

Bausteine Forschungsdatenmanagement
Empfehlungen und Erfahrungsberichte für die Praxis von
Forschungsdatenmanager*innen

Bedarfsanalyse zur Entwicklung von Software und anderen Services für Forschungsdatenmanagement

Sascha Gruscheⁱ

2025

Zitiervorschlag

Grusche, Sascha. 2025. Bedarfsanalyse zur Entwicklung von Software und anderen Services für Forschungsdatenmanagement. *Bausteine Forschungsdatenmanagement. Empfehlungen und Erfahrungsberichte für die Praxis von Forschungsdatenmanager*innen* Nr. 2/2025: S. 1-20. DOI: [10.17192/bfdm.2025.2.8737](https://doi.org/10.17192/bfdm.2025.2.8737).

Dieser Beitrag steht unter einer
[Creative Commons Namensnennung 4.0 International Lizenz \(CC BY 4.0\)](https://creativecommons.org/licenses/by/4.0/).

ⁱTechnische Universität München. ORCID: [0000-0003-0588-6886](https://orcid.org/0000-0003-0588-6886)

Abstract

Beteiligte im Datenzyklus brauchen Software für Forschungsdatenmanagement (FDM), um Daten FAIR zu machen. Die Entwicklung bedarfsgerechter FDM-Software erfordert ein tiefes Nutzerverständnis. Zur Erfassung von Erfahrungen und Erwartungen im Datenzyklus wurden 15 Interviews mit Data Stewards und Forschenden durchgeführt. Die Befragten erwarten von FDM-Software für heiße Forschungsdaten vor allem Annotationsfunktionen, abgestuftes Rechtemanagement, An- und Einbindung weiterer Software sowie Automation. Anhand der kategorisierten Interview-Aussagen wurden drei Personas modelliert (Dr. Steward, Gloria Giver, Tom Taker) und persona-spezifische User Stories als Vorformen von Software-Anforderungen formuliert. Die Interview-Aussagen und Personas sind nützlich für die Entwicklung von FDM-Software und weiteren FDM-Services wie Informations-Materialien, Schulungen und Beratungen.

1 Einleitung

Beim Forschungsdatenmanagement (FDM) besteht eine wesentliche Herausforderung darin, die FAIR-Prinzipien anzuwenden, also Daten auffindbar (findable), abrufbar (accessible), anschlussfähig (interoperable) und anwendbar (reusable) zu machen (Wilkinson et al., 2016).

Dies erfordert konzeptionelle und technische Unterstützung im gesamten Datenzyklus, also beim Planen, Erzeugen, Verarbeiten, Speichern, Teilen und Nachnutzen von Daten und Metadaten sowie beim Optimieren von FDM (Ball, 2012; Wolf & Leppla, 2020b). FDM-Software ist hierbei von zentraler Bedeutung, denn sie dient als technisches Mittel, um das Konzept FAIRer Daten umzusetzen.

Damit Beteiligte in einem Forschungsprojekt ihre Daten gemeinsam FAIR machen können, konzipiert die Technische Universität München (TUM) eine kollaborative FDM-Software für heiße Daten (also unfertige, noch nicht veröffentlichte Daten (vgl. Axmann et al., 2021)).

Derartige Tools sind bereits verfügbar, zum Beispiel Resource Space, Research Space, Globus sowie Sync+Share. Solche Tools sind jedoch für TUM-Zwecke ungeeignet. Nachteilig sind fehlende Open Source, uni-externe Speicherung, unflexible Metadatenfelder, oder komplizierte Bedienung. Aufgrund dieser Nachteile entwickelt die TUM eine eigene Software.

Die Entwicklung bedarfsgemäßer Software erfordert eine Bedarfsanalyse. Bedarfe hängen von lokalen Aufgaben, Vorgaben und Angeboten ab. Somit sollte eine Bedarfsanalyse standortspezifisch sein.

An vielen deutschen Einrichtungen wurden bereits FDM-Bedarfsanalysen durchgeführt. Allerdings ging es dabei nicht um kollaborative Software für heiße Daten, sondern um

Datenmanagementpläne (Kim et al., 2023), elektronische Laborbücher (Marutschke et al., 2020; Pietsch & Dees, 2022) und Repositorien (Blumesberger, 2023; Matuszkiewicz, 2022; Schneider & Landwehr, 2021; Strötgen, 2019). Zudem waren TUM-Bedarfe bisher unbekannt.

Angesichts dieser Forschungslücke war das Ziel, Bedarfe von FDM-Beteiligten an der TUM zu erheben, vor allem Software-Bedarfe beim Managen heißer Daten.

Deshalb stellten wir drei Forschungsfragen:

1. Welche Bedarfe herrschen im Datenzyklus?
2. Zu welchen Rollen gehören die Bedarfe?
3. Welche Software-Funktionen brauchen die Rollen, um heiße Daten zu managen?

Um diese Fragen zu beantworten, haben wir Interviews durchgeführt (Abschnitt 2). Dadurch haben wir folgendes herausgefunden: Zentral sind Bedarfe nach Annotation, abgestuften Rechtemanagement, An- und Einbindung weiterer Software sowie Automation (Abschnitt 3.1). Die Bedarfe gehören zu koordinierenden, datengebenden und datennehmenden Rollen (Abschnitt 3.2). Die Rollen brauchen komplementäre Software-Funktionen (Abschnitt 3.3). Die Ergebnisse helfen beim Entwickeln von FDM-Software und anderen FDM-Services (Abschnitt 4).

2 Methoden

2.1 Analyse von Interview-Aussagen

Um Bedarfe im Datenzyklus zu ermitteln, wurden Interviews qualitativ analysiert.

Zur Vorbereitung der Interviews wurden Leitfragen zu sechs Bedarfs-Aspekten formuliert: Probleme/Aufgaben, bisherige Lösungen, Erfreuliches, Unerfreuliches, Wichtiges und Lösungs-Ideen. FDM-Beteiligte wurden per Email gefragt, ob ein Interview über FDM-Software möglich wäre. TUM-intern wurden alle TUM-Schools kontaktiert, um verschiedene Fachrichtungen abzudecken. TUM-extern wurden Personen in FDM-unterstützenden Stellen kontaktiert, um Nutzerbedarfe aus Anbietersicht zu ermitteln.

Insgesamt wurden 15 einstündige Interviews an der TUM oder per zoom durchgeführt.

Auf Seite der interviewenden Personen gab es pro Gespräch eine leitende Person (LP) und mindestens eine unterstützende Person (UP). Fragen wurden von LP gestellt und durch UP ergänzt. Fragen und Antworten wurden von UP in Stichpunkten protokolliert. Die Interviewer*innen waren Dr. Sascha Grusche (8x LP, 3x UP), Dr. Manuel Hora (1x UP), Katja Kessler (3x LP, 1x UP), Francis Kiefer (4x UP), Dr. Flavio Montiel-Montoya (5x UP), Dr. Sebastian Öder (4x LP, 4x UP) und Dr. Christine Wolter (1x UP) von der TUM-Bibliothek.

