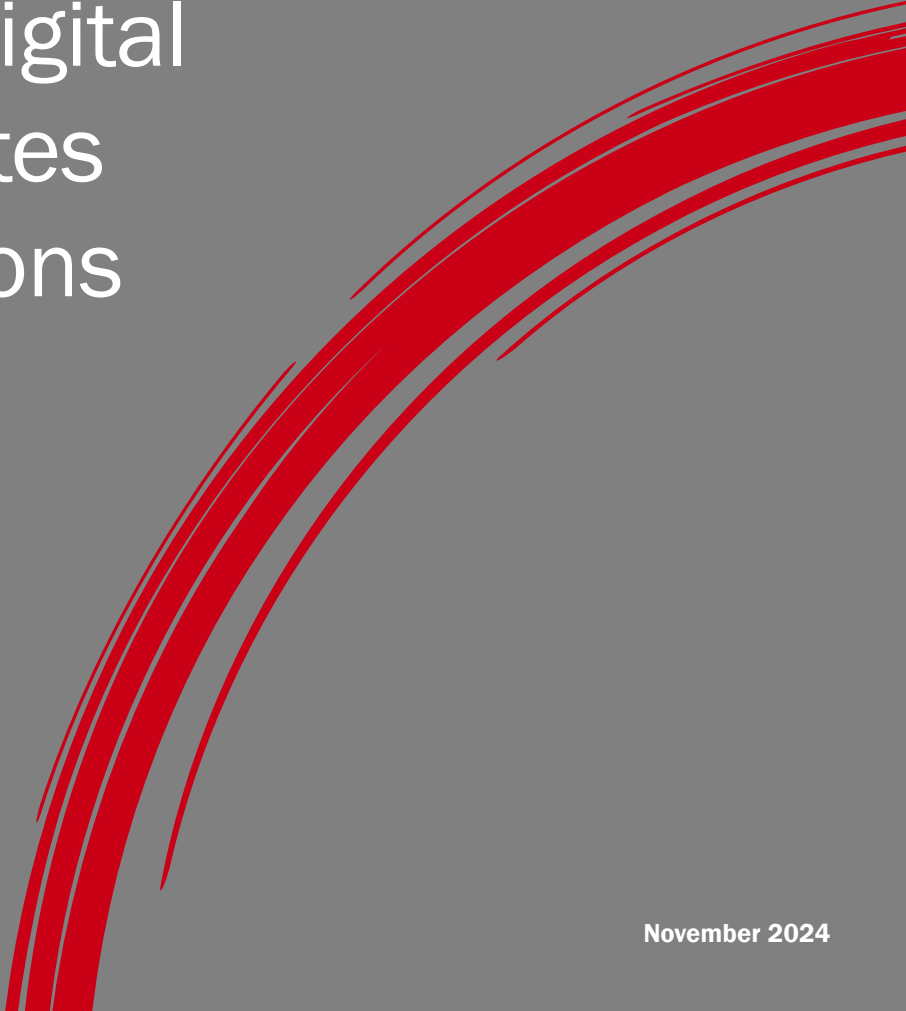


Connected Excellence in Research

Strategy Report

Towards a Shared AI Strategy for European Digital Science Institutes and Organisations



Authors and Speakers

Authors:

- Joost Geurts, Inria
- Peter Kunz, ERCIM
- Han La Poutré, CWI
- Marie-Hélène Pautrat, Inria,

Speakers:

- Benoît Sagot, Inria
- Cecile Huet, head of unit AI and Robotics, DG Connect, European Commission
- Liviu Stirbat, head of unit AI in and Science, DG Research and Innovation, European Commission
- Damien Gratadour, Centre national de la recherche scientifique, Observatoire de Paris
- Dong Nguyen, Utrecht University
- Mathieu Acher, INSA/Inria/CNRS/IRISA
- Seán Ó hÉigeartaigh, Centre for the Future of Intelligence, University of Cambridge
- Haris Papageorgiou, Institute for Language and Speech Processing, ILSP/R.C. Athena
- Anastasios (Tassos) Roussos, ICS-FORTH
- Abdallah El Ali, CWI

Contents

Executive Summary	4
Introduction	5
The 2024 ERCIM Visionary Event Overview	5
Current State of GenAI and LLMs	6
European Landscape on LLMs and GenAI	6
Impact on and Science	6
Broader Impacts and Future Perspectives	7
The Impact of Gen AI and LLMs on Society	8
Panel Discussion Reflections	9
Recommendations for an Institutional AI Strategy	11
Key Highlights	11
Recommendations	12
Annex I	13
Program	13
Annex II	14
Speakers' Bios	14

Executive Summary

This report summarizes the discussions and findings from the 2024 ERCIM Visionary Event "Challenges and Opportunities of Foundational Models and Generative AI (GenAI) for Science and Society," held on 16 April 2024 in Brussels. Experts from across Europe convened to examine the rapidly evolving landscape of GenAI and Large Language Models (LLMs), their impact on science and society, and the role of European institutions in this context.

