

Briefing Note. Chat GPT

Main Title:	Briefing Note chatIPT		
Date:	14/03/2023	Release Maturity:	Draft/Live/Final
Author:	Allen Woods		
Owner:	Allen Woods		
Client:			
Version:	1.50		
CADMID	Concept		
Line of Development	Doctrine		
Organisation			
Release Classification	U		
Document ID/Number			

Supporting Links and Documents

Ser	Description	Location
1	Academic Integrity	Left click here
2	AI Bottleneck – Computing Power	Left click here
3	Alexander Hanff – The Register	Left click here
4	An infinite number of monkeys.	Left click here
5	Bing “chatbot” goes insane NYT	Left click here
6	chaptGPT In a Corporate World	Left click here
7	chatGPT and explicit disinformation	Left click here
8	Common crawl repository	Left click here
9	Common crawl, crawl frequency	Left click here
10	Conceptual Geometry – Route Efficiency	Left click here
11	EU AI Act	Left click here
12	EU AI Act – Proposal	Left click here
13	EU AI Liability Directive	Left click here
14	EU Digital Markets Act	Left click here
15	EU Digital Services Act	Left click here
16	EU E Privacy Directive	Left click here
17	EU GDPR	Left click here
18	EU NIS 2	Left click here
19	Github – ChatIPT resources	Left click here
20	Language Models “Know what they know”	Left click here
21	Large Language Model Architecture	Left click here
22	Linkedin post chatIPT Data sources	Left click here
23	Maintaining Large Language Models	Left click here
24	Microsoft AI Supercomputer	Left click here
25	MS Lobotomised Bing Chat...	Left click here
26	Multi Modal LLM's	Left click here
27	Noam Chomsky – “High Tech Plagiarism”	Left click here
28	Open AI chatGPT “fun facts”	Left click here
29	Open AI Content Policy	Left click here
30	Open AI Copyright Infringement	Left click here
31	Open AI FAQ	Left click here
32	Open AI Introducing chatGPT	Left click here
33	Open AI Statistics	Left click here
34	Open AI Terms and Conditions	Left click here
35	Open AI. GPT 4 Release, a note of caution	Left click here
36	OpenAI – GPT 4	Left click here
37	OpenAI Board Level Organisation	Left click here
38	OpenAI Privacy Policy	Left click here
39	OpenAI profile	Left click here
40	OpenAI Research links	Left click here
41	Paper – The Unfairness of Machine Learning	Left click here
42	Perception and LLM's	Left click here
43	Project Bibliotik	Left click here
44	Project Gutenberg	Left click here
45	Reasoning in Large Language Models	Left click here
46	Solarwinds Briefing Note	Left click here
47	The “size” of Wikipedia	Left click here
48	The Lithium Experiment	Left click here
49	The UK Post Office Horizon Enquiry	Left click here
50	The US Marines Don't Do AI...	Left click here
51	Training Large Language Models	Left click here
52	Tricking chatGPT	Left click here
53	Using ChatGPT for generating Email	Left click here
54	Wolfram Application of a Maths Engine	Left click here
55	Wolfram chatGPT Functionality	Left click here

Contents

Supporting Links and Documents.....	2
The Purpose Of This Document	4
Scope	5
Caveats.....	5
On The Links.....	5
The Authors Experience	5
Open AI	7
ChatGPT	8
Serious Intellectual Review.....	8
Large Language Models – A Little Theory.....	9
The Training Data	10
Stored Data... Once Trained.....	11
Not Just English.. ..	12
Retraining.....	12
Managing The Volume Of Enquiries	13
Enquiry Response Effort	13
Short Term Memory Issues	13
Management of the Exercise	13
OK, So what Does This Big Computer Look Like?	16
Where Is It?.....	16
Legal and Compliance Issues?.....	18
Defamation, Intellectual Property, Copyright and More... ..	19
OpenAI Terms and Conditions	20
Other Risks and Issues.... ..	22
Are Large Language Models Intelligent?	24
What Next?	25
A Bit of An Assessment of Risk	25
Conclusions	26
Anything Else To Raise?	27
Recommendations	28

The Purpose Of This Document

chatGPT was launched on to an unexpected world in January 2023. It is fair to say that its reception has been one of overwhelming excitement that probably took its inventors by surprise. The world it seems, has been given its “babel fish”. However, as is often the case, in IT, there is no accounting for people and what they do when given access to new toys, chatGPT is no exception. With great enthusiasm, end users, in various trade and professional disciplines have been experimenting with it for their own purposes. The results the author has read, have been, to say the least, questionable, some degree of “fitness for purpose” being demonstrated, but many, many causes of concern too.

One of the problems in assessing the suitability of the tool, being that commentary is “here today, gone tomorrow” which is part of the ephemeral nature of the web. However, as is the case, there is a growing concern about what chatGPT does and how it works and a series of potential risks that have reared their collective heads. The risks are diverse, from Noam Chomsky’s warning of plagiarism to the exploitation of the tool to support the construction of sophisticated malware. This document provides, hopefully, a collated list of links to significant pieces that set out the nature of the risks in one place.

The key driver being a LinkedIn Post which can be read [here](#), which, in the space of a few days, was viewed over 120,000 times and resulted in extending the authors contact list considerably. If there was one single criticism of this whole chatGPT exercise, it is that there is not enough information, of an independent nature, in one place that people can study at their leisure. This is an attempt to fill that real need.

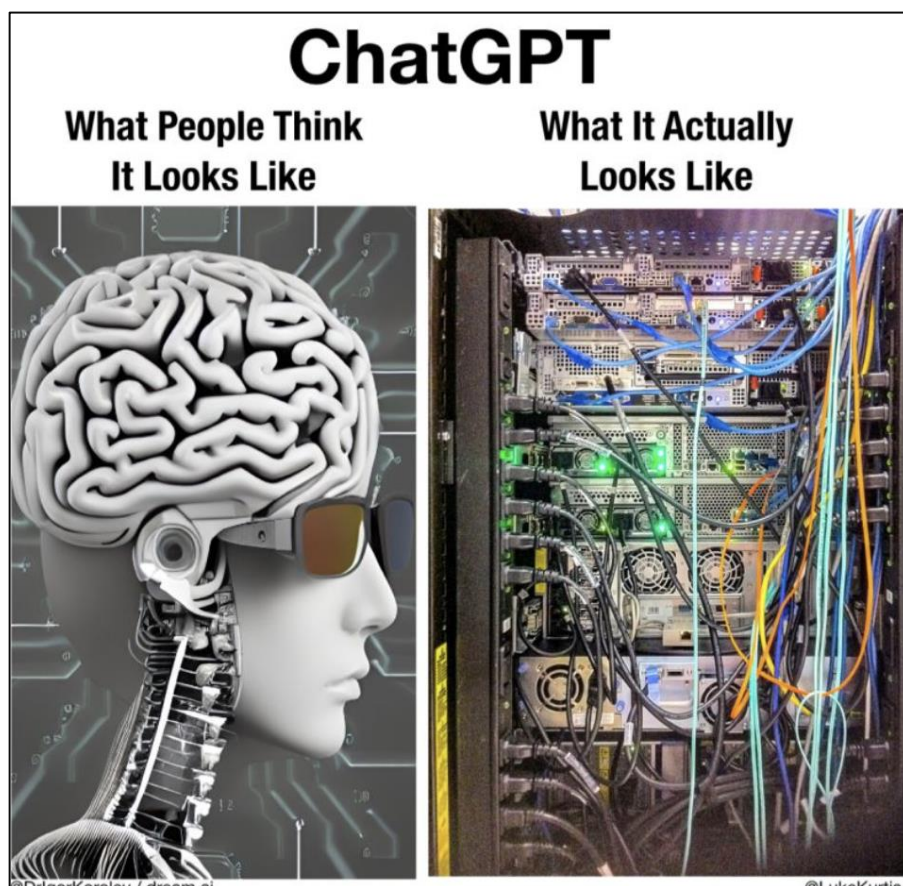


Figure 1 - Just about sums things up!

Scope

This scope of this document is set out below:

- To provide a collated list of key links. One of the major issues in trying to figure out how chatGPT does what it does is that articles and posts appear and disappear at lightning speed.
- To inform. The world of IT is prone to type, chatGPT being no exception.
- To provide the means for the reader, once informed, to make a balanced judgement on the nature of LLM's in general and the OpenAI GPT series of products in particular.
- To serve as a warning. In short, "beware of geeks bearing gifts" chatGPT and similar tools are not "JFDI" exercises and need careful consideration before adoption.

The author's background is ex military, if tasked with reviewing any product, a briefing note, of this kind, would be produced as a matter of routine as "back pocket stuff" to brief senior people on the pro's and cons of any product. The aim, to inform decision takers through the provision of verifiable evidence.

Caveats

This document has been put together by one person. Readers may disagree with the nature of the selected links and comments, that is their choice, however, all of the links in this document have been read by the author. It is recommended that those who agree or disagree with the content do the same carefully.

This version contains additional links and has been restructured to an extent. The main reason for that is that as time has passed, more details have emerged about the nature of Large Language models and how they work and about the effort involved in making chatGPT and GPT4 work, noting to that there have been reported enhancements to the tool set that address a number of functional shortfalls, the introduction of Stephen Wolframs Maths Engine through an API being an example.

On The Links

The opening page contains a list of 55 links from various sources that the author has read. Some from eminent scientists of world repute others from less august folk. Some written by the author. All in the authors view, making valid comment and observation on the nature of chatGPT in particular and Large Language Models (LLM) generally. The aim of the link page is to provide a consolidated list so that those reading this can, if they wish, research further independently of what is written in the body of the text. Bear in mind however, that there are other links, embedded in the body of the text, particularly in respect of LLM theory, that may be of use to the reader.

The aim is that each link is a gateway of its own kind, into a much, much wider world. For completeness, links to key legislation in the EU have been included, as have links to the OpenAI Terms and Conditions, FAQ and the company content policy. It is unwise to ignore both sets of documents as both will have an impact on how the chatGPT tool is operated and perhaps more importantly, paid for by end users (which means anyone really who uses the OpenAI product set for any reason).

The supporting notes are the authors view of the nature of LLM's in general using chatGPT for illustrative purposes. A phrase the author likes to use is "pathfinder not gospel" meaning they are what the author thinks, opinion with the encouragement that readers should follow the links to form your own view.

