
Plan Overview

A Data Management Plan created using DMPonline

Title: What makes baby cries impossible to ignore?

Creator:Andrey Anikin

Principal Investigator: Andrey Anikin

Data Manager: Andrey Anikin

Project Administrator: Andrey Anikin

Affiliation: Lund University

Funder: Swedish Research Council

Template: Swedish Research Council Template

ORCID iD: 0000-0002-1250-8261

Project abstract:

How can crying babies and barking dogs be harder to ignore than specially designed sirens and fire alarms? The key may lie in their temporal structure, cleverly optimized for capturing listeners' attention, and in a particularly irregular, unpredictably changing voice quality. In this project I will test this hypothesis, which requires both methodological and experimental work with long sequences of human nonverbal vocalizations and animal calls. Helped by a qualified research assistant, I will improve the existing algorithms for segmenting and analyzing the temporal structure of vocalization sequences, linking temporal irregularities at different time scales to their bottom-up salience (the ability to involuntarily attract attention). The second task is to develop the first algorithm for automatically detecting different irregularities in voice production known as nonlinear vocal phenomena. Coupled with experimental manipulation of nonlinearities, this will make it possible to investigate their role in making harsh-sounding vocalizations, such as intense baby cries, so salient and distressing to listeners. The project will deliver important methodological innovations and theoretical insights into human and animal vocal communication. It also has immediate practical applications: a better understanding of why some sounds are so distracting and annoying is crucial to managing our acoustic environments.

ID: 149067

Start date: 01-01-2024

End date: 31-12-2026

Last modified: 26-04-2024

Grant number / URL: 2023-00850

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

What makes baby cries impossible to ignore?

General Information

Project Title

What makes baby cries impossible to ignore?

Project Leader

Andrey Anikin

Registration number/corresponding

2023-00850

Version

1.0

Date

2024.04.26

Description of data - reuse of existing data and/or production of new data

How will data be collected, created or reused?

The data for this project consist of the following:

- audio material of different types, including computer-generated sounds, voice recordings obtained from publicly available sources (e.g., YouTube and other online resources) or recorded for the project, and resynthesized or otherwise manipulated voice recordings;
- datasets containing the results of acoustic analysis of the audio material and of online perceptual playback experiments;
- scripts for data preparation, processing, and analysis written in R, javascript, python, etc.

All these data are shared without restraints as they are included in online supplementary materials for each published paper and deposited in public depositories (normally osf.io) and in the institutional archive.

What types of data will be created and/or collected, in terms of data format and amount/volume of data?

The audio will be stored as .wav and .mp3 files, the datasets as .csv and .RDS files, and scripts as .R files. The uncompressed audio files in the .wav format require more storage space, but this uncompressed format ensures optimal quality for acoustic analysis. The compressed .mp3 format is also provided for ease of use.

Documentation and data quality

How will the material be documented and described, with associated metadata relating to structure, standards and format for descriptions of the content, collection method, etc.?

Each published paper has an associated project page on OSF and a dedicated folder in the institutional data archive. These all follow a similar structure, with separate folders for raw data, which usually include the following: audio, datasets (with a README file explaining the variables in each dataset), and analysis scripts (again, with a README file explaining what each script does, as well as with extensive comments within each script).

How will data quality be safeguarded and documented (for example repeated measurements, validation of data input, etc.)?

All included data is part of peer-reviewed publications.

Storage and backup

How is storage and backup of data and metadata safeguarded during the research process?

In addition to being stored on research laptops and external hard drives, it is deposited on OSF (a separate project per study) and in the institutional archive.

How is data security and controlled access to data safeguarded, in relation to the handling of sensitive data and personal data, for example?

There is no sensitive or personal data associated with the project. All participants are anonymous because data collection takes place online (prolific.com).

Legal and ethical aspects

How is data handling according to legal requirements safeguarded, e.g. in terms of handling of personal data, confidentiality and intellectual property rights?

There is no sensitive or personal data associated with the project. All participants are anonymous because data collection takes place online (prolific.com). No copyrighted material is used in the project.

How is correct data handling according to ethical aspects safeguarded?

Because the collected data does not include sensitive information or personal identities, it is freely shared in its entirety with the research community.

Accessibility and long-term storage

How, when and where will research data or information about data (metadata) be made accessible? Are there any conditions, embargoes and limitations on the access to and reuse of data to be considered?

All data are permanently stored in two publicly available depositories: OSF (osf.io, one project per published study) and the institutional archive at Lund University.

In what way is long-term storage safeguarded, and by whom? How will the selection of data for long-term storage be made?

All data necessary for understanding and replicating each study is made available (experimental stimuli, participants' responses, scripts for processing and analysis, etc.).

Will specific systems, software, source code or other types of services be necessary in order to understand, partake of or use/analyse data in the long term?

The audio, datasets, and scripts are stored in standard, widely supported file formats (wav, mp3, csv, R, plain text). All software necessary for processing the data and replicating each experiment is listed in the scripts, with the versions of key libraries.

How will the use of unique and persistent identifiers, such as a Digital Object Identifier (DOI), be safeguarded?

Each published study and its associated OSF repository are assigned unique, permanent DOIs.

Responsibility and resources

Who is responsible for data management and (possibly) supports the work with this while the research project is in progress? Who is responsible for data management, ongoing management and long-term storage after the research project has ended?

Andrey Anikin and Lund University are responsible for data management and storage.

What resources (costs, labour input or other) will be required for data management (including storage, back-up, provision of access and processing for long-term storage)? What resources will be needed to ensure that data fulfil the FAIR principles?

A project folder for permanent archiving is provided by Lund University IT services.