Auf Seite der interviewten Personen gab es 2 Habilitierte (zuständig für FDM bei sich und anderen); 4 Post-Docs (1 zuständig für eigenes FDM, 1 für eigenes FDM und IT, 2 für FDM bei anderen); 4 Promovierende (zuständig für eigenes FDM); sowie 5 Personen mit Studienabschluss (1 zuständig für IT, 1 für Technik, 3 für FDM bei anderen). Von jeder TUM-School wurde mindestens eine Person interviewt (1x Computation, Information and Technology; 4x Engineering and Design; 1x Life Sciences; 1x Management; 3x Medicine and Health; 1x Natural Sciences; 1x Social Sciences and Technology). Von TUM-externen Stellen wurden drei Personen interviewt.

Um die Interview-Protokolle qualitativ zu analysieren, wurde die Methode von Lazar et al. (2017) genutzt, denn sie fokussiert auf Maschine-Mensch-Interaktion. Die Aussagen-Kodierung kann deduktiv anhand einer bestehenden Theorie oder induktiv im Sinne von Grounded Theory erfolgen. Um den Datenzyklus als theoretischen Rahmen zu nutzen und tiefere Einsichten auf Grundlage der Interview-Daten zu gewinnen, wurden deduktive und induktive Kodierung kombiniert. Dieser hybride Ansatz wurde durch Schritte für Datenschutz und Nachnutzbarkeit ergänzt.

Im ersten Schritt wurden die Interview-Daten für basalen Datenschutz pseudonymisiert. Hierfür wurden Namen durch Zahlen ersetzt. Im zweiten Schritt wurden die Interview-Aussagen deduktiv kategorisiert. Hierzu wurden die Aussagen gemäß Bedarfs-Aspekten farbig markiert: Problem/Aufgabe (P) grau, bisherige Lösung (B) magenta, Erfreuliches (E) grün, Unerfreuliches (U) rot, Wichtiges (W) gelb, Lösungs-Idee (L) cyan. Im dritten Schritt wurden zwecks Rückverfolgbarkeit Kennungen zugewiesen. Zum Beispiel wurde hinter die zweite Aussage von Person 10 über bisherige Lösungen die Kennung "[= 10_B_002]" eingefügt. Im vierten Schritt wurde für stärkeren Datenschutz und Prägnanz hinter jede Kennung eine Paraphrase gesetzt. Im fünften Schritt wurden die Aussagen gemäß Datenphasen deduktiv kategorisiert. Hierfür wurde hinter jede Paraphrase ein Kürzel gesetzt: D1 (Planung), D2 (Erzeugung von Daten und Metadaten), D3 (Verarbeitung), D4 (Speicherung), D5 (Teilen), D6 (Nachnutzung), oder D7 (FDM-Optimierung), vgl. Ball (2012) und Wolf & Leppla (2020b). Im sechsten Schritt wurden die Aussagen übersichtshalber sortiert. Hierzu wurde jede Paraphrase mitsamt Kennung und Phasen-Kürzel in ein vorläufiges Kategoriensystem übertragen (mit Ebene 1 für Datenphasen und Ebene 2 für Bedarfs-Aspekte). Im siebten Schritt wurden die Aussagen zwecks stärkerer Ordnung induktiv kategorisiert. Dafür wurden die Aussagen aktionsorientiert geclustert, z.B. "Aufgaben beim Schützen personenbezogener Daten". Die induktiven Kategorien wurden unter den deduktiven Ebenen angeordnet. Im achten Schritt wurden die Interview-Daten für maximalen Datenschutz anonymisiert, indem alle Pseudonyme durch „xx“ ersetzt wurden.

2.2 Erstellung von Personas

Um Rollen und zugehörige Bedarfe zu modellieren, wurden anhand der Interview-Aussagen Personas erstellt. Eine Persona ist eine hypothetische Person aus einer Bedarfsgruppe (Salminen et al., 2020).

Eine Unterscheidung von Bedarfsgruppen gemäß demografischen Variablen (Aoyama, 2005, 2007) wurde nicht durchgeführt, weil die FDM-Zuständigkeiten der Befragten unabhängig vom akademischen Grad, Alter und Geschlecht waren. Stattdessen wurden Bedarfsgruppen anhand von Aufgabenbereichen unterschieden: Den Phasen D7 und D1 wurde eine koordinierende Persona zugeordnet, den Phasen D2 bis D5 eine datengebende Persona, und der Phase D6 eine datennehmende Persona.

Für die Persona-Beschreibungen wurden Interview-Aussagen zu Aufgaben, Unerfreulichem und Erfreulichem aufgegriffen, da sich eine Persona durch so genannte Jobs, Pains und Gains auszeichnet (West & Di Nardo, 2016). Hierbei wurden ausschließlich Aussagen über die zugewiesenen Phasen verwendet. Mit Fokus auf die FAIR-Prinzipien (Wilkinson et al., 2016) wurden Aussagen über das Finden, Abrufen, Integrieren und richtige Nutzen von Daten ausgewählt. Um Empathie mit den Personas zu fördern, wurden fiktive persönliche Merkmale hinzugefügt (vgl. Salminen et al., 2020).

2.3 Formulierung von User Stories

Um rollengemäße Software-Funktionen zu beschreiben, wurden anhand der Interview-Aussagen User Stories formuliert. Eine User Story ist eine Vorform von Software-Anforderung: „Als [Persona] möchte ich [eine Software-Funktion] nutzen, um [ein Ziel] zu erreichen“ (vgl. Wautelet et al., 2014).

Ziele wurden aus Aussagen zu Aufgaben und Wichtigem abgeleitet; Software-Funktionen aus Aussagen zu bisherigen Lösungen und Lösungs-Ideen. Gleichzeitig wurden Aussagen zu Unerfreulichem und Erfreulichem berücksichtigt, da jede Software-Funktion bei der Job-Erledigung als so genannter Pain Reliever oder Gain Creator dienen soll (Osterwalder et al., 2015). Bei der Aussagen-Auswahl wurde auf die FAIR-Prinzipien (Wilkinson et al., 2016) fokussiert.

3 Ergebnisse

3.1 Kategorisierte Interview-Aussagen

Die Interview-Analyse ergab ein Kategoriensystem mit ca. 1580 Aussagen (Grusche, 2023), vgl. Abb. 1.

Die kategorisierten Interview-Aussagen beleuchten pro Datenphase verschiedene Bedarfs-Aspekte bei einzelnen Aktionen. Im Folgenden tragen wir jene Aussagen zusammen, die einen gegebenen Bedarf unter möglichst vielen Aspekten betrachten, mit Fokus auf die FAIR-Prinzipien.

