

# Artificial Intelligence: Past, Present, and Future(?)

M. Vidyasagar FRS

Distinguished Professor, IIT Hyderabad

[m.vidyasagar@iith.ac.in](mailto:m.vidyasagar@iith.ac.in)

[https://people.iith.ac.in/m\\_vidyasagar/index.html](https://people.iith.ac.in/m_vidyasagar/index.html)

AI symposium for Science and Engineering  
IIT Hyderabad, 21 July 2026

# Outline

- ① Terminology
- ② Brief Historical Overview: 1949–2017
- ③ Recent Developments: 2017–Date
- ④ Current Situation
- ⑤ The Future
  - Achieving AI Sovereignty
  - Impact on Education
  - Impact on Research
  - Impact on Employment

# Outline

- 1 Terminology
- 2 Brief Historical Overview: 1949–2017
- 3 Recent Developments: 2017–Date
- 4 Current Situation
- 5 The Future
  - Achieving AI Sovereignty
  - Impact on Education
  - Impact on Research
  - Impact on Employment

# Some Terminology

- What do the terms AI (Artificial Intelligence), AGI (Artificial General Intelligence), ML (Machine Learning) mean, and how are they related?
- *My Definitions:*
  - The ability carry out large scale *numerical* computations does not, by itself, count as AI!
  - AI refers to enabling computers to carry out tasks that humans find easy, but machines (previously) found difficult. Examples include speech recognition, pattern recognition.
  - AGI refers to developing *one utility* that can perform *multiple tasks* after proper training.
  - Examples include DeepMind's "AlphaZero," which can play Chess, Go, and Shugi (Japanese Chess).

# Some Terminology: Others

- Machine Learning (ML) refers to a set of algorithms (mathematical procedures, and/or their implementations into computer programs) that *permit the realization* of AI or AGI.
- Thus AI and AGI are the *destinations*, while ML is the *path towards* the destination(s).

*Caution:* These are *my* definitions, not universal.

## Some More Terminology: Gen AI & LLMs

- Further terminology: Generative AI tools, Sovereign AI models, Large Language Models (LLMs), open-source models, foundational models, fine-tuned models etc.
- What do these terms mean, and how are they related to each other?
- A **Generative AI** tool is one that can accept previously unseen inputs (prompts) and produces accurate (or at least, plausible) output.
- A **Large Language Model (LLM)** is a utility that accepts a “context” (a sequence of “tokens”) to predict the next “token.” It is probabilistic in nature, and aims to produce high-quality output that sounds “human.”

## Some More Terminology: Open Source Software

- **Open source software** is any software that is freely downloadable.
- Once you have downloaded the software, it is yours forever. There cannot be any “kill switches” in the software.
- The user is allowed to make modification in the software for *internal* use. But if the modified software is shared with *some* users outside the organization, then it must be put back into the public domain.
- It is strictly forbidden to include open source software inside a commercial product; but this is an “honor system.”

## Some More Terminology (Cont'd)

- A **foundational model** is an LLM that is trained “from scratch” on a corpus of tokens, which have been created and curated by the persons building the LLM.
- We can think of a token as a word, though often a token is actually smaller.
- A **fine-tuned model** is one that starts with an open source model, and then adjusted to suit local requirements.
- A **sovereign** AI tool is one that is entirely owned by India.
- In this talk, I focus only on the AI *software*, not the hardware on which it is executed.



# Outline

- 1 Terminology
- 2 Brief Historical Overview: 1949–2017**
- 3 Recent Developments: 2017–Date
- 4 Current Situation
- 5 The Future
  - Achieving AI Sovereignty
  - Impact on Education
  - Impact on Research
  - Impact on Employment

# Trivia Quiz

Who coined the phrase “Artificial Intelligence”, and when?

# Trivia Quiz

Who coined the phrase “Artificial Intelligence”, and when?

The phrase was coined by Stanford Professor John McCarthy in

# Trivia Quiz

Who coined the phrase “Artificial Intelligence”, and when?

The phrase was coined by Stanford Professor John McCarthy in 1955!

# Trivia Quiz

Who coined the phrase “Artificial Intelligence”, and when?

The phrase was coined by Stanford Professor John McCarthy in 1955!

