



WHITEPAPER

ARTIFICIAL INTELLIGENCE IN RESEARCH DATA MANAGEMENT

Expert Perspectives, Emerging Practices, and Recommendations

AUGUST 2026

ARTIFICIAL INTELLIGENCE IN RESEARCH DATA MANAGEMENT

Expert Perspectives, Emerging Practices, and Recommendations

Authors:

Ariana Dongus and Gereon Rahnfeld

INDEX

Summary 4

1. Introduction..... 5

2. Scope and Approach 6

3. AI as Part of Everyday Research Practice..... 7

4. Main Application Areas in Research Data Management..... 8

5. AI and Research Software Engineering 10

6. Benefits and Opportunities 11

7. Risks, Limits, and Tensions 12

8. Transparency, Reproducibility, and Documentation 15

9. Skills and Infrastructures Needed for Responsible Adoption 16

10. Epistemic Changes and Shifts in Knowledge Production..... 17

11. Recommendations 19

Summary

This white paper examines the role of artificial intelligence (AI) in research data management (RDM) on the basis of five exploratory expert interviews conducted in the NFDI¹ context. Across the interview, corpus interviewees agree that AI is already embedded in everyday research practice. Rather than replacing scholarly interpretation itself, AI is applied to help in support-oriented tasks, such as metadata enrichment, information extraction, retrieval support, documentation assistance, and bounded software-related work.

Having AI perform these tasks leads to a range of benefits as well as tensions. While the utilisation of AI in RDM, for example, increases the efficiency in repetitive and labour-intensive tasks, improves discoverability, semantic richness and structured reuse as well as supports multilinguality and cross-domain interoperability, it also causes opacity and incomplete traceability, a validation burden, as well as legal issues linked to data sensitivity and licensing conditions, amongst others.

Both benefits and tensions make new approaches in RDM necessary. Firstly, they necessitate an expansion of documenting practices in order to increase transparency and to strengthen reproducibility. This should not be taken to imply that existing norms in documentation practices have to be reinvented. Rather, they ought to be extended to include prompting, model choices, interaction histories etc. Secondly, there is a need for new skills and infrastructures. While AI systems should become sharable, governable, privacy-aware, and flexible, methodological and technical literacy is required in order to enable different actors to evaluate and correct AI. Finally, AI does not just change how research tasks are performed, but also how knowledge itself is produced, formulated, circulated, and evaluated. There is a need to further engage with what these epistemic changes as well as changes in the production of knowledge imply for RDM.

The white paper argues that the most promising path for AI in RDM lies neither in blanket adoption nor blanket rejection, but in selective and well-governed integration. Against this background, the white paper offers six recommendations: 1. prioritise AI where current value is clearest; 2. establish human-in-the-loop review as a baseline principle; 3. develop shared guidance for documenting AI use in data management and software engineering workflows; 4. invest in governable and, where appropriate, locally controllable AI infrastructure; 5. treat AI competence as a collective training issue; and 6. maintain methodological pluralism rather than treating large generative models as the default answer to every task.

¹ NFDI is the short form for Nationale Forschungsdateninfrastruktur (English: National Research Data Infrastructure). It is an initiative in Germany whose aim is to collect and link research data systematically from all scientific disciplines and make it permanently accessible. For more information visit <https://www.nfdi.de/?lang=en>.

1. Introduction

Artificial intelligence (AI), and especially the rapid diffusion of large language models, has become a visible and increasingly consequential part of academic work. What only recently appeared as an experimental add-on has, in many contexts, already become embedded in everyday scholarly routines: searching for literature, summarizing information, drafting text, generating code, enriching metadata, and navigating large and heterogeneous information spaces. The launch of specific AI applications intended to support scientific research, like Claude Science², illustrates this development well. The question is therefore no longer whether AI has entered research practice, but rather: how it is doing so, where its use is genuinely valuable, and under what conditions it can be integrated responsibly.

This question is especially relevant in contexts like research data management (RDM), where technical feasibility, scholarly standards, disciplinary diversity, and public responsibility intersect. Here, AI matters less as an abstract trend than as a practical issue: can it make data more findable, reusable, and interoperable; reduce repetitive metadata work; improve access to complex repositories and knowledge systems; and support better software practices where time and technical resources are limited?

