

Data Designer: Empowering Data Preparation, Transformation, ARD Creation and Summarization for Insights and Decision Making.

Mukul Mittal, Sanofi, Morristown, NJ, USA

Sanjeev Rathore, Quantzig, Bengaluru, India

Andre Couturier, Sanofi, Morristown, NJ, USA

ABSTRACT:

Data Designer is intuitive **R Shiny® application**, built to streamline and accelerate data transformation, as well as to summarize trials data. The tool supports **seamless ingestion** of datasets from various file formats e.g., **SAS, XPT, CSV, Excel, Parquet**, and **JSON**.

At its core, Data Designer empowers users to create **Analysis Results Datasets (ARDs)** using *{card}* package by leveraging analysis results Metadata specs. These ARDs can be further summarized using *{gtsummary}* to generate publication ready clinical tables. Application also offers support for multilingual outputs and includes features for generating **mock tables** from dummy data as well.

For customized reporting, users can create summary outputs from the ADaM datasets using *{rtables}*, either through predefined templates or custom logic, with outputs exportable in RTF and DOCX formats. Titles and footnotes are dynamically populated from reflow to ensure consistency/compliance, and R code is generated to support reproducibility and traceability.

INTRODUCTION:

Data import, cleaning, wrangling, and transformation are foundational steps of clinical trials data flow, ensuring that collected information is accurate, consistent, reliable and ready to generate insights that support safety profiling and efficacy evaluation for operational decision making and regulatory submissions. Together, these steps create a robust process that turns complex clinical information into simplified reliable evidence.

As we all know, clinical trial data analysis is governed by regulatory standards and internal reporting conventions. To support the generation of Analysis-ready-datasets (ARDs) and Tables-listing-figures (TLFs), Sanofi has recently adopted new standards: the CDISC ARS guidelines.

To support our statistical programming teams with this transition ART (Analysis and Reporting Technologies) team has developed Data Designer; a **R Shiny®** based application that streamlines data ingestion, transformation, and reporting. The application enables seamless handling of industry-standard data formats, including SAS, XPT, CSV, Excel, Parquet, and JSON. The data can be filtered on multiple conditions, new variables can be derived, and missing values can be imputed using multiple methods. Datasets can be merged according to different methods using multiple key variables. Furthermore, the application generates R code, ensuring traceability and reproducibility.

Currently, the application supports ARD generation for key domains such as Demographics, Medical History, and Adverse Events using predefined TLF configurations. These configurations automatically drive variable mapping, analysis specifications, and display attributes. Users can generate stratified summaries, frequency-based adverse event tables, header counts, overall columns, and “Any Event” rows consistent with ARS (Analysis results specifications). Clinical tables can be produced using *{gtsummary}* R package, either directly from ARDs or through ARS-aligned templates. Outputs are generated with titles, footnotes, and orientations, populated from Analysis Output Metadata (AOM) and reflow metadata, ensuring consistency across all deliverables.

Technically, Data Designer is developed using a modular R architecture built on the Shiny framework. Core dataset derivation and manipulation rely on *{dplyr}* package functions, including *select()*, *mutate()*, *filter()*, *arrange()*, and *summarize()*. Reporting outputs leverage *{cards}* and *{gtsummary}* packages to meet both standard and custom reporting needs. The application generates fully reproducible R scripts, including standardized header capturing tool i.e., version, author, and study information. This technical foundation enables Data Designer to deliver ARS-compliant, reproducible, and submission-ready clinical outputs.

METHODOLOGY AND RESULTS:

MULTI-FORMAT DATA SUPPORT:

Data analysis begins with the essential step of importing data into the application. To create a comprehensive master dataset, users often require multiple raw datasets. Data Designer is specifically designed to support this requirement. Users can seamlessly import multiple datasets of different formats through an intuitive interface (Fig. 1).

To ensure efficient handling of various file formats, different R packages are used, including *{readr}* for CSV files, *{openxlsx}* for Excel files, *{haven}* for SAS and XPT files, *{jsonlite}* for JSON files, and *{arrow}* for Parquet files. This tailored approach enhances the application's versatility and enables users to work with diverse datasets efficiently.

