

LES MODÈLES DE LANGUE USAGES ET ENJEUX SOCIÉTAUX

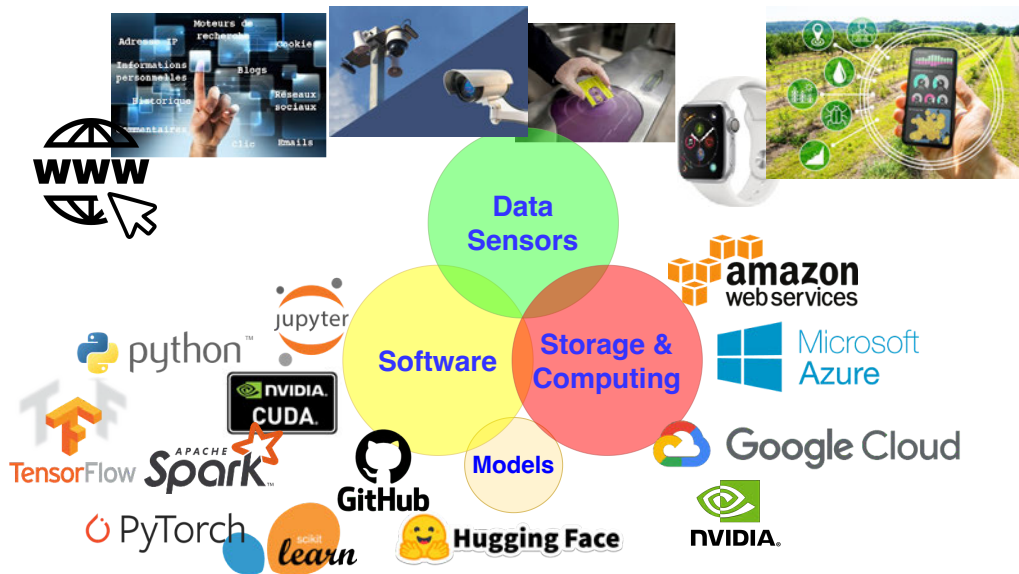
Mercredi 8 Octobre 2025
Séminaire CNRS, ANF-TDM-IA 2025

Vincent Guigue
<https://vguigue.github.io>

INTRODUCTION



The Ingredients of Machine-Learning



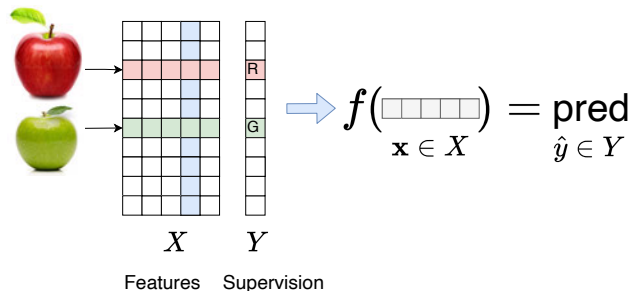


From tabular data to text

■ Tabular data

- Fixed dimension
- Continuous values

⇒ A perfect playground for machine learning



■ Textual data

- Various length
- Discrete values

⇒ Complex for machine learning

This new iPhone, what a marvel

An iPhone, What a scam!

Half the price is for the logo

Apple once again proves that perfection can be sold

How do we turn this text data into a table?

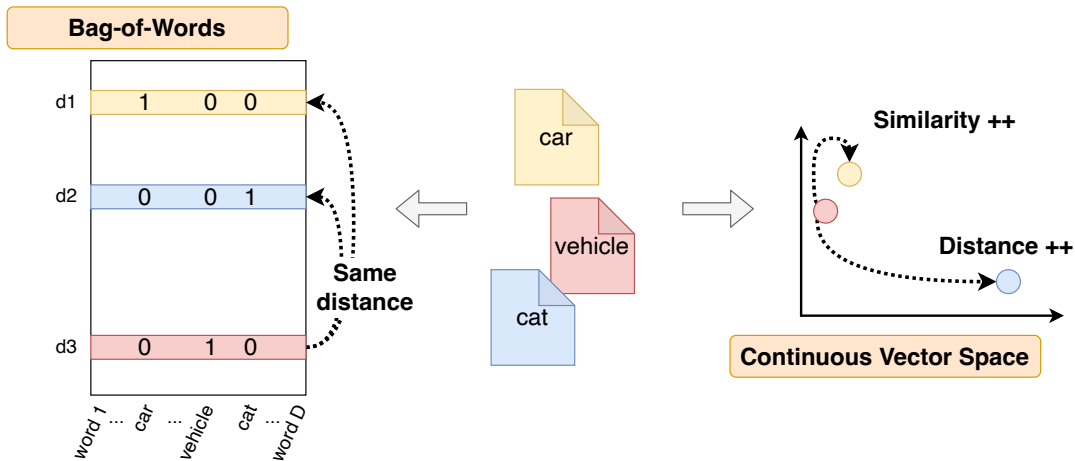




Deep/Representation Learning for Text Data

From Bag of Words to Vector Representations

[2008, 2013, 2016]



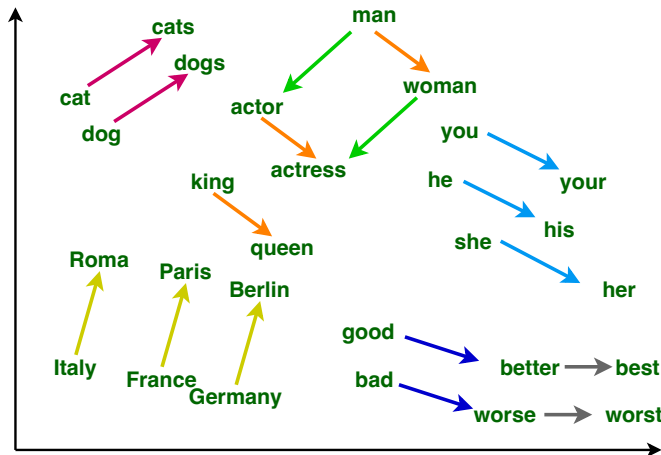
LeCun, Y., Bengio, Y., Hinton, G. (2015). [Deep learning](#). Nature, 521(7553), 436-444.



Deep/Representation Learning for Text Data

From Bag of Words to Vector Representations

[2008, 2013, 2016]



- **Semantic Space:**
similar meanings
 $\Leftrightarrow$
close positions
- **Structured Space:**
grammatical regularities,
basic knowledge, ...

Distributed representations of words and phrases and their compositionality, [Mikolov et al. NeurIPS 2013](#)



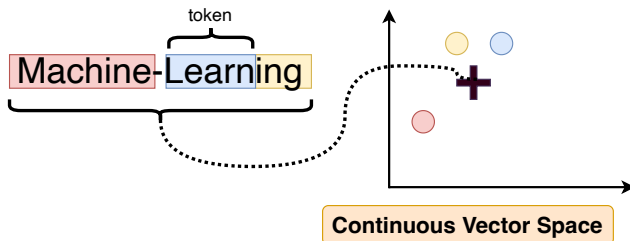
Deep/Representation Learning for Text Data

From Bag of Words to Vector Representations

[2008, 2013, 2016]

From Words to Tokens

Word Piece statistical split



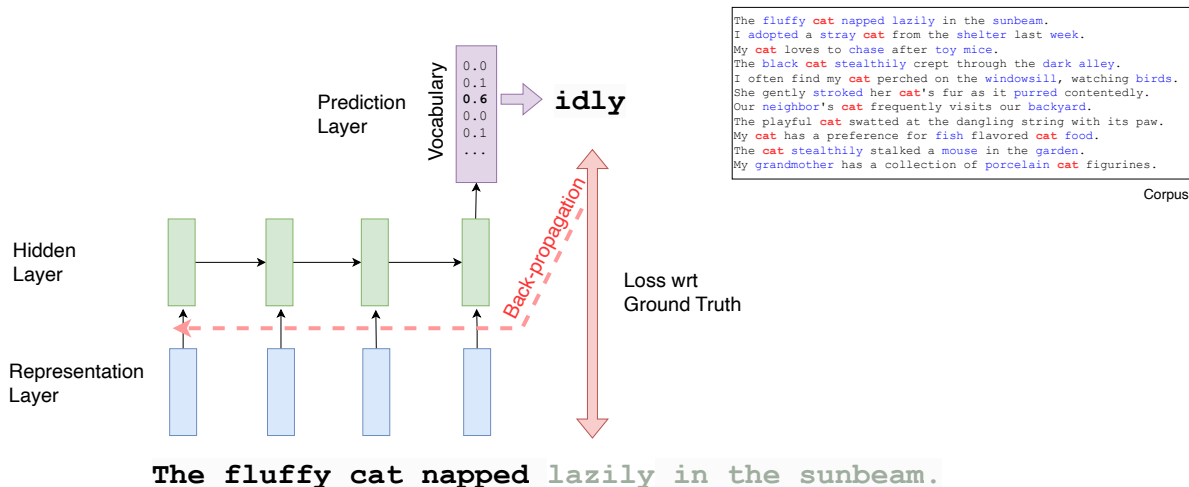
- Representation of unknown words
- Adaptation to technical domains
- Resistance to spelling errors