The Key Points discussed by the experts include:

- Technological advancements in LLMs and GenAI
- European political landscape and initiatives
- Impact on sciences
- Broader impacts and future perspectives
- Impact on society
- Panel reflections

The report concludes with a set of High-Level Recommendations, such as:

- Promote balanced and objective discussions
- Develop organisational oversight and monitoring
- Facilitate multidisciplinary exchange

GenAI and LLMs are transforming science and society, presenting significant opportunities alongside inherent challenges. European digital science institutes and organizations are well-positioned to contribute to responsible development and use of these technologies. By promoting balanced discussions, implementing oversight mechanisms, and fostering multidisciplinary collaboration, institutions can harness the benefits of GenAI while mitigating its risks. These efforts are essential to ensure ethical alignment, maintain competitiveness, and maximize positive impacts on society and science.

We invite readers to explore the full report for a comprehensive understanding of the discussions, insights, and detailed analyses that underpin these key points and recommendations.

Introduction

GenAI is undergoing a revolutionary phase of development profoundly impacting both society and the scientific community. While it serves as powerful enabler for innovation and advancement, it also presents inherent challenges that must be addressed responsibly. Developing a realistic understanding of the transformative nature of GenAI is crucial for leveraging its benefits while mitigating potential risks. .

As a network of European centres of excellence in digital technology, the ERCIM institutes are well positioned to contribute to this important discussion at both national and European levels. This report aims to assist institutes and organizations in developing an organizational understanding and strategy to promote the responsible use and development of GenAI.

Based on the 2024 ERCIM Visionary event “Challenges and opportunities of Foundational Models and GenAI for science and society”, this document aims to help institutes and organisations to develop an organisational understanding and possibly a strategy, to promote the responsible use and development of GenAI.

The 2024 ERCIM Visionary Event Overview

Experts from across Europe gathered on 16 April 2024 in Brussels at the Maison Irène et Frédéric Joliot Curie to discuss the rapidly evolving landscape of GenAI and large language models (LLMs) at the ERCIM visionary event titled “Challenges and Opportunities of Foundational Models and GenAI for Science and Society.” The event featured focused sessions covering the current state of affairs, an overview on the European political landscape, and the impact of LLMs on the sciences and society. The day concluded with a panel discussion on the role and expectations for the digital sciences in addressing the opportunities and challenges posed by LLMs and GenAI.

The following experts were invited to give presentations at the event:

- Benoît Sagot, Inria
- Cecile Huet, head of unit AI and Robotics, DG Connect, European Commission
- Liviu Stirbat, head of unit AI in and Science, DG Research and Innovation, European Commission
- Damien Gratadour, Centre national de la recherche scientifique, Observatoire de Paris
- Dong Nguyen, Utrecht University
- Mathieu Acher, INSA/Inria/CNRS/IRISA
- Seán Ó hÉigeartaigh, Centre for the Future of Intelligence, University of Cambridge
- Haris Papageorgiou, Institute for Language and Speech Processing, ILSP/R.C. Athena
- Anastasios (Tassos) Roussos, ICS-FORTH
- Abdallah El Ali, CWI.

In this document, we, organisers and reporters of the event, present the discussions, opinions, and statements made during this event. We emphasize that this is not a scientific document, but rather our view on and impression of the event.

The report is structured in two parts. The first part focuses on aspects of LLMs and AI at the EU level, as well as specific and technical perspectives on their applications. The second part presents broader views on the impact of AI on society and science as a whole, along with reflections from the panel discussion. These subsequently give rise to several key highlights and high-level recommendations distilled by the authors.

Current State of GenAI and LLMs

Benoît Sagot of Inria opened the event with an overview of advancements in LLMs, focusing on the transformative effects of Transformer architectures. There is a trend towards larger models. Additionally, there is a need for balancing data diversity with volume. Finding an optimal balance is an active topic of research, not only from a performance point of view, but also considering computational resources and limiting their impact on carbon emissions. A lack of diversity in the datasets leads to representational biases in the trained models. Languages (and their cultural context) are already very unevenly represented on the internet, LLM-generated content further reduces data diversity. Ultimately this could lead to LLMs producing unreliable output if iteratively retrained on generated data. There are also concerns about the proprietary nature of training data, legal challenges, environmental impacts, and biases in these models that could reinforce societal stereotypes.