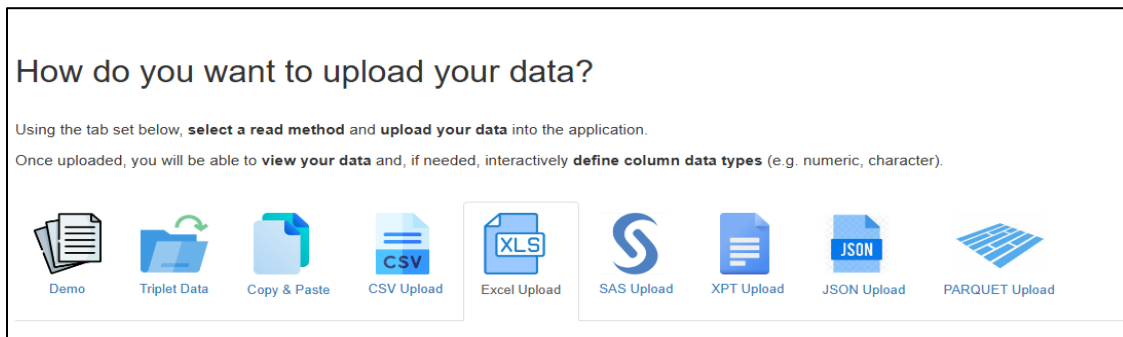


Figure 1: Multi-Format data Ingestion

ADVANCE DATA SUBSETTING:

Enhanced filtering functionality allows users to define multiple filters, each composed of one or more logical conditions, enabling complex subsetting scenarios without manual programming. In addition, it also enables value comparisons between columns, supporting flexible dataset subsetting based on inter-variable dependencies (e.g. TRTSDT > TRTEDT). Specialized functions for **POSIXct** and **POSIXlt** classes enable accurate filtering of *date and date-time variables*. These capabilities support controlled and reproducible subsetting. We have used different functions from *{dplyr}* package like *select()* and *filter()* to select the variables and filter the data based on one or multiple conditions(s).

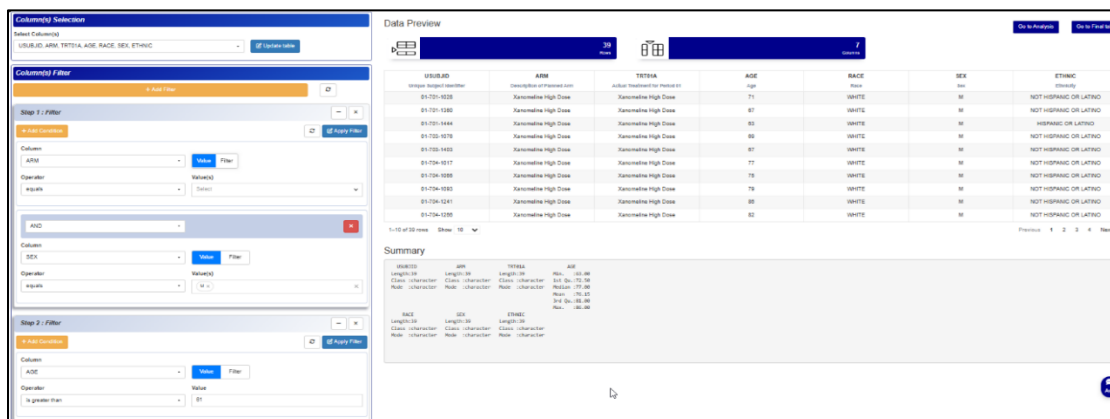


Figure 2: Advanced Data Subsetting

MISSING VALUE IMPUTATION:

In clinical trials, collected data may have missing value due to patient dropout, measurement errors or absence at some time point for some subjects. There are many methods to deal with missing data. In this application, we utilized predefined imputation strategies e.g., mean, median, mode, KNN (K Nearest Neighbors), and last/first observation carried forward, or by removing incomplete observations when appropriate. It reduces the need for manual intervention during data preparation.