The Authors Experience

Readers are invited to review the authors "LinkedIn" profile available [here](#). While not an AI expert he has some experience in developing software to support the user of words as a data source – see the "Document Management Case Study" on the profile a key article in respect of the significance of a corporate dictionary and the effort involved in building such a thing. There is also a short description of

an exercise carried out to support the supply of Health and Safety documentation, called the “Lithium Experiment”, involving the combination and alignment of a library of some 30,000 documents from various suppliers of hazardous material as part of an inventory of 17,000,000 items, in order to detect and minimise a serious health risk.

Additionally, the author has been asked to provide commentary on the nature of AI, but focussing on facial recognition, the case study for which can be read [here](#).

This is the second exercise of this kind, one of collation of hyperlinks to key documents, that the author has written and posted, the first focussing on the Solarwinds incursion, which can be read [here](#).

Trying To Put It All Into Context

This document is aimed at trying to explain what a large language model is and some of the issues to consider when thinking about using one. As a first step exercise, to try and put some context into what an LLM does, a neural network and tokenisation consider this as an exercise to try and illustrate just what is going on. Let's assume you are in an office and in the office are some reference books. The purpose of the reference books is immaterial as it is the book as an information source that is the key concept.

1. From the reference books, pick one. Preferably one with a minimum of 100 pages.
2. Open it not to read, but to consider as an information source.
3. Note the identifying details like the author name, the publisher name, the ISBN, table of contents etc.
4. Having noted then, look at the content. Consider it not as subject matter but instead as an extensive collection of words.
5. If at all possible, carry out a word count of unique words in the book. The result will be the basis for calculating something called the token count.
6. Whilst identifying unique words, keep a running total of all the words.
7. Whilst doing the counts, keep track of the time required to do both and how complicated it may or not be to identify the uniqueness of the words in the text.
8. Once both of the counts have been completed, multiple the unique word count by the total word count.
9. The result will give you an approximated first token count as a simple relationship between uniqueness and appearance counts will be established.
10. Now, take other books from the reference library and do the same exercise for each book..
11. Once complete, consider any dependencies between books in the reference library, where dependencies accrue, multiply the token count of each book by the token count of any associated books. What this will do will give an approximate estimate of the number of relationships between words across multiple subject areas.
12. Then, do the same exercise for each and every document that has a significant content relationship with and and/or all books.
13. Then multiply the document token counts by the reference book token counts. What this will do is to give a ball park figure of the number of relationships between word content across a broad spectrum of subject matters of interest.
14. Whilst doing the content counts, note and count spelling mistakes, antonyms, synonyms and keeping a count of the number of relationships between each class of textual anomaly and

multiply the word counts of each anomaly by the total token count. This will give a very approximate count of how many comparison tests between properly formed content words and those that are errors

All the time keeping track of the time taken to do steps one to 14, noting too how many people, other than yourself need to be involved in the various counting exercises.

The result will be an approximate of the complexity of any resulting neural network that may be considered being used for LLM purposes. Try to work out the probability of each and every word appearing in any document to be created using the original reference material.

Getting exasperated? The exercise set out above is quite complicated, but in essence is the kind of thing an LLM does as part of its data induction and subsequent “transformation”. Note too that it is likely the number of relationships between words in the reference library might vary considerably from word to word.

Bear in mind that the kind of word count review set out above does not include more abstract issues like the nature of meaning and grammatical use of each word. Nor does the exercise give any indication of contextual complexity in such a way that when an LLM is responding to a query there is some understanding of “fit in a sentence” that may draw a response from multiple subject matter areas of expertise for each and every word being produced to answer any given question.

Open AI

In order to understand the nature of chatGPT it is necessary to understand a little of how the product came about. In respect of that, the reader's attention is drawn to this profile on the company who built the product from Wikipedia which is available [here](#).

In support of the Wikipedia profile, the image below includes some key facts on OpenAI

About OpenAI: Company Overview	
OpenAI was launched in 2015 in California, USA. One of the founders of OpenAI was Elon Musk, the current CEO of Twitter!	
Now, let's have a look at a key data about OpenAI:	
Key Stats	Key Info
Launch Date	December 11, 2015
Head Quarter	Pioneer Building, San Francisco, California, US
Founder(s):	Elon Musk (retreated), Sam Altman, Ilya Sutskever, Greg Brockman, Wojciech Zaremba, John Schulman
OpenAI Products	DALL-E, GPT-3, GPT-2, OpenAI Gym
Industry	Artificial intelligence
Business Type	Private (Currently listed as a For-Profit Type of Business)
Valuation	USD 20 Billion (2022)
Total Funding	USD 1 Billion (4 Rounds)
Number of Employees	335+ Employees (2022)
OpenAI Monthly Active Users (MAU)	21.1M+ (2022)
Official Website	openai.com

Figure 2 - Open AI Key Facts.

The full profile can be read [here](#). It should be noted that the content of the Wikipedia piece is like others reporting on the company as an operating capability in respect of detail.

ChatGPT

ChatGPT is a form of software called a “Large Language Model” (LLM) by the Artificial Intelligence (AI) community. In nutshell, the focus of LLM's is on the nature of language and its use, from a grammatical perspective, to converse with, typically, human beings in as realistic a manner as possible. The following links describe the “nature of the beast”, the first being general introduction to the product.

- A general introduction of the product is available [here](#).
- A second explanation of how chatGPT works, in more layman terms is available [here](#).
- A more detailed critique from Stephen Wolfram a renowned mathematician is available [here](#).
- A second more scientific explanation is available [here](#)
- As matter of balance, an assessment of the product, from Wikipedia, including comments by OpenAI senior staff, can be found [here](#).

It is also useful to understand more of what chatGPT has difficulty with as a piece of software, one such issue is its ability to “do” mathematics which would appear to be limited (its focus being on “language” rather than calculation per se). The difficulty with maths being explained by Stephen Wolfram, [here](#).

Serious Intellectual Review

It is also worthwhile to review some comment from computer scientists of repute.. For consideration:

Marvin Minsky – On the modelling of the nervous system available [here](#)

Noam Chomsky – On the nature of plagiarism [here](#)

Igor Korolev – A decision flow chart, available [here](#)

Steven Wolfram – LLM's are not designed to do math [here](#)

Felix Hovespian - This gentleman regularly raises issues about the nature of information management. [This](#) on what makes a “good programmer”..

Kelvin Low – On the legal implications of using chatGPT which can be read [here](#)

Max Tegmark – MIT professor on “ai super intelligence” a video, [here](#)

Terence Tao – World famous mathematician, an assessment can be read [here](#)

The authors interpretation of the comments above are:

LLM's are not intelligent. The focus of them is to solve a very complex problem, how language, as it is written and spoken can be replicated by software means. There are three general stages:

- Ingestion of data in huge volumes in order to build a model of language and how it is used conversationally.
- Transformation of the data, to “train” a model with the transformation adopting a sophisticated pattern recognition approach based on “Recurrent Neural Networks” on which to derive a dictionary and a set of relationships between words such that it can be inferred, with accuracy,

how sentences may be composed. The nature of how words are connected to each other being in a form of tokenisation.

- Finally, when questioned or queried, the provision of the means to calculate a possible valid response based on sophisticated probability testing with the aim of responding in the most contextually sound way possible, but with the constraint that any response will be constrained by the scope of any training data set.

It should be noted that the term “Large Language Model” is specific in that the focus of an LLM is grammatical and syntactical construction of sentences based on inference drawn from frequency of word use in any given subject area. LLM's are not capable of “understanding” matters like meaning or conceptual issues like “truth” or “lies” the way they respond to questions is on a “best bet” basis which in turn is dependent on the availability of a considerable volume of data in order to improve the inferential capabilities of the model. Large data sets are therefore essential for the construction of a viable model.

Large Language Models – A Little Theory

At this point, there is a need to understand the nature of large language model architecture, as illustrated in figure 2 and described in detail [here](#) .

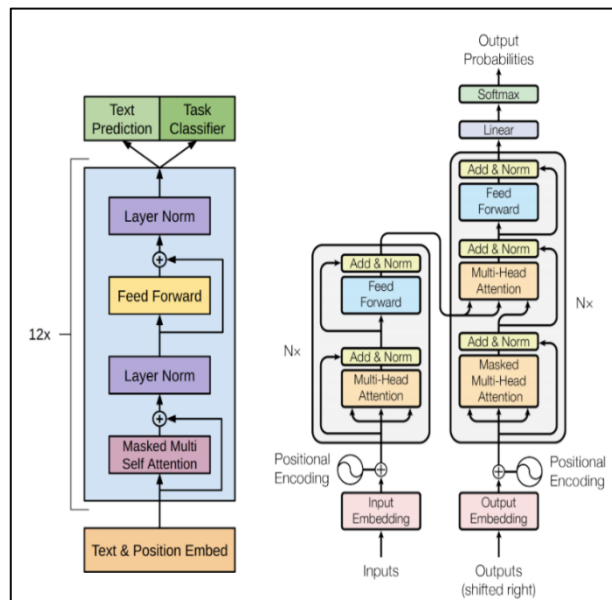


Figure 3 - LLM Architecture

The videos listed below add additional detail to the explanation of the mechanics of LLM's set out in the link in the previous sentence.

- Large Language Models From Scratch, can be viewed [here](#)
- Large Language Models Part 2 can be viewed [here](#)
- A video explaining “Recurrent Neural Networks” on which transformation of language to “tokens” can be viewed [here](#).

There are of course many other videos available on YouTube of a similar nature, the samples listed above are the ones that make the most sense to the author.

A further aid to how LLM models are constructed can be found by searching for “wordnet” and following the links to Princeton University in the US. Once there readers will find details of a series of experiments in the field of LLM development in general and the development of an associated dictionary or lexicon

deserves serious study be found [here](#). As an observation, it is worth noting that there is a major indication of the cost of creating an LLM in that Princeton was obliged to stop work on the programme of work for lack of funds.

The Training Data

What it is hoped has become clear, is that in order for an LLM to work, there is a need to use lots of data. A paper from Cornell University in the US, which can be read [here](#) explains in some detail the need for significant quantities of data for LLM training purposes. A detailed examination of the datasets used for training purposes by OpenAI can be read [here](#). A video that explains how chatGPT is trained can be viewed [here](#).

In the image below, there is a list of the major training datasets used by OpenAI to build their current LLM set up. It is worth taking note that the “[common crawl](#)” programme contains some 13Pb of web content taken from across the world from many, many web sites as a spidering exercise of a similar nature to that executed by Google and other major search engines as a matter of routine. In short, ingesting training data and tokenising it is a major processing exercise. It also worth taking into account that in the case “common crawl” (at least), a decision was taken to stop ingesting data for training purposes in 2021. That ingestion in part or in whole ceased in 2021, suggests that the data used for testing training of the model itself and subsequently deployed in the live platform is neither complete, or current.

