

AI in Action

RDA Global Workshops
February 2025

Table of Contents

Executive Summary	3
Introduction	5
Why This Report?	5
Background and Context	5
AI in Action February 2025 sessions	6
Objectives and format	6
Diversity of topics and speakers	7
Session registration and participation	7
Session Summaries	8
Session 1: 13 Feb 2025 (Europe/Americas)	8
Towards Reproducible AI Research in Life Sciences: Best Practices with DOME- Gavin Farrell	8
Role of Generative AI in Research Software Engineering: Opportunities, Challenges, and Outlook- Carlos Utrilla Guerrero	9
Health data for AI: Insights from a PhD's logbook- Giulia de Luca	9
itwinai: Large-Scale ML on HPC for Digital Twins in Science- Matteo Bunino	10
AI Tools for Research: A Tried-and-Tested Workflow in Digital Humanities- Saloni Chaudhary	11
Discussion summary	11
Session 2: 18 Feb 2025 (Asia/Oceania)	13
Research Use of LLMs: Beyond the Chatbot- Brian Ballsun-Stanton	13
Leveraging AI for Research Classification- Amir Aryani	13
Healthy Connections: An AI-Driven Mobile Medi-Kit for Remote Health Equity- Susannah Soon	14
Discussion Summary	15
Feedback and Next Steps	16
Acknowledgement	17
About the Research Data Alliance	18
Annex: Speaker bios	19

Executive Summary

AI in Action – How Researchers Leverage AI (February 2025)

The *AI in Action* workshops, conducted by the Research Data Alliance (RDA) in February 2025, spotlighted practical applications of Artificial Intelligence (AI) in scientific research. Following the success of the RDA's State of AI Pilot Workshops in 2024, which identified key challenges in AI adoption—including reproducibility, data governance, and ethical considerations—these workshops aimed to demonstrate real-world solutions and community-driven approaches to responsibly leveraging AI.

Context and Purpose

AI is transforming research methodologies across disciplines, from biomedical sciences to digital humanities. Yet, alongside AI's promise lies the challenge of ensuring transparency, reproducibility, and responsible governance. These workshops brought together global researchers, technologists, and policymakers to share experiences and practical applications of AI in research settings, emphasizing FAIR (Findable, Accessible, Interoperable, Reusable) principles and ethical AI adoption.

Workshop Format and Participation

Two sessions were held to cover global time zones:

- **13 February 2025 (Europe, Americas, Africa):** 144 registrants, 65 live participants.
- **18 February 2025 (Asia, Oceania):** 122 registrants, 59 live participants.

A total of **266 registrants from 36 countries** signed up for the workshops. Sessions featured diverse speakers from Europe, Asia, Australia, and America, covering biosciences, health research, digital humanities, and research infrastructure.

Highlights of the Sessions

Session 1 (Europe/Americas) Key Themes:

- **Reproducible AI Research:** Gavin Farrell introduced DOME guidelines for transparent AI research in life sciences.
- **Generative AI in Research Software:** Carlos Utrilla Guerrero discussed AI-driven automation in research software engineering.

- **Medical AI Reproducibility:** Giulia de Luca emphasized challenges in reproducing AI models for lung cancer screening.
- **Digital Twins & HPC:** Matteo Bunino showcased AI in large-scale simulations for science.
- **AI in Digital Humanities:** Saloni Chaudhary demonstrated AI tools for analyzing research collaboration networks.

Session 2 (Asia/Oceania) Key Themes:

- **Beyond Chatbots:** Brian Ballsun-Stanton explored advanced uses of LLMs in research aggregation and academic writing.
- **AI-driven Research Classification:** Amir Aryani presented Australia's National PID Graph for mapping research outputs.
- **Remote Health Equity:** Susannah Soon introduced the Medi-Kit project, using AI diagnostics and telehealth to support Indigenous communities.

Discussion and Insights

The sessions fostered vibrant discussions on:

- The gap between FAIR principles and actual AI reproducibility.
- Ethical and privacy considerations, especially for health-related AI applications.
- Responsible AI use in research classification.
- Data/model drift, validation, and community standards for AI governance.

Participants highlighted the need for stronger AI reporting guidelines, model auditability, and community-led best practices to bridge policy and practice gaps.

Feedback and Next Steps

Post-session feedback indicated strong interest in continuing this dialogue. Participants appreciated the diversity of speakers and practical focus, requesting future sessions to delve deeper into specific AI research methodologies, ethical frameworks, and underrepresented fields like humanities and social sciences.

The RDA will incorporate these insights into its ongoing AI strategy, to kickstart various community initiatives and across its global research community.

In line with the workshops' AI focus, this executive summary was generated using a combination of AI tools, including Claude Sonnet 3.5 and ChatGPT-4o, with the prompt: "Generate an executive summary of the following report." The models were used with privacy-preserving configurations ensuring that no data provided to the models was stored, retained, or used for further training. The summary was generated and reviewed for accuracy and completeness by Hilary Hanahoe and Shalini Kurapati. Both were present during the workshops and confirmed that the AI-generated summary accurately reflects the key themes and content of the report.

Introduction

Why This Report?