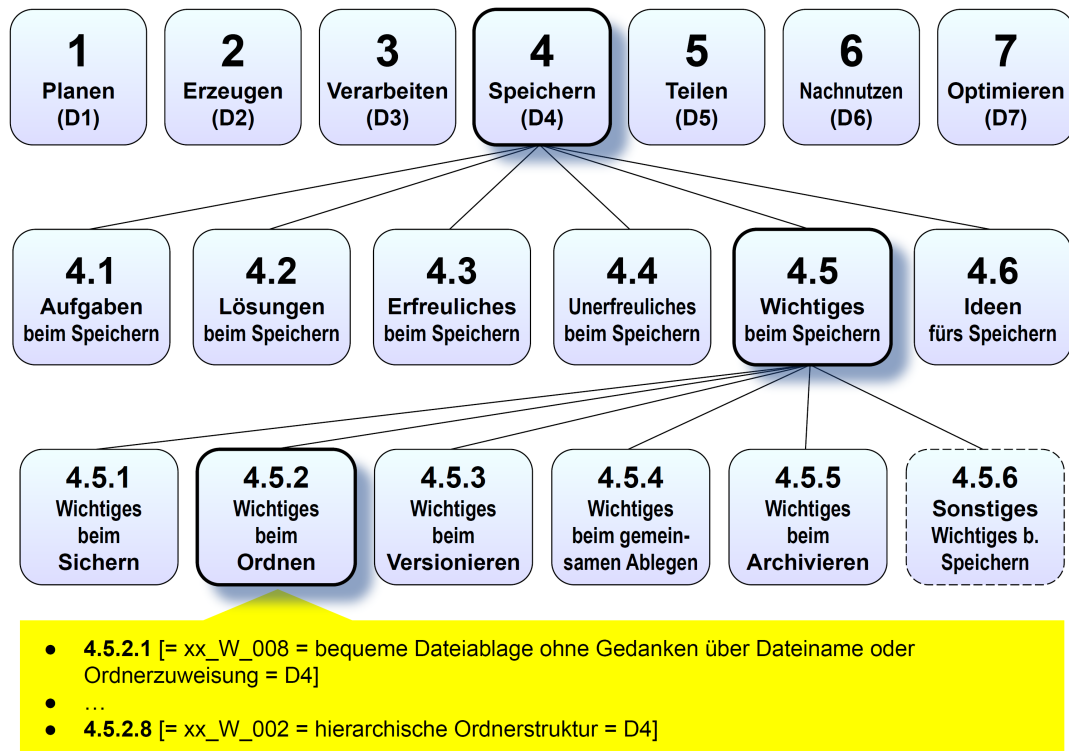


Abbildung 1: Das Kategoriensystem mit Interview-Aussagen über FDM (hier auszugsweise gezeigt) ist als Datenpublikation (Grusche, 2023) verfügbar.

Am Ende jedes Absatzes zitieren wir die Aussagen-Nummern aus dem Kategoriensystem. Ein Asterisk (*) markiert, welche Aussagen auch aus anderen Studien bekannt sind.

3.1.1 Planen

Beim Einplanen von FDM-Software ist der Bedarf der Forschenden zu berücksichtigen. Deshalb planen Verantwortliche für verschiedene Projekte meist verschiedene FDM-Systeme ein. Open-Source-Software ist besonders attraktiv. Ärgerlich ist es, wenn jedes Projekt seine eigene FDM-Software kaufen muss. In jedem Fall ist es wichtig, dass die FDM-Software flexibel ist. Daraus erwächst die Nachfrage nach einer FDM-Software, die aus austauschbaren Komponenten besteht. (1.1.5.8*; 1.2.3.2; 1.3.1.1*; 1.4.1.2; 1.5.1.10*; 1.6.1.5)

Beim Planen der Speicher lautet die Aufgabe, eine leistungsfähige Speicher-Infrastruktur zu finden und somit die Datenvorhaltung zu verantworten. Hierfür werden meist mehrere Server eingeplant. Erfreulich ist es, wenn die FDM-Software verschiedene Speicher unterstützt. Unerfreulich ist es, wenn Server außerhalb der universitären Infrastruktur liegen. Wichtig ist eine Speicherauswahl hinsichtlich Volumen, Zugriff,

Flexibilität, Performanz und Speicherdauer. Im Idealfall kommen nur uni-eigene Server zum Einsatz. (1.1.3.3 & 1.1.3.4*; 1.2.3.1*; 1.3.3.1*; 1.4.3.1*; 1.5.3.3 - 1.5.3.7*; 1.6.3.2*)

Beim Planen der Metadaten sind insbesondere Pflichtfelder zu definieren. Hierfür kommt in manchen Projekten ein Metadaten-Template zum Einsatz. Besonders nützlich sind fachspezifische Metadatenfelder. Enttäuschend ist es, wenn die FDM-Software nur einen Metadatenstandard zulässt. Entscheidend ist die Kombination generischer und spezifischer Standards. Daher soll die FDM-Software projektspezifische Metadatenfelder ermöglichen. (1.1.2.1; 1.2.2.5*; 1.3.2.2*; 1.4.2.8; 1.5.2.7*; 1.6.2.5)

Beim Planen von Datenstrukturen ist zum Beispiel die Struktur von Tabellen vorzugeben. Hierfür werden typischerweise die Spaltenüberschriften definiert. Bei solchen Strukturvorgaben geht es um Standardisierung in der Forschung. Eine allgemeine Lösungs-Idee besteht darin, dass die FDM-Software Vorlagen für Projekte, Prozesse und Experimente liefert. (1.1.1.3; 1.2.1.1; 1.5.6.1*; 1.6.6.2)

3.1.2 Erzeugen

Beim Datensammeln wird oft ein Laborbuch geführt. An der TUM wird hierfür das Labbook der TUM Workbench genutzt. Erfreulich ist, dass das Labbook in die FDM-Software integriert ist. Eine universitätsspezifische Insellösung ist jedoch unerwünscht. Nichtsdestotrotz ist ein elektronisches Labornotizbuch (ELN) wichtig. Eine denkbare Lösung besteht darin, fertige Protokolle in der FDM-Software abzulegen. (2.1.1.11*; 2.2.5.1; 2.3.5.1; 2.4.5.9; 2.5.5.1*; 2.6.5.1)

Die erzeugten Metadaten müssen in dafür vorgesehene Felder eingetragen werden. Als Hilfsmittel dienen zum Beispiel Listen bestehender Eingaben, der DataCite-Metadatengenerator, oder JSON-Dateien. Eine tabellarische Darstellung von Metadaten erscheint besonders sinnvoll. Wenn Metadaten in die Datei-Eigenschaften eingetragen werden, gehen sie beim Wechsel des Betriebssystems manchmal verloren. Entscheidend ist jedoch eine zeitsparende Metadaten-Eingabe. Deshalb besteht der Wunsch nach automatischer Feldbefüllung in der FDM-Software, zum Beispiel durch Parser, JSON-Import oder Bulk-Zuweisung. (2.1.3.8*; 2.2.4.7 & 2.2.4.15 & 2.2.4.5; 2.3.2.3; 4.4.2.3; 2.5.2.2*; 2.6.3.8 & 2.6.3.7 & 2.6.3.4)

Eine häufige Aufgabe besteht darin, personenbezogene Daten zu sammeln und vorher eine Einwilligung einzuholen. Hierbei orientieren sich Forschende gern an Dokumenten über standardisierte Abläufe. (2.1.2.1* & 2.1.2.3*; 2.2.3.2)