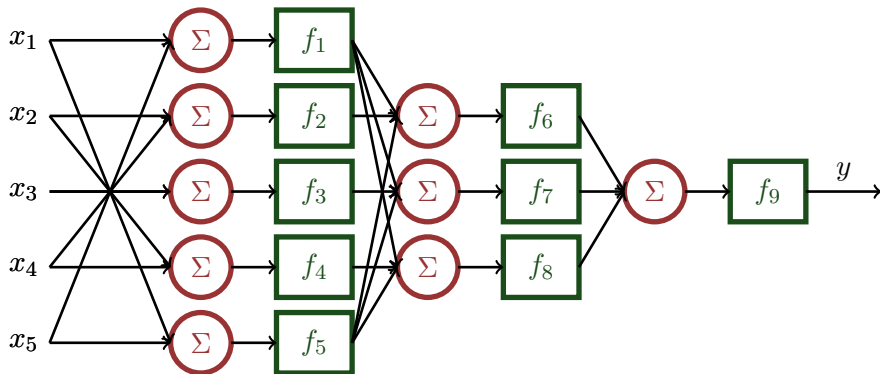
AI is an *old* subject!

The idea of using “computers” to process non-numerical data started along with the advent of the digital computer ENIAC in 1949.

# Neural Networks: 1986–2012

- The current “workhorse” of AI utilities are *neural networks*.
- Introduced in 1962(!) as “perceptrons”
- Reinvented as “MLPNs (Multi-Layer Perceptron Networks)” in 1986 (depiction on next slide)
- “Deep” neural networks invented in 2012
- Today’s networks are still larger

# Neural Networks: Depiction



Each  $f_i$  is a “neuron,” i.e., a map from  $\mathbb{R}$  to  $\mathbb{R}$ . Not all neurons need to be the same.

# Advent of Deep Neural Networks: 2012

- Until the early 2010s, neural networks were “shallow” – a few hidden layers, and a few hundred neurons in all.
- That changed dramatically in 2012.

---

## ImageNet Classification with Deep Convolutional Neural Networks

---

**Alex Krizhevsky**  
University of Toronto  
kriz@cs.utoronto.ca

**Ilya Sutskever**  
University of Toronto  
ilya@cs.utoronto.ca

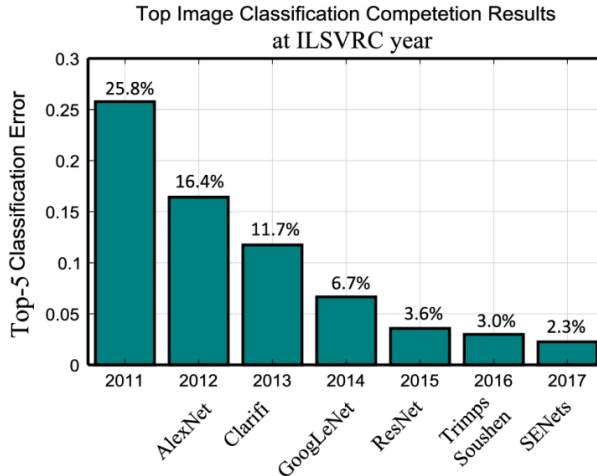
**Geoffrey E. Hinton**  
University of Toronto  
hinton@cs.utoronto.ca



# Large-Scale Visual Recognition Challenge (LSVRC)

- A “hand-curated” collection of 1.3 million images, grouped into 1,000 classes.
- Provided a benchmark for comparing different image processing algorithms.
- Prior to 2012, the typically best error rates were  $\sim 25\%$ .
- The network in 2012 had 650,000 neurons and 60 million parameters.
- Deep Neural Networks (DNNs) have today thousands of hidden layers and hundreds of millions of neurons.
- The impact of DNNs on the LSVRC are shown on the next slide.

# Impact of Deep Neural Networks: LSVRC



## A Bit of Trivia

The NN in the 2012 paper used an NVIDIA GPU.

*Trivia Question:* What was the original target market of NVIDIA?

## A Bit of Trivia

The NN in the 2012 paper used an NVIDIA GPU.

*Trivia Question:* What was the original target market of NVIDIA?

*Gaming!*

The credit for turning NVIDIA from a gaming company to an AI company can perhaps be given to this paper.

# Outline

- 1 Terminology
- 2 Brief Historical Overview: 1949–2017
- 3 Recent Developments: 2017–Date
- 4 Current Situation
- 5 The Future
  - Achieving AI Sovereignty
  - Impact on Education
  - Impact on Research
  - Impact on Employment

# Shortcomings of Existing NN Architectures

- Due to the temporal nature, they process “one word at a time.”
  - No scope for parallelization.
  - Will fail on highly inflected languages (such as Indic), where word order is not crucial.