This report provides a structured synthesis of the AI in Action workshop sessions conducted in February 2025 by the Research Data Alliance, including:

- Context and objectives of the sessions.
- Summaries of presentations and discussions from both sessions.
- Insights from Q&A and live chat interactions.
- Key resources shared by speakers and participants, including AI tools, research projects, and policy frameworks.
- Participant feedback and next steps to guide future AI initiatives within the RDA community.

Session recordings and slides are available on the [RDA website](#).

Background and Context

Artificial Intelligence (AI) is rapidly transforming research across disciplines, from biomedicine and environmental sciences to digital humanities and engineering. A 2025 Nature survey of more than 7000 researchers across disciplines highlighted that researchers are increasingly integrating AI into their work, with more than 50% expecting AI to play a major role in scientific discovery within the next five years¹.

AI is being used to automate data analysis, accelerate pattern recognition, and enhance research workflows, helping scientists process vast datasets more efficiently than ever before. The integration of large language models (LLMs), generative AI, and Machine Learning (ML) algorithms into scientific research is unlocking new possibilities for knowledge extraction, hypothesis generation, and even experimental design. As AI becomes increasingly embedded in research, ensuring transparency, reproducibility, and responsible data governance has become a grand challenge.

¹ <https://www.nature.com/articles/d41586-025-00343-5>

Recognising these challenges, the Research Data Alliance (RDA) has identified AI as a strategic topic of global importance. AI's data-intensive nature and far-reaching impact across research disciplines present key challenges in reproducibility, FAIR data governance, bias mitigation, and sustainability, all of which demand coordinated global solutions and standards. As a global organisation with over 15,000 members across 150+ countries, RDA's strength lies in its ability to connect researchers, policymakers, and infrastructure providers to embrace the AI opportunities responsibly and address the grand challenges posed by AI for the research community in an open and community driven way.

As an initial step, RDA launched the [State of AI Pilot Workshops \(September 2024\)](#), engaging key RDA stakeholders to assess the most pressing challenges and state of AI adoption within the global RDA community context. These discussions focused on data readiness as a foundation for reproducibility, the complexities of sensitive data management, the sustainability of AI infrastructure, and the gap between AI technological advancements and policy frameworks. The workshops underscored the need for clear metadata standards, ethical oversight, and improved interoperability between AI tools and research workflows.

Following feedback from these workshops, participants emphasised the need to move beyond discussions of challenges and focus on real-world AI applications in research. As a direct follow-up, RDA organised the AI in Action workshops (February 2025) to showcase practical implementations of AI, explore community-driven solutions, and highlight how researchers are addressing key barriers in AI adoption.

AI in Action February 2025 sessions

The *"AI in Action – How Researchers Leverage AI"* workshops, held on 13 and 18 February 2025, were organised by the RDA for the entire global community.

Objectives and format

The main objectives of the AI in Action workshop sessions were to:

- Demonstrate practical AI implementations across diverse research domains.
- Explore solutions to key challenges identified in the pilot workshops, including reproducibility, FAIR data governance, and ethical AI adoption.
- Foster interdisciplinary dialogue by bringing together researchers, technologists, and policymakers to share experiences and insights.
- Promote best practices and community-driven solutions, ensuring AI is used transparently, responsibly, and effectively.

The workshops were structured around presentations, Q&A sessions, and open discussions, allowing participants to engage directly with speakers. All RDA members had the chance to send in their interest to speak at either session. Two sessions were held to accommodate different time zones: 13 February for Europe, the Americas, and Africa and 18 February for Asia and Oceania, ensuring broad global participation. Speakers were chosen based on their relevance to the session topics.

Diversity of topics and speakers

The AI in Action workshops featured a diverse range of topics and speakers, reflecting the broad application of AI across research disciplines and global regions. Presenters represented institutions across Europe, Asia, Australia, covering areas such as biosciences, digital humanities, health research, environmental sciences, and research infrastructure.

Discussions showcased AI's role in high-performance computing, digital twins, and large-scale simulations in physics, as well as AI-assisted diagnostics in medicine, AI-driven metadata classification for open science, and AI-supported research discovery in digital humanities. Speaker biographies can be found in a dedicated section at the end of the report.

Session registration and participation

The workshops attracted strong global interest, with 266 registered participants from 36 countries. The Europe, Americas, and Africa-friendly session had 144 registrants from 23 countries, while the Asia and Oceania-friendly session saw 122 registrations from 21 countries. The diversity of participants across RDA global regions is illustrated in Figure 1.



Figure 1: Registered participant diversity across the RDA global regions for the Feb 13 and Feb 18 workshops

Attendance was high across both sessions, with 65 live participants in the Europe/Americas session and 59 in the Asia/Oceania session. Engagement was particularly strong in live discussions, chat interactions, and Q&A, where participants raised critical questions about AI transparency, reproducibility, model validation, and ethical AI frameworks. The following section provides structured summaries of both sessions.