This white paper offers a grounded synthesis of several expert perspectives on three connected questions concerning the application of AI in the field of RDM: where is AI currently being used in RDM and Research Software Engineering (RSE)? What practical benefits and recurring tensions are visible across these use contexts? And what kinds of infrastructures, competencies, and governance practices appear necessary for responsible and effective adoption?

The white paper supports neither uncritical embrace nor wholesale rejection. Instead, it argues that AI is already beneficial in focused, assistance-oriented tasks, although its application is not without risks. To balance both benefits and risks, and to accommodate the needs which result from them, the white paper offers six recommendations concerning the utilisation of AI in RDM:

1. prioritise AI where current value is clearest;
2. establish human-in-the-loop review as a baseline principle;
3. develop shared guidance for documenting AI use in data management and software engineering workflows;
4. invest in governable and, where appropriate, locally controllable AI infrastructure;
5. treat AI competence as a collective training issue; and
6. maintain methodological pluralism rather than treating large generative models as the default answer to every task.

² <https://www.technologyreview.com/2026/06/30/1139987/claude-science-is-anthropics-newest-flagship-product/> (accessed: 16 July 2026)

2. Scope and Approach

This white paper has been created in the context of the “National Research Data Infrastructure for and with Computer Science” (NFDIxCs)³ project. It is based on five exploratory expert interviews conducted with six interviewees between October 2025 and February 2026. All of the interviewees belong to NFDI, each of them being part of one or several NFDI consortia. The white paper is therefore informed by NFDI and its different participants rather than being representative of the project in general.

Methodologically, the paper treats the interviews as an exploratory qualitative corpus rather than as isolated testimonies. The analysis is organised thematically rather than interview-by-interview, with the aim of identifying recurring patterns, converging judgements, productive differences in emphasis, and especially the tensions that emerge across the material. The goal is not statistical representativeness but a structured overview of the current situation as seen from expert practice across several related environments.

The small but analytically rich corpus of interviews brings together perspectives from several research-infrastructure settings in which AI is relevant in different but overlapping ways: scholarly search and exploration, language data and corpora, metadata production and interoperability, graph-based and structured information access, code generation and research-software support, and semantic extraction from scientific literature. Taken together, the material captures a range of environments in which AI intersects with data-intensive work, technical workflows, and the challenge of making large and complex information spaces more usable.

The structure of the white paper is guided by the following subtopics: the current role of AI in research practice, its use in research data management, its impact on research software and modes of knowledge production, the benefits as well as tensions created by its application in these fields, questions of trust, reproducibility, and transparency, as well as future expectations regarding skills and infrastructures. The material therefore captures both concrete applications and the broader conditions under which AI is becoming embedded in scholarly work. The white paper ends by formulating six recommendations regarding AI in RDM.

³ <https://nfdixcs.org/>

3. AI as Part of Everyday Research Practice

One of the clearest findings across the interview corpus is that AI is already part of everyday research practice. This does not mean that research has become fully AI-driven, nor that all scholarly tasks are being delegated to generative models. It does mean, however, that AI has become normalised as a layer of practical assistance in routine and semi-routine activities such as paraphrasing, summarisation, search support, metadata support, extraction from text, code generation, data-format transformation, and conversational access to information that would otherwise require more laborious navigation.

“When I write a paper, I might use it for drafting. I might have ideas of how the paper should be structured and what the topics are. But I might be stuck with this blank page thing. And what I find helpful is having some AI tool write drafts and then iterate on them.”⁴

This normalisation is especially evident where large volumes of textual material, descriptive information, or technical documentation must be handled. In such settings, the use of both classical Natural Language Processing (NLP) tools and newer Large Language Model (LLM) – based capabilities is framed less as optional innovation than as a practical necessity for making corpora, archives, and related resources searchable, visible, and reusable at scale.

A similar pattern appears in discussions of structured information and portal-like environments. Here, AI is explored not primarily as a drafting tool, but as a means of improving access to semantically organised information systems whose formal representations, query languages, or internal logic would otherwise create high barriers for many users. Natural-language interaction is valued because it can lower those barriers without abandoning structured organisations.

Several interviews also suggest that AI is reshaping the sequence of technical work itself. In software-related contexts, users increasingly begin with AI-generated suggestions and then revise them, rather than always producing a first draft manually. In search and exploratory contexts, they move more often from result-list navigation to asking systems to summarise, compare, or extract salient points directly.