European Landscape on LLMs and GenAI

Cecile Huet from DG Connect, European Commission, emphasized the transformative nature of AI technology on European society and industry. The European Commission identifies five key ingredients for a successful European AI ecosystem: data access, algorithms, investment, skills, and processors. The “European Innovation Package” promises to boost the EU’s AI capabilities with a four-billion-euro investment. It includes an EU strategy for developing European AI (i.e. AI factories) that is built around access to HPC infrastructure, high-quality data, algorithms, and talents. Furthermore, the GenAI4EU initiative has a focus on 14 industrial ecosystems to stimulate the uptake of GenAI technology. Development and adoption of AI algorithms are supported through the ALT-EDIC multi-country project, Testing and Experimentation Facilities (TEFs), regulatory sandboxes, and the EIC accelerator to support startups. Finally, data access is established through common European dataspace that develop a single European market for data.

Liviu Stirbat from DG Research and Innovation, European Commission, presented the activities and goals of the European Commission’s “AI in Science” unit. AI has transformative potential in scientific research. Integrating AI in science can accelerate scientific discovery and increase the productivity and quality of science, but misuse also poses risks to the integrity of the scientific method. He pointed out that, based on bibliometric analysis, China is accelerating the adoption of AI in scientific disciplines. Large corporations like Microsoft and Google are also making significant contributions to basic research, including fields like mathematics through initiatives like Google DeepMind. A point of concern is the ownership of scientific data controlled by private companies, which poses challenges for scientists needing access to datasets. Likewise, proprietary archives of scientific publications can theoretically be used to generate new research findings with GenAI. There is a need for an EU policy framework that optimizes AI use in scientific settings while ensuring competitiveness on a global scale. The European Commission’s objectives are to 1) Accelerate AI uptake by scientists in the EU, and 2) Monitor and steer the impact of AI on the scientific process. These include a policy brief titled “AI in Science”, the publication of “Guidelines on the Responsible use of Generated AI in Research”, and the setup of a Scientific Advice mechanism on AI in Science (see <https://scientificadvice.eu/>).

Impact on and Science

Damien Gratadour of the Observatoire de Paris highlighted that Europe is at the forefront in the field of astronomy. The construction of the European Extremely Large Telescope in Chile is designed to facilitate major scientific breakthroughs. Important there is Adaptive Optics (AO), a technology that enhances the performance of telescopes by dynamically correcting the distortion of incoming light. This system uses thousands of actuators to control the shape of the incoming wavefront in real time, with latency under one millisecond (a stable time to solution is critical). This technology is used in applications such as planet hunting and observing distant celestial objects. GenAI is used for building advanced models for interpreting noisy observations using giant astronomical telescopes. Important potentials are in fields like planet hunting and observing distant celestial objects. Moreover, Generative Adversarial Networks (GANs) show promise in achieving super-resolution and have the ability to extrapolate data effectively. Reconstructing a long exposure (involving multiple

stochastic processes) can benefit from super-resolution, however they can sometimes produce inaccurate instant results that make it unsuited for optimizing operations. Thus, several thought-provoking questions about the reliability and ethical considerations of using GenAI in science can be posed. The use of GenAI improves results on average, but cannot be trusted on a specific measurement.

Dong Nguyen of Utrecht University pointed out that the focus of Natural Language Processing (NLP) typically are on task-oriented applications such as automatic summarization, part-of-speech tagging, and spam detection. NLP, particularly through the use of LLMs, can address broader, more complex questions within social sciences, such as the prevalence of hate speech, effectiveness of different interventions, public opinion dynamics, and changes in the meaning of words over time. Translating social and cultural concepts into measurable quantities is an area where LLMs can play a significant role. One methodology for fine-tuning LLMs can involve annotating a small dataset, training a classifier on this data, and then applying the model to larger datasets. Three challenges in integrating LLMs into social science research can be identified:

1. Determining the correct answers, as NLP models often require ground truth data for validation.
2. Assessing whether LLMs are providing the right answers for the right reasons, which is crucial for their reliability.
3. Comparing the performance and outcomes of LLMs directly with human capabilities.

While LLMs hold transformative potential for scientific research, their integration into social sciences must be handled carefully. Machine learning models may exploit certain spurious patterns in the data that are difficult to identify and, therefore, cannot be trusted without verification. Nevertheless, LLMs will transform science sooner or later. Therefore it is important to understand the transformation and to possibly shape it. In that context interdisciplinary collaboration is advised.

Mathieu Acher from INSA/Inria/CNRS/IRISA discussed the implications of LLMs in software engineering, highlighting their role in code generation and optimization. He noted that a parallel between software and AI exists. Ultimately, AI is software. Several studies report that the same data analyzed with different software can lead to different results. Many of the challenges identified for AI (e.g. trustworthiness) apply to software more generally. Having said that, LLMs offer many opportunities to facilitate software development. LLMs can make advanced software development tools more accessible to a broader range of developers. Furthermore, by automating routine tasks, LLMs can increase productivity and ensure the reliability of the software produced. Additionally, LLMs enable developers to explore different variations of software more efficiently, enhancing innovation and customization in software solutions.