Session Summaries

Session 1: 13 Feb 2025 (Europe/Americas)

Towards Reproducible AI Research in Life Sciences: Best Practices with DOME- Gavin Farrell

Gavin highlighted the role of Machine learning in rapidly advancing life sciences research, driven by improved hardware (GPUs) and the availability of high-quality FAIR data. One of the most transformative AI achievements in biology is DeepMind's AlphaFold, which has solved the long-standing challenge of protein structure prediction. Applications of ML in the life sciences are broad, including drug discovery, where AI models support drug design, screening, and bioimaging, where AI-driven tools contribute to cancer detection. Large EU-funded projects such as AI4Life exemplify AI's role in biomedical imaging.

Despite these advancements, a reproducibility crisis exists in ML-based biomedical research. Many journals lack standardised ML reporting requirements, and peer reviewers often lack expertise in machine learning, making it difficult to assess and reproduce results. Gavin Farrell underscored the need for best practices in documenting and validating ML experiments to ensure transparency and reproducibility. He introduced the DOME (Dataset Optimisation, Model Evaluation) framework, originally proposed by Walsh et al. in Nature Methods (2021), which provides standardized guidelines for reporting ML experiments. DOME ensures transparency through four key aspects: Data (provenance, splits), Optimisation (algorithm selection, parameter tuning), Model (interpretability, software stack), and Evaluation (performance metrics, comparative analysis).

The DOME framework has led to the creation of a DOME Registry, a curated database of ML publications that adhere to DOME guidelines. The registry provides annotations to assess the reproducibility of ML papers. Gavin demonstrated the registry's website and adoption metrics, highlighting 117 active users and 214 recorded ML studies. He emphasized best practices for research reproducibility, such as the use of persistent identifiers (DOIs) for datasets and models, containerisation for reproducible computing environments, and the sharing of code and results on open platforms like Zenodo.

Notable links & references from Gavin's talk (checked and last accessed on 15 March 2025):

- AlphaFold by DeepMind: [AI-driven protein structure prediction](#)
- AI4Life (Euro-Biolmaging Initiative): [AI for biomedical imaging](#)
- DOME Recommendations: [DOME ML reporting framework](#)
- Nature Methods Paper (Walsh et al., 2021): [DOI: 10.1038/s41592-021-01205-4](#)
- FAIRsharing Entry for DOME: [FAIRsharing ID 6198](#)

Role of Generative AI in Research Software Engineering: Opportunities, Challenges, and Outlook- Carlos Utrilla Guerrero

Carlos Utrilla Guerrero explored the impact of Generative AI (GenAI) in research software engineering (RSE), particularly at TU Delft, where AI-driven automation is being integrated into research workflows. GenAI tools like large language models (LLMs) can generate code, provide documentation assistance, and serve as interactive tutors for researchers. The primary objective of using GenAI in RSE is to improve productivity and automation while ensuring maintainability and reproducibility.

Beyond productivity gains, sustainability remains a challenge for AI-generated research software. Carlos discussed best practices for validating and auditing AI-generated code, emphasising the need for human-in-the-loop oversight, robust versioning of AI models, and adherence to FAIR principles for research software. He shared insights from GenAI4RS, a training initiative at TU Delft that organises hands-on workshops where researchers experiment with LLM-generated code and assess its reliability. A Spring Symposium on AI in Education was held, featuring interactive sessions to gather community perceptions of AI in research software.

Carlos highlighted the importance of community engagement and shared resources in AI-assisted research software. He referenced The Carpentries' initiative on teaching with LLMs, where collaborative note-taking and discussions on best practices are encouraged. Key challenges include trust and verification of AI-generated outputs, ethical concerns related to bias, and the need for clear audit trails of AI-generated research artifacts. Moving forward, Carlos insisted that RSE practices must align with the principles and recommendations of community groups like RDA's FAIR for Research Software (FAIR4RS) whose objective is to address similar challenges.

Notable links & references from Carlos' talk (checked and last accessed on 15 March 2025):

- GenAI Training Events: [Delft AI Spring Symposium 2024](#)
- The Carpentries Blog: [Teaching with LLMs](#)
- FAIR4RS Working Group: [FAIR4 Research Software](#) RDA group
- Zenodo Community for RSE: [RSE Community Zenodo](#)

Health data for AI: Insights from a PhD's logbook- Giulia de Luca

Giulia Raffaella De Luca's research focuses on applying AI to lung cancer screening, leveraging the National Lung Screening Trial (NLST) dataset. Lung cancer remains one of the leading causes of cancer deaths worldwide, and AI offers the potential to enhance early detection by assisting radiologists in identifying risks and nodules more consistently. However, she encountered significant challenges in accessing and standardising data, given the heterogeneity of scans from multiple centres and manufacturers. Preprocessing involved extensive format conversion and normalisation to ensure AI readiness of the datasets.

A core aspect of her work was reproducibility. She attempted to replicate findings from multiple open-source AI models for lung CT analysis, including risk models for predicting lung cancer and cardiovascular risk from a single chest CTscan. Reproducing these models proved difficult due to missing datasets, undefined computational environments, and incomplete methodological details. To address this, she employed containerisation using Docker/Singularity, ensuring that all dependencies were captured, and models could be reliably rerun.

Her key takeaway: medical AI must be rigorously validated through testing on independent datasets. Simply publishing an AI model is insufficient—openness requires adhering to FAIR principles to allow true reproducibility. She advocates for sharing datasets, code, and computing environments to enable AI models to be trusted and translated into clinical practice. Her study contributes to guidelines on publishing AI research transparently, aligning with initiatives like MIDRC and the RDA Health Data Interest Group.