3.1.3 Verarbeiten

Vor der Auswertung müssen Daten manchmal erst zusammengeführt werden. Hierfür nutzen Forschende mitunter eigene Python-Skripte. Eine Python-basierte FDM-Software ist daher von Vorteil. Unhandlich wird es aber, wenn eine Datei durch Zusammenfügen

von Daten immer größer wird. Datenvernetzung ist jedoch von großer Bedeutung. Ein beliebter Ansatz besteht darin, Daten miteinander zu verlinken. (3.4.2*; 3.2.1.2; 3.3.1.1; 3.4.2.1; 3.5.4.1*; 3.6.2.1)

Die Schlüsselaufgabe beim Datenverarbeiten ist die Datenauswertung. Dazu nutzen Forschende unter anderem Softwarepakete im Netzlaufwerk. Eine performante grafische Darstellung der Daten ist hierbei vorteilhaft. Eine Software mit unnötig komplexer Datenverarbeitung ist jedoch frustrierend. Besonders wertvoll wäre eine Komplettlösung für die Datenverarbeitung. Mögliche Lösungen bestehen darin, Jupyter anzubinden oder Analyse-Tools in die FDM-Software zu integrieren. (3.1.3.1*; 3.2.2.2*; 3.3.3.1*; 3.4.6.2; 3.5.1.1; 3.6.1.1 & 3.6.1.2*)

Über die Einzelauswertungen hinaus müssen Daten auch miteinander verglichen werden. Dafür nutzen Forschende gern Visualisierungen. Interessant wird es, wenn Zusammenhänge zwischen Daten erkennbar werden, auch projektübergreifend. Kritisch wird es, wenn Daten verwechselt werden. Beziehungen zwischen Ergebnissen sind von hoher Relevanz. Denkbare Ansätze bestehen darin, Metadaten heranzuziehen und Künstliche Intelligenz einzusetzen. (3.3.5.1*; 3.2.3.1*; 3.3.5.1; 3.4.5.2; 3.5.5.2*; 3.6.3.1 & 3.6.4.3)

Personenbezogene Daten werden meist pseudonymisiert oder anonymisiert. Hierbei werden mitunter Pseudonymisierungs-Software und randomisierte Schlüssel genutzt. (3.1.5.2 & 3.1.5.6*; 3.2.4.1 & 3.2.4.2)

3.1.4 Speichern

Damit Daten auffindbar sind, müssen sie beim Speichern geordnet werden. Hierzu dient oft eine Ordnerstruktur gemäß relevanten Parametern. Hierbei ist eine hierarchische Struktur praktisch. Allerdings bezeichnen manche diesen Ordner-Ansatz als veraltet. Wichtig ist eine bequeme Dateiablage ohne Gedanken über Dateinamen oder Ordnerzuweisungen. Ein Vorschlag lautet, Dateien nach Metadaten zu sortieren, ähnlich wie bei iTunes oder WinAmp. (4.1.2.5; 4.2.2.13; 4.3.2.3*; 4.4.2.2; 4.5.2.1; 4.6.2.6* & 4.6.2.2)

Damit Daten abrufbar bleiben, müssen Sicherungskopien erstellt werden. Die Sicherung erfolgt unter anderem auf externen Festplatten und Servern des Rechenzentrums. Beruhigend ist es, wenn die Originaldaten gesichert sind, denn das Neueinspeisen von verlorenen Daten ist mühsam. Deshalb ist die Vorbeugung von Datenverlust unverzichtbar. (4.1.1.1*; 4.2.1.7; 4.2.1.4*; 4.3.1.1; 4.4.1.1; 4.5.1.4)

Zur Interoperabilität gehört es, Dateien zu versionieren. Zur Versionsverwaltung nutzen Forschende teilweise Git. Zufriedenstellend ist es, wenn die Versionierung gelingt. Verwirrend ist es, wenn originale und importierte Dateien nicht synchronisiert sind. Eine funktionierende Versionierung ist unabdingbar. Für eine einfache Versionskontrolle erscheint es sinnvoll, zunächst nur ausgewählte Dateien zu versionieren. (4.1.3.1*; 4.2.4.37*; 4.3.3.1; 4.4.3.2; 4.5.3.2*; 4.6.3.2)

Damit Daten im Team nachnutzbar sind, sind sie in eine gemeinsame Ablage hochzuladen. Als Ablage nutzen Forschende zum Beispiel das Network Attached Storage vom Lehrstuhl oder Sync+Share. Angenehm ist es, wenn der Datentransfer zwischen Speichern reibungslos abläuft. Lästig sind Upload-Probleme. Unkompliziertes Hochladen ist entscheidend. Als Vorbild gilt Dropbox. (4.1.4.1*; 4.2.4.15 & 4.2.4.16; 4.3.4.2*; 4.4.4.2*; 4.5.4.2*; 4.6.4.3)

3.1.5 Teilen

Zunächst müssen beim Teilen Zugriffsrechte vergeben werden. Beliebt sind die Berechtigungssysteme in Sync+Share und Globus Online. Bequem ist es, wenn sich Dateien via Link teilen lassen. Die Zugriffsverwaltung kann komplex werden, sobald Externe im Projekt mitarbeiten. Daher ist ein flexibles Berechtigungssystem notwendig. Eine Lösungs-Idee lautet, Berechtigungen auf Ordner Ebene zu vergeben. (5.1.6.5*; 5.2.2.5 & 5.2.2.7; 5.3.2.3*; 5.4.2.3; 5.5.1.1*; 5.6.2.3)

Sodann geht es oft darum, mit anderen an einer Datei zusammenzuarbeiten. Zu diesem Zweck nutzen Forschende unter anderem Google Drive, Dropbox, Sync+Share und WebDAV. Angenehm ist es, wenn die Zusammenarbeit funktioniert. Datei-Änderungen sind jedoch manchmal schwer nachvollziehbar. Zudem besteht bei Google-Diensten das Risiko von Datenmissbrauch. Die Einfachheit der gemeinsamen Bearbeitung wird jedoch geschätzt. (5.1.3.2*; 5.2.3.4 & 5.2.3.5 & 5.2.3.6 & 5.2.3.2; 5.3.3.1*; 5.4.3.3*; 5.4.3.2; 5.5.3.1*)

Manche Daten sollen auch veröffentlicht werden. An der TUM dient mediaTUM als institutionelles Repositorium. Nützlich sind Review-Links für eine Begutachtung vor Veröffentlichung. Immer wichtiger wird die Veröffentlichung von heißen Forschungsdaten. Eine Datenpublikation ist wichtig, da sie als wissenschaftliche Währung gilt. Somit besteht der Wunsch nach einer Anbindung an das institutionelle Repositorium. (5.1.5.9*; 5.2.5.3; 5.3.5.1; 5.5.5.20*; 5.5.1.4*; 5.6.5.1*)

