have captured everyone's imagination.

If AI has been around since Turing's (1950) paper 'Computing Machinery and Intelligence', proposing his now-famous imitation game as a test of machine intelligence, why has it suddenly become the buzzword of the last decade?

Two shifts quietly prepared the ground. First, decades of internet growth produced unimaginably large volumes of text data available online and in digital form to train models. Second, the gaming industry's investment in powerful Graphics Processing Units (GPUs) made the heavy-duty calculations required by DL computationally feasible. Together, these enabled DL's extraordinary ability to learn complex patterns, which in turn created the conditions for Generative AI (LeCun et al., 2015; Tiwari et al., 2018). Generative AI is that part of AI which creates long documents, images, audio, video, and code. Once these systems started producing fluent paragraphs and coherent code, the long-discussed dystopian AI future began, for many, to feel less like science fiction and rather more like next Tuesday.

The Advent of Transformers and GPT

In addition to the above two changes that actually made today's LLMs, was there something more? What happened post 2022? The short answer is a combination of architecture and scale.

In 2017, a team of researchers at Google published a paper titled Attention Is All You Need. The architecture they proposed, the Transformer, changed everything. Instead of reading text word by word in sequence, the Transformer could look at an entire passage at once and weigh the relationships between all words simultaneously. For the first time, a model could genuinely track context across long stretches of language (Vaswani et al., 2017a).

On top of the Transformer architecture emerged a new class of models, called GPT-built around a three-step idea that now underpins virtually every major language model in use today:

- **Generative:** The model does not merely classify or retrieve; it produces new text.
- **Pre-trained:** Rather than being built from scratch for each task, the model is first trained on vast corpora of text drawn from books, websites, research papers, and more, absorbing the patterns, grammar, facts, and reasoning styles embedded in human language at scale.

- **Transformer-based:** Using the attention mechanism described (Vaswani et al., 2017b) to process and generate language with genuine contextual understanding.

This recipe was first put into practice by OpenAI. In June 2018, researchers from OpenAI published *"Improving Language Understanding by Generative Pre-Training"*, introducing GPT-1 (Radford et al., 2018). The model was trained in two stages-first, unsupervised pre-training on BooksCorpus, a collection of roughly 7,000 unpublished books, simply learning to predict the next word, and then, supervised fine-tuning on much smaller labelled datasets for specific tasks such as question answering and textual entailment. GPT-2 (2019) scaled up and startled researchers with how fluently it could generate coherent text. GPT-3 (2020), with 175 billion parameters, was the research turning point-it could write, reason, summarise, and translate across tasks it had never been explicitly trained for. However, it largely remained a researcher's tool and unfamiliar to the general public.

That changed in November 2022, when OpenAI released ChatGPT-a conversational interface built on GPT-3.5. The underlying model, GPT-3.5, was not radically different from GPT-3, but what changed was the packaging. A simple chat window was introduced and the model was fine-tuned using Reinforcement Learning from Human Feedback (RLHF) to follow instructions, refuse harmful requests, and hold a conversation rather than merely complete a prompt. The impact was immediate and unprecedented-ChatGPT reached a million users within five days and roughly a hundred million within two months, making it the fastest-adopted consumer application at the time (OpenAI, 2023a).

GPT-4 was introduced in March 2023, adding image input and an enhanced reasoning ability over GPT-3.5, with GPT-4 Turbo, GPT-4o, and a separate "o-series" of step-by-step reasoning models arriving through 2023-24. These models converged into GPT-5 in August 2025-a unified system pairing fast and reasoning models behind an automatic router (OpenAI, 2025a), and was refined through GPT-5.1 to GPT-5.5, before the latest iteration, GPT-5.6, arrived in July 2026 with improvements in coding, biological research, and cybersecurity tasks (OpenAI, 2026a).

In similar lines, Anthropic's Claude is built on the transformer architecture, trained using the alignment method, Constitutional AI, which shapes model behaviour according to a set of explicit written principles rather than relying solely on human

feedback. Unlike OpenAI's GPT branding, however, Anthropic has not disclosed a distinct name for its base architecture. Claude was publicly released in March 2023, followed by Claude 2 later that year. Claude 3 (2024) introduced the tiered Opus, Sonnet, and Haiku system that continues today, refined further through Claude 4 in 2025. The most recent development came in June 2026, with the introduction of Claude Fable 5 and Claude Mythos 5, Anthropic's first models in a new tier positioned above Opus (Anthropic, 2026a). Their access, however, was briefly suspended to comply with United States export controls, before being restored selectively on July 1, 2026 (Anthropic, 2026b).