Notable links & references from Giulia's talk (checked and last accessed on 15 March 2025):

- National Lung Screening Trial (NLST): [Public Dataset from NCI](#)
- Hofmanning et al. (2020): [Lung Segmentation – Data Diversity Challenges](#)
- Sybil Model (Yale/MGH, 2022): [AI Predicting Lung Cancer Risk](#)
- DeepCAC model (Zeleznik et al, 2021): [Deep convolutional neural networks to predict cardiovascular risk from computed tomography](#)
- Tri-2DNet model (Chao et al, 2021): [Deep learning predicts cardiovascular disease risks from lung cancer screening low dose computed tomography](#)
- Code & Containers for Reproducibility: [Docker](#) | [Singularity](#)
- RDA [Health Data Interest Group](#)
- Medical Imaging and Data Resource Centre ([MIDRC](#))

itwinai: Large-Scale ML on HPC for Digital Twins in Science- Matteo Bunino

Matteo presented interTwin, an EU-funded initiative developing an open-source Digital Twin Engine that enables large-scale simulations on High-Performance Computing (HPC) infrastructure. The platform integrates AI with persistent identifiers that connect datasets, software, and computing resources.

Scientific applications showcased include AI-driven simulations in high-energy physics (CERN), gravitational wave astronomy (Virgo), and Earth observation (wildfire prediction in ML4Fires). The yProvML library facilitates reproducibility by tracking data lineage in ML experiments across HPC environments. Challenges remain in orchestrating computations across heterogeneous systems and ensuring reproducibility in distributed AI training.

Notable links & references from Matteo's talk (checked and last accessed on 15 March 2025):

- itwinai: <https://itwinai.readthedocs.io/latest/>

- interTwin Project: <https://itwinai.readthedocs.io>
- yProvML GitHub Repository: <https://github.com/HPCI-Lab/yProv>
- ML4Fires – Wildfire Prediction AI: <https://github.com/CMCC-Foundation/ML4Fires>
- Hython- [Python package for distributed hydrological models](#)

AI Tools for Research: A Tried-and-Tested Workflow in Digital Humanities- Saloni Chaudhary

Saloni Chaudhary presented her AI-driven bibliometric workflow for studying research collaboration in India. Using Web of Science (WoS), she collected metadata from academic publications and leveraged Elicit, an AI-powered research assistant, for literature discovery. To analyse patterns in research networks, she applied topic modeling (LDA) for text classification and VOSviewer to visualize co-authorship and keyword co-occurrence networks.

Her findings highlight AI's potential to accelerate literature review and uncover hidden connections in large research corpora. However, she stressed that AI outputs require careful human validation, as models can generate misleading or incomplete results. She emphasized the need for training humanities scholars in AI tools to enhance research efficiency and ensure proper interpretation of computational insights.

Notable links & references from Saloni's talk (checked and last accessed on 15 March 2025):

- Web of Science (WoS): [Web of Science Database](#)
- Elicit – AI for Literature Review: [Elicit AI Research Assistant](#)
- VOSviewer – Bibliometric Network Analysis: [VOSviewer Tool](#)
- Topic Modeling (LDA): [Gensim Python Library](#) | [MALLET for NLP](#)
- Scientometric Research & Collaboration Analysis: [CWTS Bibliometric Studies](#)

The speakers received questions orally and via chat, and we will summarise the participant interactions and discussions in the following section.

Discussion summary

The key discussion topics of the 13 Feb session included reproducibility, role and relevance of FAIR principles, trust and verification of AI outputs.

The session featured lively discussion on AI reproducibility and FAIR principles, particularly around whether the "R" in FAIR should be extended to FAIRR to include Reproducibility. Some participants opined that FAIR (Findable, Accessible, Interoperable, Reusable) ensures reproducibility by making data and code accessible, while others contended that it does not explicitly enforce methodological reproducibility in AI and machine

learning. Further discussions clarified that while FAIR principles support reproducibility, they do not guarantee it, as reproducibility also depends on software dependencies, data versioning, and documentation practices.

A related conversation ensued around "reproducibility vs. replicability", where attendees questioned whether reproducing AI results requires access to the same training data or if it is sufficient to replicate results using similar methodologies. Some argued that in medical AI, full reproducibility is often impossible due to patient data privacy laws, while others pointed out that open datasets like NLST (National Lung Screening Trial) enable at least partial validation of medical AI models.

Fairlyze, a platform developed to support the ethical implementation of AI Data Visitation (AIDV) policies in research was mentioned in discussions about AI compliance with FAIR and Reproducibility principles. Some participants sought details on how Fairlyze assesses AI model reproducibility. As a follow-up, Seonyoung Kim, an active member of RDA's Artificial Intelligence and Data Visitation ([AIDV-WG](#)) group provided more information and on behalf of her group, invited community members to become testers of FAIRlyz. The platform enables users to run quality control on datasets using OpenAI tools and contribute feedback on policy effectiveness and usability. This initiative supports responsible AI innovation by shaping AIDV frameworks that are practical and beneficial for the research community. Interested testers are encouraged to sign up and actively participate in shaping the future of AI-enabled data governance. [Link to Fairlyze and full call for testers.](#)

A notable concern raised in chat was the gap in ML reproducibility standards across biomedical journals. It was pointed out that most journals lack clear AI/ML reporting guidelines, leading to inconsistencies in how models are evaluated. Participants debated whether standard checklists, such as DOME (Dataset Optimisation, Model Evaluation), should be required by publishers and reviewers. A few attendees highlighted that certain fields, like bioimaging and drug discovery, are moving towards structured model validation protocols, but these are not yet widely enforced.