Data Designer

Welcome Page | Upload Data | Merge & Split Data | **Data Manipulation** | ARD and Clinical Summary | Downloads | Feedback | User Manual & Tutorials | Info

Dataset Selection

Select a dataset

uploaded_data_label

Submit Dataset loaded

Save Dataset

Average name for dataset

+ Save E Code Save

Data Manipulation

Cuts Operations

Select Filter Rename Set Column Labels Order Columns Remove Duplicates Group & Summarize Outlier Detection & Treatment Missing Value Imputation Sort By

Add Column

Data Reshape

Data Cleaning and Formatting

Teleneza Data Processing

+ Add Operations

Step 1 : Missing Value Imputation

Select Columns

ACDCFL_ADRIN

ACDCFL ADRIN Missing 273

Imputation method Impale missing values with Mode

ADORN ADRIN Missing 477

Imputation method Impale missing values with Mean

Reset Apply Changes

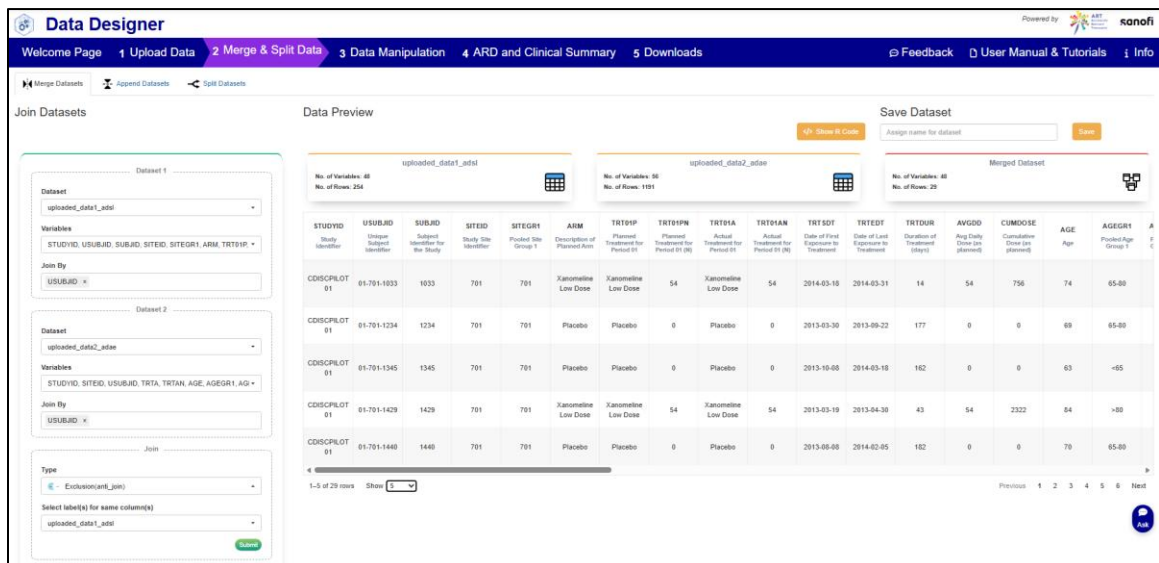
Data Preview

	SITEID	USUBJID	TRTA	TRTSDT	TRTEDT	ADURN	AEHLT	AEBCODSYS	AOCCFL
1	791	01-701-1015	Placebo	2014-01-02	2014-07-02	23.94	HLT_0017	GENERAL DISORDERS AND ADMINISTRATION SITE CONDITIONS	Y
2	791	01-701-1015	Placebo	2014-01-02	2014-07-02	23.94	HLT_0317	GENERAL DISORDERS AND ADMINISTRATION SITE CONDITIONS	Y
3	791	01-701-1015	Placebo	2014-01-02	2014-07-02	3	HLT_0148	GASTROINTESTINAL DISORDERS	Y
4	791	01-701-1023	Placebo	2012-08-05	2012-09-01	24	HLT_0204	SKIN AND SUBCUTANEOUS TISSUE DISORDERS	Y
5	791	01-701-1023	Placebo	2012-08-05	2012-09-01	23.94	HLT_0204	SKIN AND SUBCUTANEOUS TISSUE DISORDERS	Y

1-5 of 1191 rows Show 5 Previous 1 2 3 4 5 ... 239 Next

MERGING, APPENDING and SPLITTING DATASET(S):

From an implementation perspective, various functions are used from `{dplyr}` package to perform all column-wise merge operations, including `inner_join()`, `left_join()`, `right_join()`, `full_join()`, `anti_join()`, and `semi_join()`. Row-wise merging is implemented using `bind_rows()`, while dataset splitting and subsetting are achieved through `slice()` and `select()` functions.