# Transformers: 2017

## Attention Is All You Need

**Ashish Vaswani\***  
Google Brain  
avaswani@google.com

**Noam Shazeer\***  
Google Brain  
noam@google.com

**Niki Parmar\***  
Google Research  
nikip@google.com

**Jakob Uszkoreit\***  
Google Research  
usz@google.com

**Llion Jones\***  
Google Research  
llion@google.com

**Aidan N. Gomez\*<sup>†</sup>**  
University of Toronto  
aidan@cs.toronto.edu

**Łukasz Kaiser\***  
Google Brain  
lukaszkaizer@google.com

**Illia Polosukhin\*<sup>‡</sup>**  
illia.polosukhin@gmail.com

# Transformer Architecture

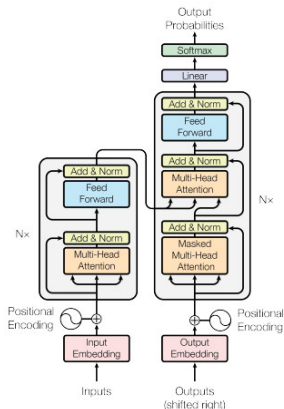


Figure 1: The Transformer - model architecture.

*Source:* Original transformer paper.



# Advantages of the Transformer Approach

- Processes *all* words at once, without regard for position; this permits massive parallelization.
- Due to this, it can cope with much longer-range contexts.
- Independence of word order makes it well-suited for Indic languages
- Positional encoding retains *some* information about word order
- Versatile architecture, applicable accross multiple domains (text, voice, images, etc.)
- Now the *de facto* architecture for implementing AI

# Statistical Language Models

- The phrase currently in vogue is “Large Language Models (LLMs).”
- Now people talk about “Small Language Models.”
- (Fortunately, thus far no one has said “Small LLMs”!)
- The issue is not the *size* of the model, but its *nature*.
- “Statistical” as opposed to “syntactic.”
- Statistical LMs are based on the *frequency of occurrence* of pairs of words (or triplets) etc.

# Methodology of Statistical LMs

- The language corpus is broken up into “tokens” (which I will treat as just words).
- The text is analyzed for the frequency of occurrence of doublets, triplets, quadruplets etc. of tokens.
- Given the preceding  $N$  tokens, we make a prediction of the *next* token using the observed frequencies.
- Thus a Statistical LM is just a *Markov process with memory*: The probability distribution of the next token is determined solely by the preceding  $N$  tokens.
- The larger we take  $N$  to be, the more accurate is our prediction.

# Operation of Statistical LMs

- Suppose our language has  $10^4$  tokens. Then there are  $10^8$  pairs of tokens,  $10^{12}$  triplets, etc. This number blows up very quickly.
- *Fortunately* the overwhelming majority of these combinations do not occur. So by using clever indexing, the memory requirements can be reduced.
- To produce more natural-sounding text output, the LM *does not* merely produce the most likely next token.
- Rather, the next token is *random*, with the same probability distribution as the observed frequency.

## Example: Syntactic vs. Statistical

*Question:* Is it correct to say “They was coming”?

- *Syntactic:* No. “They” is plural and “was” is singular. So it has to be either “They were coming” or “He was coming.”
- *Statistical:* I have parsed billions of sentences, and I rarely saw “They was”. So it is *probably* incorrect.

Statistical models have no sense of *why* something is incorrect.

But statistical models are cost- and time-efficient for building cross-lingual models where the syntaxes may vary widely (e.g., English to Indian and vice versa).

# The Importance of Context

Complete this sentence:

*The capital of India is . . . .*

*Easy!* No ML is needed, just a searchable database.

Now complete this sentence:

*Delhi is the . . . of India.*

There are many possible *valid* completions.

Knowing the preceding text gives the *context* and allows us to give an accurate answer.

## The Importance of Context (Cont'd)

Consider the possible preceding sentences:

*Mumbai is the commercial capital of India.*

This suggests

*Delhi is the **political capital** of India.*

How about

*Hyderabad is the second largest city in India by area.*

This suggests

*Delhi is the **largest city by area** in India.*

And so on.

The preceding  $N$  tokens are the “context.”

# Outline

- 1 Terminology
- 2 Brief Historical Overview: 1949–2017
- 3 Recent Developments: 2017–Date
- 4 Current Situation**
- 5 The Future
  - Achieving AI Sovereignty
  - Impact on Education
  - Impact on Research
  - Impact on Employment



# Foundational Models vs. Fine-Tuning Existing Models

- A **foundational model** is one trained *ab initio* (from scratch), usually on one's own proprietary corpus.
- Fine-tuning an existing model consists of
  - Starting with an existing (invariably open source) model, and
  - Tuning the weights of the LLM on one's proprietary corpus.