Another key discussion focused on AI model audits. Questions were raised on how research teams can document experiments to ensure future reproducibility, particularly when models rely on external APIs and third-party training datasets. Some attendees suggested using containerisation (Docker, Singularity) and environment snapshots, while others warned about long-term availability issues if dependencies are not carefully archived.

In the next section we will describe the 18 Feb 2025 session featuring participants and speakers from Asia and Oceania.

Session 2: 18 Feb 2025 (Asia/Oceania)

Research Use of LLMs: Beyond the Chatbot- Brian Ballsun-Stanton

LLMs are often seen as chatbots, but they have far greater potential in research. The discussion explored how LLMs can assist in literature review, structured summarisation, academic writing, and indexing. A structured system prompt approach was highlighted as essential for getting reliable and meaningful AI-generated content.

A case study, “Calculator for Words”, demonstrated how structured LLM prompting enabled students to synthesise insights from 160 sources, producing a well-organised bibliography with narrative explanations. This showcased LLMs' potential for structured reasoning beyond simple text generation. Transparency and documentation were key themes, with Brian advocating for publishing AI-generated outputs and system interactions to allow for verification. His team makes LLM transcripts openly available via the Open Science Framework (OSF).

The talk also addressed prompt engineering beyond the chatbot interface. He gave an example of using the Claude API for academic copy-editing, instructing the model to improve clarity while tagging its changes transparently. The session also covered ethics and policy, focusing on responsible AI use. Brian stressed the need for disclosure when using AI-generated content in research, advocating for standards similar to software citation. The discussion emphasised how institutions should create clear policies around LLM documentation, reproducibility, and auditability.

Notable links & references from Brian's talk (checked and last accessed on 15 March 2025):

- Caligula Bibliography Project :LLM-assisted research aggregation ([Zenodo record](#))
- OSF Transcripts ([osf.io/kn2aq](#), [osf.io/4b7sd](#), [osf.io/8dxj6](#))
- Book Index Generator : LLM-generated book indexes ([Denubis/Book-Index-Generator](#))
- Brian's [Zenodo resources](#)

Leveraging AI for Research Classification- Amir Aryani

Amir Aryani presented Research Link Australia (RLA), a national platform designed to classify research outputs and map collaborations using AI. A core component of RLA is the National PID Graph, which integrates Persistent Identifiers (PIDs) such as ORCIDs, DOIs, and RORs to connect researchers, institutions, and datasets. This structured knowledge graph serves as a foundation for AI-driven research classification.

RLA employs AI models for topic classification and capability mapping, automatically tagging research publications and researchers with Fields of Research (FoR) codes. To enhance accuracy, the team experimented with a large-scale language model (~22 billion parameters) for zero-shot classification. Ensuring

precision, RLA implements human-in-the-loop validation, where subject matter experts review and refine AI-generated classifications.

Validated outputs enable collaboration network analysis and expert discovery, helping policymakers and industry partners identify researchers in specialised fields. RLA also supports projects in health, climate science, and agritech, where structured research intelligence accelerates innovation. Ongoing challenges include data integration, identity disambiguation, and ensuring AI remains transparent and bias-free. Future development focuses on network embeddings for predicting research trends and refining governance frameworks for sustainable AI-driven classification.

Notable links & references from Amir's talk (checked and last accessed on 15 March 2025):

- Research Link Australia (RLA) – ardc.edu.au/project/research-link
- Persistent ID (PID) Graph – researchgraph.org
- Fields of Research (FoR) Classification – abs.gov.au/statistics/classifications/anzsrc

Healthy Connections: An AI-Driven Mobile Medi-Kit for Remote Health Equity- Susannah Soon

Susannah Soon's team introduced Healthy Connections, a project designed to bridge healthcare gaps in remote and rural areas, particularly for Indigenous communities with limited access to medical services. Their proof-of-concept solution, a mobile Medi-Kit, combines AI-driven diagnostics, medical IoT devices, and telehealth to bring clinic-level care directly to underserved populations, with a focus on chronic disease screening and prevention for cardiac, diabetes, and renal conditions

The Medi-Kit is a portable, rugged case containing sensors and diagnostic tools, including vital signs monitors, and point-of-care testing devices. A key innovation is the multilingual AI Virtual Health Assistant, which guides non-expert users, such as community health workers or patients, through medical assessments, symptom analysis, and preliminary diagnostics, real-time triage and health education. AI is used to generate health and wellbeing recommendations such as diet and exercise that patients can access on a patient companion app. The project also incorporates LLM analysis of ECG waveforms to generate reports for clinicians to assess. Computer Vision is also used to detect skin conditions such as scabies, eczema or wounds, and throat conditions Strep A. Additionally, the Medi-Kit supports telehealth via satellite or mobile networks.