Enriching word vectors with subword information. [Bojanowski et al. TACL 2017.](#)



Aggregating word representations: towards generative AI

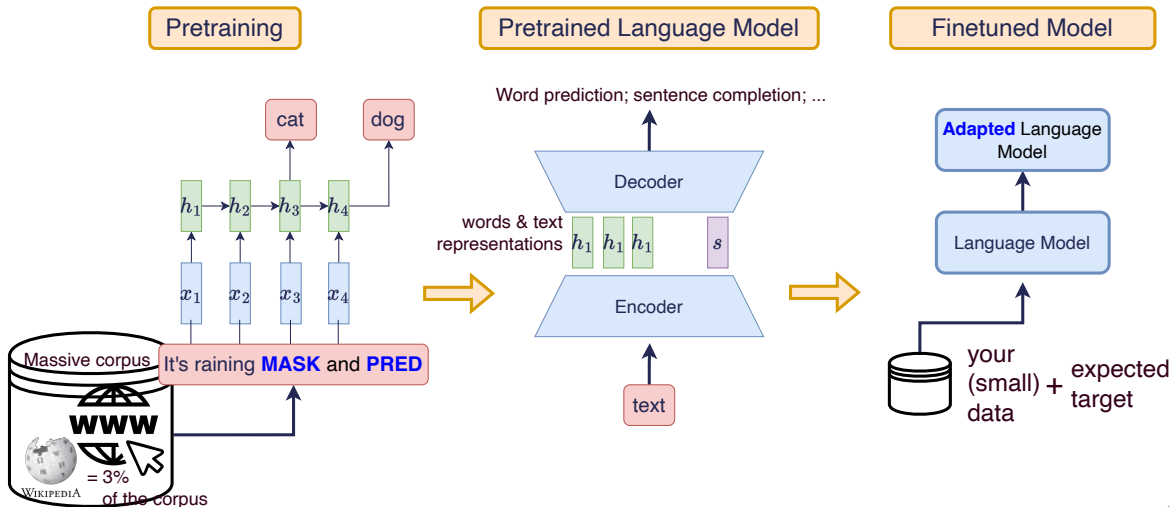
- Generation & Representation
- New way of learning word positions





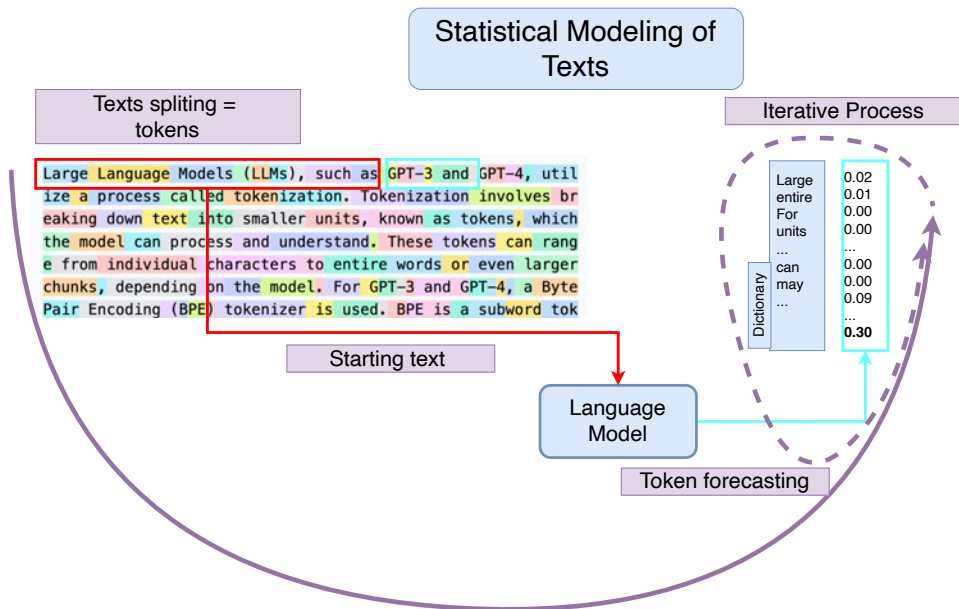
A new developpement paradigm since 2015

- Huge dataset + huge archi. $\Rightarrow$ unreasonable training cost
- Pre-trained architecture + 0-shot / finetuning





At the end of the day: a stochastic parrot :)



CHATGPT

NOVEMBER 30, 2022

1 MILLION USERS IN 5 DAYS

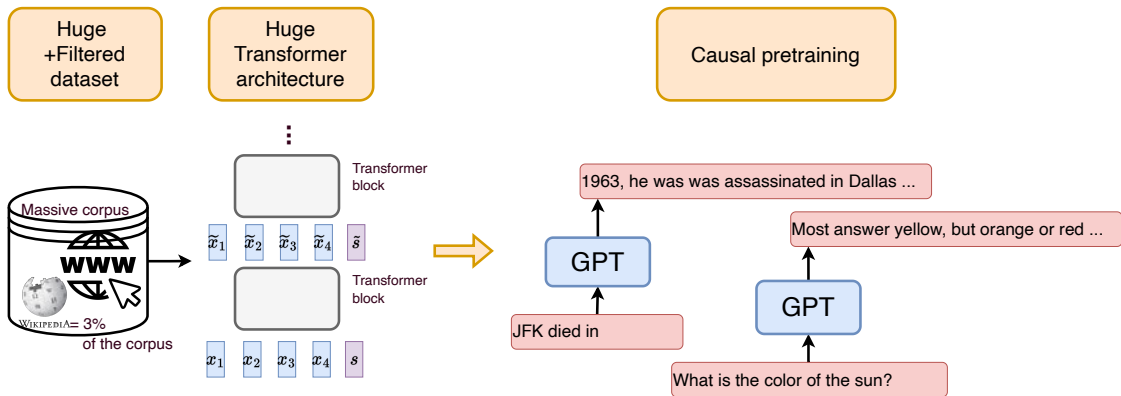
100 MILLION BY THE END OF JANUARY 2023

1.16 BILLION BY MARCH 2023



The Ingredients of chatGPT

0. Transformer + massive data (GPT)

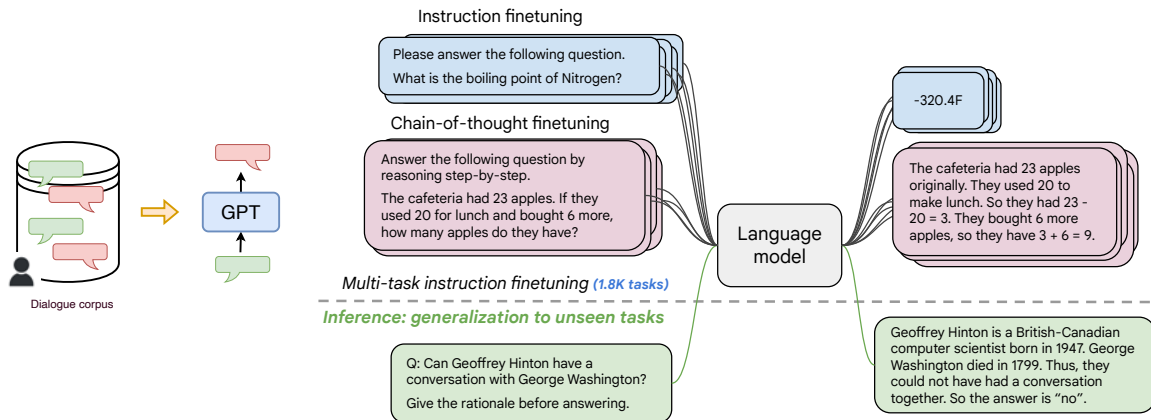


- Grammatical skills: singular/plural agreement, tense concordance
- (Parametric) Knowledge: entities, names, dates, places