Which is better for India?

## Current Large Language Models

- OpenAI's products (ChatGPT family), and Anthropic's products, are the best-known “closed” LLMs that are available only via subscription.
- There are many *notionally open source* LLMs, e.g., Google Gemma, FAIR (Facebook) LLaMA, etc.
- In recent times, a whole host of open source LLMs have come out of China, e.g., Deep Seek, Qwen, Kimi-Moonshot, etc.
- Most recently, US-based Thinking Machines has come out with its own open source LLM
- These open source utilities:
  - Nearly match the performance of “closed” LLMs
  - Are an order of magnitude cheaper to execute, mainly due to the absence of “fees for token usage”

## Foundational Models for India?

- Well-known closed models such as ChatGPT, Claude etc. are trained on  $\sim 10^{14}$  tokens.
- Creating such an enormous corpus, and then training an LLM on it, involves *huge* expenditure.
- Anthropic just *raised* \$ 85 Bn, and Google raised \$ 80 Bn, on top of what they have already spent.
- In contrast, India's AI Mission is INR  $10^{11}$ , or about USD 1 billion.
- Thus any foundational model in India will necessarily be trained on a much smaller corpus, and be focused on a *niche* application.

# Fine-Tuning Open Source LLMs

- With open source models, *only the trained model* is open, and not the *corpus* on which it was trained.
- However, the Open Source models can be “tweaked” to accommodate the additional corpus, using “low rank adaptations.”
- *Worry*: During fine-tuning, has the performance on the (unseen) original training data become degraded?
- Two remedies:
  - If the additional corpus is much smaller than the original (unseen) corpus, the degradation if any would be negligible.
  - Additionally, the tweaked models can be benchmarked against publicly available libraries of token corporuses.

# Outline

- 1 Terminology
- 2 Brief Historical Overview: 1949–2017
- 3 Recent Developments: 2017–Date
- 4 Current Situation
- 5 The Future**
  - Achieving AI Sovereignty
  - Impact on Education
  - Impact on Research
  - Impact on Employment

# Outline

- 1 Terminology
- 2 Brief Historical Overview: 1949–2017
- 3 Recent Developments: 2017–Date
- 4 Current Situation
- 5 The Future
  - Achieving AI Sovereignty
  - Impact on Education
  - Impact on Research
  - Impact on Employment

# Components of AI Sovereignty

- Ownership of the AI *utility*
- Ownership of the AI *data*
- Ownership of the AI *platform*.
- The third item is not feasible, at any time in the near future.
- The second item is achievable, thanks to India-based cloud servers (e.g., Yotta)
- The first item is also achievable, and this will be the subject of the remainder of the talk.

# Objectives of a Sovereign AI Utility for India

- It must be *sanctions-proof*:
  - No part of it should be vulnerable to sudden and arbitrary access denial, as happened with Anthropic recently.
  - It should not be vulnerable to “snooping” by an external agency.
- There are two other desirable (almost mandatory) features:
  - Debiasing and
  - Localization.



# Debiasing and Localization

## Debiasing:

- Western sources (e.g., Wikipedia) are incorrigibly biased against India.
- Our only alternative is to create our own alternative.
- *Only the biased entries* need to be purged and/or updated. It is neither desirable nor practicable to recreate it in its entirety.

## Localization:

- Enabling Indian citizens to access GoI services in Indic languages is a good first step; but it is not sufficient.
- Ideally, *the whole world's knowledge corpus* should be accessible to Indians in their own mother tongues.

## Fine-Tuning an Existing Model for India?

- Even in an “open source” model, only the *parameters* of the model are open, and not the *corpus* used in training.
- The open source model can be “fine-tuned” using *low-rank adaptations*
- The performance on the original (unknown) training corpus will not be degraded, provided the new corpus is *small*.<sup>1</sup>
- The performance of the fine-tuned model can be benchmarked against public-domain databases of tokens.

In this way, the fine-tuned model can capture *India-specific* requirements, while retaining performance on the original corpus.

<sup>1</sup>As a percentage of the size of the original. Typical LLMs have  $\sim 10^{14}$  tokens. So fine-tuning on  $\sim 10^{11}$  tokens would be fine. 



# Should India Stop Developing Foundational Models?

*Definitely not!*

There are two aspects to LLM Development: Algorithmic, and software engineering

- For LLMs with  $< 10^{10}$  parameters, python code implementing standard algorithms will run “out of the box” in e.g., PyTorch.
- Therefore algorithmic understanding is sufficient.
- However, the *impementation* of even moderate-sized LLMs also requires *software engineering* for managing memory calls, data representation (fixed-point vs. floating point) etc.
- India needs people skilled in *both* aspects.