The project has conducted extensive co-design with stakeholders including clinicians and community. The project's user-centred design reflects cultural appropriateness, incorporating Indigenous artwork on the case and considering local healthcare workflows. The team has also worked in language translation using LLMs to provide better access to medical advice and information to rural and remote patients that may speak English as their third or fourth language.

Scalability is a key focus, with ongoing model refinement for Indigenous populations, and improved privacy and security measures. The project aligns with global health equity initiatives and has the potential for deployment in disaster relief and other resource-limited settings.

Notable links & references from Susannah's talk (checked and last accessed on 15 March 2025):

- Healthy Connections Project – research.curtin.edu.au/healthy-connections
- Indigenous Health Equity & Digital Health Initiatives – aihw.gov.au

Note

Liu Chang's talk was interrupted by technology issues. While her talk could not be fully delivered, she kindly shared her [slides](#) on the topic of “Artificial Intelligence for Geographical Indications Environment & Sustainability”

Saloni Chaudary participated in both sessions as a speaker with the same content so we reported her presentation only once.

The speakers received questions orally and via chat, and we will summarise the participant interactions and discussions in the following section.

Discussion Summary

This session explored the complexities of transparency in AI research, model and data drift, sensitive data management, small vs large models, and the FAIR and CARE principles.

Insights around transparency in LLM driven research centered on how system prompts and AI-generated content should be documented. Some participants asked how researchers should acknowledge AI-generated contributions and whether there should be formal citation guidelines.

There were also concerns about AI model size and validation. There were questions regarding the use of larger vs smaller models in terms of performance and accuracy. While larger models were more generalisable there was a lot of potential for smaller models to perform well with fine tuning with good quality datasets. This also reduces computational costs and increases access to the technology. Others questioned how AI-driven research classification systems could be improved, particularly in avoiding false categorisations and ensuring proper human oversight in the process.

A key topic was data and model drift, with questions about how AI models remain accurate as data changes over time. Participants asked how often models should be retrained and what best practices exist to ensure that updates do not compromise reproducibility. There were also questions on versioning AI models, particularly when different research teams use different iterations of the same model.

The handling of sensitive data was also discussed, with participants asking how health-related AI models ensure data security and patient confidentiality. There were concerns about whether AI models trained on historical medical datasets could unintentionally reinforce outdated biases. In the context of health AI in remote

and Indigenous communities, questions were raised about how CARE principles (Collective Benefit, Authority to Control, Responsibility, Ethics) were being applied alongside FAIR. Some asked whether data sovereignty was being respected and how communities could retain control over their own health data rather than it being solely governed by external institutions.

Participants also asked about trust in AI outputs and how researchers should critically evaluate AI-generated recommendations. Some questioned whether AI suggestions in clinical settings should be taken at face value or whether strict validation steps should always be in place. Others wanted to know how bias in AI-generated insights could be detected, particularly when using pre-trained models where training data might not be fully disclosed.

These discussions highlighted concerns about ensuring AI remains transparent, ethically managed, and reproducible in research, particularly when applied to sensitive domains like health and research classification.

Following the two workshops and live feedback from participants, we also distributed a post-session survey to gather additional input. Incorporating this feedback, along with aligning it with RDA's initiative on Artificial Intelligence, we will conclude the report with key insights from the feedback and next steps.

Feedback and Next Steps

Live feedback from the chat during both sessions reflected active engagement and a keen interest in the practical applications of AI and the discussions have been reported in the session summaries of the report. In this section we analysed the feedback from the post-session surveys sent to all participants.

The feedback analysis is based on responses from 12 participants for the 13 February session and 6 participants for the 18 February session. Given the small sample sizes, we provide a qualitative assessment rather than a quantitative analysis.

The workshop was generally well received, with most participants stating that it met or exceeded their expectations. The aims and objectives were clear, and the agenda was structured effectively and provided rich content. The breadth of topics covered was appreciated, though some felt that a more focused discussion on fewer areas would have allowed for deeper engagement.

Many attendees valued the opportunity to explore different AI research practices, particularly in relation to generative AI, FAIR principles, and reproducibility. The speakers were commended for their clarity and expertise. The event facilitated meaningful networking, with attendees learning about the work and interests of others in the community.

There were calls for a more concentrated focus on specific topics in future sessions, particularly around AI's limitations, such as bias, outdated data, and ethical and privacy considerations in AI research. Many participants expressed interest in gaining practical guidance on AI research methodologies, software solutions, system architecture, and AI model implementation strategies. Some requested greater attention to AI's role in humanities and social sciences research, as these fields are often underrepresented in such discussions.

Almost all respondents expressed a desire to attend similar workshops in the future, demonstrating a strong demand for continued discussions in this space.