3.1.6 Nachnutzen

Zuerst müssen die Daten gefunden werden. Hierzu nutzen Forschende Ordnerstrukturen, Dateinamen und Metadaten. Erfreulich ist es, wenn die gesuchten Daten leicht zu finden sind. Anstrengend ist es, wenn die FDM-Software eine ungefilterte Liste von Dateien anzeigt. Entscheidend ist die Unterscheidbarkeit der Daten anhand ihrer Metadaten. Eine Metadaten-Suchfunktion wäre hierbei hilfreich. (6.1.1.1*; 6.2.1.5*; 6.3.1.1*; 6.4.1.1; 6.5.1.3*; 6.6.1.8*)

Sobald die gesuchten Daten gefunden sind, sollen sie abgerufen werden. Datenzugriffe erfolgen unter anderem über Lehrstuhl-Server, Behörden-Server, sowie die Server von TUM Workbench, REDCap und Labfolder. Praktisch ist es, wenn Daten auch von zuhause aus abrufbar sind. Frustrierend wird es, wenn man auf Datenfreigabe warten

muss. In vielen Forschungsprojekten ist es nötig, dass auch Uni-Externe Zugriff erhalten können. Eine Lösungs-Idee besteht darin, verschlüsselte Dateien mit Passwort entschlüsseln zu können. (6.1.6.1*; 6.2.3.4* & 6.2.3.6 & 6.2.3.5 & 6.2.2.1* & 6.2.3.1*; 6.3.2.1; 6.4.2.4; 6.5.2.4*; 6.6.2.1*)

Beim Nachnutzen sind meist mehrere Daten miteinander zu kombinieren. Hierfür nutzen viele TUM-Forschende bisher die Labbook-Exportfunktion der TUM Workbench, wobei alle verknüpften Dateien in einem pdf-Dokument exportiert werden. Nützlich ist es, wenn sich heruntergeladene Dateien offline bearbeiten lassen. Problematisch sind jedoch Dateiformate, die sich nicht öffnen lassen. Wichtig ist es, Daten miteinander vergleichen zu können. Die Idee einer Export-Funktion ist entsprechend beliebt. (6.1.5.2*; 6.2.5.1; 6.3.3.1*; 6.4.3.1*; 6.5.3.1; 6.6.3.1)

Die Daten sollen typischerweise weiterverarbeitet werden. Zum Beispiel dienen Inhalte aus einem digitalen Laborbuch oft als Grundlage für eine Masterarbeit. Erleichternd ist es, wenn die Daten gut dokumentiert sind. Schlecht beschriebene Daten hingegen sind ärgerlich. Für eine korrekte Weiterverarbeitung ist die Nachvollziehbarkeit der Daten unerlässlich. Eine Inhaltsübersicht über die Daten kann hierzu beitragen. (6.1.5*; 6.2.5.1; 6.3.5.1*; 6.4.4.4*; 6.5.4.4*; 6.6.4.1*)

3.1.7 Optimierung von Forschungsdatenmanagement

Manche Data Stewards begleiten Forschende als Projektleiter. An der TU München wird hierfür bisher die TUM Workbench eingesetzt. Es ist befriedigend, mit solcher Software die Arbeit der anderen zu sehen. Belastend wird es, wenn im Projektverlauf immer wieder Rückfragen und Kommunikationsschleifen nötig sind. Wichtig ist nicht nur der eigene Nachvollzug des Projektfortschritts, sondern auch dessen Dokumentation für den Projektträger. Eine mögliche Lösung besteht in der automatischen Analyse von Projekt-Aktivitäten innerhalb eines definierten Zeitintervalls. (7.1.5.7; 7.2.1.4; 7.3.3.1; 7.4.5.22; 7.5.3.2 & 7.5.3.5; 7.6.3.8)

Oft unterstützen Data Stewards die Forschenden beim Metadaten-Management. Manche Data Stewards nutzen Software, die leere Metadatenfelder unterbindet. Vorteilhaft ist es, wenn für die Forschenden kein Mehraufwand entsteht. Umso belastender ist es für Data Stewards, wenn sie die Forschenden zur Metadatenvergabe zwingen oder überreden müssen. Überzeugend ist eine FDM-Software nur, wenn sie Zeit einspart. Somit liegt es nahe, eine regelbasierte Automatisierung umzusetzen. (7.1.3.3*; 7.2.3.2; 7.3.5.4*; 7.4.5.5; 7.5.5.14*; 7.6.5.2*)

Zudem steuern Data Stewards die Datenzugriffe, meist durch Ordner mit Zugriffsverwaltung. Ideal ist es, wenn die Datenhoheit beim Data Steward liegt. Unabhängig davon besteht aber das Dilemma, Daten gleichzeitig zu schützen und zu teilen. In jedem Fall müssen Datenzugriffe projektvertragskonform sein. Als Lösung wurde ein rollenbasiertes Berechtigungssystem vorgeschlagen, bei dem jede Rolle bestimmte Arten

von Zugriffsrechten bekommt. (7.1.4.1; 7.2.4.6; 7.3.4.1; 7.4.4.3; 7.5.4.6; 7.6.4.5*, vgl. 7.6.4.1 - 7.6.4.4 und 7.6.4.6 - 7.6.4.8)

Data Stewards wollen FDM-Software bedarfsgetrieben anbieten. Dies erreichen sie vor allem durch projektspezifische APIs zum Anbinden anderer Tools und durch Absprachen im Arbeitskreis rdmuc. Erfreulich ist es, wenn Forschende mit der FDM-Software zufrieden sind. Software-Funktionen selbst zu entwickeln, ist aufwändig. Es ist jedoch wichtig, dass die FDM-Software institutionsspezifisch anpassbar ist. Vorgeschlagene Lösungen bestehen darin, eine feingranulare Software zu bieten, bereits entwickelte Tool-Features von der Community einzubinden und bei Bedarf weitere Tools anzubinden. (7.1.1.6; 7.2.2.8 & 7.2.2.10; 7.3.1.5; 7.4.2.4*; 7.5.1.25; 7.6.1.29 & 7.6.1.30 & 7.6.1.20*)

3.2 Personas

Die Synthese von Interview-Aussagen über Aufgaben, Unerfreuliches und Erfreuliches ergab drei Personas (Abb. 2 bis 4).

Die Personas repräsentieren koordinierende, datengebende und datennehmende Rollen.

Dr. Sue Steward (Abb. 2) ist federführend beim Planen und Optimieren von FDM (Phasen D1 und D7). Sie ärgert sich über die Vielzahl von Speicherorten und Metadatenstandards, das Dilemma zwischen Daten-Schützen und -Teilen, sowie fehlende Metadaten. Dr. Steward wünscht sich eine Entlastung der Forschenden, verlässliche Speicherlösungen, flexible Datenzugriffe und eine Projekt-Übersicht.

Gloria Giver (Abb. 3) ist verantwortlich für das Erzeugen, Verarbeiten, Speichern und Teilen von Daten und Metadaten (Phasen D2 bis D5). Sie hasst Datenmanipulation, aufwändige Metadaten-Zuweisung, Datenverlust und Datenmissbrauch. Ms. Giver wäre dankbar für schnelle Uploads, gezieltes Datenteilen, Nachnutzung ihrer Daten und Anerkennung für ihre Arbeit.

