

One eye open, one eye closed: How to approach masking in open-label oncology study

Abhay Kumar, GSK, Philadelphia, USA

Tanya Dubinsky, GSK, London, UK

ABSTRACT

Open-label designs are common in oncology due to the life-threatening nature of the disease. Moreover, ethical considerations make it crucial to avoid delaying or withholding effective treatment, ensuring that patients obtain the best available care. Practical challenges also arise because certain oncology treatments have noticeable side effects or specific modes of action that render blinding impractical. Additionally, open-label trials facilitate closer monitoring of side effects and treatment responses, which is vital in rapidly evolving conditions like cancer, allowing for timely adjustments when adverse effects occur. In order to maintain the integrity of the results and minimize the possibility of sponsor's bias, both real and perceived, a comprehensive data masking may need to be performed to avoid revealing any differences between drug of interest and control arm.

This paper outlines different techniques for masking the collected (SDTM) and the analysis data (ADAM). It examines various implementation challenges and presents approaches to address them. Additionally, the paper discusses strategies for managing different Case Report Form (CRF) scenarios and mitigating associated risks when developing a data masking plan. We will review various programming techniques for masking specific data elements and discuss how different study documentation e.g. protocol, aCRF can be used to identify the data elements that require masking.

INTRODUCTION

Open-label study designs are very common in oncology, primarily because the severe nature of cancer often leads patients to prefer transparency regarding their treatment. Ethical obligations also play a significant role, as delaying or withholding potentially effective therapies is a serious concern. Additionally, the distinct side effects or mechanisms of action associated with many oncology treatments can render traditional blinding impractical. Therefore, to ensure study integrity and mitigate potential bias, sponsors often implement comprehensive data masking strategies during the trial's execution. This paper will further explore various considerations and techniques for effectively masking data in such trials.

ACRONYMS

aCRF	Annotated Case Report Form
SAP	Statistical Analysis Plan
SDTM	Study Data Tabulation Model
ADAM	Analysis Data Model
TFL	Tables, Figures, and Listings

GENERAL CONSIDERATIONS

Study data masking plan: Developing a comprehensive data masking plan or specification is good practice. This document typically details the masking rationale and specific methods for each data point, developed in collaboration with the entire study team, including data management, clinical, safety, and biostatistics functions. This collaborative approach ensures that all masking requirements, data transfer mechanisms, and validation processes are meticulously planned, outlined, and agreed upon by all involved study members in advance. Particularly in oncology, where investigational drugs may be compared against one or more standard-of-care treatments, Case Report Forms (CRFs) often capture sensitive treatment-specific details in various forms e.g. Exposure, Study treatment discontinuation form, Adverse Events. In such instances, this critical information can either be left blank, or a dummy value can be added. Ideally, key stakeholders from the wider study team, such as clinicians, meticulously review all collected data elements to identify any information that could potentially reveal the study treatment. These potential unblinding scenarios require thorough examination and discussion with statistical and data management teams to reach a definitive decision on how to mask these data elements. To maintain consistency, the same masking principles applied to the collected data can also be applied to external data.

During the study conduct, after each protocol and/or CRF amendment, the treatment sensitivity plan should be revised.

To ensure the validity of the masking process, implementing a verification process is highly recommended. This typically involves running frequency statistics and performing manual inspections of the datasets designated for masking to confirm the successful implementation of variable masking. The validation programmer responsible for this task should maintain strict independence from the core study team. Their mandate is to accurately ascertain that all masking procedures have been performed correctly before the masked data is released for analysis.

Below, we will explore some study-specific masking scenarios:

TREATMENT ARMS

While some data masking approaches suggest using dummy subject IDs, this should be carefully considered for its suitability, as it may lead to difficulties in merging different data (e.g., ADL with another dataset) and result in inconsistencies that ultimately compromise the accuracy of Tables, Figures, and Listings (TFLs).

The most effective method involves disrupting the expected treatment arm ratios in randomized Open-Label trial, as defined by the study protocol. This approach eliminates bias and allows for data aggregation without revealing the true treatment allocations. For example, in a study design specifying a 2:1 protocol ratio, the new treatment assignment could be set to a different ratio, such as 2:1:1 or 3:4, or any other ratio decided by the team, with treatments scrambled among subjects.