Next steps

Given the transformative role of AI in research and its intrinsic relationship with data, and recognising the importance of valorising the expertise and engagement of the RDA community members, the RDA Secretariat has outlined the following plans of action to continue fostering meaningful dialogue and community-driven initiatives in this space. The following proposed initiatives are not limited to these and will be adapted based on evolving community needs and interests:

- Two major roundtable discussions bringing together global community experts and stakeholders, co-organised by RDA and Microsoft, have already been scheduled as part of this ongoing effort: **Data Readiness and Data-Centric AI Skills** (May 2025) and **From HPC to Quantum: Shaping the Future of Research Computing** (June 2025), with outcomes including joint white papers and community recommendations
- the Secretariat plans to propose two Birds of a Feather (BoF) sessions at IDW2025 (International data Week October 2025), based on community insights from the State of AI workshops held in September 2024 together with the feedback from the February 2025 AI in Action workshops, as well as input from the RDA-Microsoft roundtables' white papers.
- to identify the topics of the BoFs, the RDA Secretariat will deep dive in to the main focus areas identified from both workshops including Data Readiness for AI, Sensitive Data Management, Sustainable Infrastructure and Policy together with Reproducibility in AI research workflows, and the application of AI in Humanities and Social Sciences.
- a follow-up workshop series is envisaged to take place in Q4 2025 to consolidate insights, address emerging priorities, and promote best practices.
- the RDA Secretariat will collate and curate all relevant resources, outputs, and outcomes from these activities in a dedicated page on the RDA website, titled "The Value of RDA for Artificial Intelligence," to serve as a central hub for community knowledge and ongoing initiatives in this area.
- RDA AI community experts featured in the RDA events related to AI will also be added to the dedicated AI page by July 2025
- Finally, the Secretariat will facilitate a community consultation process to assess interest in forming a Community of Practice (CoP) on AI, providing an ongoing platform for collaboration and knowledge exchange.

Acknowledgements

The development and delivery of the AI in Action workshops, along with the preparation of this report, would not have been possible without the support of the outstanding RDA Secretariat team.

I would like to extend my sincere thanks to Hilary Hanahoe for her foresight and leadership in recognising the emerging relevance of AI and its intersection with the RDA community, and for initiating the design and strategic direction of this initiative.

Special thanks to Irina Hope for her logistical support, often beyond regular working hours. I am grateful to Bridget Walker for her operational expertise and, just as importantly, her invaluable emotional support.

A heartfelt thank you to Trish Radotic for her superb efforts in securing and coordinating the participation of the brilliant Oceania speakers, ensuring global representation. I would like to thank Rosie Allison for her thoughtful and efficient communication support, making it easier for the community to engage with these activities. Finally, my appreciation goes to Connie Clare for her careful review of this report and her commitment to carrying this work forward in the next phase of RDA's AI initiatives.

Most importantly, my thanks to all the speakers, participants, and community members who contributed their time, expertise, and insights to make the AI in Action workshops a meaningful and valuable experience for the global research community.

About the Research Data Alliance

The [Research Data Alliance \(RDA\)](#) was launched as a community-driven initiative in 2013 with the vision that researchers and innovators can openly share and re-use data across technologies, disciplines, and countries to address the grand challenges of society.

The RDA's mission is to build the social and technical bridges that enable that vision, accomplished through the creation, adoption and use of the social, organisational, and technical infrastructure needed to reduce barriers to data sharing and exchange. Scientists & researchers join forces with technical experts in focused Working Groups, exploratory Interest Groups and Communities of Practice. As of March 2025, the RDA is a 15000-member strong community spread across the globe.

Everyone is welcome to join the RDA as an individual member free of charge. If you're interested in getting involved at the organisational level, please explore our [membership options](#).

Annex: Speaker bios

In alphabetical order

Amir Aryani

Swinburne University of Technology, Australia

ORCID: [0000-0002-9994-1462](https://orcid.org/0000-0002-9994-1462)



Amir Aryani leads the Social Data Analytics (SoDA) Lab at Swinburne University of Technology. The lab applies data analytics techniques for insights into health and social challenges. His expertise is in data modelling, information retrieval techniques and real-time data analysis. In other capacities, Amir leads the Swinburne node for the Centre for Information Resilience (CIRES), an Australian Research Council (ARC) Industrial Transformation Training Centre funded 2021-2026.

Brian Ballsun-Stanton

Macquarie University, Australia

ORCID: <https://orcid.org/0000-0003-4932-7912>



Dr. Brian Ballsun-Stanton leads Generative AI policy, research, and education at Macquarie University's Faculty of Arts, shaping AI integration across disciplines. With a PhD in Philosophy of Data, he bridges theory and practice, advising organizations like IEEE Silicon Valley and NBNco on AI policy. He has delivered AI-enhanced learning initiatives, led multi-million-dollar research projects, and developed 70+ field data collection modules through the FAIMS Project. A recognized expert in Digital Humanities and AI policy, he has 18+ publications and is a key figure in AI-driven education and research.

Carlos Utrilla Guerrero

TU Delft, the Netherlands

ORCID: [0000-0002-9994-1462](https://orcid.org/0000-0002-9994-1462)



Carlos is a trainer at TU Delft, Netherlands, with research interests in Semantic Web and open-source, working on how to increase the understandability of research software with AI.