GPT-3 Training Data		
Dataset	# Tokens	Weight in Training Mix
Common Crawl	410 billion	60%
WebText2	19 billion	22%
Books1	12 billion	8%
Books2	55 billion	8%
Wikipedia	3 billion	3%

Figure 4 - Training data by token count

The table below gives an indication of the known data volumes of the chatGPT data sets

Ser	Source	Data Footprint	Link
1	Common Crawl	Estimated 13 Pb	Left click here
2	Wikipedia	Estimated 21.23 Gb	Left click here
3	Project Gutenberg On Line Library		Left click here
4	Blibliotik On Line Library Tracker		Left click here

However, storage footprint is not the only measure of “size”. As both sets of statistics listed in the table above indicate, the word count of both is in the billions split or used in millions of articles on the part of Wikipedia and many, many millions of web pages for all of the rest.

Where volumes are not listed, then the issue is that at any one time, no-one can be quite sure what the storage volumes might be. An article [here](#) gives an insight into the nature of the on line Gutenberg library and its management. A further complication in the data sets is how frequently both are updated, in the case of “common crawl” by the crawl itself, in the case of Wikipedia by any number of post editors and new contributors. In respect of currency of content, it is worth considering that all of the data sets listed above are open to random change of unknown reliability and completeness.

It is also worth considering another aspect of currency and its impact, that is how the source data is refreshed and how frequently that is. The reason, the source data refresh rate will drive the need, for no other reason than correctness and integrity of chatGPT responses should be of a high order of accuracy. With that in mind consider Wikipedia.

In respect of its content, Wikipedia is an on line encyclopaedia maintained and otherwise updated by an army of volunteer editors. The Wikipedia article on OpenAI, available [here](#). It provides a very good briefing page on the nature of OpenAI as a company. The revision history page, available [here](#), on reading, note the number of approved changes to just this one article in just a few weeks. Each Wikipedia piece contains a similar history page. In short, Wikipedia, all of it, is open to review and amendment on an entirely random basis. The same kind of change complexity is exhibited by the “common crawl” exercise with the updating of its repository being frequent and while planned as an exercise, is entirely random in respect of the changes that may be made by web masters to their content. The crawl history can be found [here](#).

The same kind of change frequency and randomness probably applies to all of the other data sources used to populate the openGPT model. Given that ingestion is reported to have stopped in 2021, it is safe to assume that much, if not all of the data used to train openGPT is nether complete, or up to date.

A significant scheduling issue therefore, is “time to train”. The estimates of how long it takes to train a model vary. Without the processing capability it has at its disposal, training a model using the volume of data required would be a long affair, in extremis, centuries. However, the nature of the hardware used to train indicates that the training time may be as low as SIRO 30 days.

Stored Data... Once Trained...

Next, consideration must be made, in respect of the nature of the data that chatGPT stores for whatever reason. Firstly, in terms of the data used to provide responses to queries, from OpenAI themselves, the statistics quoted suggest just 570 Gb of data is held for use for enquiry response purposes. A summary of the holding can be read [here](#). While the forint is small, the readers attention is drawn to the number of tokens listed above and the claim that the openAI dictionary contains some 16 billion individual words drawn from multiple languages (some 95 in all). Big numbers.

However, the training data is not the only source of data OpenAI have the opportunity to collect, there is of course the questions and the data end user use to train themselves that are a rich source of training data.

In respect of query responses, OpenAI through chatGPT state that:

“Because chatgpt activity helps millions of people daily in creating a topic or answering questions, the question that many people ask is does chatgpt save data? the user. Here is an answer to this question, from chatgpt itself.

As an AI language model created by OpenAI, ChatGPT does not store or save data from its interactions. The model processes each request and generates a response in real-time, without retaining any information about previous conversations or interactions.

This means that if a user asks a question and then asks the same question again at a later time, the model will not remember the previous answer and will generate a new response based on its current understanding.

In terms of data privacy, this is beneficial as it ensures that any personal or sensitive information shared with ChatGPT will not be stored or used for any future purposes. OpenAI has implemented strict data protection measures to ensure that all user interactions with ChatGPT are secure and protected.

It's important to note that while ChatGPT itself does not save data, the interactions with it may be logged by the platform or service provider hosting it. These logs could potentially be used for purposes such as improving the quality of the model or for monitoring the performance of the service. However, such

logs typically do not include the content of the conversations and are used for aggregate statistical analysis only.

In conclusion, ChatGPT does not save data from its interactions, ensuring that user privacy is protected. However, it's important to be aware that the interactions may be logged by the platform hosting the model, but these logs are typically used for aggregate statistical analysis only."

However, that chatGPT does not store source training data, suggests that in the event that a conversation stops part way through, that there will likely be continuity issues if a conversation is restarted for whatever reason. It would appear too, that there is some conflict of opinion in respect of the on-line documentation in support of chatGPT. The "frequently asked questions" page, stating clearly that conversation data will be recorded for "training" purposes. The FAQ page is available [here](#).

Not Just English..

OpenAI are reported to have trained chatGPT trained on the syntax and grammar of 95 spoken languages. A remarkable achievement in and of itself. However, that suggests 95 variations on a theme given human behaviour which generates things like patois or back slang.

There is also the issue of spelling mistakes, on the face of it a minor matter, but in the authors experience, spelling mistakes and poor grammar are potentially major issues that should be taken into account. As is "patois" and "slang. From experience chatGPT does not do "Scouse", or "Glaswegian".

Retraining

The frequency and random nature of data change and the fact that the main training data set, the "common crawl" is not complete (having stopped at 2021) implies a need for some form of retraining exercise of a regular, planned nature.

Following on from that is the need to plan an update and release strategy as inevitably, the "shape" of a model will change with revision and they will be a real risk of any answer to a query being substantially different from one release to the next.

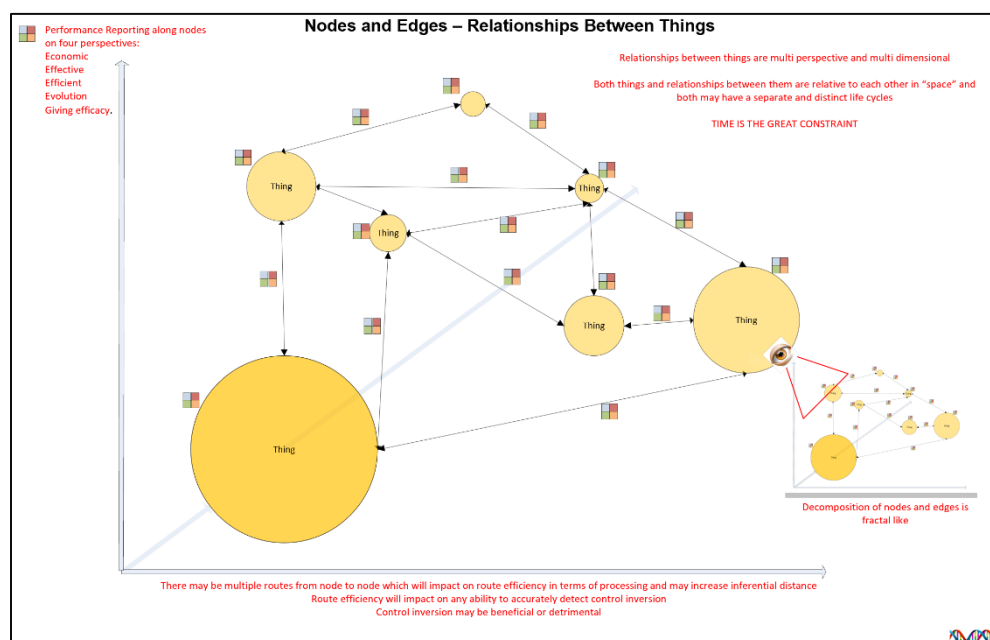


Figure 5 - Multi Dimensional relationships between things.

Figure 5 is an attempt to illustrate the nature of the structural fragility that may occur. Two further issues may arise, firstly referential integrity, as the model changes in shape, then relationships between words

will be broken or enhanced in some way and the nature of inferential distance (how words are connected to each other in a model) will change too. In short, while the aim will be to produce a model that is structurally sound in operation (it has to be) as new relationships between words are introduced, structural continuity between versions will change over time introducing a significant element of fragility over time too. Additionally in terms of the management of processing, issues of route efficiency will arise as more and more relationships and scoring considerations between words are given rise to. An overview of the complexity referred to in this section from Stanford University can be read [here](#). At the conceptual level, a video of a lecture on the nature of the kind of geometry something like an LLM gives rise to can be read [here](#).

In short, the training and retraining effort is a major exercise for any implementation of an LLM, one that has considerable responsibility given the real need to ensure the correctness and relevance of answers such a platform may give.

Managing The Volume Of Enquiries

Firstly, there is the volume of enquiries. Figure 2 presents an estimate that currently, chatGPT has 21 million plus active users. A number that is repeated in more than a few estimates of use. A more detailed set of demographics on how the chatGPT user population might look is available [here](#). The number of individual “conversations” will vary in number and complexity in respect of volume, one estimate put the number of conversations and uses at around 100,000,000 since January 2023. The sudden and rapid increasing in popularity of chatGPT is remarkable in itself. It should be noted that while there is a “free” account, use is metered by some form of token count, once the limit allocated by OpenAI is reached, then further usage is chargeable. OpenAI publish a warning to that effect [here](#). In terms of volume of traffic, chatGPT must be supported by a robust communications infrastructure.

Enquiry Response Effort

Finally, there is the need to provide to its end users a speedy response to all enquiries put to it. Given the volume of enquiries and issues like route efficiency to traverse the model on the basis of probability testing to try and make sure responses are contextually coherent, then enquiry processing is an intensive exercise of great complexity and demands the availability and use of very, very high specification machines hence the reason for Microsoft's investment in what is reported to be the 5th most powerful supercomputer in the world as at 2023 and that it has taken [some time](#) to build it. A discussion of the response processing requirements is available [here](#).