3

GROUPING AND SUMMARIZATION:

The **Group and Summarize** feature provide a flexible framework for aggregating dataset based on user-defined grouping variables, summarize variables, and selected aggregation functions. For continuous variables, it offers a range of aggregate functions, including *sum*, *mean*, *minimum*, *maximum*, *median*, *standard deviation*, *standard error*, and *quantiles*, enabling comprehensive statistical summarization. For categorical variables, *counts* are computed to summarize frequency distributions. Users can further refine output by specifying the number of decimal points for the aggregated values, ensuring consistent numerical presentation. Additionally, the feature allows custom renaming of summarized columns, providing clarity and enhancing interpretability of the results.

From a technical perspective, this functionality is implemented using the *{dplyr}* package, with core operations driven by functions such as *group_by()*, *summarise()*, and *rename()*. Aggregation computations leverage standard *{base R}* functions, including *mean()*, *min()*, *max()*, *n()* etc., ensuring efficient and reproducible summary generation.

Step 1 : Group & Summarise

Grouping Column(s) ⓘ
ARM

Summarise Column(s) ⓘ
AGE

Aggregate Function(s) ⓘ
mean median sd Quantile(s)

Quantile value(s) ⓘ
25,75

☒ Set decimal points ⓘ 0

☒ Set column name(s) ⓘ

AGE_mean Mean

AGE_median Median

AGE_sd Standard Deviation

AGE_Q25 Q25

AGE_Q75 Q75

Apply Changes

Figure 5.1: Input configuration panel

Data Preview

3 Rows 6 Columns

	ARM Description of Planned Arm	Mean	Median	Standard Deviation	Q25	Q75
1	Placebo	75	76	9	69	82
2	Xanomeline High Dose	74	76	8	71	80
3	Xanomeline Low Dose	76	78	8	71	82

1-3 of 3 rows Show 5

Previous 1 Next

Figure 5.2: Summarized output

ARD AND SUMMARY CREATION USING SANOFI ARS:

ARD creation within the *ARD & Summary (Sanofi ARS)* module is driven by Sanofi ARS-compliant TLF configurations. The application automatically loads the selected TLF specifications and applies it to generate standardized ARD. *Automatic variable mapping* is performed based on TLF specs, with the flexibility to apply manual mapping for unmapped variables. ARD generation is currently supported for some tables e.g., Demographics, Medical History, and Adverse Events.

We have extensively used the complex list structure to capture the configuration/meta data about ARS as well as *ard_stack()*, *ard_stack_hierarchical()* functions from *{cards}* package to generate the ARD.