At the same time, the corpus does not present this integration as a simple success story. Its value depends strongly on how AI is used and by whom. A recurring distinction is between using AI as support after the underlying thinking has already been done and using it as a substitute for thinking itself. That distinction matters both in writing and in coding, where the benefits of AI remain closely tied to prior expertise.

Taken together, the interviews suggest that AI is already embedded in everyday research practice, but mostly in assistive rather than fully delegated forms. The dominant pattern is one of augmentation: AI supports searching, drafting, coding, annotating, transforming, or accessing information, while human actors remain responsible for interpretation, validation, and contextual judgement.

⁴ Interview: SA_03:54.

4. Main Application Areas in Research Data Management

When asked about the current practical value of AI, the interviewees' answer suggest that it is not autonomous science or fully automated interpretation. Instead, they emphasise a narrower but highly consequential set of use cases within RDM: metadata generation and enrichment, extraction from unstructured material, improved findability and retrieval, and the facilitation of access to complex or heterogeneous data environments.

A first major application area is metadata generation and enrichment. Multiple parts of the interviews describe AI-supported metadata annotation in infrastructural terms: data deposited in repositories, archival systems, or digital research environments must be annotated in a structured and semantically rich manner if they are to remain interpretable, interoperable, and reusable across contexts. AI is valuable here because it can extract terms from textual descriptions and map them onto controlled vocabularies or ontologies, supporting more precise and machine-actionable reuse.

Other parts of the interview corpus reinforce this point from the perspective of large text collections and long-term archives. Here, metadata enrichment includes topic annotation, named entity recognition, keyword extraction, and linking to existing knowledge bases. The purpose is not only descriptive completeness, but better filtering, searching, and the construction of meaningful subcollections from otherwise unwieldy holdings.

“One thing where AI can help is about the metadata part. I think that has been always the biggest problem, that we know that we should generate metadata, but we are all too lazy to do it. I think ... LLMs, they are quite, quite helpful in producing this metadata.”⁵

A complementary observation helps explain why metadata is such a recurring topic: it has long been recognised as necessary, but often remains incomplete because researchers and institutions do not consistently produce it. From this perspective, AI is valuable because it addresses one of the chronic weaknesses of research data management – the gap between acknowledged best practice and actual practice under time and labour constraints.

A second major application area is information extraction from unstructured material, especially full text. The interviews describe concrete workflows in which AI is used to extract semantic concepts from scientific publications and turn them into structured metadata that can support comparison, linking, and discovery. Framed this way, AI addresses a clear scalability problem: manually extracting relevant aspects from large numbers of papers is labour-intensive and difficult to sustain.

This concern with extraction connects to a broader pattern in the corpus: AI is repeatedly valued where it helps transform textual abundance into more usable and better-structured forms of scholarly information. Whether through generated term sets, extracted answer elements, or summaries of text-based records, the underlying logic is similar – AI reduces the friction involved in moving from unstructured text to semantically or procedurally more tractable forms.

A third major application area is findability and retrieval. Several interviews indicate that AI can improve not only the generation of metadata, but also the practical retrieval of information from existing systems, especially where keyword-based search proves insufficient. Natural-language querying is valued because it can provide a more direct route into large information environments and support sensemaking in technically sophisticated or domain-specific settings.

⁵ Interview: AD_05:02.

Closely related is the facilitation of access to complex knowledge environments through natural-language interfaces. In infrastructure contexts, usability is not a superficial concern; it is part of whether resources can actually be used by broader scholarly communities. The interview corpus does, however, introduce an important nuance: text-based summarisation and comparison currently appear more reliable than direct interaction with graph-structured or otherwise highly formalised information.

5. AI and Research Software Engineering

The interview corpus suggests that AI is not only changing RDM but, in connection with this, RSE practice as well. However, the interviewees do not hold that AI eliminates the need for software expertise. Rather, AI can accelerate many aspects of research software work, but the extent to which this acceleration produces robust results depends heavily on prior knowledge, review capacity, and the specific kind of software task involved.

One of the clearest recurring observations is that code is increasingly being generated and then reviewed rather than written entirely from scratch. This change matters because it shifts the sequence of software production itself: the first draft of a solution may now come from the model, while the human role moves more strongly toward evaluation, correction, and contextual adaptation. The benefits are therefore real but uneven, accruing most reliably to those who already possess enough software understanding to diagnose and revise what the model produces.