Other players followed the same basic pattern. For instance, Gemini & Grok are all transformer-based architectures trained at scale, each shaped by its own choices of training data, safety philosophy, and institutional context.

It is interesting to know that the United States Patent and Trademark Office (USPTO) had refused OpenAI's application to trademark "GPT" itself, ruling it a generic term that merely describes a feature of the technology and cannot be owned by any single entity (Vaseghi, 2024). OpenAI does own, and guards closely, the proprietary model weights, training data, and underlying code, leasing access to these via their Application Programming Interface (API) and through their consumer application, ChatGPT. The term is generic; the model behind it is not.

The correct umbrella term for all of them is **LLM**. Knowing this distinction is not simply a semantic quibble. The foundation of an informed decision lies in knowing which LLM a system runs on, who built it, under what data governance terms, and with what safety architecture.

Speaking of LLMs, they are not the whole of AI either.

Generative AI Is the Parent. LLM Is Just One Child

With the proliferation and increasing usage of AI tools like ChatGPT, Gemini, and Claude, it is easy to conflate Generative AI with LLMs. They often appear together as if they are two sides of the same coin.

Generative AI refers to any AI that creates new content: text, images, music, video, code. An LLM is merely the subset that was trained with large volumes of text, and initially dealt with text generation. DALL·E (image generation), Suno (music), and Runway (video) are all Generative AI but none of them are LLMs. A GenAI system extends well beyond what a pure LLM does.

The Engine and the Car Are Not the Same Thing

ChatGPT or Copilot-which is a better LLM? To answer this, it would be crucial to think about another question-is ChatGPT really an LLM? Speaking the layman's language, the answer would be yes without hesitation. The answer is not entirely wrong, but technically, ChatGPT is not the LLM.

The LLM, in ChatGPT's case, is a model from OpenAI's GPT series. GPT-4, released in March 2023, was the model that first brought these capabilities into serious public conversation, demonstrating a significant advancement in reasoning, instruction-following, and contextual understanding (OpenAI, 2023b). GPT-4o launched in May 2024, brought multimodal capabilities, that is, the ability to process text, images, and audio within the same model. GPT-5 arrived in August 2025, and the series has since moved through GPT-5.1, GPT-5.2, and GPT-5.4. The current model, GPT-5.5, is described by OpenAI as their smartest and most intuitive yet, excelling at writing and debugging code, researching online, analysing data, and operating across tools until a task is finished (OpenAI, 2026). ChatGPT is the consumer-facing interface wrapped around whichever is the current model in this series, constituting the conversation layout, the safety filters, the memory features.

Now, returning to the first question-*ChatGPT or Copilot-which is a better LLM?* The question itself, it turns out, is the trap. Neither ChatGPT nor Copilot is an LLM, both are applications. Interestingly, Microsoft Copilot, when first launched, ran on GPT-4 and later GPT-4 Turbo through Microsoft's partnership with OpenAI. However, that is no longer the case. Copilot is now based on a combination of LLMs, including OpenAI's GPT-4 and GPT-5, as well as Anthropic's Claude Opus and Claude Sonnet, matched to the specific task at hand (Microsoft, 2024).

Multimodal Is Not the Same as Multilingual

When you hear that a model is *multimodal*, does it mean it is *multilingual*, that is, supporting multiple languages? This is one of the most common misunderstandings in AI discussions, particularly in a country as linguistically diverse as India.

The two terms describe entirely different capabilities. *Multimodal* means the model can process and generate multiple types of data-text, images, audio, video, often within one integrated system. A model can be multimodal but English-only, or multilingual but text-only. These are independent axes. Google DeepMind's

technical documentation describes *Gemini* as designed from the ground up to be natively multimodal, trained jointly across modalities rather than having capabilities bolted on as separate modules (Google DeepMind, 2023b). ChatGPT, by contrast, achieves multilingual capability in a manner closer to an extended feature rather than a foundational design choice, a distinction that matters when selecting tools for diverse citizen-facing applications.