Short Term Memory Issues

[Giving Auto-GPT Long-Term Memory with Weaviate | Weaviate - vector database](#)

[Linked In Feed](#)

Management of the Exercise

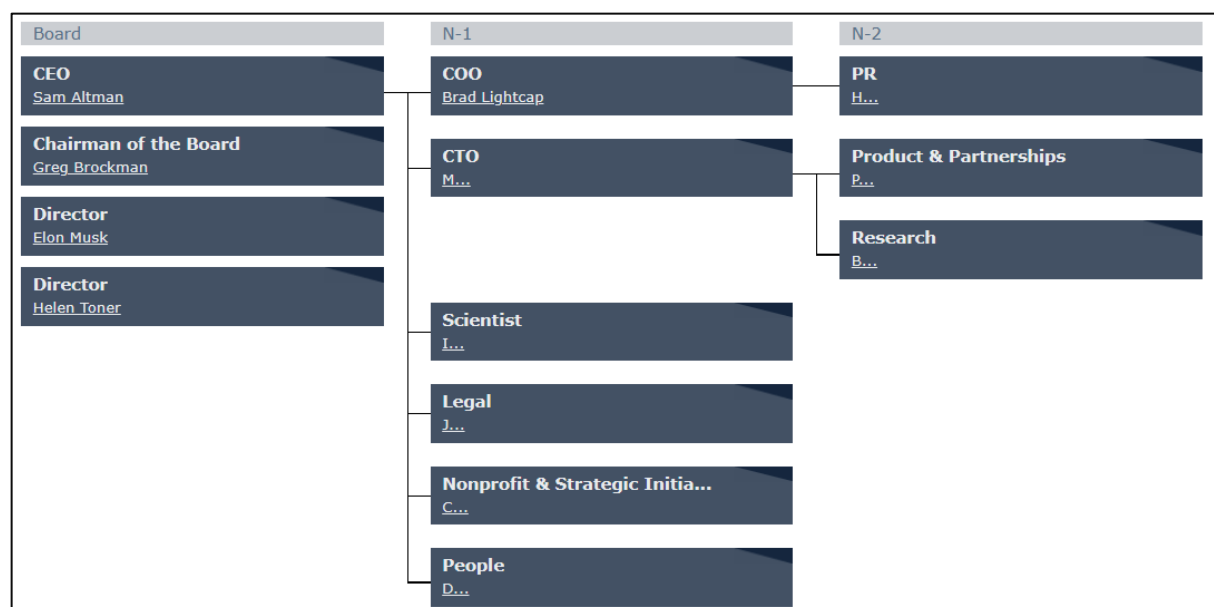


Figure 6 - OpenAI Management Structure

As indicated by the processing requirement, the nature of chatGPT requires highly sophisticated hardware resources and of course personnel. Figure 4 gives some indication of the way the company is organised. Figure 1 gives an indication of the estimated number of employees OpenAI has on its payroll. The number of employees varies depending on review but each staffing figure examined does not go above 375 people.

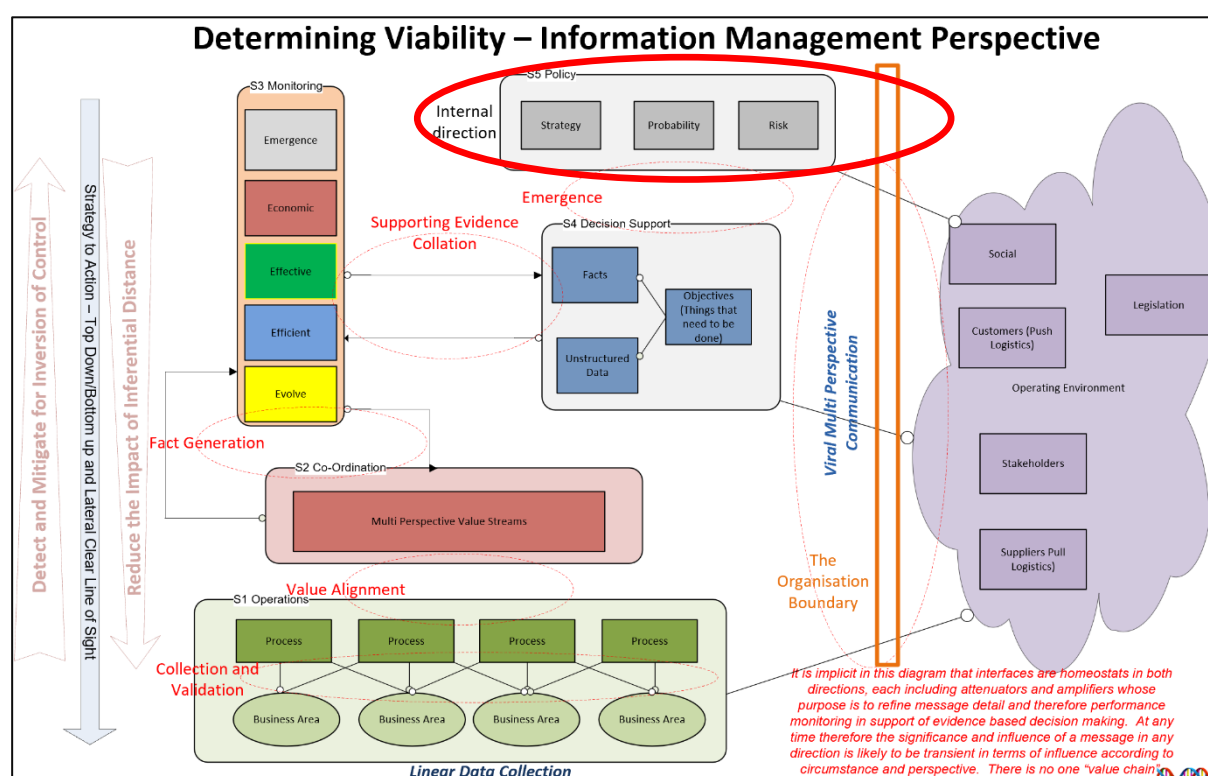


Figure 7 – It's not just "the model".

Given that clearly, the company has all of the staff management issues other companies have, as a ballpark figure it may be the case that some 20% of staff are employed on administrative and other support activities. It is also reasonable to assume that at any one time, up to 30% of staff are away from work on leave, training and any other number of other activities. The remaining 50% of available staff at any one time being involved in development and maintenance of the chatGPT product set. Bearing in mind that OpenAI has a product portfolio, each of which will have a dedicated support team.

As it is hoped is becoming apparent, chatGPT involves a major transformation exercise in respect of its model building. The end result, many millions of "tokens" and a dictionary of many millions of words, each word scored such that the probability of its appearance in any given sentence, written in response to any enquiry, can be answered in a human like way. However, that transformation exercise, architecturally speaking is one of several that are involved in almost any information management exercise. Figure 5 indicates the nature of the transition and attenuation exercises that may come in to play. The red ellipse indicates, in the authors view, the position of chatGPT as a function of dealing

with emergence. Below it are several other transitions for which the chatGPT model is dependent for its accuracy. The table below sets put, in overview, what each transition might be,

Ser	Transition Stage	Purpose
1	Validate	Data collection. All data collection should be validated for reasons of integrity of evidence further up the reporting chain
2	Alignment	The combination of multiples of data sets which will almost inevitably involve multiple forms of attenuation given that different data sets may contain the same values, but named differently and more besides.
3	Fact Generation	The generation of hard evidence to support operational reporting requirements that may impact on effectiveness and efficiency of the “organisation as a system”
4	Evidence Collation	The combination, on an architecturally sound basis of both structured and unstructured data.
5	Emergence	<i>A transition from deterministic forms of analysis to those of a “Bayesian” nature involving techniques like probability and regression testing.</i>
6	Inform	A two way transition relating to extracting data from the outside world where possible and presenting validated and verifiable information back into the outside world if the need arises.

Briefly, OpenAI and its operations are entirely dependent on the ability of many, many others (in respect of the common crawl as many web site operators of sites it spiders) to maintain in integrity of the source data it bases its modelling effort on.

The reason for raising the nature of staffing , transition and attenuation of the OpenAI source data is that while no doubt their technical teams are highly qualified and very competent, given the scale and scope of what is being attempted and associated modelling activities like ingestion, testing and so on, it is likely that staff are hard pressed. Furthermore, the complications of the management of source data, entirely outside of the control of OpenAI, suggests there is a considerable dependency, not to mention trust, on the accuracy of the work of many, many people.

Inevitably, there would have been a need for OpenAI to seek a Microsoft or a Google to take the chatGPT further. Or indeed to maintain the stability of the product architecture. Complicated software of this kind and its ongoing development are resource intensive.

OK, So what Does This Big Computer Look Like?

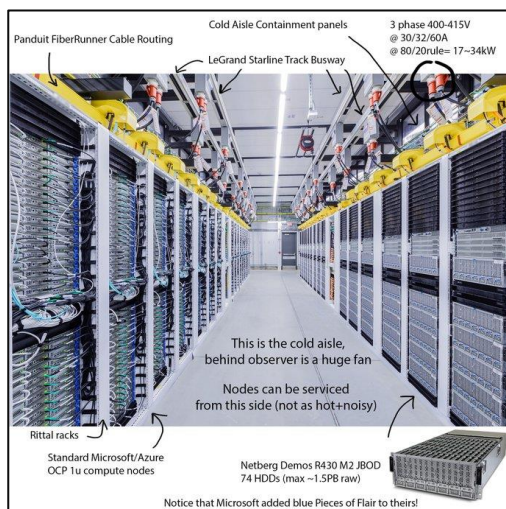


Figure 9 - An indication of physical architecture

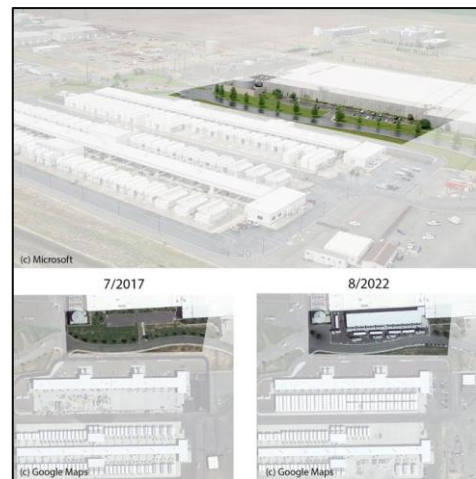


Figure 8 - And indication of physical size

There have been a number of suggestions to the effect that it is possible to “run” chatGPT et al from a laptop. That is more than open to debate. LLM’s in particular and neural networks generally are processor intensive, they require very high specification machines to support them and it is prudent therefore to get some understanding of what the machines concerned might look like and the kind of facilities needed to support them.

Note that Microsoft commissioned a machine, specific to the needs of the OpenAI in 2010. Readers can review the announcement [here](#). It would appear that Microsoft have been planning the GPT series of exercises for quite some time, far longer than many of those that have used it and claim expertise given that chatGPT rose to prominence, like a shooting star, to great excitement in January 2023.