# Advantages of Fine-Tuned LLMs for India

- Existing open source models have been developed using *huge* resources. It would be a mistake to not build on them.
- The twin objectives—debiasing, and localization—can be achieved at an affordable cost.
  - Debiasing*: Since foreign LLMs are trained on local sources, they have a lot of bias, esp. on our history, culture, religion etc.
  - Localization*: Information about many items is *simply absent*, or misleading due to a small amount of inputs.
- The code as well as data *resides on our own computers*.
- By fine-tuning open source LLMs, we can *preserve and protect* our traditional knowledge.
- Starting with an open source model and fine-tuning to achieve these objectives is, in my view, the way to go.*

# Outline

- 1 Terminology
- 2 Brief Historical Overview: 1949–2017
- 3 Recent Developments: 2017–Date
- 4 Current Situation
- 5 The Future
  - Achieving AI Sovereignty
  - **Impact on Education**
  - Impact on Research
  - Impact on Employment

## A Potentially Impactful Application

*My hope:* AI-enabled education in local languages, at the primary and secondary levels.

- Start with “hard sciences” such as maths, physics, and “factual” topics like Civics; then expand scope.
- To my mind, fine-tuning an existing LLM by debiasing, and then localizing it, is the best option.
- A one-year project of INR 60 to 100 crores would kick-start this effort, with follow-on funding.
- At the moment, there is *no GOI funding* for this.
- Private entities will have to step up (through CSR?)
- Several NGOs are doing laudable work; but they need to pool efforts.

# Ethics in Education

What should “ethics” mean in the AI era?

- ① Student A emails the exam via smart phone to a older Student B who has already taken the course. In turn B mails the solution to A who submits it as original work.
- ② Student A uses AI tools to solve homework problems

*My view (FWIW):*

- Item 1 is unethical, because if A were to join a company, B would not be there to assist.
- Item 2 is OK, because if A were to join a company, s/he would have access to similar or better AI tools.

A lot of nuances need to be thought through!

# Outline

- 1 Terminology
- 2 Brief Historical Overview: 1949–2017
- 3 Recent Developments: 2017–Date
- 4 Current Situation
- 5 The Future**
  - Achieving AI Sovereignty
  - Impact on Education
  - Impact on Research**
  - Impact on Employment



# Use of AI in Solving Research Problems in Mathematics

- In recent days, several open problems in mathematics and theoretical CS have been solved by expert researchers (sometimes repeatedly) prompting an AI utility.
- Clearly AI utilities are able not only to “scan” research papers and books, but also to “integrate” them to arrive at new approaches for solving problems.
- The problems solved thus far have been “open,” but not “significant” in my view (e.g., the Riemann hypothesis)
- Will this lead to an elimination of human researchers?

# Can AI Utilities Solve Problems in Other Domains?

- Mathematics is particularly well-suited for the use of AI tools.
- Physics less so than mathematics, and biology / chemistry less so than physics.
- Though the 2024 Nobel Prize in Chemistry went to AlphaFold, I doubt that novel drugs will emerge through AI – novel *drug candidates*, yes, but not validated drugs.

## Another Impact on “Research”

- Nowadays almost all papers submitted to ML conferences are “reviewed” by AI tools, mostly ChatGPT
- The quality of these reviews is *exceptionally poor*
- The organizers would be better off just picking accepted papers via a lottery.
- Eventually this will cause the “major” conferences to lose prestige.

# Outline

- 1 Terminology
- 2 Brief Historical Overview: 1949–2017
- 3 Recent Developments: 2017–Date
- 4 Current Situation
- 5 The Future
  - Achieving AI Sovereignty
  - Impact on Education
  - Impact on Research
  - Impact on Employment

# Impact on Global Employment

- Like all new technologies, AI will create more jobs than it destroys.
- Unlike (e.g.) the industrial revolution, the persons displaced by AI can *largely* be retrained to fill the jobs created by AI.
- *However*, the impact will *not* be uniform throughout all countries.
- I also believe that AI will *widen disparities* in incomes.
- AI engineers might become like lawyers, film stars and athletes.

## Concluding Remarks

- The advent of AI throws up more questions than answers.
- Undoubtedly AI is one of the most “disruptive” developments.
- *Remember:* I am just guessing, just like everyone else.

# Thank You

# Thank You!