Broader Impacts and Future Perspectives

Seán Ó hÉigearthaigh of the University of Cambridge stated that we are only at the start of the AI revolution. This is indeed the case, considering the variety of AI techniques developed in previous decades, the well-developed learning techniques, and especially the breakthrough of LLM techniques. Every few months, substantial technical improvements are made. So, we have increasingly capable, large, and general systems, while we don't fully understand what's going on inside them. On the positive side, they are important for economic growth and social benefits. However, current concerns include hallucinations, deepfake videos, impact on jobs, environmental impacts, biases, data privacy leaks, IP infringement, and possible future misuse for attacks. We have to anticipate drastic future developments in AI techniques and deployment, and hence drastic impacts on our society. This includes possible malicious use of AI affecting physical security, information security, and political security.

We are now heading for more general AI capabilities, centralized and/or decentralized, which can be combined into a whole. Here, three AI features can be combined, yielding serious concerns: generality of AI systems, their autonomy (agency), and the combination with strategic planning and execution. This thus combines capabilities, alignment, and control into a whole. So, we are heading for AI that can carry out all or most intellectual tasks as well as a humans, including the execution of AI research itself. Such combined AI systems can perform operational and control tasks in society. AI systems can act in the real world and many such AI systems will exist.

So, important questions are: What are we heading for? And how can we counter the negative and/or threatening effects? Which disciplines, inside as well as outside of AI can counter these, and in what way? Important international measures to take also include e.g. developing robust standards and regulations, including regulations about licensing, liability and external auditing.

Haris Papageorgiou of ILSP/R.C. Athena stated that the landscape of LLM research and development has shifted significantly from predominantly academic institutions to a more industry-dominated field over the last ten years. This transition has implications for the direction of research and the accessibility of innovations. This not only holds for the opportunities of LLM development and usage, but especially also for the threats.

Limitations include misalignment with human needs, lack of interpretability, susceptibility to hallucinations, issues with obsolete data, limited reasoning capabilities, and lacking attribution in scholarly citations. LLMs' impact on society has dual concerns: intrinsic biases (built into the models through biased input data sets or algorithmic preferences) and extrinsic harms (such as those that affect broader societal norms). In addition, there are environmental concerns, legal issues, and ethical dilemmas. So, possibly the monitoring of AI systems with AI could help to prevent abuse: being able to 'trace' AI outputs and AI-generated content, and to take action if needed.

Legal and ethical concerns include, for example, using the same LLM models in different domains, thus amplifying the (biased) impact or weaknesses. Also, LLMs can be deployed in scientific research. LLMs in Open Science can significantly accelerate scientific discovery across disciplines. Here, the (lack of) broadness and rigor of AI-supported research should be taken care of and users should be aware of that.

Possible strategies for harnessing the benefits of LLMs while mitigating their risks include source tracing, adoption of "good" data and model management practices, and reinforcing community values and norms to guide the development and use of AI technologies (by, for example, diversity, intervention, feedback, etc.). The EU AI Act as an important regulatory framework that already addresses many of the issues discussed. Important actions include mobilizing the open-source community, improving quality control on data, and involving stakeholders early in the development processes to ensure that LLMs contribute positively to society.

In conclusion, LLMs can be seen as a "sociotechnical system" LLMs should be viewed and managed not just as technical tools, but as parts of a broader sociotechnical system that encompasses ethical, cultural, and social dimensions.

The Impact of Gen AI and LLMs on Society

Anastasios Roussos of FORTH explained that the evolution of GenAI goes back to the deep learning revolution that began in 2012. Developments range from automatic captioning to text-to-image generation. The techniques of AI for generating realistic images and videos are constantly improving, e.g. with GANs. The current AI models are able to generate synthetic images and videos of unprecedented quality and realism. These developments have a broad range of applications that affect many fields, from visual arts and entertainment to politics and journalism. Text-to-image generation attracts a large audience; this leads to the democratization of artistic means, or at least the tools to be artistic: artists will focus on more creative and genuine creations. Also, the ability to create virtual humans include interactive tutors and trainers, movie dubbing, human-computer interaction (HCI) and chatbots, and natural communication across different languages and cultures.

Concerns and ethical considerations include potential job loss for artists or performers, copyright issues, environmental impact, and AI-driven content creation in the style of a living artist/performer. Also, AI systems reflect the biases in the datasets that they were trained on, and could be mistakenly interpreted by the general public as completely objective; thus possibly leading to reinforcing stereotypes. E.g. 18 AI-generated images of scientists are all male, whereas in reality 43% of the scientists are female.