The team may apply this step to either the subject-level analysis dataset or the demographics dataset (SDTM.DM), selecting the option that best fits the study scenario. Consider a scenario where some subjects receive intravenous injections in one treatment arm while others receive oral dosing in a control arm; scrambling subjects in a different ratio addresses these data elements without causing the TFLs to be completely misaligned or out of context. This method helps to minimize significant logical inconsistencies that can arise when CRFs are designed to accommodate various treatment schemes. The Appendix below contains an SAS/R code snippet designed to scramble treatment values, which can be customized based on the ratio the team decides to use.

The following snippet illustrates the dataset's appearance after scrambling, reflecting a change in the treatment ratio from 1:1 to 2:1:1. It may be carried out by programming team, using a "Randomization Flag" (e.g., 'Y' or 'N') macro variable set up at study set up level to appropriately mask treatment codes:

Subject Identifier	Original Treatment Before Masking is applied	Masked Treatment Values after Treatment Scrambling in between Subjects
Study_A_001	Study Drug A	Control Drug B
Study_A_002	Control Drug B	Study Drug A
Study_A_003	Study Drug A	Control Drug B
Study_A_004	Study Drug A	Control Drug B
Study_A_005	Control Drug B	Study Drug A
Study_A_006	Control Drug B	Study Drug C
Study_A_007	Study Drug A	Control Drug B
Study_A_008	Control Drug B	Study Drug A
Study_A_009	Control Drug B	Study Drug C
Study_A_010	Study Drug A	Control Drug B

Some of the specific exposure-related scenarios are listed below:

STUDY DRUG ADMINISTRATION

Study Drug Administration panels can be the most complicated to handle. The critical exposure information could either be masked by replacing it with dummy values or by completely removing the values. Removing the values, however, would further complicate scenarios when generating statistical summaries, especially when aggregated numeric values such as dose intensity, cumulative dose, or dose frequency are desired. Therefore, using consistent dummy values is the more efficient approach.

Dose Levels and Frequencies: Scheduled and actual current cycle dose levels (mapped to ECDOSE, EXDOSE), along with dosing frequencies and cycle patterns i.e. Dose value being consistently captured across all treatment arms, yet its specific value or configuration differs for each treatment arm. For instance, the dosing and treatment schedule is a common variable, but for Treatment A it might be 1000 mg daily, while for Treatment B it could be 300 mg every alternate day. Differences in these elements across treatment arms could unintentionally reveal actual treatments, so these data points may need to be masked.

Drug Administration Method: To prevent unblinding, even the method of drug administration (e.g., oral versus infusion) may require masking. For instance, if an experimental drug is infused while a comparator is administered orally, this difference could inadvertently reveal treatment assignments if not properly addressed. Similarly, certain variables are collected only for a specific treatment arm – known as treatment-specific or conditionally collected variables. An example would be a question like "Was the dose level changed from previous cycle?", which might only be applicable and collected for patients receiving the investigational drug via infusion. In all such cases, these data elements need masking to avoid inadvertently disclosing treatment assignments

Reasons for Dose Change or any treatment related to free text fields: Reasons for dose changes, particularly if captured via free text fields in the CRF, require careful consideration due to their highly sensitive nature. These reasons often correlate with significant events, e.g., subject missed a particular kit vial dosing which was only being collected for study treatment, not control arm and leaving these details unmasked poses a high risk of unblinding actual treatment, especially since sites may inadvertently include revealing comments. These numerous data elements within the system could inadvertently reveal a patient's assigned treatment. Therefore, it is strongly recommended to mask all such elements to prevent the database from disclosing the specific treatment each patient is receiving.

ADVERSE EVENTS

When dealing with adverse event panels, especially in studies involving combination drug regimens where the relatedness of adverse events to each specific drug is individually tracked, special care is needed. In such situations, variables like AERELn, AEACNn, and other similar fields might be masked using dummy values. It's also crucial to remember that even minor pieces of information might be stored in the SUPPAE dataset; therefore, those specific details in the SUPPAE dataset should also be considered for masking.

Additionally, other information such as Adverse Event (AE) narratives and comments can inadvertently reveal unblinding information. For example, if an injection site adverse event is reported, and only the experimental drug is administered via infusion, this detail could expose a subject's actual treatment. Such unblinding introduces bias and can undermine the credibility of study results; consequently, this type of information should also be masked.