Gavin Farrell

University of Padova and Elixir, Italy

ORCID: [0000-0001-5166-8551](https://orcid.org/0000-0001-5166-8551)



Gavin is a PhD candidate at the University of Padova and member of ELIXIR Italy. He is researching machine learning (ML) best practices in the life sciences to advance transparent and reproducible ML. He has a strong background in the life sciences with a BSc in biotechnology and an MSc in bioinformatics from the University of Galway, Ireland. After his initial studies he worked at the ELIXIR Europe Hub in Cambridge, UK - an organization focused on aligning pan-European bioinformatics efforts through collaborative projects such as the Genomic Data Infrastructure. He has been involved in international initiatives working to advance open science like the European Open Science Cloud and ELIXIR Research Data Alliance Focus Group.

Giulia Raffaella De Luca

Alma Mater Studiorum – University of Bologna, Italy

ORCID: [0009-0002-8156-4221](https://orcid.org/0009-0002-8156-4221)



Giulia graduated with a master's degree in Biomedical Engineering, specialised in 'Innovative Technologies in Diagnostics and Therapy' from Alma Mater Studiorum - University of Bologna in March 2023. Giulia is PhD student in the Health and Technology programme at the Department of Electronic and Information Engineering at the same university, and her research project focuses on the integration of artificial intelligence in lung cancer screening. When she's not at her PC, you'll find her baking, playing basketball, or walking by the seaside—while also keeping up with her annual reading challenge. She is also an active member of several IEEE student chapters.

Matteo Bunino

CERN, Switzerland

ORCID: [0009-0008-5100-9300](https://orcid.org/0009-0008-5100-9300)



Matteo is a Fellow for Artificial Intelligence at CERN openlab, where he co-leads a task within interTwin, an EU-funded project aimed at developing a framework for advanced AI workflows in science. His work focuses on scaling PyTorch and TensorFlow workflows for large-scale computing infrastructure (such as HPC systems) while integrating hyperparameter optimization through Ray Tune. Alongside technical contributions, he plays a key role in managing the project's work plan and strategy.

Liu Chang

Institute of Geographical Sciences and Natural Resources Research, China

ORCID: Unavailable



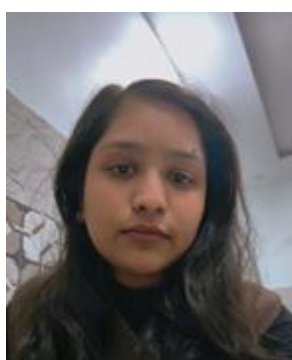
Dr. Chuang Liu is a Professor and Principal Scientist at the Institute of Geographical Sciences and Natural Resources Research, Chinese Academy of Sciences (IGSNRR/CAS) and Director of GCdataPR (World Data System). She also leads several global initiatives, including the FAO OCOP Organizing Group (Asia-Pacific) and CODATA Task Group on GIES. A pioneer in scientific data sharing, she led a CODATA international group (2002-2022), published 300+ digital datasets, and received multiple awards, including the CODATA Prize (2008). Holding a Ph.D. from Peking University,

she has held academic positions in China, Canada, and the U.S., contributing significantly to geographical data science and global environmental research.

Saloni Chaudhary

University of Delhi, India

ORCID: [0000-0002-5998-5794](https://orcid.org/0000-0002-5998-5794)



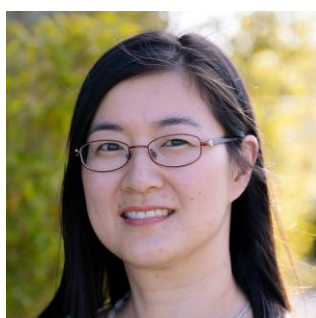
Saloni Chaudhary is an academic and researcher dedicated to the evolving landscape of Library and Information Science. She currently serves as an Assistant Librarian at the University of Delhi and holds a Ph.D. in Library and Information Science from Banaras Hindu University, Varanasi.

With a strong interdisciplinary background, she completed her undergraduate studies in Ancient Indian History, Culture, and Archaeology at Banaras Hindu University. Her research interests lie at the intersection of Scientometrics, Digital Literacy, and Digital Humanities, where she explores the impact of digital advancements on knowledge systems.

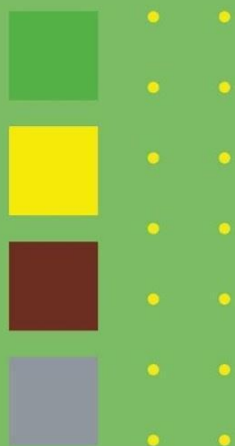
Susannah Soon

Curtin University, Australia

ORCID: <https://orcid.org/0000-0001-5091-8221>



Susannah is a Computing academic at Curtin University delivering multiple multi-disciplinary demand-driven research projects. She is an expert in AI, digital health, human-computer interaction, and product design and development. Susannah is the Academic Lead for the award-winning, high-profile, flagship Healthy Connections project developing a mobile Medi-Kit for health equity in the Pilbara. The Medi-Kit is AI-enabled and satellite-connected, allowing clinicians to take it on-Country to provide preventative health screening for chronic conditions in remote communities. In 2025, phase II of the project demonstrates immediate impact, with the project's 4WD mobile clinic visiting remote communities in the Pilbara demonstrating the Medi-Kit's potential.



research data sharing without barriers

rd-alliance.org