To address the challenges posed by GenAI, countermeasures could include:

- The implementation of tools to detect and mitigate biases in AI datasets, like those related to race, nationality, and other identities;
- Enhancing public media literacy to combat the risks posed by deep fakes and fake news;
- Education to enhance public awareness about the capabilities and risks of AI in image and video synthesis;
- Tools that reveal biases in available AI-generated images;
- -Tools for, e.g, diversifying training data, bias detection and mitigation, user feedback incorporation, etc.

Last but not least, there is an urgent need for scientists in the digital fields to be more engaged and proactive regarding the relevant ethical considerations. Big conferences now require authors to discuss the potential negative societal implications. Furthermore, interdisciplinary ethics-related courses integrated into the relevant STEM degrees should become much more common.

Abdallah El Ali of CWI stated that GenAI and Large Whatever Models (LWMs) are transforming research cycles in HCI (Human Computer Interaction). These technologies can reshape interactions between humans and computers as well as the research methodologies used to study them.

This also raises the question of the perception of AI by humans: Can AI be perceived as a reliable source of information, for example in health contexts? Is AI perceived as trustworthy as humans? And what about labelling news as human-made or AI-made? Or whether the news is true or fake? Current results show that these can be counter-intuitive. Also, behavioural and physiological responses differ when individuals interact with human-generated versus AI-generated news snippets.

Article 50 of the EU AI Act already has important practical implications, particularly for AI transparency and accountability. Clear labeling (who, what, when, where, why, and how) is very important in AI-generated content to already ensure that users can understand the origins and credibility of the information (i.e., explainable AI).

Human perception of AI is important in fostering effective collaboration between AI systems and their human users. It is therefore important to design AI systems that are not only technically proficient but also ethically aligned and transparent in order to support positive human-AI encounters.

Panel Discussion Reflections

Panelists: Liviu Stirbat, Dong Nguyen, Damien Gratadour, Seán Ó hÉigeartaigh, Haris Papageorgiou; moderators: Joost Geurts and Han La Poutré

Key Themes Discussed

Interaction Between Systems

The panelists discussed the challenges associated with systems interacting with one another. There is a need for robust systems capable of handling increasing volumes of AI-generated information without compromising data integrity or functionality. Further, there are concerns regarding the powerful combination of increasing generality of AI systems, their autonomy, and their planning and execution capabilities.

Quality and Transparency of AI-Generated Information

Concerns were raised about the integration of AI-generated content with human-generated information, particularly the risk of diminishing the overall quality of information. Panelists stressed the importance of transparency and the selection of high-quality data for AI training to prevent the degradation of information quality. An example given was Google's proactive measures to remove AI-generated web pages from search results to maintain data integrity.

Impact of GenAI on the Web

A critical question addressed was the irreversible nature of misinformation, noise, and spam once it has been disseminated online. This is a problematic consequence of machine-translated content derived from low-quality text and needs solutions to mitigate the spread of such 'polluted' information.

Need for Multidisciplinarity

The importance of multidisciplinary approaches in enhancing AI development and application is a recurring theme. There is a necessity for collaboration between AI specialists and professionals from other fields, to enrich AI models and ensure their applicability across various domains. Additionally, issues of gender equality and geographical diversity in AI were pointed out as essential for fostering comprehensive and inclusive AI solutions.

Efficiency and Safety Concerns

Questions were raised about the efficiency and potential dangers of AI, particularly in contexts like quality control and automatic fact-checking. The complexity of determining which facts can be reliably checked by AI was noted, along with the potential of AI to contribute to more resource-efficient processes in the future.

Formal Methods and Open Science

The feasibility of implementing formal methods or proof systems in AI was debated, with references to current experiments such as Inria's work on formal verification with LLMs. Further, the open science paradigm could benefit from AI.

Recommendations for Future AI Strategies

Recommendations for advancing AI research included fostering an environment that supports extensive collaboration and exchange, both within and across disciplines. It was felt that AI research institutes should adopt long-term strategic planning, emphasizing the need for effective talent management, resource allocation and pooling, broad exchange of knowledge and data, and infrastructure enhancement to ensure that AI technologies are leveraged ethically and effectively in scientific research.

Recommendations for an Institutional AI Strategy

The authors present several key highlights and high-level recommendations, as distilled from the event.

Key Highlights

LLMs are important for economic growth and social benefits. They can, for example, increase productivity and customization of software. Also, AI-powered virtual humans can be valuable as e.g. interactive tutors, in HCI, and for natural communication across different languages and cultures. Text-to-image generation leads to a democratization of artistic tools, while artists can focus on more genuine creations.

