
Plan Overview

A Data Management Plan created using DMPonline

Title: Executive Functioning and Diabetes Management Outcomes in Youth With Type 1 Diabetes: A Systematic Review

Creator: Renan Goksal

Principal Investigator: Renan Goksal

Data Manager: Renan Goksal

Project Administrator: Giulia Bassi, Ramona Cardillo, Daniela Di Riso

Contributor: Silvana Zaffani, Sofia Fasoli, Claudio Maffeis, Carolina Pizzato

Affiliation: University of Padua

Funder: Science Europe

Template: Science Europe Template

ORCID iD: 0000-0001-6894-0523

Project abstract:

Type 1 diabetes mellitus (T1DM) is the most common endocrine condition diagnosed in childhood or adolescence, and requires following a strict daily routine, including frequent blood glucose monitoring, insulin administration, carbohydrate tracking and regular physical activity to keep the blood glucose levels within the target range (i.e., 70–180 mg/dL) and prevent complications. These self-management tasks rely on executive functions (EF), and prior reviews suggest an EF–diabetes management link, but they mainly use cross-sectional designs, questionnaire-based measures, and adolescent/young adult samples, leaving uncertainty about objective adherence indicators, performance-based EF assessment, and how early the EF–management association is evident, issues central to pediatric practice and prevention.

This project aims to conduct a systematic review of etiology and risk factors (with the possibility of a meta-analysis) examining the association between EF (assessed with performance based tasks and/or rating scales) and (a) diabetes self-management / adherence and (b) glycemic outcomes (HbA1C and CGM-derived metrics) in youth with T1D aged 0-18 years. We will examine potential differences in associations by informant (parent-report or self-report), EF assessment modality (performance based tasks or rating scales), adherence indicators (rating scales or objective self-monitoring of blood glucose via meter downloads) and EF domain/construct (inhibition, working memory, cognitive flexibility vs composite scores).

We will follow the JBI manual for Evidence Synthesis (Aromataris et al., 2024), and in particular the 7th chapter on Systematic reviews of etiology and risk factors (Moola et al., 2020).

We adhere to the checklist for transparent reporting of a systematic reviews: the Preferred

Reporting Items for Systematic Reviews and Meta-Analyses (PRISMA) (Page et al. 2021) and the knowledge synthesis principles outlined by Tricco et al., (2011).

PubMed, Web of Science, APA/PsycInfo, Cinahl are the selected databases for our search. Additional articles will be identified through a snowballing approach, including backward snowballing of reference lists of relevant reviews and included studies, and forward snowballing using Google Scholar. Data Management will employ Zotero, Covidence, Excel, Drive for business, with potential quantitative synthesis using R and Rstudio. In alignment with open science principles, the protocol will be preregistered in the Open Science Framework (OSF) repository.

ID: 197472

Start date: 11-11-2025

End date: 11-11-2026

Last modified: 27-04-2026

Grant number / URL: n/a

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

Executive Functioning and Diabetes Management Outcomes in Youth With Type 1 Diabetes: A Systematic Review

Data description and collection or re-use of existing data

How will new data be collected or produced and/or how will existing data be re-used?

No new primary (human participant) data will be collected. Existing data will be re-used by identifying eligible published quantitative observational studies (cross-sectional, cohort/longitudinal, case-control; theses/dissertations if full methods and results are available) that examine associations between executive functioning (EF) and diabetes self-management/adherence and/or glycemic outcomes in youth with Type 1 Diabetes (0–18 years). Studies will be retrieved from PubMed, Web of Science, APA PsycINFO, and CINAHL and supplemented by backward/forward citation chasing (Google Scholar).

Eligible studies will be organized in a extraction dataset on Excel, containing study metadata, design, sample characteristics and variables of interest, including: EF measurement details (performance-based vs rating scales; EF domains/constructs; informant), adherence/self-management indicators (questionnaire/interview vs objective behavioral indicators such as SMBG frequency/meter downloads), glycemic outcomes (HbA1c and CGM-derived metrics such as TIR and variability indices where available), statistical methods, covariates, and reported association estimates/effect sizes. Study quality / risk of bias will be assessed through National Institutes of Health (NIH), National Heart, Lung, and Blood Institute (NHLBI) tools (Observational Cohort & Cross-Sectional checklist and the Case-Control checklist, as applicable).