The Ingredients of chatGPT

1. Dialogue + Tasks



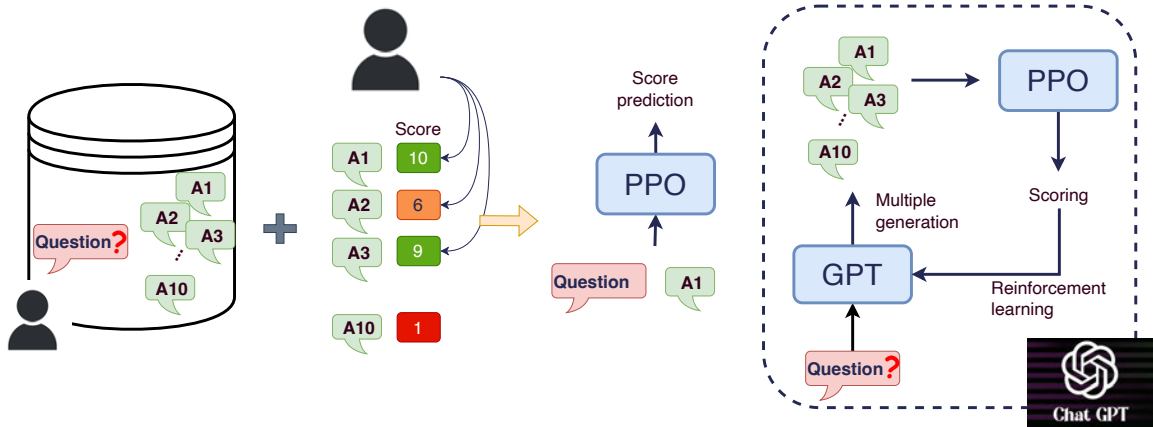
■ Very clean data

Data generated/validated/ranked by humans



The Ingredients of chatGPT

3. Instructions + answer ranking



- Database created by humans
- Response improvement

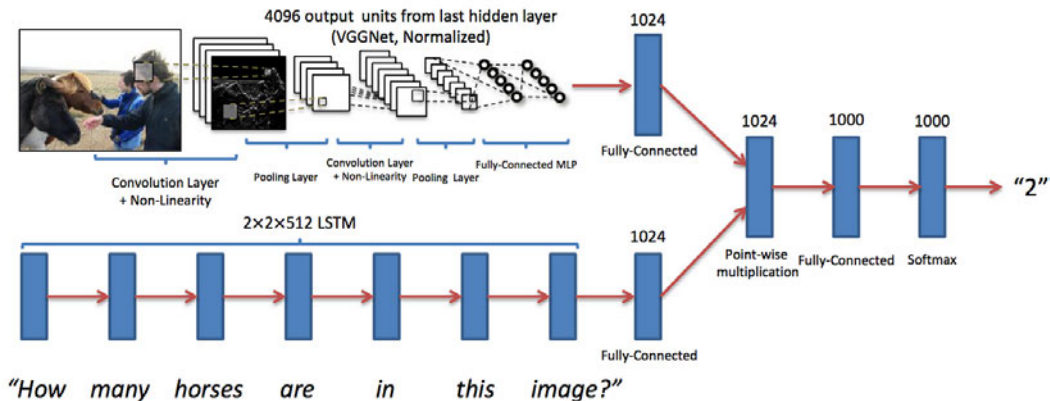
- ... Also a way to avoid critical topics = censorship



GPT4 & Multimodality

Merging information from text & image. **Learning** to exploit information jointly

The example of VQA: visual question answering



⇒ Backpropagate the error ⇒ modify word representations + image analysis



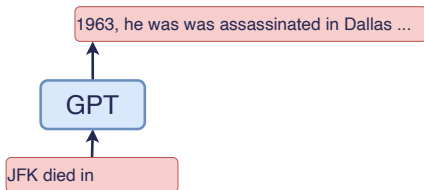
VQA: Visual Question Answering, arXiv, 2016, A. Agrawal et al.

MACHINE LEARNING LIMITS

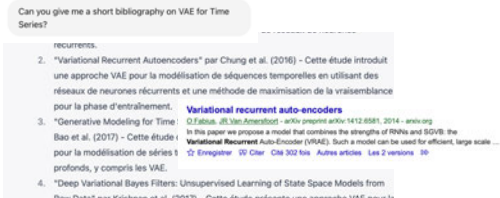


chatGPT and the relationship with truth

- 1 **Likelihood** = grammar, agreement, tense concordance, logical sequences...
⇒ Repeated knowledge
- 2 Predict the most **plausible** word...
⇒ produces **hallucinations**
- 3 **Offline** functioning
- 4 chatGPT $\neq$ **knowledge graphs**
- 5 Brilliant answers...
And silly mistakes!
+ we cannot predict the errors



Example: producing a bibliography

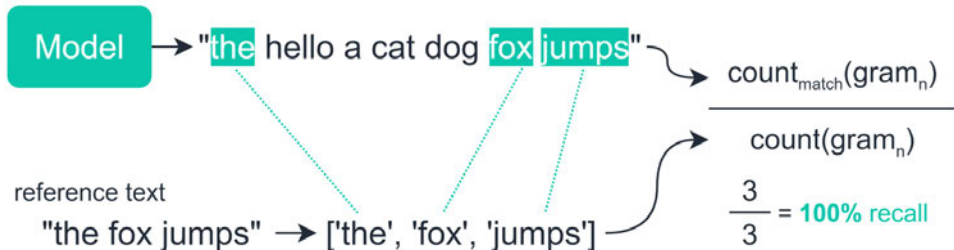




Generative AI: how to evaluate performance?

The critical point today

- How to evaluate against ground truth?
- How to evaluate system confidence / plausibility of generation?

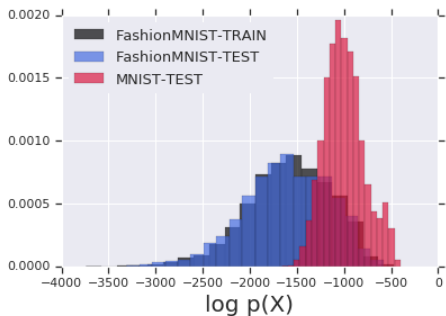




Generative AI: how to evaluate performance?

The critical point today

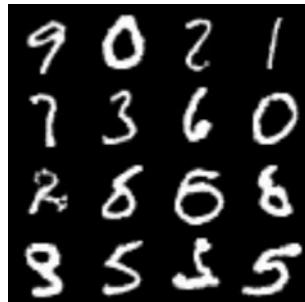
- How to evaluate against ground truth?
- How to evaluate system confidence / plausibility of generation?



Plausibility



Train



Test



Do Large Language Models Know What They Don't Know?, Yin et al. , ACL, 2023

Do Deep Generative Models Know What They Don't Know?, Nalisnick et al. , ICLR, 2019



Stability/predictability

- Difficult to bound a behavior
- Impossible to predict good/bad answers

⇒ Little/no use in video games



how old is Obama



Barack Obama was born on August 4, 1961, making him 61 years old as of February 2, 2023.





Stability/predictability

- Difficult to bound a behavior
- Impossible to predict good/bad answers

⇒ Little/no use in video games



how old is obama?



As of 2021, Barack Obama was born on August 4, 1961, so he is 60 years old.

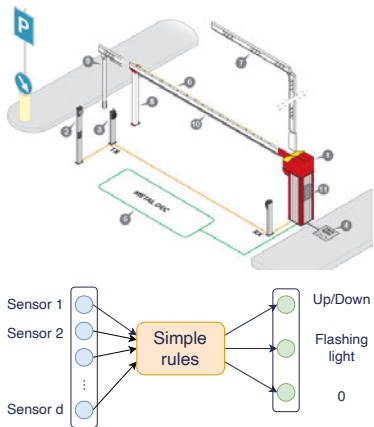


and today?