LLMs should be viewed and managed not just as technical tools, but as parts of a broader “sociotechnical system” that encompasses ethical, cultural, and social dimensions. Also, human perception of AI is important in fostering effective collaboration between AI systems and their human users.

LLM research and development has shifted significantly from predominantly academia to industry over the last ten years. In the EU, a European Innovation Package promises to boost the EU’s AI capabilities with a four billion Euro investment.

There is a need for an EU policy framework that optimizes AI use in science while ensuring competitiveness on a global scale. Integrating AI in science can accelerate scientific productivity and quality, but misuse poses a risk to scientific integrity. China is already accelerating the adoption of AI in science. Some challenges in integrating LLMs into (social science) research are, for example, determining the correct answers, as often ground truth data is required for validation, and assessing whether LLMs are providing the right answers for the right reasons.

Direct concerns about LLMs include the proprietary nature of training data, legal challenges, environmental impacts, and biases in data and models that could reinforce societal stereotypes. Additional concerns include physical security, information security, and political security, such as hallucinations, use in propaganda, deepfake videos, impact on jobs, environmental impacts, biases, data privacy leaks, IP infringement, and possible misuse for attacks.

Lack of diversity in data sets leads to a biases in the computed models by LLMs. Also, LLM-generated content further reduces data diversity, ultimately leading to LLMs producing unreliable outputs if iteratively retrained on generated data. Furthermore, machine learning models may exploit certain spurious patterns in data that are difficult to identify, and therefore cannot be trusted without verification.

AI capabilities will become more general and can be combined into a whole, centralized and/or decentralized. Three AI features that yield serious concerns when combined are: generality of AI systems, their autonomy, and the combination with planning and execution capabilities in society.

An important question is: What are we heading for? AI that can carry out intellectual tasks as well as humans, including the execution of AI research itself. How can we counter the negative and/or threatening effects? Which (sub)disciplines inside as well as outside of AI can counter these and in what way?

Possible strategies for harnessing the benefits of LLMs while mitigating their risks include source tracing, adopting “good” data and model management practices, reinforcing community values and norms to guide the development and use of AI, and tools that reveal biases in AI-generated models. Also, scientists in computer science need to be ethically engaged and collaborate with other disciplines.

Recommendations

Promote a balanced and an objective discussion grounded in scientific expertise and consensus

LLMs and GenAI have the potential to significantly transform modern society and industry. Yet, it is too early to develop a realistic understanding of the impact of the technology, and the domains and sectors that will be most affected. Additionally, the technology has also serious shortcomings and challenges in terms of trustworthiness, environmental impact and sovereignty. The trade-off between advantages and challenges is not obvious and requires a level of understanding of the technology. In order to develop a realistic, long-term understanding on the transformative impact of the technology, a balanced objective discussion should be promoted among scientists within institutes and organisations, but also with partners outside of these (e.g. industry and policy).

Develop a means for organisational oversight and monitoring to help understand trends and undesirable phenomena

With exceptions, LLMs and GenAI concern scientists who are active in the development of the technology, but also scientists who may be experimenting with or using the technology in the context of their own research. The choice of a particular technology or model has socio-economic consequences that may not be initially obvious at time of deployment. This includes ensuring access to technology, data, and infrastructure enabling the research, as well as safeguarding intellectual property obtained from the research (e.g., a system may enrich itself through interactions, or data it is exposed to). This impacts individual researchers, but also research communities (e.g., through collaboration using particular technologies or platforms), and ultimately the full value chain (e.g., innovation resulting from research that depends on proprietary models). In order to detect and quantify undesirable phenomena early on, some level of institutional oversight or monitoring may be helpful.

Set up a means for (institutional) exchange with other, non-digital sciences (e.g local institutes, universities, and organisations)

LLMs and GenAI have an impact on other sciences and disciplines. On the level of scientists, there is a need to facilitate dialogue with other disciplines to integrate this technology in their research while understanding its strengths and limitations. This may include means to support multidisciplinary research, or setting up training programmes. On the level of institute management it may include raising awareness of capacity development and technological dependencies.

