
Plan Overview

A Data Management Plan created using DMPonline

Title: Mental Imagery Interviews

Creator: Reshanne Reeder

Principal Investigator: Dr Reshanne Reeder

Data Manager: Dr Reshanne Reeder

Project Administrator: Dr Reshanne Reeder

Affiliation: Edge Hill University

Template: DCC Template

ORCID iD: 0000-0002-8525-5285

Project abstract:

Visual mental imagery is the mental simulation of visual sensory information (i.e., the “mind’s eye”). It has been long known that there is a spectrum of mental imagery abilities, including subjective vividness, precision, and control of images (Kosslyn et al., 1984). Vividness is dominantly investigated, perhaps due to the accessibility (including free, easy, and efficient implementation) of the Vividness of Visual Imagery Questionnaire (VVIQ; Marks, 1973). The VVIQ has supplemented thousands of studies, usually as a control measure for various visual perception and memory tasks. In the past few years, since “imagery extremes” began to be investigated (Zeman et al., 2020), it has been used to classify individuals as having a “blind” mind’s eye (aphantasia) and more recently as having an atypically vivid mind’s eye (hyperphantasia). Using the VVIQ, recent studies have reported a consistent, yet arbitrary, classification for different scores: scores between 16 (i.e., rating your imagery vividness as all 1s) and 23 (more than half of all ratings are 1s) are classified as aphantasia; scores between 75 and 80 (80 being the maximum score) are classified as hyperphantasia; and scores between 24 and 74 are considered to be within the typical range of imagery vividness (Dance et al., 2021; Milton et al., 2021). There are two major problems with this current classification system: firstly, the VVIQ was never intended to categorize individuals with imagery extremes; it was originally created to measure the connection between self-reported imagery vividness and visual memory performance (Marks, 1973). In a 100+ page critique of the VVIQ in 1995, McKelvie concluded that “[t]he reliability and validity coefficients may be sufficient for research purposes, but they do not justify the VVIQ as a diagnostic instrument for personal assessment” (p.87, McKelvie, 1995). Vividness was also pointed out as one of many dimensions of imagery that could impact cognitive performance, and it remains unclear whether imagery vividness should be used to categorize imagery ability (the ability to generate imagery or not). Importantly, no study of imagery extremes, to date, has found compelling cognitive performance differences between aphantasia and typical imagery (tasks: visuo-spatial working memory, pattern and verbal recognition memory, mental rotation; Pounder et al., 2021) or even hyperphantasia (tasks: visual feature judgments, visual memory, mental rotation, face and word recognition; Milton et al., 2021). There is no known cognitive task that people with imagery can do that people with aphantasia cannot do; or vice versa. It is important to determine whether the noisy results and repeated lack of a

difference between aphantasia and imagery (or even hyperphantasia) are due to imagery ability actually having no effect on these tasks; alternatively, it is possible that the most common measure used to determine imagery ability (the VVIQ) is inadequate for this purpose and leads to the misclassification of individuals. It is clear that for the future of imagery extremes research, we must explore different techniques of classifying this subjective experience. Here, we propose starting with an overhaul of the current classification methods by introducing a standardized, personalized classification system.

ID: 84810

Start date: 18-10-2021

End date: 31-12-2022

Last modified: 24-09-2021

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

Mental Imagery Interviews

Data Collection

What data will you collect or create?

In collaboration with a citizen scientist, whose work involves interviewing individuals about their mental imagery experiences on an almost-daily basis, I have come up with a 5-10 minute interview process that relies on three main sections: the interviewee describing an active visual mental image of a sunset to the interviewer, the interviewer describing different mental imagery abilities to the interviewee, then both coming to a decision about the interviewees visual mental imagery ability. The interviews will be transcribed (and audio recorded, provided the interviewee gives consent to this), and analyzed to test the hypothesis that people with different imagery abilities will use different consistent language in their visual descriptions, which will provide clues to the interviewer about their imagery ability. This will then allow the interviewer to classify the participant, with participant agreement about the interpretation of their interview. Notes and audio recording will be taken on a tablet (in-person), or on the researcher's personal computer (online). Notes and audio recordings will be stored on password-protected computers and Cloud storage for an indefinite time, to which only the researchers will have access. There is no existing data that will be reused.

How will the data be collected or created?

Via online or in-person interviews. Interviews will be short (~5 minutes for the necessary sections, but longer if the interviewee would like to provide more in-depth information about their mental imagery experiences). All consent forms and anonymized data (notes in a text file, audio recording in an .mp3) will be stored on a dedicated Cloud directory on OneDrive, shared only among the researchers. Each folder within the directory will pertain to an anonymized participant, labeled with a participant ID (e.g., "F2F_1.mp3", "F2F_1.txt", "F2F_1_consent.pdf"). "F2F" denotes that the interview took place face-to-face, whereas the prefix "ONL" will denote an interview conducted online. Online interviews will be performed via Zoom, and if participants agree, the interview will be audio recorded. Video will be switched off prior to recording, to maximize anonymity. Text notes will be taken during the interview process.