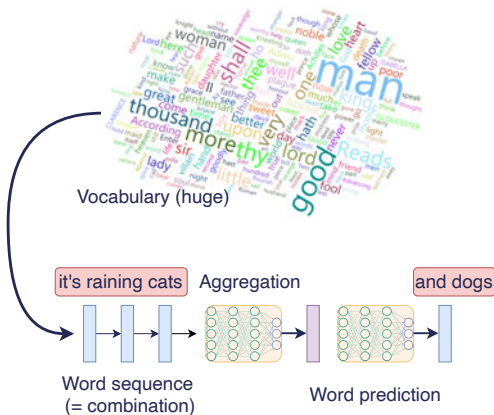




Explainability... And complexity



- Simple system
- Exhaustive testing of inputs/outputs
- Predictable & explainable



- Large dimension
- Complex non-linear combinations
- Non-predictable & non-explainable



Explainability... And complexity

Interpretability vs Post-hoc Explanation

Neural networks = **non-interpretable** (almost always)

too many combinations to anticipate

Neural networks = **explainable a posteriori** (almost always)



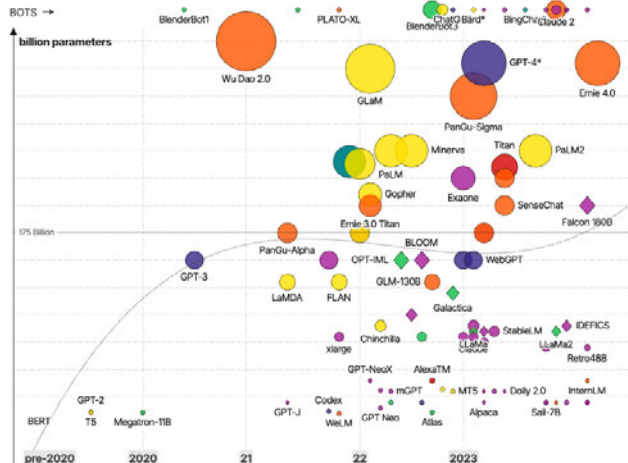
[Uber Accident, 2018]

- Simple system
- Exhaustive testing of inputs/outputs
- **Predictable & explainable**
- Large dimension
- Complex non-linear combinations
- **Non-predictable & non-explainable**



size = no. of parameters open-access

● Amazon-owned ● Chinese ● Google ● Meta / Facebook ● Microsoft ● OpenAI ● Other



* = parameters undisclosed // see the data

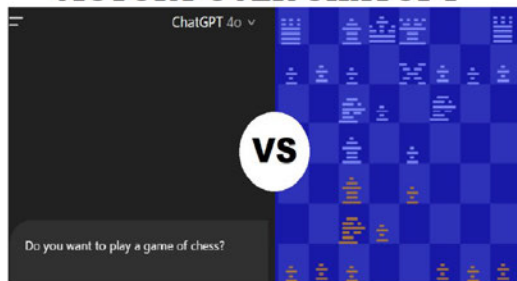
1998	LeNet-5	= 0.06M
2011	Senna	= 7.3M
2012	AlexNet	= 60M
2017	Transformer	= 65M / 210M
2018	ELMo	= 94M
2018	BERT	= 110M / 340M
2019	GPT2	= 1,500M
2020	GPT3	= 175,000M
2025	Llama-4	= 2,000,000M



Everything beyond the LLM's capabilities/training

- Simple calculations
(multiplication, division)
- Generating n -syllable animal names
(in progress)
- Playing chess
- Follow (complex) causal reasoning
- ...

ATARI 2600 SCORES STUNNING VICTORY OVER CHATGPT

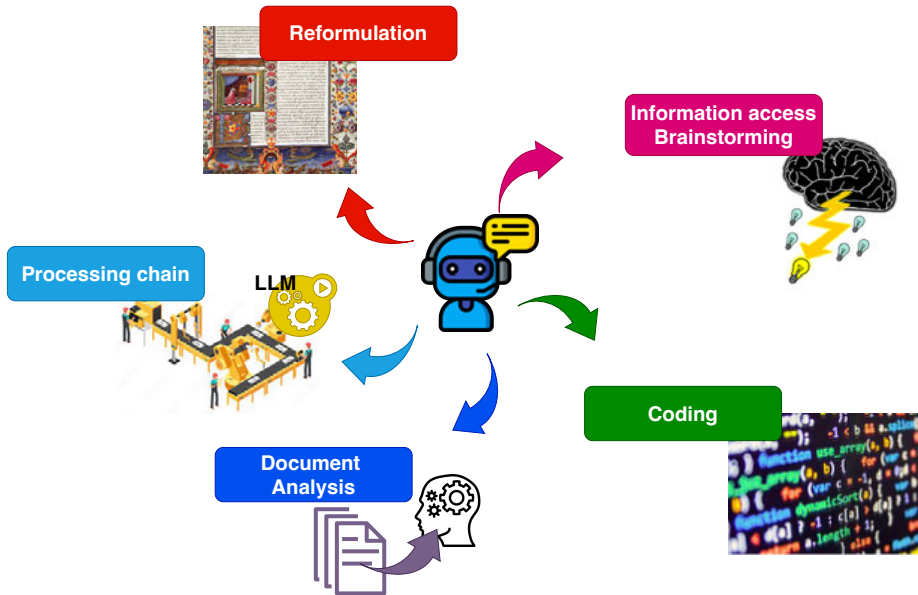


***WHEN YOU UNDERESTIMATE A 1977 CHESS ENGINE...
AND IT HUMBLER YOU IN FRONT OF THE WHOLE INTERNET***

LARGE LANGUAGE MODELS USES



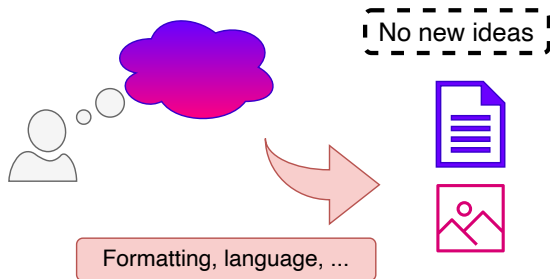
Key uses in 5 pictures





(1) Formatting information

A fantastic tool for
formatting



■ Personal assistant

- Standard letters, recommendation letters, cover letters, termination letters
- Translations

■ Meeting **reports**

- Formatting notes

■ **Writing** scientific articles

- Writing ideas, in French, in English

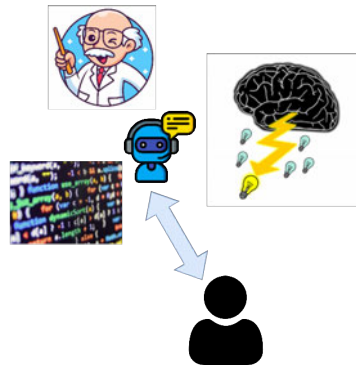
No new information ⇒ just writing, improving, translating, cleaning up, ...



(2) Brainstorming / Course Planning / Statistics Review

- **Find** inspiration [writer's block syndrome]
- **Organize** ideas quickly
- **Avoid omissions** / increase confidency
- **Search** in a targeted way, adapted to one's needs
- **Answer** student questions (24/7)
- **Partner** in research, test/enrich ideas

⇒ Impressive answers, sometimes incomplete or partially incorrect... But often useful



- In which areas are LLMs reliable?
- What are the risks for primary information sources?
- What societal risks for information?



(3) Coding: Different Tools, Different Levels

- Providing solutions to exercises
- Learning to code or getting back into it
 - New languages, new approaches (ML?)
 - Benefit from explanations...

But how to handle mistakes?

- Help with a library [*getting started*]
- Faster coding



GitHub
Copilot



- What about copyrights?
 - What impact on future code processing?
- How to adapt teaching methods?
- How many calls are needed for code completion?

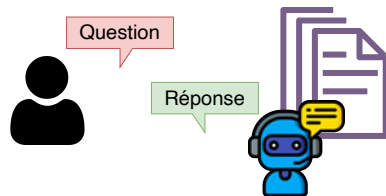
What about the carbon footprint?
- What is the risk of error propagation?

```
sentiments write_sql.go parse_expenses.py addresses.rb
1 import datetime
2
3 def parse_expenses(expenses_string):
4     """Parse the list of expenses and return the list of triples (date,
5     ignore lines starting with #.
6     Parse the date using datetime.
7     Example expenses_string:
8         2016-01-02 -34.01 USD
9         2016-01-03 2.59 DKK
10        2016-01-03 -2.72 EUR
11    """
12    expenses = []
13    for line in expenses_string.splitlines():
14        if line.startswith("#"):
15            continue
```



(4) Document Analysis

- Summarizing documents / articles
- Dialoguing with a document database
- Assistance in writing reviews
- FAQs, internal support services within companies
- Technology watch
- Generating quizzes from lecture notes



NotebookLM

Think **Smarter**,
Not Harder

Try NotebookLM

- Will articles still be read in the future?
 - Should we make our articles NotebookLM-proof?
- How to save time while remaining honest and ethical?