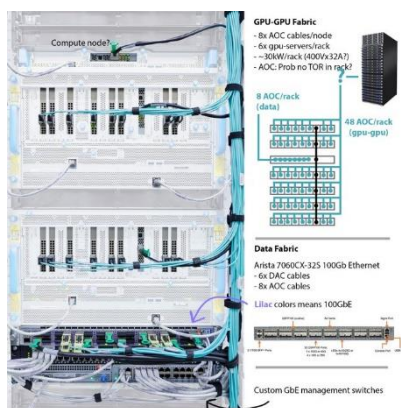


Figure 10 - A bit more detail

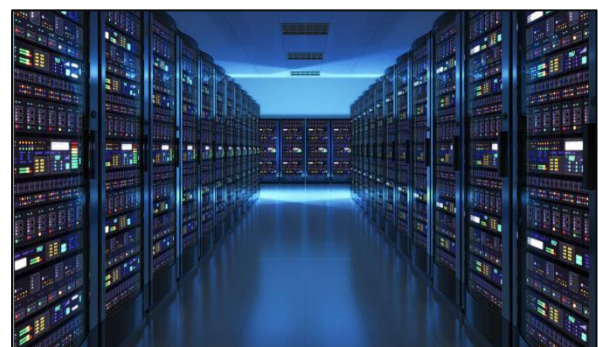


Figure 11 - Slick front end view

The level of investment in this exercise is, it is suggested, significant and the idea of being able to “run” this kind of tool from a laptop, even one with the highest level of specification, may be a tad optimistic. It follows therefore that these kinds of large computers are not “stand alone” and are supported by the kind of infrastructure illustrated above. But that raises another set of issues related to resourcing power and other utility support, which prompts the question where does the supercomputer and its supporting infrastructure sit in the world?

Where Is It?

Quite by chance, the author came across a piece mentioning the town of Quincy in the US State of Washington. The image below is a satellite picture of the location of the town and its close proximity to two very large hydro electric power generation facilities in the US.

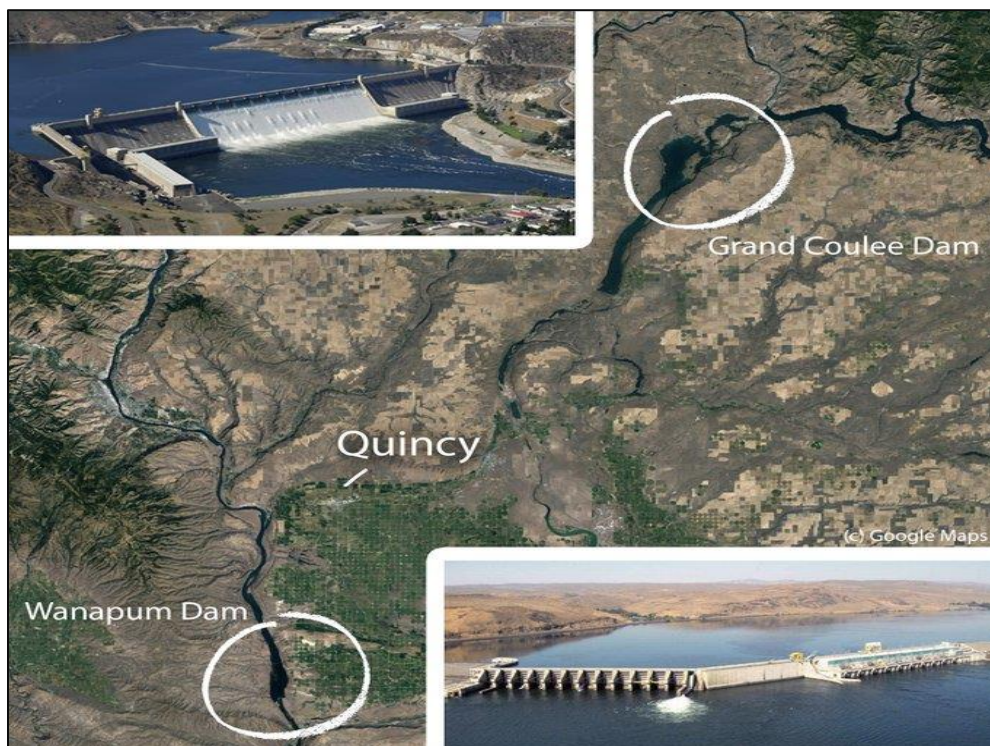


Figure 12 - Where is chatGPT located?

The area of Quincy is very attractive to owners and operators of large data centres as can be seen in figure 13.



Figure 13 - Not just Microsoft

That of course raises issues of any potential environmental impact which, with increasing frequency is being raised as a matter of concern as indicated by the document available [here](#). Needless to say, the impact of carbon emissions too is coming under review, a sample paper can be read [here](#).

It should be noted to that the OpenAI headquarters office, the Pioneer building, 3180 18th St, San Francisco, United States, is a separate and distinct facility, with the OpenAI team of 375 being supported by significant effort on the part of Microsoft.



Figure 14 - OpenAI HQ.

Which, of course, raises issues of which legal jurisdiction does the OpenAI LLM enterprise, as a whole, come under.

Legal and Compliance Issues?

The principles on which data and privacy protection are not new. One of the original expressions of operating concepts can be read [here](#).

What follows is not, by any stretch of the imagination a detailed view of the nature of legislation that may apply, the following paragraphs are raised for discussion. Above all, when considering any legislative issues, the fact that the OpenAI service is based in the US and therefore subject to US law should not be overlooked. The table below being a very small sample of the Federal law that may apply. Nor should the Sta laws of both Washington State and California be ignored, the Californian Consumer Privacy Act ([CCPA](#)) being model piece of legislation for many other US States. The state of Washington being one of a growing number of US states launching its own privacy legislation which can be read [here](#).

Ser	Legislation	Document
1	The Cloud Act	Left click here
2	FISA	Left click here
3	Sarbanes Oxley	Left click here

The US CIO has published a comprehensive guide to US Information technology law which can be found [here](#).

Nor does the compliance complication end with “just” US regulation/legislation. With an estimated 100 million people and organisations using the “free” chatGPT version and an estimated 21 million people as paid subscribers, it is the case that as a tool used internationally, it is not just US law that applies to the use of the GPT series of tools.

The table below, is a list of key EU legislation that may impact on the way LLM's like chatGPT may be governed and policed. Guidance should be sought on the applicability of each, taking into account that it is no longer just the GDPR that matters with each of the 7 Acts listed covering various aspects of IT operations.

Ser	Legislation	Document
1	EU AI Act	Left click here
2	EU AI Liability Directive	Left click here
3	EU Digital Markets Act	Left click here
4	EU Digital Services Act	Left click here
5	EU E Privacy Directive	Left click here
6	EU GDPR	Left click here
7	EU NIS 2	Left click here

The GDPR perhaps being the most well-known, but the others will in due course have an impact on the use of LLM's. For a discussion on the nature of the complex legal risks that may be involved, the readers attention is directed [here](#).

But the legislation, just in the EU, that may also apply includes the various acts and directives listed below.

1. The NIS 2 Directive
2. The European Cyber Resilience Act
3. The Digital Operational Resilience Act (DORA)
4. The Critical Entities Resilience Directive (CER)
5. The Digital Services Act (DSA)
6. The Digital Markets Act (DMA)
7. The European Health Data Space (EHDS)
8. The European Chips Act
9. The European Data Act
10. European Data Governance Act (DGA)
11. The Artificial Intelligence Act
12. The European ePrivacy Regulation
13. The European Cyber Defence Policy
14. The Strategic Compass of the European Union
15. The EU Cyber Diplomacy Toolbox

With each of them, both in the US and EU, being policed by a number (which seems to be growing) of authorities. The US FTC having received a complaint about the use of the chatGPT product, the details of which can be read [here](#), and moves by the Italian Data Protection Authorities to ban the use of chatGPT as reported by the UK BBC, which can be read [here](#). In the authors view, the regulatory actions being taken as of 31/03/2023 are the first of many to come.

Then of course, there is the rest of the world, with many countries establishing various forms of regulation in the field of IT. In short, get involved in IT, in any multinational effort, and multinational jurisdiction's apply. Users of the OpenAI tools have a real need to understand international law.

Defamation, Intellectual Property, Copyright and More...

While IT law is complex, in respect of the way the OpenAI tools appear to be being used, arguably generalist and akin to the "Babel Fish" of Douglas Adams fame, it is not just IT law that may apply and may be pursued in court, pretty much anywhere really. Consider..

At least one of the training datasets used by OpenAI for training chatGPT comes from web sites owned and operated by both private individuals and corporate bodies. Searching the internet for guidance on the nature of copyright protection for such content. The following articles and advice indicate the nature of copyright protection of web content:

- Copyright – A web developers guide available [here](#).
- Copyright – A legal summary from a leading legal practice in the UK is available [here](#).

- A reported test case, can be read about [here](#).
- Samsung, the Korean Telecoms giant has already fallen foul of copyright issues which can be read about [here](#).
- Defamation. While copyright is significant, given that an LLM is at the final stage in response to queries, a matter of probability testing, it is the case that quite frequently it gets things wrong. One such example is a case, raised in Australia related to possible defamation which can be read about [here](#).

There is one more consideration in respect of the data used to train the product and that is that the original data is disconnected from the model once it has been trained. That suggests that while the use of words, expressed as tokens, may be a viable basis on which to construct realistic sentences based on word value as a matter of grammatical context, there is no concept of “meaning” in respect of the way words are used, no conception of context in the nature of the output. There is therefore, still a need for end users to validate and verify output, to determine its accuracy. It is suggested that the disconnect set out here is why there are many reports of issues like inaccurate accreditation of key reference documents or even the mentioning of “reference” material that simply does not exist. There are of course other issues related to context which might cause concern, amongst other things, responses to enquiries, placed by individuals, have included declarations that the enquirer has died. An example is available [here](#).

In short, the legal situation surrounding the OpenAI product set is open to question. is open to interpretation and legal advice might be a prudent consideration given that it may well be the case that a web site operated by the reader may well have been used to train chatGPT.

OpenAI Terms and Conditions

Of course, as the issues set out above are raised, OpenAI, rightly acts to address concerns, firstly to reassure and secondly to reduce their liability risk.

The text below, warning of the nature of business risk associated with the Open AI terms and conditions, or indeed any other terms and conditions come to that), was posted on LinkedIn by the author of this document is reproduced below (the full text of the original post can be read [here](#)):

“In the mean time, an old soldiers tale, a while ago, I took part in a due diligence exercise reviewing the licence terms of a number of platforms, each with a high six figure annual maintenance charge. I built a model forecasting licence fee accrual over a 5 year period. The model was not far out when, a few years from construction, the license auditors came to call. The hit was a lot of money. I have to say that one of the things I learned was that the majors, GAFAM's if you like, were not in the least bit phased about demanding a license audit should they see the need. So.. A warning.