Documentation and Metadata

What documentation and metadata will accompany the data?

All data will be accompanied by a README.txt file, that will be updated as necessary during and after data collection has ended. Text data will be analyzed and quantified according to patterns in the interviews. For example, we will create a table of counts for the number of individuals who describe a sunset with language that denotes low imagery such as "simple" and "faint". No direct or paraphrased quotes will be uploaded to a public domain without the express consent of the interviewee. Quantified, anonymized count data (NOT audio files, as individuals may be identified) will be uploaded to the Open Science Framework (OSF).

Ethics and Legal Compliance

How will you manage any ethical issues?

Ethical approval will be obtained from the Edge Hill University Science Research Ethics Committee (SREC) prior to data collection. Participants will be identifiable by their email address on the consent form, which will be kept separate from their interview transcript. Consent forms and interview transcripts will be linked by a non-identifying participant ID, e.g., "F2F_1". Audio data will be stored only on password-protected folders and a Cloud drive, and will be accessible only to the researchers. Participants will receive a Participant Information Sheet and Consent Form prior to their interview, and will be provided a Debriefing Form and Data Consent Withdrawal Form following their interview.

How will you manage copyright and Intellectual Property Rights (IPR) issues?

All data are owned by Edge Hill University. The text data will be coded, anonymized, and made available on OSF. Any paraphrased or

direct quotes the researchers wish to be made public will require the express consent of the interviewee and this consent will be written and documented. Files of audio recordings will not be shared publicly due to potential participant identification.

Storage and Backup

How will the data be stored and backed up during the research?

All consent forms and anonymized data (notes in a text file, audio recording in an .mp3) will be stored on a dedicated Cloud directory on OneDrive, shared only among the researchers. EHU has sufficient storage capacity for these interviews. The PI will be responsible for the Cloud storage. All researchers will be responsible for the individual interviews they perform, and will need to back them up to the Cloud as soon as possible after the interview (on the same day of the interview).

How will you manage access and security?

The OneDrive storage directory will be shared with the researchers via e-mail access. The researchers include two interns who will help with face-to-face interviews, and a third-party collaborator who will help with online interviews. Interviews that are collected via tablet or Zoom will be uploaded to the Cloud immediately after the interview (optimal), and no later than the end of the day of the interview. This also ensures that notes taken during the interview will not be edited or changed at a later date.

Selection and Preservation

Which data are of long-term value and should be retained, shared, and/or preserved?

Text and audio files will be retained indefinitely for future use. The interviews will contain many details that may not be immediately of interest to the proposed study, but may be of interest at a later date. Researchers will have indefinite access to the interviews for research purposes, and participants will be informed as such.

What is the long-term preservation plan for the dataset?

The raw text and audio files will be stored on a dedicated directory on OneDrive, which will not be made public. Coded, anonymized, quantified data will be uploaded to OSF for public use. Audio files will not be uploaded to OSF. There is no charge for OneDrive storage within the EHU system. There is no charge for OSF storage. The cost of time and effort for preserving the data will be put on the PI, who is paid to do this as part of her job.

Data Sharing

How will you share the data?

The raw text and audio files will be stored on a dedicated directory on OneDrive, which will not be made public. Coded, anonymized, quantified data will be uploaded to OSF for public use. Audio files will not be uploaded to OSF. If another researcher outside of the original team would like access to audio files or direct quotes or interview transcripts, express written permission from the interviewee will be sought. The data that is to be made public will be made available after data collection has concluded, once the manuscript is posted as a preprint or submitted to a journal for peer review. All data on OSF is provided a DOI.

Are any restrictions on data sharing required?

Audio files and full transcripts with direct quotes will not be uploaded to OSF, as they may contain identifying information. If another researcher outside of the original team would like access to audio files or direct quotes or interview transcripts, express written permission from the interviewee will be sought. For data that will ultimately be made public, we will not post it online until data collection has concluded, once the manuscript is posted as a preprint or submitted to a journal for peer review.

Responsibilities and Resources

Who will be responsible for data management?

The PI will be responsible for data management, including training the researchers on good data management practices. All researchers will be responsible for the individual interviews they perform, and will need to back them up to the Cloud as soon as possible after the interview (on the same day of the interview). The PI will be responsible for reminding researchers to upload their data to the Cloud, and to follow-up with any missing data as soon as possible.

What resources will you require to deliver your plan?

I will train researchers on the data management plan and what is required of them in handling interview transcripts, notes, and recordings. I require no software or hardware that is not already available from EHU. There are no charges for the repositories used by the researchers.