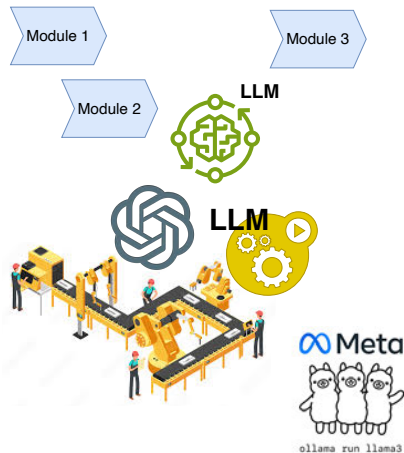


(5) LLM in a Production Pipeline / Agentic AI

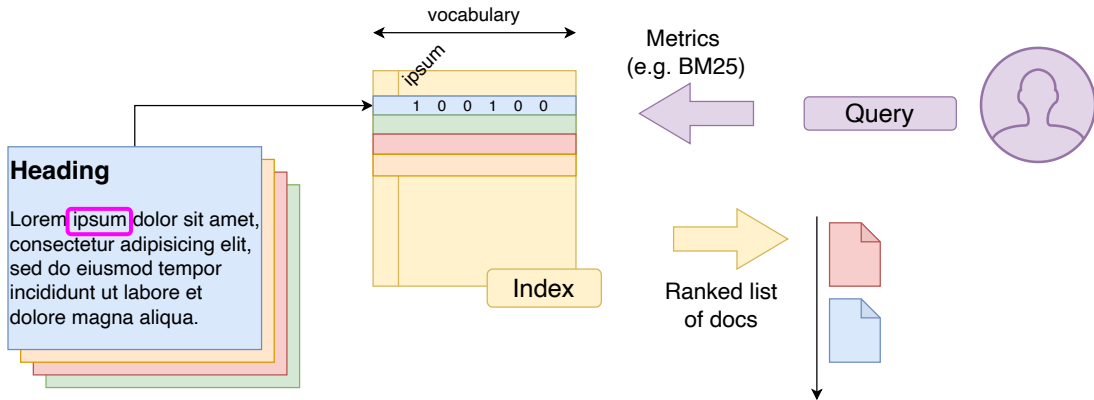
- Run LLM locally
- Extract knowledge
- Generate examples to train a model
[Teacher/student - distillation]
- Generate variants of examples ↗ ↗ increase dataset size

[Data augmentation]

⇒ Integrate the LLM into a processing pipeline
= little/less supervision = **Agentic AI**

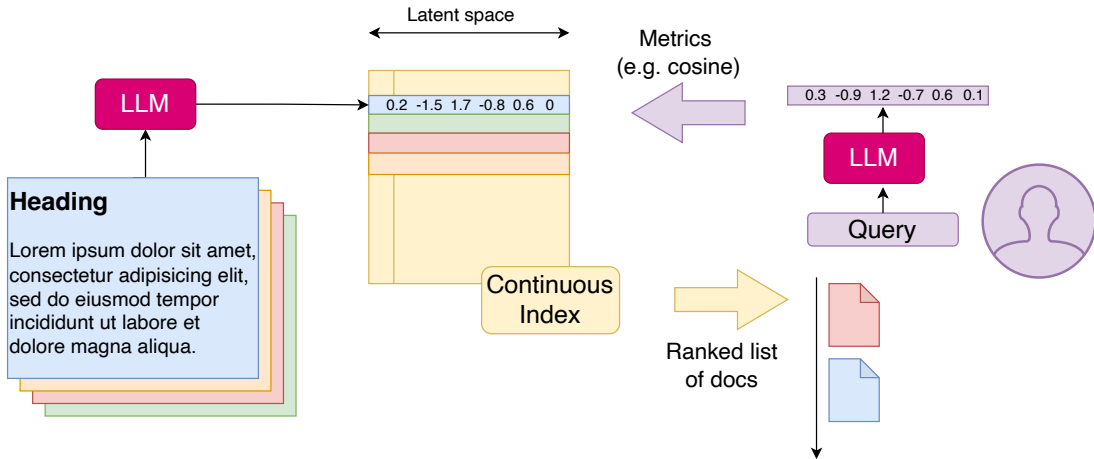


- How much does it cost? (\$ + CO₂) Need for GPUs?
- How good are open-weight models?
- How to build multiple agents?



