

LMU Open Science Center: Trainings for research organisations

Below you can find typical lectures and workshops that we offer for research organisations. We are happy to discuss your specific requirements and can offer consultation or develop and conduct customised training material.

General rate: 150 € per hour of teaching + 1h preparation time (150 €) for each existing session. Additional costs for traveling might apply. Some workshops need to be done by two instructors which might increase rates.

Fundamental basics

ID	Content	Title	Hours	Prerequisite
1	Lecture	The replicability crisis and the open science movement	1.5 - 2	
2	Lecture	Credible research with preregistration and reproducible workflows	1	
3	Lecture	Data sharing	1	
4	Lecture	Data anonymity	1.5	3
5	Lecture	How to publish open access for free – OA, preprints, postprints	1.5	
6	R Workshop	Introduction to R	3	
7	R Workshops	Version control with Git and GitHub (1 , 2)	3	2, 6
8	R Workshop	Literate programming for reproducible analyses with Quarto	3	2, 6
9	R Workshop	Package management with renv	1	6, 7
10	R Workshop	Bibliography management with Zotero	1	(8)
11	R Workshop	Reproducible code publishing	2	6, 7, 8, 9
12	Workshop	Preregistration, why not, how (slides)	2.5	2
13	R Workshop	Introduction to simulation of data and data-analyses (e.g. for preregistration)	2	2
14	Workshop	FAIR research data management	2.5	
15	R Workshop	Data validation and documentation.	1	14
16	Workshop	How to use the Open Science Framework (slides)	1.5	

Advanced or discipline specific

ID	Content	Title	Hours	Prerequisite
16	R workshop	Advanced power analyses with simulations (lecture + tutorial + facilitated work on own projects)	8	13
17	R Workshop	How to anonymize sensitive quantitative data	2.5	4
18	R Workshop	How to create synthetic datasets	2.5	4, 17

Possible customised workshops

Customised workshops can also be conducted at the same hourly rate, but with significantly more preparation time for content creation. Typically, a 3h workshop necessitates between 6 (mainly discussion-based workshop) and 40h (e.g. creation of a self-paced tutorial) of preparation.

Which session should I choose for my audience?

If you are a scientific coordinator, you might ask which session is most suitable or beneficial for your specific audience. Here are some reasons to request a specific session:
(The numbering refers to the ID in the tables above).

1. Replicability crisis lecture: In scientific fields where the replicability crisis is still rather unheard of, this lecture describes the current problem with scientific methodology and sets the stage for the need to adopt open research practices.

2. Credible research lecture: This introduces concrete ways to improve the attendees' own research and make it as credible as possible, through planning studies (including statistical planning where relevant) ahead of collecting data, to avoid biases in analyses, and through creating reproducible computational workflows. It serves as an introduction to many of the workshops we offer.

3. Data sharing lecture: This session explains why we should share data, the fundamentals of data sharing practicalities, the FAIR principles and their role in making science “as open as possible, and as closed as necessary.” It addresses the balancing act between different values at play when considering the openness of data and how to navigate that path.

4. Data anonymity lecture: Which kinds of research data need to be protected? How can the risk of re-identification be assessed? How can data be anonymized? What are easy steps to enhance procedural safety in the use of sensible data? Solutions that allow sharing sensible data (e.g. different access levels, such as scientific use files) and how to make them anonymous are presented.

5. Open Access lecture: This lecture covers the practical aspect of open access, the different shades and prices of open access, how to comply with funder requirements, what preprints and postprints are, and how researchers can retain rights to their published work.

6. R: This workshop is an introduction to the programming language R, used to wrangle, process, analyse, and visualize data in a reproducible manner.

7. Version control: Version control is an essential computing skill for any researchers using code/programming languages to e.g. analyse their data. It allows researchers to keep track of changes made to a script (e.g. in R, Python, LaTeX, etc) without having to save multiple versions of a file, and enables easy collaboration on coding projects

8. Literate programming: This workshop teaches how to create reproducible reports (or manuscript), using Quarto (similar to RMarkdown) to alternate plain language with R code chunks such that, if anything changes in the data processing or data analyses, the results, tables and figures, get automatically updated.

9. Package manager: renv is an R-package manager, i.e. a tool that assists in making your R computing environment reproducible by tracking the packages and package versions used in your projects. A project managed with {renv} can be reproduced at later dates and on different machines.

10. Bibliographic references manager: Zotero is the most recommended free and open source reference manager: one can easily insert citations into texts and render the document with an appropriately formatted reference list. Zotero is well-integrated into Quarto, RStudio, Microsoft Word (and Google Docs).

11. Code publishing: After adopting a programming language (R), a version control system (Git), a literate programming language (Quarto), a package manager (renv) and a bibliographic reference manager (Zotero), one still needs to provide, for each code repository, a codebook, a README file, a license, and get a DOI, for instance on Zenodo. This workshop goes through these final steps.

12. Preregistration: Preregistration is the practice by which researchers specify elements of the planned work in a dedicated registry (possibly under an embargo) before observing the outcomes of the work. Examples include description of the planned approach for a qualitative study, and the data analysis plan for a quantitative study. This is a practice that can prevent unintentional p-hacking through confirmation and hindsight biases. More and more journals across disciplines require them, in the form of registered reports.

13. Data simulation: Simulating data and data analyses in R allows researchers to plan their preregistration with confidence, understand the statistical properties of their data, and run power analyses. In this workshop, one learns to create 'ideal' datasets in R.

14. FAIR data management: This session covers data management essentials: organisation (folder structure, file format and names), documentation (metadata, codebook, README file), Publication (data availability statement, license, persistent identifier, repository), and Data Management Plans (DMP).



15. Data documentation and validation: This workshop covers how to create data dictionaries (or codebook) and validate your data before starting data analyses. It also provide guidance on how to create data documentation with R tools ahead of data sharing.

16. the OSF: The Open Science Framework is a free, open source web application that connects and supports research workflows and tools. Researchers use OSF to collaborate, document, archive, share, and register research projects, materials, and data.

17. Data simulations for power analyses: This power analyses workshop covers basic concepts and ways to perform power analyses for simple and complex designs and statistical tests. It also teaches how to run customised simulations of data to perform power analyses for unusual designs.

18. Data anonymity workshop: This workshop is relevant for fields of research collecting data on human participants. It covers data protection and security, GDPR laws, and how to anonymize a dataset.

19. Synthetic dataset: This session teaches how to create synthetic dataset in R. Synthetic data are drawn from a model fitted to the original data and are used to replace some or all original values in order to protect from individual reidentification while preserving relationships between variables. An analysis run on the synthetic data should provide approximately the same answers that would have been obtained if the analysis were run on the original data.