This distinction is especially important in research settings, where software often emerges under constraints very different from those of production environments. The interview corpus suggests that research software is frequently more fragile, less thoroughly maintained, and less systematically aligned with established engineering standards. From that perspective, AI can offer a meaningful benefit by helping with routine coding tasks, baseline quality improvements, and the conversion of one-off research code into something more accessible, traceable, and reusable.

The interviews also show that AI is particularly helpful where software work is repetitive, procedural, or transformation-oriented. Concrete examples include generating scripts for data-format conversion, enrichment, or reproducibility-related reconstruction, as well as assistance with software-publication requirements such as citation files and repository preparation.

These are significant use cases because they connect AI not only to coding convenience, but to the broader ecosystem of research software quality and documentation.

Another pattern in the interviews is that AI-assisted coding affects not only software artefacts themselves, but also the documentary and communicative practices around them. AI is used for validation support, phrasing assistance, and the generation of commit descriptions or other forms of metadata for code. This widens the significance of AI in software contexts beyond code generation alone.

“For software developers that are really on a more senior level, they could use it much more efficiently. But being more on a junior developer level, sometimes it was overwhelming and I just couldn’t grasp the idea.”⁶

The result is a dual conclusion: AI can clearly increase productivity in research software engineering contexts, but those gains remain expertise-dependent. Where conceptual and practical knowledge is weak, AI may speed up output while reducing understanding. That is particularly risky in research environments, where software is often part of how results are produced, interpreted, and reproduced.

In short, AI in RSE currently appears most promising as a means of acceleration under conditions of competence. It can help researchers and infrastructure teams move faster, fill in routine gaps, and improve the baseline quality of code-related work, but its value remains tied to a human capacity to understand, test, and revise what has been generated.

⁶ Interview: PS_12:35.

6. Benefits and Opportunities

Although the interviews are often cautious in tone, they are by no means pessimistic. The corpus identifies a range of areas in which AI already provides meaningful benefits, but these benefits are rarely framed in abstract or celebratory terms. AI is valued where it addresses specific forms of friction in RDM: repetitive tasks, overwhelming quantities of text, barriers to access, weaknesses in metadata practice, or the low-quality software conditions under which many research projects operate.

A first major benefit is efficiency in repetitive and labour-intensive tasks. Once datasets, corpora, or documentation demands reach a certain scale, some form of automated support becomes necessary. What newer AI approaches add is greater flexibility and, in some cases, improved precision for tasks such as clustering, tagging, annotation, keyword extraction, or minor drafting support.

A second major benefit is the improvement of discoverability, semantic richness, and structured reuse. Across the interview corpus, AI is valuable because it helps transform incomplete metadata into richer, more searchable, and more interoperable forms. This includes term extraction, ontology mapping, topic annotation, named entity recognition, entity linking, and keyword extraction. The benefit is therefore infrastructural as much as local: AI helps create the conditions under which research materials can be found, connected, interpreted, and reused more effectively.

“I think we’ll eventually reach a point where we can have conversations with datasets.”⁷

A third benefit concerns access to complex information environments. Natural-language interaction can lower barriers for users who are not fluent in formal query systems, graph standards, or technically dense portal structures. In infrastructure settings, that matters because a technically rich system may still fail if it remains too difficult for its intended communities to use productively.

A fourth benefit is support for multilinguality and cross-domain interoperability. The interview corpus suggests that AI-supported metadata work can help reduce linguistic and disciplinary barriers by extracting, translating, and mapping relevant terms in ways that remain semantically anchored.

A fifth and somewhat different benefit lies in the support of ordinary scholarly practice, even where AI is used only in modest ways. Examples include acronym generation, grammatical correction, paraphrasing, presentation support, and other small reductions of friction that, cumulatively, can still matter in everyday academic work.

⁷ Interview: FB_47:34 (translated by the author).

7. Risks, Limits, and Tensions

These positive findings are offset by the tensions created by the application of AI in RDM. Across the interviews, concerns about AI are not presented simply as occasional mistakes or familiar references to hallucinations. Rather, they point to a more substantial set of risks involving infrastructure, governance, epistemic standards, and the changing conditions of academic work.

One of the strongest and most explicit concerns is dependence on proprietary providers. Several interviewees describe this as a fundamental infrastructural challenge. The most capable models are often proprietary, leaving publicly oriented infrastructure projects dependent on private companies whose future decisions, technical changes, and value alignments remain outside the control of the research community. This is not only a matter of convenience or cost. It raises deeper questions about whether infrastructures intended to serve open and public-oriented scholarship should become dependent on systems whose design logic and long-term governance are opaque. The issue is therefore not merely technical dependence, but also the question of alignment: whether the assumptions, incentives, and priorities embedded in large commercial systems are compatible with the values of open scientific infrastructure.