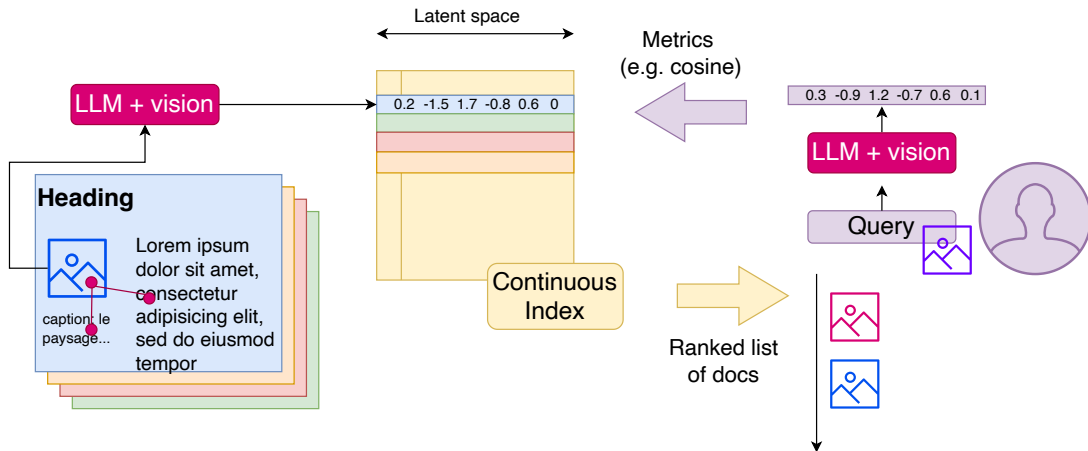


LLM vs Information Retrieval



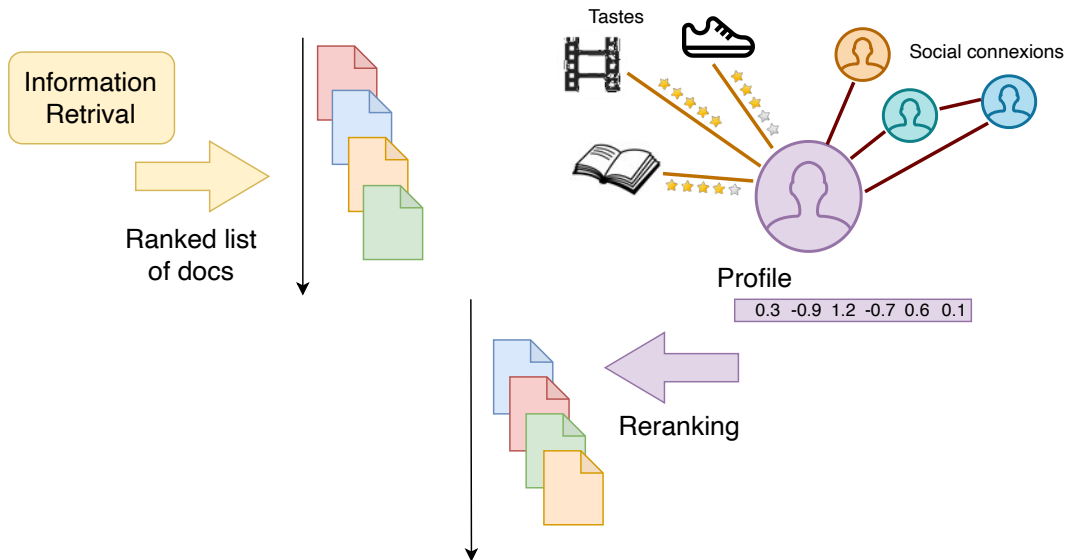


LLM vs Information Retrieval





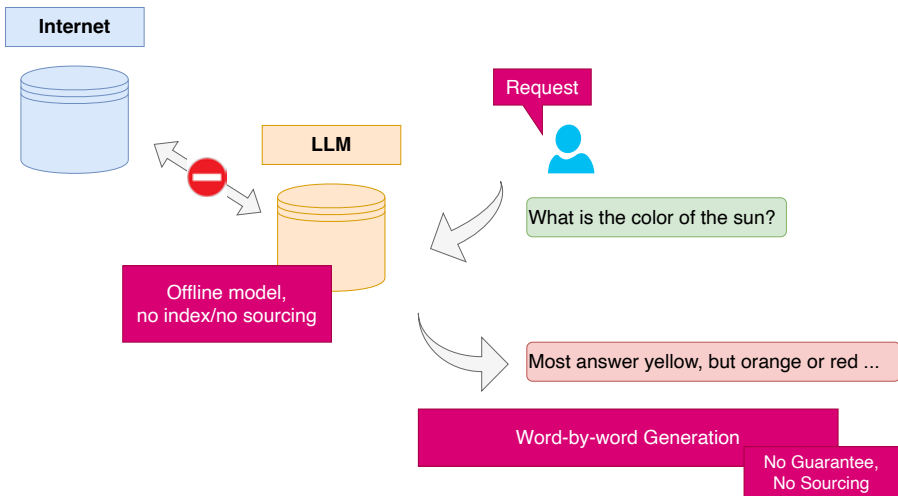
LLM vs Information Retrieval





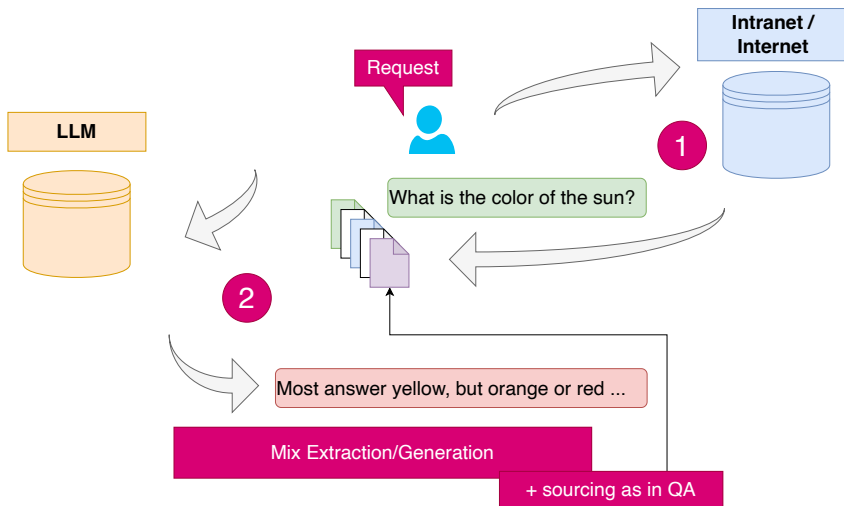
LLMs $\Rightarrow$ RAG : parametric memory vs Info. Extraction

- Asking for information from ChatGPT... A surprising use!
- But is it reasonable? [Real Open Question (!)]





LLMs $\Rightarrow$ RAG : parametric memory vs Info. Extraction



- RAG: Retrieval Augmented Generation
- (Current) limit on input size (2k then 32k tokens)



Language Handling

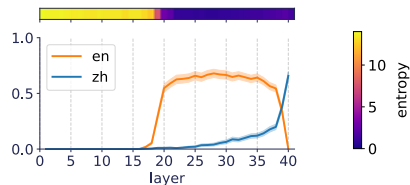
- Language models are (mostly) multilingual:

⇒ Think in the language you are most comfortable with

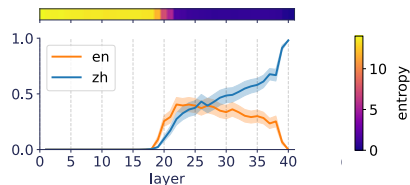
⇒ Ask for answers in the target language

[Wendler et al. 2024] Do Llamas Work in English?
On the Latent Language of Multilingual Transformers

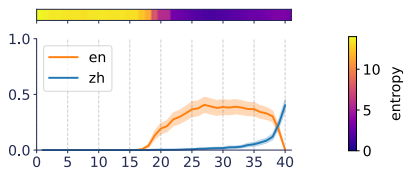
(a) Translation task



(b) Repetition task



(c) Cloze task



RISKS



Typology of AI Risks in NLP (L. Weidinger)



Discrimination, exclusion and toxicity

Harms that arise from the language model producing discriminatory and exclusionary speech.



Information hazards

Harms that arise from the language model leaking or inferring true sensitive information.



Misinformation harms

Harms that arise from the language model producing false or misleading information.



Malicious uses

Harms that arise from actors using the language model to intentionally cause harm.



Human-computer interaction harms

Harms that arise from users overly trusting the language model, or treating it as human-like.



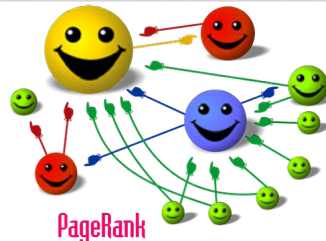
Automation, access and environmental harms

Harms that arise from environmental or downstream economic impacts of the language model.



Access to Information

- Access to dangerous/forbidden information
 - +Personal data
 - Right to be forgotten (GDPR)
- Information authorities
 - Nature: unconsciously, image = truth
 - Source: newspapers, social media, ...
 - Volume: number of variants, citations (pagerank)
- Text generation: harassment...
- Risk of anthropomorphizing the algorithm
 - Distinguishing human from machine





Machine Learning & Bias



Mustache, Triangular Ears, Fur
Texture

Cat



Over 40 years old, white,
clean-shaven, suit

Senior Executive

Bias in the data $\Rightarrow$ bias in the responses

Machine learning is based on extracting statistical biases...

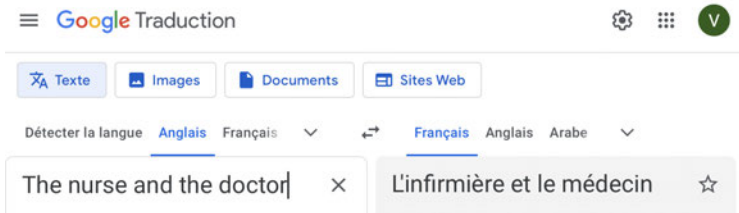
$\Rightarrow$ Fighting bias = manually adjusting the algorithm



Machine Learning & Bias



Stereotypes from *Pleated Jeans*



- Gender choice
- Skin color
- Posture
- ...

Bias in the data $\Rightarrow$ bias in the responses

Machine learning is based on extracting statistical biases...

$\Rightarrow$ Fighting bias = manually adjusting the algorithm



Bias Correction & Editorial Line

Bias Correction:

- Selection of specific data, rebalancing
- Censorship of certain information
- Censorship of algorithm results

⇒ Editorial work...

Done by whom?

- Domain experts / specifications
- Engineers, during algorithm design
- Ethics group, during result validation
- Communication group / user response

⇒ What legitimacy? What transparency? What effectiveness?





Machine learning is never neutral

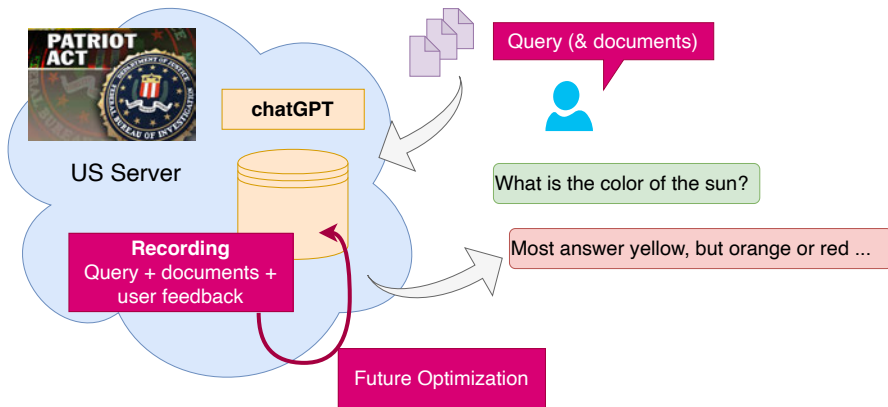
- 1 Data selection
 - Sources, balance, filtering
- 2 Data transformation
 - Information selection, combination
- 3 Prior knowledge
 - Balance, loss, a priori, operator choices...
- 4 Output filtering
 - Post processing
 - Censorship, redirection, ...