Early Generative AI systems reinforced the specialist model-one system for text, another for images, another for video or music. Later, all-rounder models emerged, that is, tools that handle text, images, data analysis, and code within a single interface. India's BharatGen, the government-funded multimodal LLM, is being designed to be precisely this-handling 22 Indian languages across text, audio, and visual data types simultaneously, though it remains under development (PIB, 2024).

Within this landscape, two distinct types of tools have emerged that are worth telling apart. As discussed above, *Multimodal models*, such as Gemini or advanced GPT versions, handle multiple input and output types directly within the model itself. The other type of tool-*Aggregators*, such as Perplexity in its early form and ChatLLM, presently, do not do all the reasoning themselves. Instead, they search, collate, and present answers by calling upon multiple underlying sources and models, functioning more as intelligent curators than independent reasoners. A useful parallel is Google's own evolution: it began as a search engine pointing users to other websites, and progressively built its own browser, maps, mail, and operating system. Tools that started as aggregators have since developed proprietary models of their own, and the lines between search engine, answer engine, LLM, and aggregator is now meaningfully blurred.

For someone evaluating any AI tool, three questions cut through the noise: *Is this product doing its own reasoning? Is it primarily aggregating others' outputs? Or is it a hybrid of both?* The answer shapes everything from data security assessment to accountability in the event of an error, and it begins with not confusing what a tool processes with what languages it speaks.

The Model Does Not Learn When You Use It-Teaching and Training Are Not the Same

A belief, even among technically curious people, is persistent: *By using ChatGPT or similar tools, I am training the models immediately.* It is, however, a technically inaccurate belief. When one types a prompt and receives a response, the model is performing inference, that is, running a fixed set of mathematical weights forward

to generate a response. Those weights do not change mid-session. Corrections within a session shape that particular exchange by providing context, but they do not permanently alter the model.

Training, in contrast, the process through which the model actually learns, is a separate, computationally expensive procedure that happens offline, requiring enormous GPU resources, and is carried out periodically by the company, not triggered by individual query. OpenAI's technical documentation draws a clear line between the training phase and the inference phase (OpenAI, 2023c). Repeated corrections in a conversation do not teach the AI anything permanently, they only provide context for that session's responses. This is precisely why companies periodically release new versions: they gather curated data, run large offline training jobs, adjust the architecture and safety layers, and release the result as a new version which is a snapshot of the improved model. GPT-3.5, GPT-4, GPT-4o, GPT-5 are the output of deliberate, resource-intensive training runs conducted by OpenAI's engineering teams over years.

Prompt Engineering Is a Planning Skill, not a Coding Skill

If a model is fixed and a user cannot change it, what can they actually do to get better results? The answer is prompt engineering, and the second word is where misconceptions multiply fast. Being well versed with prompt engineering does not necessarily mean expertise in Python, JavaScript, or any programming knowledge. It is simply the ability to craft clear, structured, context-rich instructions that consistently extract accurate and useful responses from an AI chatbot.

Think of it as the difference between asking a very smart but literal-minded colleague "write me something about the budget" versus "draft a 300-word executive summary of the Q2 defence budget for a non-technical audience, using formal Hindi with an English glossary." The model is the same. The outputs are very different.

One important technique worth naming here is Meta Prompting-the practice of asking the AI itself to generate the best prompt for a given task. Rather than crafting the prompt entirely by oneself, the user describes the goal to the AI and asks it to design the most effective instruction set. This is increasingly common in complex government use cases and is good for standardised outputs.

A question that naturally arises in practice is *why does the same prompt give different outputs across different tools-sometimes even between ChatGPT and Copilot, which broadly draw on OpenAI models?* The answer lies not only in the underlying LLM

but in also everything layered on top of it. Each application adds its own hidden system instructions about tone, safety, and priorities. Some tools retrieve live web data or internal documents before generating a response while others rely solely on the base model. The way previous messages are packaged and sent back to the model varies by product, and what one interface is permitted to say, another may suppress or rephrase entirely.

When two users compare outputs from different tools, they are not comparing models, they are comparing product decisions, safety philosophies, and retrieval configurations built around those models.