Tom Taker (Abb. 4) ist zuständig für das Nachnutzen von Daten und Metadaten (Phase 6). Dabei frustrieren ihn Zugriffsbeschränkungen, unpassende Dateiformate, unklare Nutzungsrechte und unzureichende Dokumentation. Mr. Taker erwartet zielführende Suchmöglichkeiten, kompatible Datensätze und reichhaltige Metadaten.

Details zu den phasenspezifischen Bedarfen liefert das Kategoriensystem (Grusche, 2023).

Dr. Sue Steward

Ich wache über die Daten!



Jobs: Was ich tun muss

Datenmanagement koordinieren

S/J1: **Projekt** planen & überwachen

S/J2: **Speicher** definieren

S/J3: **Struktur** von Daten vorgeben

S/J4: Metadaten **standardisieren**

S/J5: **Datenzugriff** regeln

S/J6: **Daten sichern**

Pains: Was mich ärgert

S/P1: Es gibt so viele unterschiedliche **Speicherorte!**

S/P2: Es gibt so viele unterschiedliche **Metadatenstandards!**

S/P3: Ich soll Daten zugleich **schützen und teilen!?**

S/P4: Forschende vergeben **Metadaten** ungern oder zu spät

Gains: Was mich erfreut

S/G1: **Ich entlaste die Forschenden** beim Organisatorischen

S/G2: **Sichere & langfristige Speicherlösung**

S/G3: **Datenzugriff** jederzeit und von überall möglich

S/G4: **Überblick** über das Forschungsprojekt

Alter: 40

Job: PostDoc (als Data Steward)

Beziehungsstatus: Verheiratet

Hobbies: Gospelchor leiten, Sightseeing, mit meinen 2 Kindern spielen

Liebblings-Essen: Saisonale Bio-Küche

Lebens-Motto: Ordnung ist das halbe Leben!

vgl. mit den Interview-Aussagen: Jobs 7.1.5.7; 1.1.3.3; 1.1.1.3; 1.1.2.1; 7.1.4.1; 1.1.3.4. Pains 1.4.3.1; 1.4.2.8; 7.4.4.3; 7.4.5.5. Gains 7.3.5.4; 1.5.3.4 - 1.5.3.7; 1.3.3.1 & 7.3.4.1; 7.3.3.1.

Abbildung 2: Die Persona Dr. Sue Steward repräsentiert eine koordinierende Rolle. Fotoquelle mit CC0-Lizenz: (Jira, o. D.).

Gloria Giver

Ich teile gern meine Daten,
aber vergebe ungern Metadaten!



Jobs: Was ich tun muss

Daten & Metadaten teilen

G/J1: Daten **speichern**

G/J2: Daten **hochladen**

G/J3: **Metadaten** zuweisen

G/J4: Daten dem Team oder der Welt **zur Verfügung stellen**

Pains: Was mich ärgert

G/P1: Wenn mir andere in meinen Daten **rumpfuschen**!

G/P2: Dass Metadatenvergabe so **zeitaufwändig** ist!

G/P3: Wenn ich Daten **verliere**!

G/P4: Wenn andere meine Daten **missbrauchen**!

Gains: Was mich erfreut

G/G1: **Schneller Upload** von großen Datenmengen

G/G2: Einfaches, gezieltes **Datenteilen**

G/G3: **Nachnutzung** meiner Daten

G/G4: **Anerkennung** für meine Arbeit

Alter: 30

Job: Wissenschaftliche Mitarbeiterin

Beziehungsstatus: Verlobt

Hobbies: Malen, Kochen, Häkeln

Liebblings-Essen: Selbstgemachter Obstsalat

Lebens-Motto: Glück verdoppelt sich, wenn man es teilt

vgl. mit den Interview-Aussagen: Jobs 4.1.1.1; 4.1.4.1; 2.1.3.8; 5.1.6.5 & 5.1.5.9. Pains 5.4.3.3; 2.4.4.9; 4.4.1.1; 5.4.3.2. Gains 4.3.4.2 & 4.5.4.2; 5.3.2.3; 4.5.4.3; 5.5.1.4

Abbildung 3: Die Persona Gloria Giver repräsentiert eine datengebende Rolle. Fotoquelle mit CC0-Lizenz: (U.S., [o.D.](#)).

Tom Taker



Ich will die Daten nutzen,
aber dafür brauche ich Metadaten!

Jobs: Was ich tun muss

Daten (nach)nutzen
 T/J1: Daten **finden**
 T/J2: Daten **runterladen**
 T/J3: Daten **kombinieren**
 T/J4: Daten **verstehen**

Pains: Was mich ärgert

T/P1: Wenn der **Datenzugriff** verweigert wird!
 T/P2: Wenn das **Dateiformat** nicht passt!
 T/P3: Wenn die **Nutzungsrechte** unklar sind!
 T/P4: Wenn **Entstehung & Bedeutung** der Daten unklar sind!

Gains: Was mich erfreut

T/G1: Wenn ich die gesuchten Daten **finde**
 T/G2: Inhaltlich und technisch **kompatible Datensätze**
 T/G3: **Aussagekräftige Metadaten**

Alter: 20
 Job: Studentische Hilfskraft
 Beziehungsstatus: Single
 Hobbies: Bierdeckel sammeln, Online-Shopping, Binge-Watching
 Lieblings-Essen: Fast Food
 Lebens-Motto: You only live once (YOLO)

vgl. mit den Interview-Aussagen: Jobs 6.1.1.1; 6.1.6.1; 6.1.5.2; 6.1.5. Pains 6.4.2.4; 6.4.3.1; 6.4.4.4; vgl. 6.5.5.3. Gains 6.3.1.1; vgl. 6.4.3.1 & 6.5.3.1; 6.5.1.3

Abbildung 4: Die Persona Tom Taker repräsentiert eine datennehmende Rolle. Fotoquelle mit CC0-Lizenz: (Moloney, o. D.).

3.3 User Stories

Die Synthese von Interview-Aussagen über Aufgaben und Wichtiges sowie bisherige Lösungen und Lösungs-Ideen ergab 12 exemplarische User Stories (Abb. 5).