The kerfuffle about OpenAI et al started in January. Many may know that MS has commissioned a new "supercomputer" to support its LLM initiative. Indeed, under the banner GPT 4, the next iteration, the complexity of the processing demands such a machine. Have a look at this:

"Microsoft's New Super Computer"

Note the date of publication (tiny bit of text, top left of the screen). It kind of implies Microsoft, quite rightly, have been planning this for some time now. Over the years I have been in IT, one of the things that has impressed me about MS, is that nothing they do is casual. The occasional FUBAR, but nevertheless..

Which brings me to a warning. While MS have a plan, I suspect many playing with chatGPT do not. Fair enough, its nice to play, but nevertheless, the planning capability the article is about is what you are really up against.

Your next step, should be to plan too. Surely? As a starter for 10, my experience tells me you should

all be reading the OpenAI T&C without going much further. They can be found [here](#).

"Open AI Terms of Use"

I note the legal community seem to think that they could possibly use it to help draft documents. There are others doing similar things for their professional circumstances. Me, I think that is all bonkers. Why?

Read the T&C. Just para's 7 and 8 for the moment, but the rest, soon after. But do so carefully. Line by line. Take into account the jurisdiction OpenAI must abide by. Think it through. In her book, Zuboff, rightly in my view, refers to these kinds of terms as "sadistic". Note too, that given the working relationship OpenAI have with MS, then the MS terms will come in to play too. Think about the implications for your business intel..."

While the original intent of the OpenAI experiment may have been to make the internet easier to use, inevitably, with the costs of something like an OpenAI project being very expensive, there is a need to fund day to day operations, in short, OpenAI is a business not a benevolent institution.

In the same thread, the following advice is offered related to business case preparation:

"PS. You should also read carefully how they intend to charge the service out. I read it all, given that it is to be hosted on Azure.. I had a wry smile to myself.

in my day, part of the prep for adoption would be to produce a business case, including a financial forecast of cost and ROI. quite fiddly that. Then of course it had to be approved which in my world could be traumatic because I had to prove everything I wrote.

For this particular application, which IS processor intensive in a way few other things are, running it to "train" let alone anything else, may turn out to be an expensive business. Furthermore, you will find, if you try to work out what the cost of using chatGPT is, it is difficult business because one of the charging mechanisms is based on "tokens" and their use. Which kind of implies you know what "tokenisation" refers to in an LLM world (please note, it is not at all the same as other token forms you see being talked about) and how many your use might need....

Geeks eh?

You have been warned."

Bear in mind too, the framing of the terms and conditions are a fast-changing thing. The original and noble intent was to release the technology under the "open source" banner. "open source" does not mean free of responsibility or liability however and it is typical of the IT industry and the majors in particular that they seek to reduce their liabilities through the medium of T&C which are in fact a matter of contract. Not unexpectedly, the GPT series of tools have been refocused away from "open source" and are now commercial, proprietary technology and article explaining the move can be read [here](#).

Two pieces in support of the need to review, carefully, the relevant terms of business can be read [here](#). And a slide deck, available [here](#), that reads may also wish to consider.

And finally, as the hype subsides and more consider reviews take place, the kinds of accumulated legal risk touched on here are giving rise to concern from insurers, with recommendations that the use of OpenAI tools may constitute an unacceptable professional indemnity risk. And example of which can be read [here](#).

Some warnings:

- [Professor Says High Risk of GPT-4 Being Used for 'Dangerous Chemistry' \(businessinsider.com\)](#)
- [ChatGPT and LLMs: what's the risk - NCSC.GOV.UK](#)
- [Artificial intelligence | NIST](#)

- [Brace Yourself for a Tidal Wave of ChatGPT Email Scams | WIRED](#)
- [OpenAI threatened with landmark defamation lawsuit over ChatGPT false claims \[Updated\] | Ars Technica](#)
- [Samsung workers made a major error by using ChatGPT | TechRadar](#)

“Caveat Emptor” applies with a vengeance.

Other Risks and Issues....

Finally, it should be noted that there have been concerns and issues raised in respect of the exploitation of chatGPT for nefarious purposes of a criminal and unethical kind. One authoritative paper on this issue can be found [here](#).

“Side Issues” Galore!

Ser	Subject	Link
1	Snapchat Scandal	Left click here
2	OpenAI and Bias	Left click here
3	Cyber Security Risks	Left click here
4	Open AI CTO Interview	Left click here

The readers attention is brought to the table above. Each item in the list presents commentary on the numerous ways chatGPT are open to abuse. The last is an interview given by the CTO of OpenAI that explains that the chances of abuse are not small and that “society should adapt to software of this kind”.

The author is of the opposite opinion. It is a responsibility both legally and perhaps more importantly ethically, to ensure that what is delivered into the open market is fit for purpose and safe to use. Read each of these and draw your own conclusions, but in the authors view, each represents considerable risk to the integrity of reader own information management efforts.

Site Profile OpenAI.Com

Ser	Test	Result
1	BuiltWith	Left click here
2	URL Scan IO	Left click here
3	Security Headers	Left click here
4	DNS Visibility	Left click here
5	IP Location	Left click here

The table above presents links to several tools, the author recommends the use of, the purpose of which is to test the client-side impact of a web site or page and to get an idea of the nature of the security and information management state of a web site. The tools listed are part of an extensive client-side impact review template which can be read [here](#). Annex A of the template contains the details of the toolkit itself, which lists 24 tools, free to use, that people may use to execute such a review themselves. Each tool is simple to operate. Operation is a matter of cutting and pasting a domain name into a text box clicking a single search button and waiting a few minutes for the results. The of course to review them.

The five links in the table above point directly to the review results for the OpenAI web site.

Why do the checks?

If a site host is in the US and your company is in the EU, chances are it will fall foul of EU judgements like Schrems 2 if data is collected by components from third parties, FROM YOUR DEVICE OR YOUR CUSTOMERS DEVICES. The DNS and security headers in order to get a general idea of site safety

for your end users. The other two to determine what a web page is made up of and the order they are requested from wherever they are requested from. The I/O list, it represents multiples of conversations between your site and your visitors you are not privy to as a site operator.

In the case of the OpenAI web site, given that there are tracking components installed, that deliver their payload client side, it is likely that other organisations are aware of a visitor presence on their site before Open AI might be. For one component in particular, Google Analytics, four EU Data Protection Authorities have implemented a ban. A detailed description of why the ban is inevitable, given the terms of the EU ECJ judgement known as “Schrems 2”, can be read [here](#).

Given that in order to make use of the tool for training models for corporate use, requires significant amounts of data, reviewing an company web site, where a complex data transfer is envisaged, is prudent and recommended.

B i a s

One of the major issues artificial intelligence applications of all kinds, have an issue with is the maintenance of objectivity, bias being the manifestation of a loss of that key element of information integrity. Some time ago, the author put together a briefing note and bias in relation to the use of facial recognition software which can be read [here](#). One of the points the briefing note raises is that at the time of writing, facial recognition software is in the IT equivalent of a “tank/anti tank” arms race as attempts are made to thwart such tools one way or another. The bias risk is slightly different in LLM’s as the CEO of OpenAI indicates in a recent interview available [here](#). For whatever reason, OpenAI built in a form of bias with the aim of reducing the risk of skewing responses on the basis of political persuasion. Needless to say that has raised concerns.

That aside, bias starts at the point of data collection. Those who collect data do so for their own reasons, which can be almost anything. There are limits of observation in the scope of data collection efforts for no other reason than nobody can see everything to record it. That means that bias starts much earlier than the point of data transfer into something like an LLM.

The problem for LLM’s in respect of bias and its detection, is that in order to train, there is a need for vast quantities of data and, as time goes by, to repeat such exercises periodically and retrain the model. The need to retrain raising the following issues:

- It is likely that with retraining, the number of tokens will change
- With the changes in token count will come changes to the number of probability tests that may need to be executed for sentence construction purposes
- With that will come changes to the wording of answers to questions repeated to, say, confirm an answer from a previous conversation.

Given that the number and frequency of changes to training data, which are also entirely random and cannot be controlled by those training models, it may well be the case that bias is in the eye of the beholder and while OpenAI may be able to reduce the impact of bias that they have introduced, the company has no means whatsoever to reduce external impact for no other reason than they do not control, in any meaningful way, the completeness of the training data.

Bias therefore becomes endemic.

C o d e G e n e r a t i o n

A further issue is that chatGPT has been trained on some computer programming languages and how they can be used. A description of the kinds of risk this capability offers can be read [here](#). From the authors perspective, one of the more dangerous examples of the art of the possible is described [here](#). The key acronym in the example is “DLL” which refers to a form of compiled code used extensively in the Microsoft Windows world.

Of all of the issues chatGPT has generated, in the authors view it is code generation that should be causing concern in the wider information management world. The example above, focussing on the

generation of code that can be used to create a DLL, being of particular concern an example of something similar in nature, but extremely sophisticated, is the “Solarwinds” incursion which can be read about [here](#). It would appear that one of the unintended consequences of chatGPT has that it seems to have democratised the means to create sophisticated code in that in order to generate the code, all that needs to happen is that those needing such a capability need to ask the right kind of question and chatGPT does the hard work in seconds.

Study, carefully, what the researchers in this field of activity are doing, the democratisation of hacking is, not maybe, is a major security risk all by itself.

Tricksters! How the hell did that happen?

Unfortunately, because chatGPT gives the illusion of human like response, there are words used to describe the nature of its responses that imply sentience. Disturbingly, there have recently been posts claiming the “singularity” is at hand. On the contrary, chatGPT has no concept of what lying is or truth. What is happening once data has been ingested and “tokenised” or “scored” is highly sophisticated probability testing with the aim that words combined into sentences at the point of interaction are as good a match as possible. The abuses of OpenAI intent with chatGPT, all very laudable, are becoming more and more frequent and more and more ingenious. Some of the methods identified to, in effect, “trick” chatGPT (it has no concept of subterfuge and how to deal with it either), can be found [here](#).

Dependency

Finally, a sub plot of any software use is the nature of dependency of the end user on this or that product. The dependency starts the moment a piece of software is used for the first time and rises with use. Typically, there are two reasons why the dependency is exposed:

- Termination of contract, for whatever reason
- A Mergers and Acquisition exercise of some sort in which multiple software portfolios of software need to be integrated for whatever reason.