⇒ Choices that influence algorithm results





Data Leak(s): different security levels



- Transfer of sensitive data
- Exploitation of data by OpenAI (or others)
- Data leakage in future models



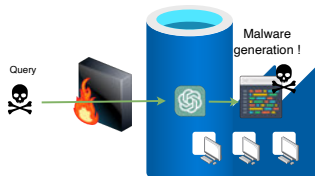
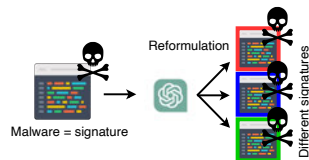
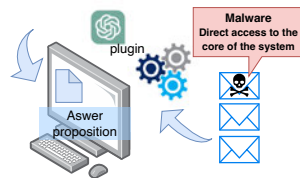
Data Leak(s): different security levels

Level 1: Commercial tools, free to use	Variable licenses (depending on the companies and subject to change over time). Uncertain data protection, risk to personal data. <i>chatGPT, Mistral, Perplexity, ...</i>
Level 2: Commercial tools, paid licence	Strong contractual guarantees. Risks associated with the <i>Patriot Act</i> . Possible to enforce non-storage of queries. <i>chatGPT, Mistral, Perplexity, ...</i>
Level 3: Local dev., Commercial tools & paid licence ++	+ Negotiation on the server location/data security. <i>Microsoft Azur, Mistral, AWS, Aristote, Ragarenn...</i>
Level 4: Local use	Use of a locally operated LLM, with no data transferred over the web. <i>HuggingFace, Ollama, ...</i>



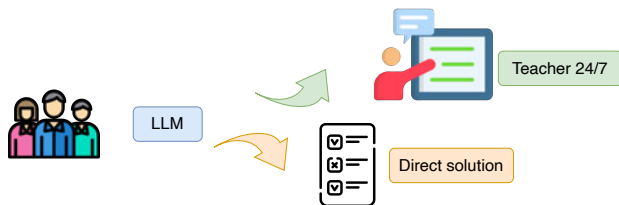
Security Issues

- Plug-ins $\Rightarrow$ Often significant security vulnerabilities for users
 - Email access / transfer of sensitive information etc...
- Management issues for companies
 - Securing (very) large files
- Increased opportunities for malware signatures
 - $\approx$ software rephrasing
- New problems!
 - Direct malware generation

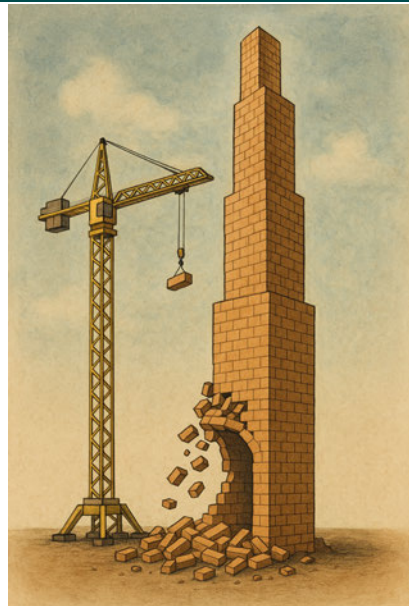


Educational Challenges

- Redefine our **educational priorities**, subject by subject, as we did with Wikipedia/calculator/...
 - Accept the **decline of certain skills**
- Train students in the use of LLMs, while managing to temporarily prohibit their use



- Learn to **recognize LLM-generated content**, use detection tools.



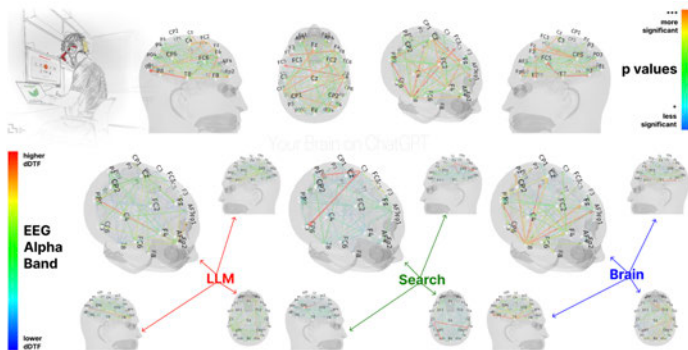
Decline / Evolution of Cognitive skills

Our brain will evolve with these new tools...

What is the scope of these transformations? What will be the consequences?

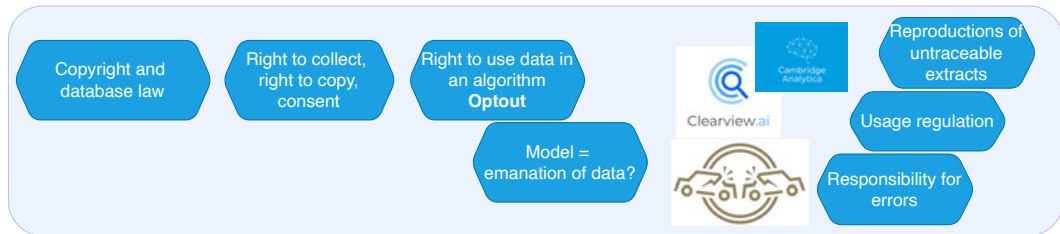
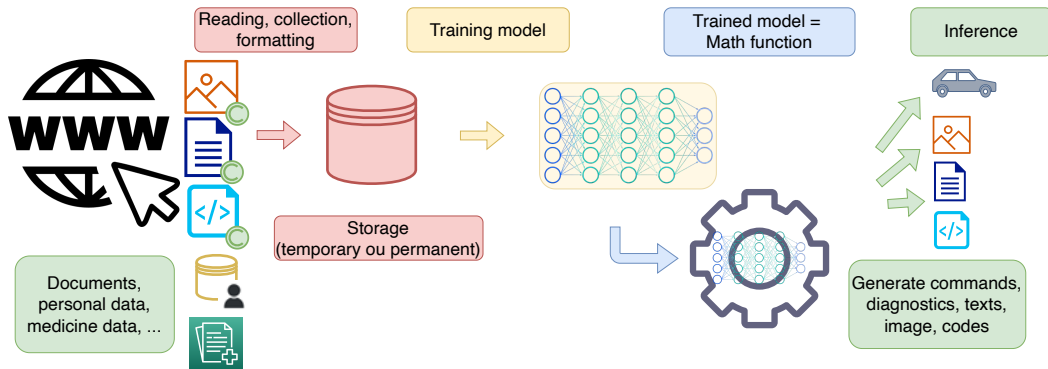
■ Education sciences and psychology had conjectured it...

cognitive sciences have measured it





Legal Risks/Questions



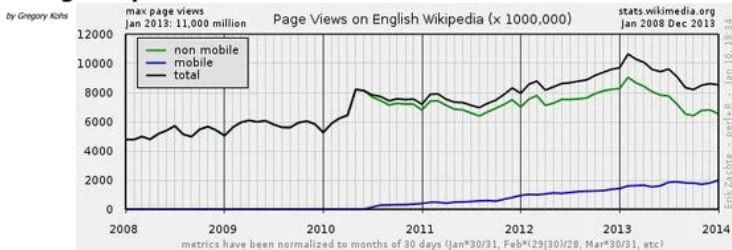


Economic Questions

- Funding/Advertising $\Leftrightarrow$ **visits** by internet users
- Google knowledge graph (2012) $\Rightarrow$ fewer visits, less revenue
- chatGPT = encoding web information... $\Rightarrow$ much fewer visits?

$\Rightarrow$ What **business model** for information sources with chatGPT?

Google's Knowledge Graph Boxes: killing Wikipedia?



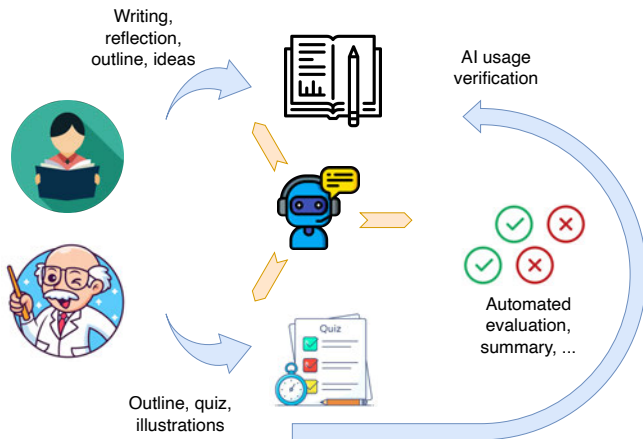
$\Rightarrow$ Who does **benefit from the feedback?** [StackOverflow]



Risks of AI Generalization

AI everywhere =
loss of meaning?

- In the educational domain
- Transposition to HR
- To project-based funding systems





How to approach the ethics question?

Medicine

- 1 **Autonomy:** the patient must be able to make informed decisions.
- 2 **Beneficence:** obligation to do good, in the interest of patients.
- 3 **Non-maleficence:** avoid causing harm, assess risks and benefits.
- 4 **Equality:** fairness in the distribution of health resources and care.
- 5 **Confidentiality:** confidentiality of patient information.
- 6 **Truth and transparency:** provide honest, complete, and understandable information.
- 7 **Informed consent:** obtain the free and informed consent of patients.
- 8 **Respect for human dignity:** treat all patients with respect and dignity.