Which Tool for Which Task: A Comparative Snapshot

With multiple AI tools now available, each optimised for different strengths, a natural question follows: *which one should be used, and for what?* Rather than treating any single platform as a universal solution, the table below (Table 1) offers a task-based comparison, matching common professional requirements with the tool best suited to each, along with a viable alternative and the underlying reasoning.

Table 1

AI Tool Mapping with Professional Requirements

Action Required	**Best AI Tool	**Second Best	Why
Email/Letter Writing	ChatGPT (GPT-4oor later)	Microsoft Copilot	Versatility, nuanced tone-matching and drafting context-aware replies directly in email clients.
Deep Research	ChatGPT (OpenAI o3)	PerplexityPro (quick)	Acts as an autonomous agent to browse, cross-reference, and structure comprehensive reports.
Large Document Summarization	Claude (Fable5/ Sonnet 4.6)	Gemini2.5 Flash Lite	Processes massive 1M-token PDFs with high accuracy for extracting key points faithfully.
Real-Time Trends Tracking	Grok	Perplexity Pro	Has real-time pipeline access to X (Twitter) for immediate social and breaking news.

Indian Govt Laws/Circulars	Perplexity Pro	Kanoon AI	Uses real-time legal RAG to cite the exact.gov.in URLs instead of hallucinating legal clauses.
Brainstorming Ideas	GPT-5.2	Claude Opus 4.8	Generates a wide diversity of divergent, creative ideas from a single prompt.
Creating Podcasts/ Audio summaries	Google NotebookLM	ElevenLabs	Instantly converts uploaded text/PDFs into hyper-realistic, two-host audio podcast discussions.
Image Generation	Midjourney v7	-	Known for artistic, conceptual quality and aesthetic image creation.
Video Generation	Google Veo 3.1	Sora 2	Delivers high cinematic fidelity for AI-generated video outputs.
Code Generation	Claude Code (Anthropic)	Cursor IDE / GitHub Copilot	Operates autonomously using loop engineering to write, debug, and commit entire workflows.
PPT/Presentations	Gamma.ai	PlusAI / Claude	Generates complete, beautifully designed web-native or PowerPoint decks from simple text prompts instantly.

***This comparison does not constitute an official policy or endorsement of any AI tool by the author; it reflects the author's subjective experience at the time of writing, in a landscape that is evolving rapidly. Readers should verify current tool-specific information independently and exercise their own judgement in selecting and using AI tools, in accordance with their department's requirements and applicable data security guidelines.*

Agentic AI: The Shift from Answering to Doing

Commonly used chatbots like ChatGPT, Gemini, Claude describes AI that waits for the user to give a prompt and then responds. That is the current dominant paradigm. But *Agentic AI*, a term flooding technology headlines, describes something fundamentally different.

An agentic AI system is given a *goal*, not a question, and then it autonomously plans, executes, and iterates across multiple steps and tools to achieve it. It could search the web, draft an email, book a calendar slot, and file a form, all without being asked step by step. The analogy is the difference between a very smart intern who needs detailed instructions for each task, and a capable deputy who takes a brief and delivers results.

Agentic AI is still nascent but rapidly maturing. Its implications for government workflows—from grievance redressal to end-to-end public service delivery, are significant. While agentic AI creates numerous opportunities, it also poses risks that seemed hypothetical at a point of time.

The Invisible Connector: API

What makes integration across models, applications, and interfaces possible? The answer lies in a technology called API. The most common example is the click "Login with Google" on a website. By doing so, the site does not build its own authentication system, instead, it uses Google's API to verify user identity. When Zomato shows a map of nearby restaurants, it is not building mapping software, it is calling Google Maps through an API.

In the AI context, when Microsoft Copilot sometimes produces results similar to ChatGPT, it is because Copilot is calling OpenAI's GPT API. The product is Microsoft's, the engine, accessed through a standardised connector, is OpenAI's. A parallel from Indian governance would be the Aadhaar authentication in the *Jeevan Pramaan app*, or in AEBAS, is another result of API operation.

Within the AI ecosystem, knowledge of APIs allows one to critically examine key questions. When two portals, two AI agents, or two different software systems need to interact, even if coded in entirely different programming languages, the API is the standardised connector that makes this possible. But that connector carries risk: *Is the vendor building its own model, or simply calling someone else's via API, and if so, where does the data go when that call is made? What happens if the upstream provider changes its pricing, revises its terms, or limits access?*

Thinking Mode Is Not Always Better

A frequent temptation, while using tools such as Gemini or ChatGPT, is to always switch on *Deep Thinking* or *Research Mode*. The reasoning: *if the AI thinks harder, the answer will be better?* Not necessarily.