Annex I

Program

Program ERCIM Visionary Event - “Challenges and Opportunities of Foundational Models and Generative AI for Science and Society”, 16 April 2024, Brussels.¹

Venue Address:

Maison Irène et Frédéric Joliot Curie
Rue du Trône 100 / Troon 100
Brussels
Metro Trône

9:00 – 9:05 Welcome

9:05 – 9:45 State of play in Large Language Models (LLM) and Generative AI (technology perspective)

- Benoît Sagot, Inria – *Language models and their training data: challenges and perspectives.*

9:45 – 10:30 European political landscape on LLMs and GenAI

- Cecile Huet, head of unit AI and Robotics, DG Connect, European Commission
- Liviu Stirbat, head of unit AI and Science, DG Research and Innovation, European Commission

10:30 – 10:45 Coffee

10:45 -12:00 Impact of Gen AI and LLM on the Sciences (Tool or Threat)

- Damien Gratadour, Centre national de la recherche scientifique, Observatoire de Paris – *Building new brains for giant astronomical telescopes with generative AI*
- Dong Nguyen, Utrecht University – *LLMs and the social sciences: One answer to a million questions?*
- Mathieu Acher, INSA/Inria/CNRS/IRISA (DiverSE team) – *LLM for Software Engineering*

12:00 – 13:00 Lunch

13:00 - 13:40: Generative AI and Future Perspectives

- Seán Ó hÉigeartaigh, Centre for the Future of Intelligence, University of Cambridge – *Frontier AI models: Scaling, risks and governance*

13:40 – 14:55 Impact of Gen AI and LLM on Society (Tool or Threat)

- Haris Papageorgiou, Institute for Language and Speech Processing, ILSP/R.C. Athena – *Large Language Models as Infrastructure for Open Science*
- Anastasios (Tassos) Roussos, Institute of Computer Science (ICS), Hellas (FORTH) – *Assessing the Impact of Generative AI for Image and Video Synthesis*
- Abdallah El Ali, Centrum Wiskunde & Informatica (CWI) – *Fake Realities: Toward Transparent and Trustworthy Human-AI Interaction*

14:55 – 15:10 Coffee

15:10 – 16:25 Panel and open debate : Which challenges and expectations do we see for the digital sciences (as response or proactive to the above)

- Liviu Stirbat, Dong Nguyen, Damien Gratadour, Seán Ó hÉigeartaigh, Haris Papageorgiou

16:25 – 16:30 Concluding remarks

16:30: Wrap up and drinks

¹ The Visionary Event is sponsored by CWI and Inria.

Annex II

Speakers' Bios

Mathieu Acher is Professor at University of Rennes (INSA) and Inria (DiverSE team), France. His research focuses on modelling, reverse engineering, and learning (deep) variability of software-intensive systems. Beyond its applicability, his research is original in combining software engineering and artificial intelligence techniques (symbolic reasoning, machine learning, generative AI). He is the author of more than 150 peer-reviewed publications in international journals and conferences. His work has received Most Influential Paper Award (SLE'19) and Best Paper Awards (SPLC'13, ICPE'19, SPLC'21, ICSR'22, MODELS'23). Since 2021, he is a junior research fellow at Institut Universitaire de France (IUF). He's co-leading the "Inria défi" LLM4Code. More information: <https://mathieuacher.com/>

Abdallah "Abdo" El Ali is a research scientist in Human Computer Interaction at Centrum Wiskunde & Informatica (CWI) in Amsterdam. He leads the research area on Affective Interactive Systems, where he combines advances in sensing and actuation technologies, eXtended Reality, and Artificial Intelligence to measure, infer, and augment human cognitive, affective, and social interactions. He is also affiliated with the AI, Media, and Democracy Lab (<https://www.aim4dem.nl/>), where he leads Human-AI Interaction research focusing on AI transparency in media. He is also part of the executive board for CHI NL (<https://chinerland.nl/>), the Dutch ACM SIGCHI Chapter. Website: <https://abdoelali.com>

Damien Gratadour is a Senior Research Scientist at Observatoire de Paris, CNRS. Damien holds a PhD in Observational Astronomy from Université Paris-Diderot (2005). He has been an Adaptive Optics (AO) fellow, responsible for the last stages of commissioning of the Altair AO system on the Gemini North Telescope in Hawaii (2006) ; and an Instrument Scientist (2007-2008), for GeMS, the Gemini MCAO System, a facility featuring 6 Laser guide stars. Since 2008, at Observatoire de Paris - PSL, Damien has been leading an original research program on high performance numerical techniques for astronomy including modeling, signal processing and instrumentation for large telescopes. He has been the P.I. of several large programs at national and European levels targeting AO Real-Time Controllers for giant optical telescopes with emerging computing technologies. Since 2021, with France officially joining SKAO, he is also getting strongly involved in the French effort dedicated to the construction of this giant radio-telescope. In particular, he is currently the inaugural head of ECLAT, a joint laboratory between CNRS, INRIA and Atos/Eviden, as a long-term support structure federating resources from academic and industrial teams that will engage in the R&D work for the French contribution to the SKA.