	Als Dr. Steward möchte ich ...	Als Ms. Giver möchte ich ...	Als Mr. Taker möchte ich ...
F	ein Metadaten-Template erstellen, um Metadaten zu standardisieren. (1.6.2.10 & 1.1.2.1)	Metadaten mehreren Dateien gleichzeitig zuweisen, um standardisierte Metadatenfelder schnell auszufüllen. (2.6.3.4 & 2.1.3.8)	flexible, metadatenbasierte Suchfunktionen nutzen, um Daten zu finden und voneinander zu unterscheiden. (6.6.1.8 & 6.1.1.1)
A	Forschenden verschiedene Rollen zuweisen, um Datenzugriffe vorgabenkonform zu steuern. (7.6.4.5 & 7.1.4.1)	meinen Ordnern ausgewählte Personen zuweisen, um ihnen den Zugriff auf bestimmte Dateien zu erlauben. (5.6.2.3 & 5.1.6.5)	ein Passwort zur Datei-Entschlüsselung bekommen, um auf Daten zuzugreifen. (6.6.2.1 & 6.1.6.1)
I	vorzugsweise uni-eigene Server anbinden, um eine sichere Speicher-Infrastruktur herzustellen. (1.6.3.2 & 1.1.3.3)	JSON-Dateien importieren, um standardisierte Metadatenfelder zeitsparend auszufüllen. (2.6.3.7 & 2.1.3.8)	Daten exportieren, um sie miteinander zu vergleichen. (6.6.3.1 & 6.1.5.2)
R	Automatisierungs-Regeln definieren, um Forschende bei der standardisierten Metadatenvergabe zu entlasten. (7.6.5.2 & 7.1.3.3)	einen Parser auf Daten anwenden, um standardisierte Metadatenfelder zeitsparend auszufüllen. (2.6.3.8 & 2.1.3.8)	eine Inhaltsübersicht mit standardisierten Metadaten sehen, um Daten vor der Weiterverarbeitung richtig zu deuten. (6.6.4.1 & 6.1.5)

Abbildung 5: Exemplarische User Stories handeln von drei Personas, die komplementäre Software-Funktionen brauchen, um ihre Daten gemeinsam FAIR zu managen.

In den User Stories betrachten die drei Personas jedes FAIR-Prinzip aus drei Perspek-

tiven: Sie brauchen Software-Funktionen zum Koordinieren, Erstellen und Nutzen von standardisierten Metadaten, Zugriffsrechten, Schnittstellen und automatisch annotierten Daten.

4 Diskussion

Durch Interviews haben wir herausgefunden, welche Bedarfe im Datenzyklus herrschen, zu welchen Rollen die Bedarfe gehören, und welche Software-Funktionen die Rollen brauchen, um heiße Daten zu managen. Die Erkenntnisse haben wir in Form von kategorisierten Interview-Aussagen, Personas und User Stories dargestellt.

4.1 Zusammenfassung der Ergebnisse

Zentral sind Bedarfe nach Annotation, abgestufter Rechtemanagement, An- und Einbindung weiterer Software sowie Automation. Die Bedarfe gehören zu koordinierenden, datengebenden und datennehmenden Rollen. Diese Rollen brauchen komplementäre Software-Funktionen zum Koordinieren, Erstellen und Nutzen von standardisierten Metadaten, rollengemäßen Zugriffsrechten, Software-Schnittstellen und automatisch annotierten Daten.

4.2 Einordnung der Ergebnisse

Ähnlich wie in anderen Studien äußerten Beteiligte im Datenzyklus Bedarfe nach **Annotation** (vgl. Blumtritt & Helling, 2019; Iglezakis & Schembera, 2018; Matuszkiewicz, 2022; Schneider & Landwehr, 2021; Wolf & Leppla, 2020a), **abgestufter Rechtemanagement** (vgl. Marutschke et al., 2020; Pietsch & Dees, 2022; Schneider & Landwehr, 2021), **Anbindung oder Einbindung weiterer Software** (vgl. Kim et al., 2023; Pietsch & Dees, 2022; Queckbörner, 2019), sowie **Automation** (vgl. Kim et al., 2023; Pietsch & Dees, 2022). Mit unseren Interviews haben wir aber auch neue Bedarfs-Aspekte beleuchtet, vor allem Emotionen, Ideen und lokale Lösungen.

Die Rollen unserer Personas ähneln den datenkuratierenden, -produzierenden und -konsumierenden Rollen gemäß Wolf & Leppla (2020a). Unsere eigenen Rollenbezeichnungen betonen jedoch das Zwischenmenschliche: Geben und Nehmen werden koordiniert. Zudem haben wir die Rollen mit Bedarfen verknüpft. Dies ermöglicht eine bedarfsgemäße Software-Entwicklung.

Unsere User Stories ähneln denen aus anderen Studien (Neuroth et al., 2018; von Rumel et al., 2024). Allerdings haben wir Bezüge zu den FAIR-Prinzipien hergestellt. Dies ist entscheidend für das Entwickeln von Software, die FAIRe Daten fördert.

4.3 Verwendung zur Software-Entwicklung

Die Interview-Ergebnisse helfen beim Entwickeln von FDM-Software für heiße und kalte Daten an der TUM und anderswo, dank der Vielfalt der Befragten und der Offenheit der Fragen.

Wir empfehlen folgendes Vorgehen: 1. Im Kategoriensystem markieren, auf welche Aufgaben die Software abzielt. 2. Pro Aufgabe eine User Story schreiben. 3. Pro Persona eine User-Rolle mit zugehörigen Rechten definieren. 4. User Stories in Anforderungen übersetzen. 5. Module mit Schnittstellen (vor allem zu DMP-Tools und Repositorien) programmieren. 6. Ein rollenspezifisches User Interface designen.

4.4 Verwendung zur Entwicklung anderer FDM-Services

Die kategorisierten Interview-Aussagen helfen auch beim Entwickeln anderer FDM-Services. Autor*innen von Informationstexten können anhand der Aussagen zu Aufgaben und Wichtigem beschreiben, was im Datenzyklus zu tun und zu beachten ist. FDM-Trainer*innen können Aussagen zu Unerfreulichem und Erfreulichem aufgreifen, um durch Horror Stories und Erfolgsgeschichten zum FDM zu motivieren. FDM-Berater*innen können Lösungen und Lösungs-Ideen für typische Probleme ablesen. Anhand der Personas lassen sich zielgruppengemäße Angebote entwickeln.

4.5 Forschungsdesiderate

Wir haben nur studierte Personen befragt. Welche Bedarfe Studierende haben, könnte durch weitere Interviews ermittelt werden.

Die Persona-Steckbriefe enthalten datenbasierte Hypothesen zu Rollen und Bedarfen. Welche dieser Rollen und Bedarfe auf einzelne Personen zutreffen, könnte anhand von Fragebögen quantitativ analysiert werden.

Die User Stories sind datenbasierte Hypothesen zu Software-Bedarfen. Ob die neue Software den Bedarfen gerecht wird, sollte in Akzeptanzbefragungen untersucht werden.