Extended reasoning modes are designed for genuinely complex tasks and makes sense for multi-step policy analysis, resolving contradictory legal clauses, structuring a budget with intricate dependencies.

For simpler tasks such as translating a sentence, summarising bullet points, or drafting a routine email, deep thinking mode adds no value, over-complicates the output, and for users on limited plans, needlessly consumes paid tokens or free query limits.

There is also a resource argument that extends well beyond individual usage. Deep reasoning models consume significantly more computational energy than standard responses, and AI data centres already account for a growing share of global electricity consumption and greenhouse effect. Using maximum-intensity modes for trivial queries is both wasteful and environmentally costly.

This anecdote points to a tension that the world has not yet resolved. While Generative AI holds genuine promise for pro-active and better governance, public service delivery, and human productivity, its energy appetite is neither small nor static. The GenAI-energy conundrum-the growing collision between AI's computational demands and global climate commitments is not a problem for the distant future. It is here now, and it deserves a place in every serious conversation about AI adoption at scale.

When AI Confidently Lies

Even the best prompt cannot prevent what is perhaps the most consequential characteristic of LLMs: *hallucination*. These models do not retrieve facts from a database, instead, they predict the most statistically likely sequence of words given a context. When they lack relevant training data, they generate plausible-sounding text, and plausible is not the same as true. The danger is that the output reads with total confidence.

The practical solution developed for this problem is RAG, where the AI is anchored to a specific, verified document vault rather than its general training data. Instead of guessing, it is constrained to answer only from uploaded policy documents, circulars, or regulations, dramatically reducing hallucination for domain-specific deployments. *NotebookLM*, Google's research assistant tool, operationalises this principle as a closed RAG system. Within this tool, responses are generated solely from documents the user uploads, which significantly reduces the risk of hallucination.

For someone working with specific Acts, circulars, or departmental guidelines, this approach might considerably be more reliable than an open-ended query to a general-purpose LLM.

AI Bias

The bulk of data on which major LLMs are trained originates from English-language,

Western-centric sources, thereby skewing the model's understanding of what is normal, neutral, or desirable toward a particular cultural frame. The consequences for India are specific and documented. Khandelwal et al. (2024) introduced the Indian-BhED dataset to evaluate caste and religion-based stereotypes, dimensions largely absent from standard fairness benchmarks, and found that popular LLMs frequently reproduced harmful stereotypes related to caste and religion, even more so than for attributes like gender or race that are commonly studied in the West. Gender bias is equally evident, with studies showing that models are more likely to generate male-associated names and roles, reinforcing occupational stereotypes (Gallegos et al., 2023). For a country as socially stratified and linguistically diverse as India, this is a governance risk. An AI tool used for citizen-facing services or grievance categorisation that reproduces caste, gender, or religious stereotypes is not merely inaccurate, it is institutionally harmful.

Will AI Replace Humans?

The future cannot be predicted, but important lessons can be drawn from the past. Since the Industrial Revolution, there have been massive protests and claims that machines will replace humans, so much so that the *Luddites* went ahead and damaged the mechanical equipment, the power looms, of their era.

Recent public statements by AI industry leaders, including the CEO of Anthropic, have highlighted concerns regarding the potential risks associated with increasingly capable AI systems. These discussions have renewed debates on AI governance, safety, and the need for appropriate regulatory safeguards.

The question of mass-scale unemployment driven by AI deserves honest engagement, not dismissal. What history suggests is that technology displaces some jobs and creates others. The question is whether the transition is managed equitably and whether reskilling systems keep pace with displacement. For most, the implication is immediate: *AI literacy is not optional*. It is the new baseline competency for almost all walks of life.

Conclusion

The beliefs, myths, questions, and terms introduced through this article carries significant implications for individuals navigating AI tools, departments involved in public service delivery, and policy-level discussions on AI governance architecture. Understanding the nuances that ML is not the same as AI, that ChatGPT is not an

LLM but uses one, that Generative AI includes far more than language models, that automation is not automatically AI, and that hallucination is a structural feature rather than a correctable bug, directly shape how decisions are made about data security, public service, financial management and institutional accountability.