A related concern, albeit from a more operational perspective, appears in several interviews. The issue is not only one of political or normative dependency, but also the practical difficulty of using external commercial systems in contexts where data sensitivity, licensing conditions, or institutional expectations require more control over processing. In some cases, data must remain within specific legal or national jurisdictions, or be deleted after processing, which makes reliance on commercial services problematic. In other cases, smaller, self-hostable models are seen as especially relevant in privacy-sensitive settings. Together, these observations suggest that questions of infrastructure sovereignty and operational control are likely to become more, not less, important as AI becomes more deeply embedded in RDM workflows.

“When work is taken off your hands as a researcher, you’re usually happy at first. But actually, what should follow is a thorough check to see if everything is correct. ... And, in my view, there’s a high risk that this doesn’t always happen.”⁸

A third recurring tension concerns validation burden and the persistence of human review. Several interviews converge on the idea that AI reduces one kind of labour while generating another. Tasks such as metadata annotation, extraction, and semantic mapping may become increasingly automatable, but the need for careful checking remains. The implicit risk is obvious: if AI makes such work fast enough to scale, institutions may be tempted to accept outputs without investing proportionately in review. Yet the same features that make AI attractive at scale also make unchecked adoption risky, particularly where semantic precision, documentation quality, or interoperability matter. The problem, then, is not that AI cannot support annotation or data-related processing. It is that automation may create incentives to under-resource the validation that still has to follow.

A closely related point emerges in connection with software and data workflows more broadly. The interviews suggest that, where AI is used in research processes, it is preferable for it to generate reviewable and inspectable outputs rather than opaque direct transformations. For example, code that can be read, debugged, rerun, and shared is fundamentally different from a black-box output that directly alters data or produces results without leaving a transparent procedural trace. This distinction is methodologically significant. Reviewable code or documented transformation steps preserve the possibility of inspection and correction, whereas direct AI outputs may be much easier to accept uncritically. The broader implication is that, where AI is

⁸ Interview: FB_15:54 (translated by the author).

integrated into research workflows, its contribution should be made as inspectable as possible, especially when it affects data processing, analysis, or technical implementation.

A fifth major concern is skill erosion through overdelegation. This issue appears repeatedly, both in relation to writing and to coding. Several interviews stress that writing is not merely a mechanical act of text production, but part of scholarly thinking itself. Delegating too much of that process to AI therefore risks weakening the researcher's own capacity to formulate, reflect, and judge. A comparable concern appears in discussions of code generation, where the benefits of AI are repeatedly linked to prior expertise. Those with strong foundations can use AI productively, while those without such foundations may become dependent on outputs they do not fully understand. In both cases, the risk is not simply incorrect output. It is the gradual relocation of intellectual work from a reflective human practice to a pattern of acceptance and post hoc checking that may never fully recover the lost understanding.

A sixth concern is opacity and incomplete traceability, especially in generative and proprietary systems. Several interviews emphasise that responsible scholarly use requires more than plausible-looking results. It requires at least some means of tracing the conditions under which outputs were produced. Yet the interview corpus also makes clear that existing practice is still provisional and might have limits. Current responses include documenting prompts, taking screenshots, recording settings, and repeating runs in order to report a range rather than a single exact value. These are pragmatic and useful strategies, but they also illustrate the current incompleteness of available standards. The research community is still developing documentation practices around systems that were not originally designed with scholarly traceability as a primary objective.

The same problem also appears at the level of model transparency more broadly. Even where more open models are becoming available, the interviewees note that they often still do not provide adequate documentation of training data, behaviour, or design choices. This makes it difficult to evaluate models not only in terms of performance, but in terms of provenance, embedded assumptions, and the implications of their training conditions. From the perspective of scientific practice, this is a serious limitation: if a model becomes increasingly involved in search, annotation, transformation, or drafting, then the inability to understand how it was built and what shaped its behaviour becomes more than a technical inconvenience. It becomes a constraint on accountability.