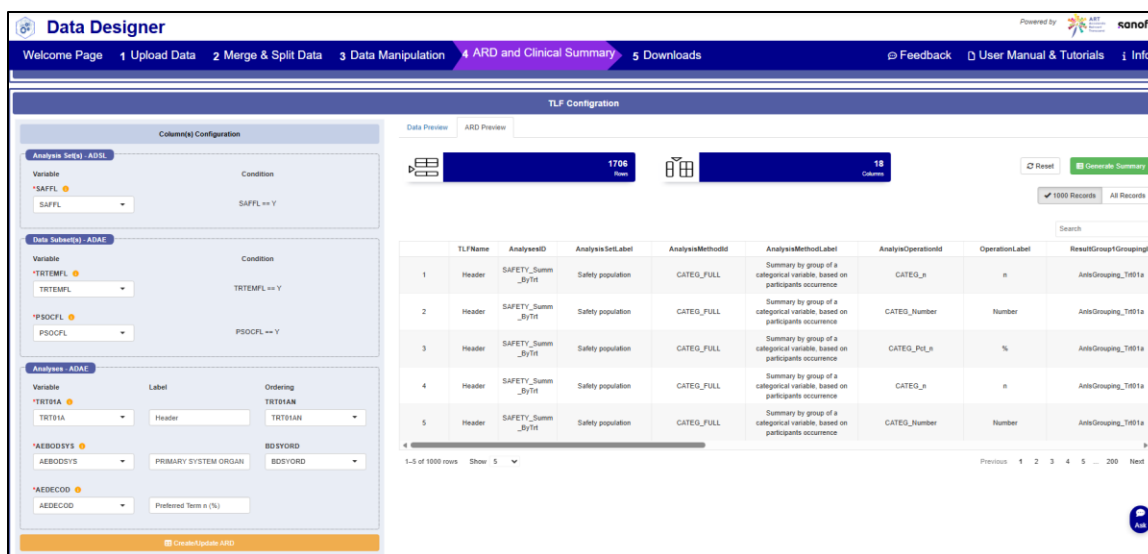


Figure 6.1: ARD Creation

Summary creation is tightly integrated with ARD, enabling one-click summary output directly from the generated ARD. The module applies ARS-defined analysis logic and sorting rules to produce compliant clinical tables. *Display attributes* such as titles, footnotes, and table orientation are automatically populated from AOM, ensuring consistency with approved reporting standards. Users can download summaries in DOCX and HTML formats. This seamless transition from ARD to summary outputs significantly reduces time and effort.

At the backend, we are using several functions from `{gtsummary}` package like `tbl_ard_summary()`, `tbl_ard_hierarchical()`, `modify_table_styling()`, `modify_header()` etc. to generate summary reports.

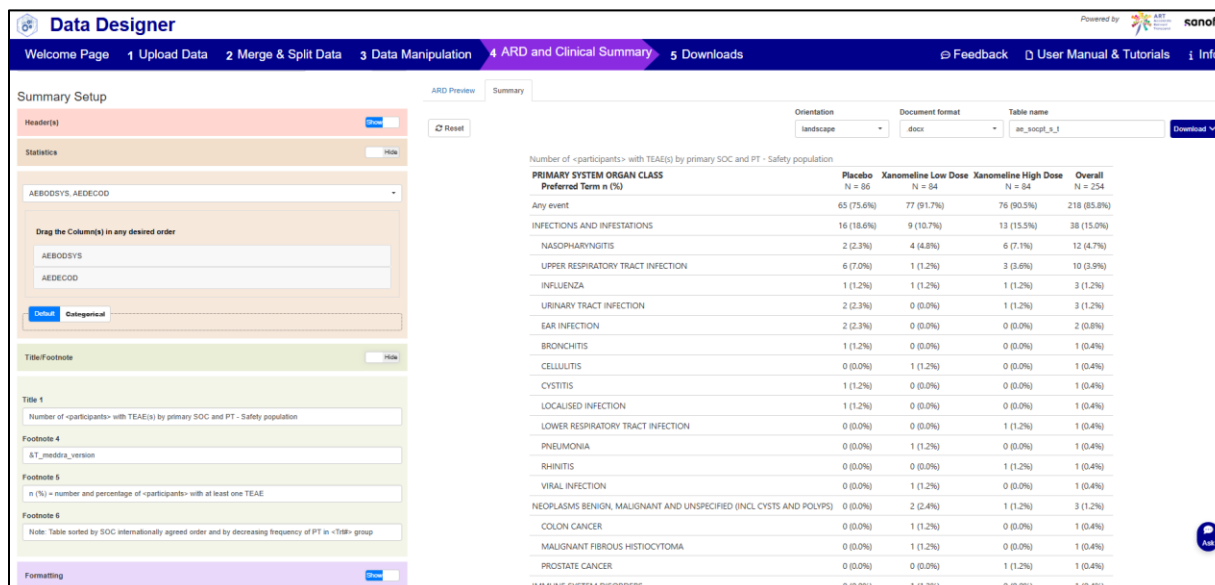


Figure 6.2: Summary creation using ARD

Overall, the ARD & Summary outputs functionality provides an **end-to-end, metadata-driven workflow** for standardized clinical reporting. By integrating TLF configuration, variable mapping, ARD creation, and summary generation within a single module, ensures traceability and reproducibility. Automation minimizes variability while preserving flexibility for study-specific needs. Collectively, this capability strengthens analytical consistency and operational efficiency across clinical studies.