Literaturverzeichnis

Aoyama, M. (2005). Persona-and-scenario based requirements engineering for software embedded in digital consumer products. *13th IEEE International Conference on Requirements Engineering (RE'05)*, 85–94.
<https://doi.org/10.1109/RE.2005.50>

- Aoyama, M. (2007). Persona-scenario-goal methodology for user-centered requirements engineering. *15th IEEE International Requirements Engineering Conference (RE 2007)*, 185–194. <https://doi.org/10.1109/RE.2007.50>
- Axtmann, A., Bach, F., Bauer, J., Blessing, A., Bönisch, T., Buck, N., Gauza, H., Hess, J., Holz, A., Jung, K., et al. (2021). Kriterien für die Auswahl einer Softwarelösung für den Betrieb eines Repositoriums für Forschungsdaten. *Bausteine Forschungsdatenmanagement*, (3), 14–26. <https://doi.org/10.17192/bfdm.2021.3.8348>
- Ball, A. (2012). Review of data management lifecycle models [Publisher: University of Bath]. <https://purehost.bath.ac.uk/ws/portalfiles/portal/206543/redm1rep120110ab10.pdf>
- Blumesberger, S. (2023). FAIRe Forschungsdaten in den Geisteswissenschaften: Umfrage über Aufbereitung und Archivierung von Daten. *o-bib. Das offene Bibliotheksjournal/Herausgeber VDB*, 10(4), 1–10.
- Blumtritt, J., & Helling, P. (2019). Umfragen und Analyse von Beratungsgesprächen als strategische Wegweiser. *Bausteine Forschungsdatenmanagement*, (2), 96–103. <https://doi.org/10.17192/bfdm.2019.2.8165>
- Grusche, S. (2023). Kategoriensystem mit Interview-Aussagen über Forschungsdatenmanagement. <https://doi.org/10.14459/2023MP1730898>
- Iglezakis, D., & Schembera, B. (2018). Anforderungen der Ingenieurwissenschaften an das Forschungsdatenmanagement der Universität Stuttgart - Ergebnisse der Bedarfsanalyse des Projektes DIPL-ING. *o-bib. Das offene Bibliotheksjournal/Herausgeber VDB*, 5(3), 46–60. <https://doi.org/10.5282/o-bib/2018H3S46-60>
- Jira. (o. D.). Business people in a video call meeting.
- Kim, S.-Y., Hillemacher, S., Decker, S., Rumpe, B., & Geisler, S. (2023). Designing and Implementing Practicable Data Management Plans in Large-Scale Projects. *Bausteine Forschungsdatenmanagement*, (3), 1–12. <https://doi.org/10.17192/bfdm.2023.3.8571>
- Lazar, J., Feng, J. H., & Hochheiser, H. (2017). Chapter 11 - Analyzing qualitative data. In J. Lazar, J. H. Feng & H. Hochheiser (Hrsg.), *Research Methods in Human Computer Interaction (Second Edition)* (S. 299–327). Morgan Kaufmann. <https://doi.org/https://doi.org/10.1016/B978-0-12-805390-4.00011-X>
- Marutschke, C., Fuhrmans, M., Jagusch, G., & Freund, J. (2020). Elektronische Laborbücher an der TU Darmstadt: Beispiel für ein strategisches Vorgehen. *Bausteine Forschungsdatenmanagement*, (2), 65–79. <https://doi.org/10.17192/bfdm.2020.2.8282>
- Matuszkiewicz, K. (2022). Forschungsdatenmanagement in der Medienwissenschaft. Eine Auswertung von qualitativen Interviews zur Bedarfsermittlung für die Gestaltung eines medienwissenschaftlichen Forschungsdatenrepositoriums [Publisher: Philipps-Universität Marburg]. *Bausteine Forschungsdatenmanagement*, 5(2), 1–14. <https://doi.org/10.17192/bfdm.2022.2.8433>

- Moloney, M. (o. D.). Laptop work free stock image. Zugriff 7. Januar 2024 unter <https://stocksnap.io/photo/laptop-work-FP2WQB568R>
- Neuroth, H., Engelhardt, C., Klar, J., Ludwig, J., & Enke, H. (2018). Aktives Forschungsdatenmanagement [Publisher: De Gruyter]. *ABI Technik*, 38(1), 55–64.
- Osterwalder, A., Pigneur, Y., Bernarda, G., & Smith, A. (2015). *Value proposition design: How to create products and services customers want* (Bd. 2). John Wiley & Sons. https://www.academia.edu/download/35014387/Value_Proposition_Design.pdf
- Pietsch, A. M., & Dees, W. (2022). Bedarfserhebung zu elektronischen Laborbüchern an der JLU Gießen: Ergebnisbericht zur Umfrage.
- Queckbörner, B. (2019). Forschungsdaten und Forschungsdatenmanagement in der Geschichtswissenschaft. *Berliner Handreichungen zur Bibliotheks- und Informationswissenschaft*, (441), 1–88. <https://edoc.hu-berlin.de/server/api/core/bitstreams/34ed83e2-d143-4536-8eac-cb903d9d586a/content>
- Salminen, J., Santos, J. M., Kwak, H., An, J., Jung, S.-g., & Jansen, B. J. (2020). Persona perception scale: development and exploratory validation of an instrument for evaluating individuals' perceptions of personas [Publisher: Elsevier]. *International Journal of Human-Computer Studies*, 141, 102437. <https://doi.org/10.1016/j.ijhcs.2020.102437>
- Schneider, G., & Landwehr, M. (2021). Forschungsdatenrepositorium mit einem kooperativen Betriebsmodell: Fallstudie. *o-bib. Das offene Bibliotheksjournal/Herausgeber VDB*, 8(4), 1–11.
- Strötgen, R. (2019). Bedarfsplanung eines institutionellen Repositoriums für Forschungsdaten. *Bausteine Forschungsdatenmanagement*, (2), 112–120. <https://doi.org/10.17192/bfdm.2019.2.8106>
- U.S., D. o. E. (o. D.). ORAU summer program at ORNL 2014 Oak Ridge. Zugriff 7. Januar 2024 unter <https://www.rawpixel.com/image/8736408/orau-summer-program-ornl-2014-oak-ridge>
- von Rummel, P., Keller, C., & Fricke, F. (2024). NFDI4Objects:: Warum brauchen wir eine gemeinsame Forschungsdateninfrastruktur für die materiellen Hinterlassenschaften der Menschheitsgeschichte? *Archäologische Informationen*, 47, 63–72.
- Wautelet, Y., Heng, S., Kolp, M., & Mirbel, I. (2014). Unifying and extending user story models. *Advanced Information Systems Engineering: 26th International Conference, CAiSE 2014, Thessaloniki, Greece, June 16-20, 2014. Proceedings* 26, 211–225. https://doi.org/10.1007/978-3-319-07881-6_15
- West, S., & Di Nardo, S. (2016). Creating product-service system opportunities for small and medium size firms using service design tools [Publisher: Elsevier]. *Procedia CIRP*, 47, 96–101. <https://doi.org/10.1016/j.procir.2016.03.218>
- Wilkinson, M. D., Dumontier, M., Aalbersberg, I. J., Appleton, G., Axton, M., Baak, A., Blomberg, N., Boiten, J.-W., da Silva Santos, L. B., Bourne, P. E., et al. (2016). The FAIR Guiding Principles for scientific data management and stewardship

[Publisher: Nature Publishing Group]. *Scientific data*, 3(1), 1–9.

<https://doi.org/10.1038/sdata.2016.18>

Wolf, A. H., & Leppla, C. (2020a). FDM-Services einer Universitätsbibliothek zur Entfaltung von Regenbogenqualitäten im Forschungsprozess. *o-bib. Das offene Bibliotheksjournal/Herausgeber VDB*, 7(1), 1–16.

Wolf, A. H., & Leppla, C. (2020b). Harmonisierung von Datenlebenszyklus-Modellen: Nutzung von Synergien für optimierte Anwendungen im FDM. *Bausteine Forschungsdatenmanagement*, (2), 1–19.

<https://doi.org/10.17192/bfdm.2020.2.8281>