“Because LLMs can do everything and anything ... I kind of now shifted to skipping the user in the loop and replaced it to which I've called artificial data. So I feel like artificial data ... is the highest risk for the whole thing, right? You know, the tendency to just generate everything artificially because it's cheap.”⁹

A seventh and broader concern is that AI may contribute to the acceleration of academic production itself. Although this theme is less uniformly developed across the interviews than metadata or software quality, it appears strongly enough to matter. Parts of the interview material point towards a productivity dynamic in which AI makes it easier to produce texts, applications, reviews, and technical artefacts more quickly, while institutions struggle to absorb the resulting increase in volume. Even where the interviews do not formulate a full theory of systemic acceleration, the concern is visible: when both production and filtering become increasingly AI-mediated, the underlying issue is no longer only tool use, but the changing tempo and density of academic work. This concern is especially salient in contexts where limited human attention remains the bottleneck, regardless of how quickly text or code can be generated.

⁹ Interview: AD_23:14.

Finally, the interviews also point to more specific practical limits. Some interviews note that the use of large models for very extensive datasets can be costly in energy terms and may still require substantial evaluation, making careful tool selection essential. Others report that text-based summarisation and comparison currently work better than direct interaction with graph-structured or otherwise highly formalised information. These observations are important because they add a necessary degree of realism: AI is not equally suitable for all data types, scales, or infrastructural configurations.

Taken together, the risks and tensions identified in the corpus are substantial. They suggest that the central problem is not simply that AI can be wrong. It is that, under the pressure of convenience and productivity, research environments may normalise forms of work that are less transparent, less reviewable, more dependent, and potentially less skill-forming than the practices they partially replace. These risks point directly to the need for stronger documentation and reproducibility practices.

8. Transparency, Reproducibility, and Documentation

One of the strongest recurring themes in the interview corpus is that the growing use of AI in research data management and research software requires a corresponding expansion of documentation practices. Existing standards of transparency and reproducibility remain essential, but they now have to be extended to prompts, model choices, interaction histories, automatic annotations, and the distinction between machine-generated and human-generated contributions.

“I think it would be very interesting in the future to see. Okay, there was an article, they published something and the conclusion was completely wrong. What went wrong? And if you have a conversation history somehow that shows the prompts, how did you prompt the LLM? I think that could be helpful.”¹⁰

The interviews suggest that this challenge is already being addressed in practice, but only in partial and still evolving ways. Several parts of the interviews describe efforts to mark when annotations or other outputs were created automatically rather than manually, yet they also make clear that such practices are not fully standardised across infrastructures.

One particularly important distinction emerging from the corpus concerns the difference between documenting outputs and documenting processes. In generative settings, exact reproducibility may not always be achievable, but meaningful traceability can still be improved by recording prompts, settings where available, screenshots or output captures, repeated runs, and the points at which human review intervened.

The material also offers a concrete methodological principle for research software and data-related transformations: where possible, AI should generate reviewable procedural artefacts rather than uninspectable direct outputs. Code that can be inspected, debugged, rerun, and shared preserves the possibility of understanding and correction in a way that black-box transformations do not.

Another important pattern concerns disclosure and attribution. Several interviews indicate that transparency requires explicitly stating where AI tools have been used, what role they played, and where human verification took place. This is not only a compliance issue; it is also a way of preserving distinctions that may become difficult to reconstruct later if they are not recorded at the moment of production.

Documentation alone does not solve the problem of opacity, especially where adequate information about training data, model construction, or behavioural characteristics is missing. Yet the absence of complete model transparency makes workflow-level documentation even more important. If model-level transparency remains incomplete, then the conditions of use, the outputs generated, and the points of human intervention should be documented as carefully as possible.

Hence, the interviews suggest that responsible AI use in RDM requires an extension of existing norms rather than a reinvention of them. Good scientific practice still depends on transparency, justification, and the possibility of scrutiny; what changes is that those requirements now have to be applied more systematically to AI-supported processes as well.

¹⁰ Interview: AO_27:03.

9. Skills and Infrastructures Needed for Responsible Adoption

If the earlier sections identify both concrete opportunities and substantial tensions, the next question is what kinds of conditions are needed to make AI use in RDM effective and responsible. Here, the interviews offer a coherent answer: responsible adoption depends on both suitable infrastructures and the cultivation of skills that allow users to evaluate, correct, and situate AI outputs rather than merely consume them. On the infrastructure side, one of the clearest themes is the need for shared and governable access to AI systems. Several interviews imply that it is neither efficient nor desirable for every project or institution to solve model access, integration, and governance independently. More collective provision can improve efficiency, clarify responsibility, and reduce the proliferation of ad hoc or poorly documented solutions.

