
Plan Overview

A Data Management Plan created using DMPonline

Title: The application and validation of artificial intelligence to automate the planning of assisted conception protocols

Creator: Jessica Wilkinson

Principal Investigator: Jessica Wilkinson

Data Manager: Jessica Wilkinson

Project Administrator: Jessica Wilkinson

Contributor: Kanishka Gogna, Tom Bamford

Affiliation: University of Birmingham

Template: UoB short template

Project abstract:

Currently, there is a broad range of practices of planning assisted conception protocols both within clinics and between clinicians. This allows for a great disparity in treatment of patients with subfertility and thus patient outcomes. Artificial Intelligence is a new and exciting field in reproductive medicine and has been shown to be successful, especially in embryo selection. If an artificial intelligence model was available that could take a host of different patient demographic and clinical factors into account and predict the protocol that would produce the most optimal response, we postulate that it could personalise the treatment for individuals but standardise the care across specific patient groups and improve outcomes.

ID: 160432

Start date: 05-08-2024

End date: 03-08-2027

Last modified: 09-06-2025

Copyright information:

The above plan creator(s) have agreed that others may use as much of the text of this plan as they would like in their own plans, and customise it as necessary. You do not need to credit the creator(s) as the source of the language used, but using any of the plan's text does not imply that the creator(s) endorse, or have any relationship to, your project or proposal

The application and validation of artificial intelligence to automate the planning of assisted conception protocols

Data description

What types of data will be used or created?

I will use already published data in the form of plain text prints for my systematic reviews.

I will then create tabular data, experimental measures, graphs and computational data. Formats will include .xlsx, .csv, .dta, .rtf, .pdf, .docx, .html

How will the data be structured and documented?

Folders with adequate naming will be used. Datasets will be named appropriately and according to the data type, date in YYYY-MM-DD format. Versioning will be used when appropriate as follows:

1.0 = original document

1.1 = minor revisions made

1.2 = further minor revisions

2.0 = substantive changes

Any changes made to files will be recorded in a readme file.

Data storage and archiving

How will your data be stored and backed up?

The University of Birmingham provides a Research Data Store (RDS); access to the RDS is restricted to project members. Backup copies of data are taken on a daily basis and data is stored in separate buildings from the live data. The RDS has a backup and retention policy on how it looks after the data including archiving of primary data here :

<https://intranet.birmingham.ac.uk/it/teams/infrastructure/research/bear/research-data-service/RDS/BackupRetentionPolicy.aspx>

Data may also be stored within the CARE Fertility electronic patient record system (clinic information system or CIS) in accordance to the Data Protection Act 2018. Extracted data will be compiled onto one spreadsheet and stored on the CARE Fertility server. Patient data will be anonymised. Both the computer and spreadsheet will be encrypted with a password.

Is any of the data of (ethically or commercially) sensitive nature? If so, how do you ensure the data are protected accordingly?

Data will be stored securely and anonymously in the University of Birmingham RDS and only shared with research group members. Our processes will comply with the University's Data Protection Policy.

Data will also be stored within the CARE Fertility electronic patient record system (clinic information system or CIS) in accordance to the Data Protection Act 2018. Extracted data will be compiled onto one spreadsheet and stored on the CARE Fertility server. Patient data will be anonymised. Both the computer and spreadsheet will be encrypted with a password.

Where will your data be archived in the long term?

At the publication of a paper, a subset of the data that underpins the paper will be transferred to the UoB BEAR Archive. Once transferred the data will be set to read-only to prevent any inadvertent additions or deletions of the dataset, Any changes will result in a new dataset, which will be archived separately. The BEAR Archive solution has been created to be highly resilient

and is located at two data centers in two different sites, with a backup placed in a third site. Data will be stored for 10 years, should access to the data be requested within a 10 year period, the 10 year clock is then reset from the point of last access. After the 10 year period the data will be deleted.

Data sharing

Which data will you share, and under which conditions? How will you make the data available to others?

Data will be shared through the University of Birmingham's eData repository (<https://edata.bham.ac.uk/>) which makes the datasets discoverable through search engines like Google. eData uses Dublin Core as a metadata standard and the minimum metadata provided for published datasets will cover amongst others title, type of data, creators, publication date and related publications.