It is at the point where either/or raise their heads that those wishing to move platform or integrate platforms find out just what they own. Which, usually, is very little.

Any selection of any product, including AI, is an architectural issue and should be treated as such. Arguably, unlike many dependency issues, the use of the kind of tool chatGPT represents is of particular significance because the purpose of such tools is to exploit data more effectively for business intelligence purposes and end users do not own the output. Instead Microsoft and OpenAI reserve the right to use any experiments for their own purposes.

Are Large Language Models Intelligent?

The hype surrounding chatGPT in particular has been a surprise, to the author at least. Some of the claims related to the GPT series being more than a little curious many raising the idea that the “singularity” will be reached or that there are “sparks” of an early intelligence exhibited by LLM’s. in particular. Such claims are, in the authors view, open to debate, but as with so much in information technology, it makes sense to do a bit of reading on the nature of intelligence, certainly that of machine learning. As door openers, the two links below may be useful.

[System 1 and System 2 Thinking - The Decision Lab](#)

[Narrow AI vs General AI vs Super AI \(whyofai.com\)](#)

As to the nature of the human brain (bearing in mind it is Mk1 Human beings who invent tools like an LLM, it might also be sensible to learn something of how the human brain might work. The following books are highly recommended:

[“Master and His Emissary”](#)

[“The Matter With Things”](#)

All by the renowned psychiatrist Iain McGhilchrist. Both titles requiring careful study, “The Matter With Things” in particular being comprehensive in their detail and well worth the study.

The authors response to the question prompting this section is a resounding “No”.

What Next?

chatGPT has generated a considerable amount of interest since January 2023 however, at the moment, its focus is purely on text and its manipulation. The next step under the banner “[multi modal Large Language Models](#)” with the aim of introducing a greater degree of “[perception](#)” into the OpenAI product portfolio has been announced under the banner “[GPT 4](#)”.

However, it should be noted that while there is great excitement at GPT – 4 reaching the level of the second coming, a note of caution from the CEO of OpenAI himself can be read [here](#).

It is worth considering the potential impact of the introduction of future updates as each will mean retraining of some kind. It is likely that in retraining, more tokens will be generated as more relationships between words are identified in multiple languages. As the token number increases, at the point of end user interaction it is likely that there may be degradation of service over time given that multiples of users will “ask” chatGPT things in multiples of human spoken languages. As a consequence, deploying LLM’s may well be a complicated business that by itself requires some considered and [nuanced planning](#).

With that in mind and it is suggested a growing awareness of the nature of legal risk, the CEO of OpenAI has released a couple of interviews suggesting the need for a pause. [This](#) on a delay to the release of GPT 5 and a further interview suggesting that the further development of LLM’s may be a redundant idea and that they have run their course which can be read [here](#).

A Bit of An Assessment of Risk

One of the capabilities on offer is for third parties to train the LLM engine with the aim of building their own models. The product has an API which can be read about [here](#). However, when considering training a model, there are a number of issues that it would be prudent to examine:

- Pricing – the openAI price structure can be read [here](#).
- Volume of data. Bear in mind that chatGPT required huge volumes of data to train it to the position it is in.
- Legal sensitivity. The chatGPT platform is a third party, furthermore a third party that is subject to US legislation.
- Commercial Sensitivity. In respect of that there is a clear need to understand the nature of any agreement signed up to use the chatGPT in the way intended in respect of the use of models to improve the overarching capabilities of the product.
- How long it takes to train a local model and how frequently it may need retraining.

Over the past few weeks, while looking into how much it costs to train an LLM “from the ground up” has revealed some very large numbers. The estimate contained [here](#) at US \$5 million being typical. It should be born in mind too that to train a model requires huge volumes of data as the training data figures illustrated on page 9 indicate. It should be noted that given nothing is static, that it is likely that over time, a model will become increasingly irrelevant and there will be a need to retrain. Retraining will need to be planned, but it is worth considering that the training cost will be a recurring cost.

Then of course, there is the need for a dedicated and high specification device capable of supporting the kind of complex processing an LLM requires. #Key of course, in respect of determining if such an exercise is viable, is the return on investment.

The calculations related to pricing, through life, are almost certainly complex and in any event are difficult to do unless there is some concept of the nature of “transformation”, what a “token” is and the nature of the technology architecture or stack as it drives the specification of requirement for a machine capable of meeting the need.

Tread carefully. Processing of the nature indicated by LLM is expensive. Those costs will be passed on.

Of course, it is more than possible to build your own Large Language Model as indicated by the explanation of how to go about it that can be read [here](#). A paper on the nature of the optimal training effort can be read [here](#) with a supporting article [here](#) and [here](#) on LLM training “laws”.

However, it is suggested that contemplating such an exercise should give rise to determining its viability in respect of the delivery of any business advantage with the sample questions below being indicative of the kind that should be asked before contemplating “having a go”.

- Do we have enough data to train a viable model (bearing in mind just one of OpenAI's data sets is 13Pb)?
- Do we have a viable reason to put in the training effort?
- How much will the training effort cost?
- What is the commercial risk (given that in training, the makers of these things reserve the right to use the training effort of others to their benefit)?
- How frequently will the initial model need to be updated and what are the associated retraining costs?
- Do we have anyone (by that it will mean more than one person) who even knows how to train in the most efficient way in respect of achieving the business aim?
- Once trained, what is the nature of the associated effort along mutually supportive lines of development (TEPIDOIL for example)?
- How is ROI determined?
- How do we test the model for viability as a business asset?
- How much does using the model cost?
- Are there any legislative compliance implications?

And so on. In short, executing an investment appraisal instead of JFDI.

But there are other considerations in respect of how to approach the kind of problem an LLM seeks to address. Consider for a moment that the average corporate vocabulary is approximately 140k words. The number of documents held may vary of course, but the size of the vocabulary should be a determining factor in assessing whether or not the kind of modelling technique that an LLM applies is worth the effort. In the past, the author built a corporate dictionary, supported by a document “spider” the aim of which was to catalogue document content in a way that established a relationship between words, appearance in documents and the location of document files in a formal folder structure. The nature of the exercise can be read about [here](#).

In short, there are options. 7P's apply, with a vengeance.

Anything Else To Raise?

Some time ago, to great fanfare and excitement, Google announced an AI product called "[Tensorflow](#)". There was, as ever, great excitement in the world of IT. However, the excitement soon wore off, a recent review of existing machine learning platforms seems to indicate that Tensorflow has fallen out of flavour and can be read [here](#). Readers may be aware of the IBM machine known as "Watson" which drew similar excitement, Watson and its history can be read about [here](#). Watson too has not achieved the status claimed for it. Given that chatGPT is already raising concerns and its makers are seeking some understanding of its shortcomings, it may well be that it goes the same way as its illustrious predecessors.

Take the time to read [this](#). With someone in the AI crowd claiming "it's social media". It seems that with the best ethical will in the world, others, of more evil intent are exploiting OpenAI in ways that are morally and ethically bankrupt. It would seem the concept of "guardrails" and their effectiveness is somewhat suspect.

And with that in mind, readers should note that a number of major players in the AI field have published an open letter advising a pause in tool development of AI tools in order to work out what the impact of the unrestrained use of such tools might be. The letter can be read [here](#).

US White House AI Bill of Rights available [here](#)

As an observation, it would appear that increasingly, something like an LLM that seeks to be generalist in nature across multiple languages and jurisdictions might have been a mistake. Far more sensible, from the authors viewpoint, to specialise on a specific subject area for which there is high quality high volume subject matter that is freely available. It would appear that with the launch of the computer security co-pilot announced by Microsoft, that the need to specialise has been noted and acted on. One of the major factors that render such a thing as more practicable is that Microsoft have considerable expertise in this field themselves with supporting data and documentation they own.

Conclusions

The aim of this document is to address something the author sees as a major problem in respect of the great excitement chatGPT has generated. Disturbingly, recently there have even been claims that the "[singularity](#)" has come a big step closer. In the authors view, that is unlikely, however the evidence to argue that is spread to the four winds electronically speaking and there is a need to collate sufficient evidence to bring about informed decision making, which is what this document tries to do. That there is a need for an exercise of this kind is indicated by post on LinkedIn available [here](#) which has seen more than 120k views in a very short space of time that illustrates the idea that many, many people do what to know more about how LLM's in general and chatGPT in particular, work.

By any standard, what has been achieved with chatGPT is remarkable, it can indeed mimic human response to enquiries in a manner that on first glance, reflects human language and how it is used in day to day conversation. However, on closer inspection, things are not as precise and accurate as it would seem. The nature of risk in its use is becoming apparent, from Noam Chomsky's "plagiarism" remarks to the recent speech by the OpenAI [CTO](#) warning of the kinds of limitations the product has.

The author of this document has more than a few reservations on what chatGPT is capable of and its limitations, the more he has looked (written about above), the more one of the most appropriate comments, made many, many years ago by the US comedian Bob Newhart called "[An infinite number of monkeys](#)" bears listening to as a matter of caution.

In the authors view, that OpenAI have tried to be something akin to a "Babelfish" of "Hitchhikers Guide to the Galaxy" fame, it might have been more sensible to have released something into the wild that demonstrate, in the authors view, considerable immaturity and structural fragility. In short, this tool has no concept of "truth" or "lies" and is ethically agnostic. What it does is to score words in respect of syntactical and grammatical context to determine on the balance of probability which word will come next in any given sentence it constructs as output to any question or request put to it. However, if Microsoft achieve their aim to use an LLM as a basis for search enhancement, then that will be a game changer. In the authors view there are some major hurdles to overcome, in particular being in a position

to properly verify and validate the training data. Once and if that is achieved however, then Microsoft will have an extremely powerful capability.

Nevertheless, at the time of writing, chatGPT, it is not, in any meaningful way, intelligent in a human sense. Perhaps the greatest danger it represents is that an enormous amount of people seem to think it is.

Recommendations

The key recommendation? Read and study the links and their content if considering using chatGPT for anything really. At the moment the focus of the tool is the replication of sentence and paragraph construction as a pre-requisite to maintaining conversation of a human kind.

In the EU GDPR, its article [17](#) (and [22](#)) place some very tight technical requirements related to the protection of the right to privacy of the individual in respect of the nature of automated processing and the right to erasure. Both are key articles in respect of information and data management. Associated with them is the concept of a "Data Subject Access Request" (DSAR) based on the idea that an end user can ask to see any data held on an individual and request removal.

On linked in, the "[Nightmare Letter](#)" was offered as a means of testing the ability of an organisation to meet some key compliance requirements associated with a DSAR that is recommended to be used as the basis for a training exercise related to demonstrating and proving compliance with GDPR like regulation.