A second infrastructural requirement concerns local or self-hostable options, especially where privacy, data sensitivity, or long-term control matter. The corpus repeatedly suggests that responsible AI infrastructures should include models that can be run in-house or under institutionally controlled conditions, at least for some classes of use case.

A third infrastructural issue is the need for interfaces and integration modes that support different forms of scholarly work. Web interfaces alone are not enough. In RDM, the value of AI often lies in its incorporation into broader workflows for annotation, transformation, retrieval, and tool development, which makes programmatic access equally important.

On the side of skills, the interviews are equally clear that responsible AI use cannot be reduced to prompt literacy. Several interviews stress the need for a background understanding of the relevant methods, including more classical approaches in the field. Users who know only how to invoke the newest models but lack a basic understanding of earlier methods, evaluation criteria, or task-specific quality standards may be poorly positioned to judge whether outputs are actually good.

“If I really delegate almost all activities to the tool, ... I will kind of lose skills.”¹¹

The same principle appears in a broader epistemic form in discussions of domain knowledge. A central skill is the ability to recognise when AI is wrong, misleading, too generic, or simply not saying what the researcher intends to say. From this perspective, the critical skill is not merely using AI well, but correcting it well.

Another strong skill-related theme is the need to address knowledge gaps about limitations. Training needs to do more than teach people how to access or prompt models. It also has to explain where AI is likely to be unreliable, how grounding changes reliability, what should be double-checked, and why plausible language is not the same as trustworthy content.

Finally, the interviews suggest that educational and infrastructural resources should aim at equitable access. If only commercial premium tools are realistically available, the ability to use AI well may become tied to personal resources, institutional privilege, or uneven local support. A fragmented landscape of unequal access is likely to produce fragmented standards and inconsistent expectations in scholarly work.

Taken together, the interviews support a demanding but clear conclusion. Responsible AI adoption requires infrastructures that are shared, governable, privacy-aware, and technically flexible, as well as users who possess enough methodological, disciplinary, and technical literacy to evaluate and correct AI outputs rather than simply receive them.

¹¹ Interview: AD_15:36.

10. Epistemic Changes and Shifts in Knowledge Production

Beyond questions of workflow efficiency, metadata quality, software support, and infrastructure design, the interview corpus also points to a broader issue: AI is changing not only how research tasks are performed, but also how knowledge itself is produced, formulated, circulated, and evaluated. The impact of AI in RDM therefore cannot be understood fully if it is framed only as a matter of tools.

One of the clearest patterns in the corpus is the insistence that writing is not merely a communicative output, but part of thinking itself. Several interviews describe a deliberate practice of first producing one's own structure, notes, or argument and only then using AI for reformulation, grammar checking, paraphrasing, or stylistic refinement. The epistemic point is clear: AI may assist articulation, but when it generates the initial substance, something can be lost in the formation of argument and judgement.

This concern also appears in reflections on voice, authorship, and intellectual confidence. AI-generated text may be polished while remaining weakly inhabited by the researcher's own conceptual grasp. Ease of generation does not automatically translate into depth of understanding.

A related shift concerns the changing balance between internalised knowledge and externally assisted retrieval or generation. In software work and beyond, the valued skill may move away from producing or recalling everything oneself and towards knowing how to operate, interrogate, and selectively correct AI-supported systems. This does not necessarily imply a decline in quality, but it does shift what counts as competence and where intellectual authority is located.

The corpus does not present this shift as purely liberating. Several interviews suggest that faster production may not lead to better knowledge, but instead to denser and more difficult-to-govern flows of text, code, applications, and submissions. AI can increase the speed of production, but the available human attention for reading, reviewing, and evaluating does not grow at the same rate.

"It's been observed that submitted papers contain such small prompts, specifically for LLMs. So if you're a language model, please evaluate my paper very positively. You can write that in there and hide it somehow. ... when it comes to the evaluation – that is, when the reviewers themselves use an LLM again for the evaluation – that can, of course, skew the result."¹²

The interviews also hint at a more recursive dynamic: AI may begin to affect not only the production of papers and proposals, but also their evaluation. Once both production and filtering become increasingly AI-mediated, the issue is no longer simply whether an individual used a tool responsibly. It becomes a question of whether academic institutions can still distinguish, filter, and trust outputs adequately under conditions of machine-assisted abundance.