MOCK SUMMARY CREATION WITH MULTILINGUAL SUPPORT:

The mock summary creation allows users to generate mock shells / clinical summaries without revealing actual data. By using this feature, the structure and layout of summary tables can be shared with all team members while maintaining data confidentiality. When creating ARDs with custom logic, users can enable **“Convert to Mock ARD”** option and subsequently generate summaries from the resulting dataset. The resulting summaries display placeholder values such as “xx” instead of real aggregated numbers, ensuring sensitive data is protected. This capability facilitates review and validation of summary tables prior to final data availability. Moreover, the tool supports multi-language summary generation, enabling tables to be presented in languages such as *Chinese* or *Japanese*. Language selection is configurable at the time of summary creation, supporting *global collaboration*.

The screenshot shows the 'Data Designer' application with the 'ARD and Clinical Summary' tab selected. On the left, the 'ARD Setup' panel allows selecting base variables (AGE, AGEGR1, RACE, SEX) and analyzing variables. The 'Convert to mock ARD' checkbox is checked. The main area displays a table with 10 rows of data, including columns for group, group_level, variable, variable_level, contrast, stat_name, stat_label, stat, stat_unit, test_fun, warning, and error. The table shows statistical results for AGE and AGEGR1 across different groups and levels.

Figure 7.1: Creating mock ARD dataset using custom logic

The screenshot shows the 'Data Designer' application with the 'Summary' tab selected. On the left, the 'Summary Setup' panel allows selecting variables (AGE, AGEGR1, RACE, SEX) and formatting options. The 'Table name' field is set to 'Assign names for Summary'. The main area displays a summary table with columns for characteristic, Placebo (N=86), Xanomeline Low Dose (N=84), Xanomeline High Dose (N=84), and Total (N=254). The table shows summary statistics for AGE, AGEGR1, RACE, and SEX across different groups and levels.

Figure 7.2: Creating mock summary

Application also enables the creation of clinical summaries using *{tables}* framework from **ADaM datasets**, leveraging a **set of pre-defined, standardized templates**. The currently available templates include *Demographics* and *Adverse Event* tables. Each template includes an *information panel* detailing essential metadata, such as the required dataset, mandatory variables, and optional variables. Upon template selection, a configuration form is launched, enabling users to define *reporting parameters* and generate summaries based on the *configured variables*. Users can further map the title/footnote using the *reference list (reflist)* from a specific study to ensure alignment with study metadata and standards. The generated output can be exported in multiple formats, including HTML, RTF, PDF, and DOCX.

Clinical Summary Templates

Demographic Summary Template(s)

☒ dem_demo_r_t

Adverse Event Summary Template(s)

☐ ae_overview_pretrt_s_t ☐ ae_overview_posttrt_s_t ☐ ae_socpt_s_t

Select Template

Template Information				
Template Name	demographics 1			
Definition	Demographics and patient-subject characteristics at baseline - Randomized population			
Required Dataset	ADSL (Subject Level Analysis Dataset) is a core dataset in clinical trials that contains one record per subject. It includes key demographic, treatment, and trial participation information, such as age, sex, race, treatment group, randomization details, study site/enrolment details, and important analysis flags. ADSL serves as the foundation for linking and summarizing other subject-level analysis datasets.			
Required Variable List				
Section	Name	Description	In Worksheet	Exclude
Flow	ENR01C	1 indicates whether a subject was randomized in the clinical trial. Typically, this is a checkbox or refers with values like "Y" (yes) or missing (blank) if not randomized.	Yes	Exclude
	TRE01F	1 indicates whether the subject was treated. A subject was placed in treatment during a specified trial.	Yes	Exclude
	TRE01F01	A numeric code corresponding to the planned treatment in the same trial. It is only relevant for the "Treat?"	No	
Analysis	AGE	Subject's age—00 to 99+ age at baseline (in years). Used to describe demographics and for subgroup analyses.	Yes	
Analysis	AGE01R	Age category—into membership groups 0, <45, 45-64, 65-74, 75+ in studies on adults and missing age-based effects.	No	
Analysis	AGE01R01	Member's code for AGE01R group. Used for sorting, grouping, or programming logic.	No	
Analysis	SEX	Subject's sex—00 to 99+ biological sex (Male/Female). Used for demographic summaries and sex-based subgroup analysis.	Yes	
Analysis	RACE	Self-identified race of the subject. Important for diversity analysis and subgroup reporting.	Yes	
Analysis	RACE2N	Member's representation of RACE. Supports sorting and consistent categorization across datasets.	No	
Analysis	ETHNIC	Ethnicity of the subject (e.g., Hispanic or Latino). Assists in population characterization and subgroup analysis.	Yes	
Analysis	ETHN01R	Member's code for ETHNIC. Used to assess baseline characteristics and dosing considerations.	No	
Analysis	WEIGHT01R01	Categorical grouping of baseline weights. Facilitates stratification and comparison across weight groups.	No	
Analysis	WEIGHT01R01N	Member's code for WEIGHT01R01. Used in programming for sorting and analysis logic.	No	
Analysis	BASE01R01	Baseline BMI grouped into categories (e.g., Normal, Obese). Used to assess the impact of BMI on treatment outcomes.	Yes	
Analysis	BASE01R01N	Member's code for BASE01R01. Helps with sorting and grouping in statistical analyses.	No	