Artificial Intelligence

- 1 **Autonomy:** Humans control the process
- 2 **Beneficence:** in the interest of whom? User + GAFAM...
- 3 **Non-maleficence:** Humans + environment / sustainability / malicious uses
- 4 **Equality:** access to AI and equal opportunities
- 5 **Confidentiality:** what about the Google/Facebook business model?
- 6 **Truth and transparency:** the tragedy of modern AI
- 7 **Informed consent:** from cookies to algorithms, knowing when interacting with an AI
- 8 **Respect for human dignity:** harassment behavior/ human-machine distinction



How to approach the ethics question?

Medicine

- 1 **Autonomy:** the patient must be able to make informed decisions.
- 2 **Beneficence:** obligation to do good, in the interest of patients.
- 3 **Non-maleficence:** avoid causing harm, assess risks and benefits.
- 4 **Equality:** fairness in the distribution of health resources and care.
- 5 **Confidentiality:** confidentiality of patient information.
- 6 **Truth and transparency:** provide honest, complete, and understandable information.
- 7 **Informed consent:** obtain the free and informed consent of patients.
- 8 **Respect for human dignity:** treat all patients with respect and dignity.

Artificial Intelligence

- 1 **Autonomy:** Humans control the process
- 2 **Beneficence:** in the interest of whom? User + GAFAM...
- 3 **Non-maleficence:** Humans + environment / sustainability / malicious uses
- 4 **Equality:** access to AI and equal opportunities
- 5 **Confidentiality:** what about the Google/Facebook business model?
- 6 **Truth and transparency:** the tragedy of modern AI
- 7 **Informed consent:** from cookies to algorithms, knowing when interacting with an AI
- 8 **Respect for human dignity:** harassment behavior/ human-machine distinction

CONCLUSION



Upcoming Challenges

■ What about hallucinations?

- Should we try to reduce them or learn to live with them?
- Will LLMs improve? In what directions?
- Do LLMs make us *lose* our connection to truth? To verification?

■ Do we need small or large language models?

- How much does it cost? Is it sustainable?
- With or without fine-tuning?
- What does frugality mean in the world of LLMs?

■ When others use them... What impact does it have on me?

- Productivity (fellow researchers, coders, reviewers, ...)
- Education: managing/training *tech-savvy* students

■ Data protection... Mine and others'

- Is it reasonable to train LLMs on GitHub, Wikipedia, scientific papers, news outlets, etc.?
- How important is privacy? What are the risks when using an LLM?



Upcoming Challenges

■ What about hallucinations?

- Should we try to reduce them or learn to live with them?
- Will LLMs improve? In what directions?
- Do LLMs make us *lose* our connection to truth? To verification?

■ Do we need small or large language models?

- H The smartphone has made me an *augmented human*...
- V Will the LLM make me an *augmented researcher*?
- V ⇒ Still, have a look at NotebookLM

■ When others use them... what impact does it have on me:

- Productivity (fellow researchers, coders, reviewers, ...)
- Education: managing/training *tech-savvy* students

■ Data protection... Mine and others'

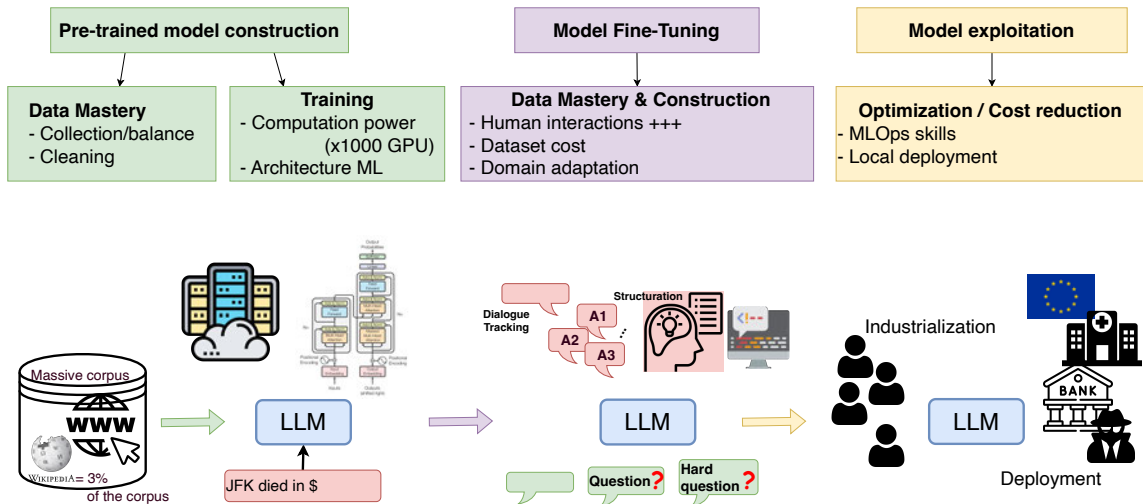
- Is it reasonable to train LLMs on GitHub, Wikipedia, scientific papers, news outlets, etc.?
- How important is privacy? What are the risks when using an LLM?



Levels of Access to Artificial Intelligence

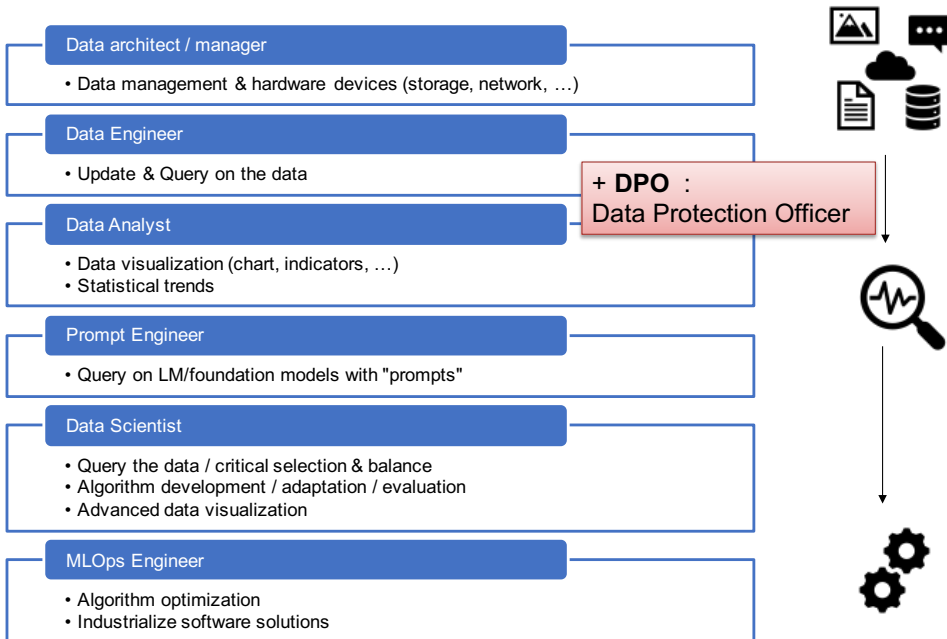
- 1 User via an interface: *chatGPT*
 - Some training is still required (2-4h)
- 2 Using Python libraries
 - Basics on protocols
 - Standard processing chains
 - Training: 1 week-3 months (ML/DL)
- 3 Tool developer
 - Adapt tools to a specific case
 - Integrate business constraints
 - Build hybrid systems (mechanistic/symbolic)
 - Mix text and images
 - Training: ≥ 1 year

Digital Sovereignty: the Entire Chain





A Multitude of Professions





Factors of Acceptability for Generative AI

1 Utilitarianism:

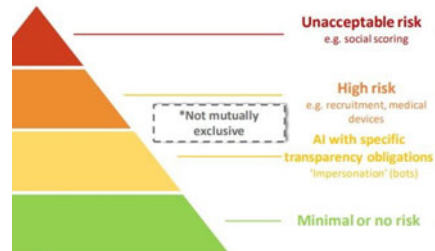
- Performance (acceptance factor of chatGPT)
- Reliability / Self-assessment

2 Non-dangerousness:

- Bias / Correction
- Transparency (editorial line, human/machine confusion)
- Reliable Implementation
- Sovereignty (?)
- Regulation (AI act)
 - Avoid dangerous applications

3 Know-how:

- Training (usage/development)





Why So Much Controversy?

- New tool [December 2022]
- + Unprecedented adoption speed [1M users in 5 days]
- Strengths and weaknesses... Poorly understood by users
 - Significant productivity gains
 - Surprising / sometimes absurd uses
 - Bias / dangerous uses / risks
- Misinterpreted feedback
 - Anthropomorphization of the algorithm and its errors
- Prohibitive cost: what economic, ecological, and societal model?