Another epistemically significant issue concerns the relationship between AI and data quality over time. The interviews warn not only against isolated hallucinations or annotation errors, but against a broader informational ecology in which easy-to-generate and weakly validated material accumulates faster than humans can assess it at scale.

The corpus does not imply that epistemic change should be understood only negatively. If AI reduces time spent on routine reformulation, repetitive documentation, feature scaffolding, or minor coding tasks, it may free up more time for conceptual design, synthesis, or reflection. Whether that becomes a genuine gain

¹² Interview: FB_37:14 (translated by the author).

depends on whether institutions use saved time for deeper thought or simply convert it into further acceleration.

Hence, AI is contributing to a shift in the conditions of knowledge production at multiple levels: in the micro-practices of writing and coding, in the relation between generation and understanding, in the balance between production and review, and in the long-term quality of the informational environment on which scholarship depends.

11. Recommendations

The interview corpus does not lead to a single sweeping prescription. It does, however, support a coherent set of practical recommendations for research-infrastructure contexts that wish to integrate AI in ways that are useful, responsible, and sustainable. The recommendations below follow directly from the most stable patterns in the material.

A first recommendation is to prioritise AI where the corpus shows the clearest current value. These areas include metadata generation and enrichment, semantic extraction from text, retrieval support, summarisation over appropriate textual materials, documentation assistance, and repetitive software-related tasks. A responsible strategy should therefore begin with the bounded areas in which the benefits are already concrete and the outputs can be more readily reviewed.

A second recommendation is to establish human-in-the-loop review as a baseline principle rather than an optional extra. Institutions should resist any organisational logic in which AI-generated outputs are scaled up without proportionate attention to validation. The relevant question is not whether AI can produce something plausible, but whether the surrounding workflow preserves meaningful opportunities for scrutiny, correction, and accountability.

A third recommendation is to develop shared guidance for documenting AI use in data and software workflows. Such guidance should include, where relevant, disclosure of AI-supported steps, documentation of prompts or interaction settings, indication of whether outputs were produced manually or automatically, preservation of reviewable code where AI contributed to analysis or transformation, and explicit marking of human verification.

A fourth recommendation is to invest in governable and, where appropriate, locally controllable AI infrastructure. The corpus is too consistent on provider dependence to ignore it. Research infrastructures should therefore avoid designing themselves into exclusive reliance on external commercial services wherever sensitive data, legal requirements, continuity of access, or scholarly accountability are at stake.

A fifth recommendation is to treat AI competence as a collective training issue rather than a matter of individual improvisation. The interviews make clear that effective use depends on background knowledge, quality awareness, critical judgement, and a clear understanding of limitations. These capacities cannot be assumed to emerge automatically from tool exposure.

A sixth recommendation is to maintain methodological pluralism rather than treating large generative models as the default answer to every task. In some contexts, classical or more specialised methods remain more efficient, more controllable, or simply better suited than large models. A mature approach asks which method is most appropriate for a given task given constraints of scale, quality, cost, energy use, and documentation burden.

Taken together, these recommendations point towards selective, governed, and methodologically disciplined integration rather than maximal automation. They suggest that the central task ahead is institutional shaping: deciding where AI belongs in RDM, how it should be documented, which infrastructures should sustain it, and which forms of competence should accompany it.

ABOUT THE GERMAN INFORMATICS SOCIETY (GI)

With more than 17,000 individual and 250 corporate members, the German Informatics Society (GI) is the largest professional society for computer science in the German-speaking world and has been representing the interests of computer scientists since 1969. With 14 specialist areas, over 30 active regional groups and countless specialist groups, the GI is a platform and mouthpiece for all disciplines in computer science.

GESELLSCHAFT FÜR INFORMATIK E. V. (GI)

Geschäftsstelle Bonn
Wissenschaftszentrum
Ahrstr. 45
53175 Bonn
Tel.: +49 228 302-145
E-Mail: bonn@gi.de

Geschäftsstelle Berlin
Weydinger Straße 14-16
10178 Berlin
Tel.: +49 30 240 09-805
Fax: +49 30 7261 566-19
E-Mail: berlin@gi.de

gs@gi.de
www.gi.de

 [/informatikradar](https://twitter.com/informatikradar)
 [/company/gesellschaft-fuer-informatik](https://www.linkedin.com/company/gesellschaft-fuer-informatik)
 [/net/gi](https://x.com/net_gi)