Study results will be synthesised narratively, with quantitative synthesis (meta-analysis) conducted where feasible.

What data (for example the kinds, formats, and volumes) will be collected or produced?

Data will be collected and created through a transparent, pre-registered systematic review process following PRISMA 2020 statement (Page et al., 2011) and the knowledge synthesis principles described by Tricco et al., (2011).

Search strategies will be developed in consultation of other researchers and similar reviews. Search queries will be documented verbatim.

Records will be exported from PubMed, Web of Science, APA PsycINFO, and CINAHL in RIS, CSV or BibTeX format and uploaded into Covidence for de-duplication.

After the automatic de-duplication in Covidence, title /abstract screening will be done, followed by full-text screening against inclusion/exclusion criteria by two independent blinded researchers.

Inconsistencies in inclusion/exclusion decisions will be discussed and resolved by consensus or through contacting a third party.

Full-text decisions and reasons for exclusion will be recorded in Covidence and extracted as CSV spreadsheets.

For included studies, data extraction will be performed in an Excel extraction sheet.

Two independent researchers will extract a small sample of studies to align extraction fields.

Extraction will use predefined categories organised in columns. All studies will be double-extracted to detect inconsistencies; disagreements will be resolved by consensus or through contacting a third researcher.

Synthesis will follow a narrative approach, with meta-analysis conducted only where appropriate; all analysis steps (if applicable) will be reproducible via scripted code.

Documentation and data quality

What metadata and documentation (for example the methodology of data collection and way of organising data) will accompany data?

Each dataset uploaded in the OSF repository will include metadata describing the project title, contributing authors, institutional affiliations, date of creation and software environment.

Accompanying documentation will include the methodological and analytical context of the review, including the search strategy, inclusion and exclusion criteria, extraction templates, coding frameworks, screening and extraction instructions. Variable definitions, abbreviations and units of measurement will be compiled in a dedicated data dictionary.

What data quality control measures will be used?

Data quality will be ensured through standardised procedures, calibration, and verification checks consistent with PRISMA reporting and the review protocol.

Screening will be conducted by two researchers independently in a blinded manner. The first screening phase will be preceded by a calibration/pilot phase to align interpretation of eligibility criteria. Title/abstract screening will then be completed independently, and disagreements will be resolved by discussion or, if needed, by a third reviewer. Full-text screening will be conducted similarly by the same two researchers, with disagreements resolved by consensus or third-reviewer adjudication.

Full-text decisions and reasons for exclusion will be recorded in Covidence and exported as CSV files to retain an audit trail. Reliability is further supported through the assessment of inter-rater reliability using Cohen's kappa, a validated measure of screening consistency (Viera & Garrett, 2005).

Data extraction will be performed using a predefined extraction form in Excel with clear coding rules and a dedicated codebook. A calibration/pilot phase on a small set of studies will be used to align extraction rules prior to full extraction. Studies will be double-extracted or independently checked to detect inconsistencies, missingness, and data entry errors.

Risk of bias / study quality will be assessed using NIH/NHLBI checklists (as appropriate to study design) and cross-checked for consistency and completeness.

Any conversions or derived values (e.g., effect-size transformations for meta-analysis, if feasible) will be documented and reproducible through scripts.

Files will be organised in a consistent folder structure (e.g., protocol/registration, search materials, full-text exclusion log, extraction, risk of bias, analysis, outputs).

Storage and backup during the research process

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

During the project, working files (search exports, reference library, screening records, extraction sheets, risk-of-bias files, and analysis scripts/outputs where applicable) will be stored on institutional managed Google Drive storage with controlled access limited to the review team. This platform provides automatic, encrypted cloud backups and complete version control, ensuring recoverability in the event of accidental deletion or system failure. Local copies of working files will also be maintained on the researchers' password - protected laptops.

Reference management and screening will be conducted using Zotero and Covidence.

The protocol and final datasets and documentation will be stored on OSF repository for long term archiving and public access.

No sensitive or personal data will be collected or stored, as the review is based solely on published studies.

How will data security and protection of sensitive data be taken care of during the research?

The project does not involve any sensitive participant data, the primary focus of security management will be to maintain the integrity of project files.