Cécile Huet is Head of the Unit "Robotics and Artificial Intelligence Innovation and Excellence" at the European Commission. This unit funds and assists beneficial robotics and AI developments within Europe. Under Horizon Europe, the unit launched the Public-Private Partnership on AI, Data and Robotics. This unit is also at the heart of the European "Ecosystem of Excellence" in AI and the "AI Innovation Package" (<https://kwz.me/hFO>), a 4 Bn € investment including the "AI Factories" and the "GenAI4EU" initiative, aiming to boost European's capacity in Generative AI. Cécile joined the unit since its creation in 2004. Previously, she worked for the industry in signal processing after a post-doc at the University of California Santa Barbara and a PhD at University of Nice Sophia Antipolis, France. In 2015, she has been selected as one of the "25 women in robotics you need to know about" (<https://kwz.me/hFk>).

Sean O hEigeartaigh is Director of the AI:Futures and Responsibility programme at the Centre for the Future of Intelligence, where he leads a team of interdisciplinary researchers working on foresight, governance and risk relating to advances in artificial intelligence. His work has covered the impacts and risks of AI in pandemic response and agriculture, malicious use and disinformation challenges associated with AI, and benchmarking and forecasting progress in AI. He has a PhD in bioinformatics.

Dong Nguyen is an assistant professor at Utrecht University (the Netherlands). She completed her Ph.D at the University of Twente and she received a master's degree from the Language Technologies Institute at Carnegie Mellon University. At Utrecht University she is heading the NLP and Society Lab.

Haris Papageorgiou is Research Director at the Institute for Language and Speech Processing (ILSP) of ATHENA Research Center (www.athenarc.gr) and co-founder of one of the ARC spinoff companies, Opix, active in the field of Policy and Business Intelligence and data analytics. He holds a Ph.D. in Computer Science and a B.Sc. in Electrical Engineering. He teaches at several postgraduate programs at Athens University of Economics and Business (aueb.gr/en) and National Kapodistrian University of Athens (en.uoa.gr) covering AI and NLP topics. He has worked as an expert in several aspects of STI (Science, Technology and Innovation) Policy analysis in the context of evaluation and assessment of R&I programmes and policies. He was the Coordinator of the Technical Committee and Technical Responsible of operating the clarin:el shared distributed infrastructure, which is the Greek part of the European CLARIN infrastructure, making language resources, technology and expertise available to the humanities and social sciences research communities at large. Haris has held Chief Scientist positions in more than 50 European and national projects for several fields including computational social and life sciences. He specializes in Artificial Intelligence, knowledge discovery and representation, machine/deep learning, web mining and Natural Language Processing where he has published more than 130 scientific papers.

Anastasios Roussos is a Principal Researcher (Associate Professor level) at the Institute of Computer Science, Foundation for Research and Technology - Hellas (FORTH), Greece. Prior to his current position, he was a Lecturer (equivalent to Assistant Professor) in Computer Science at the University of Exeter and a Fellow of the Alan Turing Institute, UK. Before these positions, he was a postdoctoral researcher at Imperial College London (ICL), University College London (UCL) and Queen Mary, University of London, UK. He has studied Electrical and Computer Engineering (PhD 2010, Dipl-Ing 2005) at the National Technical University of Athens (NTUA), Greece. Dr. Roussos' research specializes in 3D Computer Vision and Deep Learning. The focus of his current research is the development of novel methods for detailed 3D modelling, reconstruction and synthesis of real-world objects and scenes based on image data. He has more than 2,220 citations to his articles with an h-index of 22 and an i10-index of 30 (source: Google Scholar, March 2024). Homepage: <http://users.ics.forth.gr/~troussos>
Google Scholar profile: <https://scholar.google.co.uk/citations?user=Baj1CKYAAAAJ>

Benoît Sagot is a computer scientist specialised in natural language processing (NLP). He is a Senior Researcher (Directeur de Recherches) at Inria, where he heads the Inria research project ALMANACH in Paris, France. He also holds a chair in the PRAIRIE institute dedicated to artificial intelligence, and currently holds the annual chair for computer science in the Collège de France. His research focuses on language modelling, machine translation, language resource development and computational linguistics, with a focus on French in all its form and on less-resourced languages.

Liviu Știrbăț, Head of Unit for AI in Science in the DG for Research and Innovation at the European Commission. Liviu Știrbăț is the Head of Unit for AI in Science in the DG for Research and Innovation at the European Commission. He worked previously as Member of Cabinet of EVP Margrethe Vestager in charge of innovation, space, and sustainability. He was Deputy Head of EU Adaptation to Climate Change, and also Deputy Head for better regulation for Horizon Europe. Before the European Commission, he worked at the OECD and the World Bank.



2004, Route des Lucioles
Sophia Antipolis
06410 Biot
France

Tel: +33 4 9238 5010
contact@ercim.eu
<https://www.ercim.eu>