Perhaps an effective way to understand AI, in the end, is not as a destination but as a vehicle-one with a very powerful accelerator and a brake that must be applied with balance. Press only the gas pedal, and AI outruns judgement, mistaking velocity for direction, fluency for truth, and confidence for correctness. Press only the brake, and the vehicle never leaves the driveway, its potential left idling. Neither extreme serves the traveler. The real skill, for the individual as much as for the institution, lies in learning where the pedal ends and the brake begin, that is, accelerating into the tasks AI genuinely lightens, while braking firmly wherever judgement, context, and human accountability cannot be outsourced.

References

- Anthropic. (2023). Claude model card and safety documentation. <https://www.anthropic.com/research>
- Gallegos, I. O., Rossi, R. A., Barrow, J., Tanjim, M. M., Kim, S., Dernoncourt, F., Yu, T., Zhang, R., & Ahmed, N. K. (2023). Bias and fairness in large language models: A survey. *arXiv preprint arXiv:2309.00770*. <https://arxiv.org/abs/2309.00770>
- Google DeepMind. (2023). Gemini: A family of highly capable multimodal models. <https://deepmind.google/research/publications>
- Team, I. D. a. A. (2026, May 15). AI vs. Machine Learning vs. Deep Learning vs. Neural Networks. IBM. <https://www.ibm.com/think/topics/ai-vs-machine-learning-vs-deep-learning-vs-neural-networks>
- Launch of BharatGen: The first Government-funded Multimodal Large Language Model Initiative.* (2024). <https://www.pib.gov.in/PressReleasePage.aspx?PRID=2060437&=48&lang=2>
- IRPAAI. (2023). RPA vs. AI: Understanding the distinction. <https://www.irpaai.com>
- Khandelwal, K., Tonneau, M., Bean, A. M., Kirk, H. R., & Hale, S. A. (2024). Indian-BhED: A dataset for measuring India-centric biases in large language

models. In *Proceedings of the 2024 International Conference on Information Technology for Social Good* (pp. 231–239). ACM.
<https://doi.org/10.1145/3677525.3678666>

LeCun, Y., Bengio, Y., & Hinton, G. (2015). Deep learning. *Nature*, 521(7553), 436–444. <https://doi.org/10.1038/nature14539>

Microsoft. (2024). Microsoft 365 Copilot application card. Microsoft Learn.
<https://learn.microsoft.com/en-us/microsoft-365/copilot/microsoft-365-copilot-application-card>

OpenAI. (2023a). *GPT-4 technical report*. <https://platform.openai.com/docs>
OpenAI. (2023b). *Retrieval-augmented generation: Techniques and applications*. <https://platform.openai.com/docs>

OpenAI. (2026, April 23). *Introducing GPT-5.5*.
<https://openai.com/index/introducing-gpt-5-5/>

Tiwari, T., Tiwari, T., & Tiwari, S. (2018). How artificial intelligence, machine learning and deep learning are radically different. *International Journal of Advanced Research in Computer Science and Software Engineering*, 8(2), 1.

Turing, A. M. (1950). Computing machinery and intelligence. *Mind*, 59(236), 433–460. <https://doi.org/10.1093/mind/LIX.236.433>

Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, Ł., & Polosukhin, I. (2017). Attention is all you need. *Advances in Neural Information Processing Systems*, 30, 5998–6008.
<https://arxiv.org/abs/1706.03762>

Author's Profile

Shri Shashank Shekhar Agarwal is Director at the National Communications Academy-Finance, Department of Telecommunications, Government of India. A distinguished faculty member since 2024, he conducts trainings and hands-on sessions on emerging areas such as AI and data analysis, and has also contributed to implementing Mission Karmayogi principles and iGOT content creation. He mentors future leaders in data-driven policy making and digital transformation, drawing on his experience across the public and private sectors. Since 2018, he has been contributing to transformative policy and regulatory initiatives in the Department of Telecommunications. An alumnus of IIT Roorkee, he combines strong technical expertise with strategic policy insight to support telecommunications policy, regulatory compliance, technology-enabled governance, and digital infrastructure development.