Data will not be stored solely on local laptops or external drives but will be stored on institutional managed Google drive storage with access restricted to the review team permission-controlled folders.

Legal and ethical requirements, codes of conduct

If personal data are processed, how will compliance with legislation on personal data and on data security be ensured?

This project does not collect or process personal data from research participants; all data are extracted from published studies and recorded at the study level.

How will other legal issues, such as intellectual property rights and ownership, be managed? What legislation is applicable?

Copyright for the research data and documentation generated in this project will rest with all authors, University of Padova and collaborating institutions. Any publications from the project may be subject to publisher copyright transfer agreements, but the underlying datasets and supporting documentation will remain the property of the researchers.

Copyright and other Intellectual Property Rights for third-party materials (e.g., full-text articles/PDFs, tables/figures, and proprietary questionnaire items) remain with the original rights holders and will be used for internal research purposes only.

How will possible ethical issues be taken into account, and codes of conduct followed?

No significant ethical concerns are anticipated for this project. Ethical integrity in research data management will be maintained through following the principles of transparency, accuracy, and adequate attribution of all sources.

Data extracted will be study-level, aggregated information reported in the literature.

The review will follow relevant codes of conduct for responsible research and open science (e.g., transparent reporting via PRISMA), and any deviations from the registered protocol will be documented.

Data sharing and long-term preservation

How and when will data be shared? Are there possible restrictions to data sharing or embargo reasons?

All data with long-term value (e.g., final extraction dataset, data dictionary/codebook, search strategies and search dates, PRISMA flow information, risk-of-bias assessments, and analysis scripts/outputs where applicable) will be shared via the OSF project.

Metadata and documentation will ensure users can interpret and reuse the materials appropriately.

The protocol and accompanying documentation (e.g., instructions and the data management plan) will be made openly available at the time of OSF registration.

The curated extraction dataset and associated materials will be made publicly available upon completion of the review (or at manuscript submission, if an embargo is needed for journal coordination).

We do not expect any restriction on data sharing.

How will data for preservation be selected, and where will data be preserved long-term (for example a data repository or archive)?

Data for preservation will be selected based on long-term value for transparency, and reuse (i.e., final datasets and the documentation needed to interpret them). Long-term preservation will be ensured by depositing the preserved outputs in the project's OSF repository. A restricted-access backup copy will be maintained on University-managed cloud storage (e.g., Google Drive) to ensure redundancy. Materials will be retained for a minimum of 10 years after project completion.

What methods or software tools will be needed to access and use the data?

The preserved datasets will be provided in widely used, non-proprietary or broadly accessible formats (CSV for tabular data; PDF for documentation; and R scripts where applicable). Data can therefore be accessed and reused with standard spreadsheet software (e.g., Microsoft Excel, LibreOffice) and/or statistical software (e.g., R/RStudio). Documentation will be provided as PDF and/or plain text and can be opened with standard viewers.

Zotero and Covidence will be used during the review workflow, but all shareable outputs will be exported to standard formats so that access does not depend on these platforms.

How will the application of a unique and persistent identifier (such as a Digital Object Identifier (DOI)) to each data set be ensured?

The OSF registration will provide a persistent, citable record of the protocol, and the final datasets will be deposited in OSF to obtain a DOI.

Data management responsibilities and resources

Who (for example role, position, and institution) will be responsible for data management (i.e. the data steward)?

Renan Göksal (PhD student, University of Padova) will act as the data steward and will be responsible for all data management activities, including data capture, metadata and documentation production, data quality control, storage and backup during the project, and preparation of materials for archiving and data sharing. Carolina Pizzato (Master's student, University of Padova) will support these activities as needed. Before any public data release, the materials prepared for sharing will be reviewed by the project administrators Dr. Giulia Bassi (Assistant Professor, University of Padova), Dr. Ramona Cardillo (Assistant Professor, University of Padova) and Prof. Daniela Di Riso (Associate Professor, University of Padova) to ensure completeness, consistency, and appropriateness for sharing.

What resources (for example financial and time) will be dedicated to data management and ensuring that data will be FAIR (Findable, Accessible, Interoperable, Re-usable)?

The project will use existing institutional and open tools (OSF, Zotero, Covidence, Excel, R/RStudio, and University-managed cloud storage), and no additional financial resources or repository fees are anticipated.