The author is not at all sure that the GPT series of tools can meet the requirements that the Nightmare letter implies not least because OpenAI have no ownership rights or control of the training data as it currently stands. For reasons of legislative compliance, as an acceptance into service test, a similar exercise would be strongly recommended.

The name Temnit Gebhru should be one to note in the explanation of how LLM's are constructed and the implications of their use. She was sacked by Google for raising concerns about the technology and published a paper on the risk. An MIT review of the paper can be found [here](#). It is that rare thing, an authoritative insiders view of what the issues LLM's generate might be.

If asked about a recommendation when working on whether or not to adapt chatGPT for operational use, regrettably, the recommendation would be not to use it at all, but to monitor progress and maturity as it develops. The reliability of the tool is open to question as its ability to meet the compliance requirements of legislation in the EU and elsewhere.

In order to use it, there will be a need for significant amounts of data for training purposes, the use of which will, in the authors view, represent significant legal and contractual issues that at the moment cannot be accurately worked out not least because very few people will have any idea to work out a costing plan given that competitively few people understand the technology and its costing model.

There is also a clear need for an independent and objective assessment of AI generally, in the authors view, one of the best, from Stanford HAI is available [here](#)

In short, tread carefully. Treat the hype with suspicion and contempt. Heed the comments by the [CTO of OpenAI](#) herself on the product and its shortcomings.

New links for possible next version.

[LLM training pack](#)

[Google "We Have No Moat, And Neither Does OpenAI" \(semianalysis.com\)](#)

[copyright github and openai](#)

Texas courts warning https://www.linkedin.com/posts/dazzagreenwood_united-states-district-court-activity-7069452137608462336-OpEh?utm_source=share&utm_medium=member_desktop

Graeme Johnson Mata v Avianca. New York [Mata v. Avianca, Inc., 1:22-cv-01461 – CourtListener.com](#)

Graeme Johnson Mata v Avianca supporting notes: [Post](#) | [Feed](#) | [LinkedIn](#)

Apple Crack down on use – Forbes [Apple Joins A Growing List Of Companies Cracking Down On Use Of ChatGPT By Staffers—Here's Why \(forbes.com\)](#)

Techradar. Samsung crackdown on use: [Samsung workers made a major error by using ChatGPT | TechRadar](#)

Philistine. https://www.linkedin.com/posts/louis-rosenberg-025851132_johannesvermeer-activity-7069833281948569600-uMKq?utm_source=share&utm_medium=member_desktop

Prompt “injection or prompt hacking” - [Hacking LLMs with prompt injections | by Vickie Li | Jun, 2023 | Medium](#)

[Automated data analysis? - Noteable: The ChatGPT Plugin That Automates Data Analysis | by The PyCoach | May, 2023 | Artificial Corner](#)

Let's play the humanising game OpenAI seem to like playing. Basically, it lied or fantasised. Followed closely by a legal who took something at face value a liar and fantasist produced. Not clever.

Which also raises the idea if, when it gets things wrong, it “hallucinates” and openAI cannot say why it does it, then surely it also “hallucinates” when it gets things right on occasion, because by extension it is difficult to prove why it is right too.

The need for a ban, across all jurisdictions is obvious really. Surely?

Unless of course it has now become acceptable to lie to a court.....

And where does that lead? [Homepage | Post Office Horizon IT Inquiry \(postofficehorizoninquiry.org.uk\)](#)

[Complex task solving against academic papers](#)

GPT attack vectors hallucinogenic packages: <https://vulcan.io/blog/ai-hallucinations-package-risk>

GPT Tokenisation: <https://simonwillison.net/2023/Jun/8/gpt-tokenizers/>

Humour and ChatGPT: [Researchers discover that ChatGPT prefers repeating 25 jokes over and over – Ars Technica \(ampproject.org\)](#)

Transformers and “compositionality”: [\[2305.18654\] Faith and Fate: Limits of Transformers on Compositionality \(arxiv.org\)](#)

Why AI will save the world: [Why AI Will Save the World | Andreessen Horowitz \(a16z.com\)](#)

ChatGPT Plug ins: [ChatGPT Plugins Overview \(gptgames.dev\)](#)

Stanford LLM EU AI Act Compliance: [Stanford CRFM](#)

Trust assessment LLM's: [Paper page - DecodingTrust: A Comprehensive Assessment of Trustworthiness in GPT Models \(huggingface.co\)](#)

[2306.11698.pdf \(arxiv.org\)](#)

I did not think it would be wrong m'lud...: [AI: Judge sanctions lawyers over ChatGPT legal brief \(cnbc.com\)](#)

Will we run out of data?: https://www.linkedin.com/posts/joseapbello_will-we-run-out-of-data-activity-7077609943930953728-IQwI?utm_source=share&utm_medium=member_desktop

Integrating LLM's and knowledge graphs: [2306.08302.pdf \(arxiv.org\)](#)

Flooding the web: https://www.linkedin.com/posts/rohanlight_junk-websites-filled-with-ai-generated-text-activity-7079202124768690176-zKs1?utm_source=share&utm_medium=member_desktop

Flooding the web 2: [The Scourge of LCM Low-Cost Media | LinkedIn](#)

Hallucination

[What does it mean when an LLM "hallucinates" & why do LLMs hallucinate? \(labelbox.com\)](#)

[What Are LLM Hallucinations? Causes, Ethical Concern, & Prevention - Unite.AI](#)

[\[2102.00233\] An evolutionary view on the emergence of Artificial Intelligence \(arxiv.org\)](#)

[The Evolutionary Dynamics of the Artificial Intelligence Ecosystem | Strategy Science \(informs.org\)](#)

[Evolutionary algorithms and their applications to engineering problems | SpringerLink](#)

[Evolutionary approaches towards AI: past, present, and future | by Matthew Roos | Towards Data Science](#)

Prompt Engineering

[Chain-Of-Thought Prompting In LLMs | by Cobus Greyling | Medium](#)

[Generative AI Prompt Pipelines. Prompt Pipelines extend prompt... | by Cobus Greyling | Medium](#)

[Preventing LLM Hallucination With Contextual Prompt Engineering — An Example From OpenAI | by Cobus Greyling | Medium](#)

Information geometry

Introduction: ["Introduction to Information Geometry" by Frank Nielsen - YouTube](#)

Conceptual geometry: [The Geometry of Thinking, Peter Gärdenfors - YouTube](#)

Elementary Introduction to information Geometry: [An Elementary Introduction to Information Geometry - PubMed \(nih.gov\)](#)

The Royal Society. Algorithms that remember: [Algorithms that remember: model inversion attacks and data protection law | Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences \(royalsocietypublishing.org\)](#)

MIT – People training AI to train AI: [The people paid to train AI are outsourcing their work... to AI | MIT Technology Review](#)

Its not one model its... [GPT-4's Secret Has Been Revealed - by Alberto Romero \(substack.com\)](#)

Build your own? [SAI Notes #08: LLM based Chatbots to query your Private Knowledge Base. \(swirlai.com\)](#)

https://www.linkedin.com/feed/update/urn:li:activity:7078681991377768448?utm_source=share&utm_medium=member_desktop

GPT-4

◆ Mixture of Experts

◆ LLM Router

◆ Switch

◆ Semantic Routing over the

◆ Working Memory

◆ Graph of

◆ Semantic

◆ Data Agents: <https://lnkd.in/eyJuFNci>

Rumour: <https://lnkd.in/erZGyBxY>

Model: <https://lnkd.in/edXJYmbj>

Chain: https://lnkd.in/e_kNKRgc

Transformers: <https://lnkd.in/eJYdjUVM>

network: <https://lnkd.in/exPY6B5E>

Graph: <https://lnkd.in/eQF4PE27>

Thought: <https://lnkd.in/eSVusZgs>

Layer: <https://lnkd.in/efRncmk8>

[LLM Powered Autonomous Agents | Lil'Log \(lilianweng.github.io\)](#)

[Supervised Chain-Of-Thought Reasoning Mitigates LLM Hallucination | by Cobus Greyling | Jun, 2023 | Medium](#)

https://www.linkedin.com/posts/tonyseale_ai-llm-gnn-activity-7082991673189773312-pw90?utm_source=share&utm_medium=member_desktop

The Law

[AI for Law · MIT Computational Law Report](#)

Neo 4J



Content Repository: Graph Databases / LLMs / GPT / Generative AI

Post	Date
------	------

Exploring ChatGPT for Learning, Code, Data, NLP & Fun	Nov 22
Create Neo4j Database Model with ChatGPT	Jan 23
ChatGPT and Neo4j for Creating and Importing Sample Datasets	Feb 23
Doctor.ai+GPT-3+Kendra = An Ensemble Chatbot for Healthcare	Feb 23
Creating a Knowledge Graph From Video Transcripts With ChatGPT	Apr 23
Context-Aware Knowledge Graph Chatbot With GPT-4 and Neo4j	Apr 23
Generating Company Recommendations using LLMs & Knowledge Graphs	Apr 23
Integrating Neo4j into the LangChain ecosystem	Apr 23
Fine-tuning an LLM model with H2O LLM Studio to generate Cypher statements	Apr 23
Generating Cypher Queries With ChatGPT 4 on Any Graph Schema	May 23
ChatGPT Generated Star Wars Data: Let's Explore With Neo4j	May 23
LangChain has added Cypher Search	May 23
Knowledge Graphs & LLMs: Harnessing Large Language Models with Neo4j	May 23
Jump into Graph: Creating Mock Data to get started with Graph Databases	May 23
Neo4j Live: GraphGPT	May 23
OpenAI API Access - APOC Extended Documentation	Jun 23
LangChain Cypher Search: Tips & Tricks	Jun 23
Neo4j Integrations with Generative AI Features in Google Cloud Vertex AI	Jun 23
Neo4j Knowledge Graphs and Google Generative AI	Jun 23
Integrate LLM workflows with Knowledge Graph using Neo4j and APOC	Jun 23
Unifying Large Language Models and Knowledge Graphs: A Roadmap	Jun 23
Neo4j's Role in Fueling Generative AI with Graph Technology	Jun 23
Knowledge Graphs & LLMs: Fine-Tuning Vs. Retrieval-Augmented Generation	Jun 23
Knowledge Graphs & LLMs: Multi-Hop Question Answering	Jun 23
Neo4j Knowledge Graphs and Google Generative AI	Jun 23
Personal Movie Recommendation Agent with GPT4 + Neo4j	Jun 23

Others

[\[2303.18223\] A Survey of Large Language Models \(arxiv.org\)](#)

[There Is No A.I. | The New Yorker](#)