Column(s) Configuration

Header Mapping

TRT01P

AGE

AGEGR1

SEX

SEX

RACE

RACE

ETHNIC

ETHNIC

WEIGHTBL

WEIGHTBL

BMBLGR1

BMBLGR1

☐ TRT01P Order

Analysing Variable(s) Mapping

AGE

AGEGR1

SEX

SEX

RACE

RACE

ETHNIC

ETHNIC

WEIGHTBL

WEIGHTBL

BMBLGR1

BMBLGR1

AGE Label

Custom Label for AGE

AGEGR1 Label

Custom Label for AGEGR1

SEX Label

Custom Label for SEX

RACE Label

Custom Label for RACE

ETHNIC Label

Custom Label for ETHNIC

WEIGHTBL Label

Custom Label for WEIGHTBL

BMBLGR1 Label

Custom Label for BMBLGR1

Generate summary table

[illegible]

7

CONCLUSION:

Data Designer is an intuitive, user-friendly R Shiny® application that streamlines clinical data transformation, analysis and summary generation. Through its interactive interface, users can manipulate, map, and summarize data without extensive programming, while leveraging a robust R framework built on *{dplyr}*, *{cards}*, and *{gtsummary}* R packages.

The tool supports simultaneous processing of multiple input datasets from diverse formats, including CSV, Excel, SAS, XPT, JSON and Parquet. Generated datasets and summaries can be exported in multiple formats, enabling seamless integration with downstream visualization or reporting workflows. Data Designer also facilitates the creation of clinical tables directly from analysis-ready datasets (ARDs).

A key feature of Data Designer is its regenerative capability: reproducible R code is automatically generated alongside each data processing and reporting step. This promotes reusability, validation, and traceability to ensure reproducibility, and maximizes generalizability.

Future enhancements will focus on expanding support for ARS-based summary creation, enabling the generation of a final report that includes all possible summaries in a single document, as well as supporting the creation of multiple summaries in one go.

REFERENCE:

1. <https://dplyr.tidyverse.org/>
2. <https://insightsengineering.github.io/cards/main/index.html>
3. <https://www.danielsjoberg.com/gtsummary/index.html>
4. <https://www.rdocumentation.org/packages/jsonlite/versions/2.0.0>
5. <https://arrow.apache.org/docs/r/>

ACKNOWLEDGMENTS:

We gratefully acknowledge our colleagues Parth Patadia (Quantzig), Vatsal Garg (Quantzig), Ramya Pichikala (Sanofi) and Sophie Bellec (Sanofi) for their contribution and suggestions to the development of this tool.

CONTACT INFORMATION:

Mukul K Mittal
Sanofi, USA
100 Morris St,
Morristown, NJ, 07960
Email: mukul.mittal@sanofi.com

Andre Couturier
Sanofi, USA
100 Morris St,
Morristown, NJ, 07960
Email: andre.couturier@sanofi.com

*Brand and product names are trademarks of their respective companies (R Studio/SAS/Microsoft).