HEADLINE RESULTS & PRIMARY AND SECONDARY END POINTS

It is often essential to preserve sponsor from seeing unmasked aggregated statistics related to primary or secondary objectives in oncology studies during the conduct of the trial. Such statistics encompass, for example, summary Kaplan Meier analyses for Progression-Free Survival (PFS) or Overall Survival (OS), PFS derived from blinded independent central review (BICR) assessments, Objective Response Rate (ORR), Duration of Response (DOR), and Disease Control Rate (DCR) (each based on both BICR and Investigator assessment), as well as Patient-Reported Outcomes (PROs) and Quality of Life scales including relevant statistics. In all such cases special care should be taken to mask these data points and treatment scrambling between subjects would ensure that even if we generate aggregated statistic outputs, the crucial primary and secondary efficacy end points are not revealed.

PHARMACOKINETIC DATASETS

The Analysis Dataset for Pharmacokinetic Concentration (ADPC) contains concentration measurements (for example, plasma or urine) recorded per subject, analyte, and time point. The Analysis Dataset for Pharmacokinetic Parameters (ADPP) holds derived metrics such as area under the curve (AUC) and maximum concentration (Cmax) etc. Certain ADPC elements require careful masking e.g., concentration result values, BLQ/LLOQ indicators, sampling time references, and free-text comment fields can reveal treatment details if left unprotected. Given that specific Below Limit of Quantitation (BLQ) or Lower Limit of Quantitation (LLOQ) values are uniquely linked to a particular drug's analyte, revealing these thresholds could allow someone to deduce the specific analyte being measured and the actual study drug. Since this information could be reverse coded to identify the analyte and potentially the treatment, this information may need to be masked appropriately.

In some study designs, pharmacokinetic sampling and analysis are performed only for the active treatment arm and not for the control arm. In such cases, to preserve study integrity and confidentiality, an independent programming team may be engaged to develop the pharmacokinetic datasets and outputs.

CONCLUSION

In conclusion, robust data masking strategies are paramount for upholding the integrity and validity of open-label oncology trials, where patient preference for treatment knowledge, ethical considerations, and practical challenges often preclude traditional blinding. The inherent transparency of these designs necessitates meticulous planning to minimize bias, both real and perceived. Therefore, establishing comprehensive and detailed Data Masking Specification is essential. This plan, developed collaboratively with the entire study team, outlines the rationale, specific techniques, and scope of masking across all relevant data sources, from source systems to SDTM, ADaM, and TFLs.

The successful implementation of data masking relies on clearly defined procedures for execution, transfer, and validation, agreed upon by all study participants in advance. A robust verification process, conducted by an independent validation programmer, is critical to confirm the accuracy of masking before data release. While integrating data masking adds complexity to programming activities, its proper application prevents the premature revelation of treatment, maintains data traceability and reproducibility, and ensures that masked summaries and regulatory outputs are accurate and defensible. Ultimately, a well-executed data masking strategy is indispensable for safeguarding study integrity and securing acceptance by regulatory authorities.

APPENDIX: SAMPLE SAS/R CODE TO SCRAMBLE TREATMENT ARMS:

Below sample code could be used to generate and assign randomized treatment arms to subjects. It first sets up all subjects with unique IDs and sequential unit numbers. It then defines a base treatment allocation, refines it to include multiple arms, and finally uses PROC PLAN to randomly scramble and assign these treatment arms to the subjects. The resulting treatment assignments, complete with descriptive names, are then merged into the main study dataset.

```
/**Dataset to be scrambled used as one of the parameter***/
%MACRO DRAND (DATAST= );
/* Counts all rows in the dataset */

PROC SQL NOPRINT;
    SELECT COUNT(*)
    INTO : NOBS1 FROM &DATAST.;
QUIT;
%PUT The total number of rows in the dataset is: NOBS1;
```

Assign the sequential observation number to the UNIT variable. This is a crucial step for creating a unique linking key.

```
DATA AD_RAND (KEEP= USUBJID UNIT);
    SET &DATAST.;
    UNIT = _n_;
    OUTPUT;
RUN;
```

Below step creates a sequential, non-randomized assignment of TREATMENT 1 and 2 to each UNIT.

```
/* Unrandomized Design */
DATA GEN1;
DO I = 1 TO &NOBS1;
UNIT=I;
    IF (UNIT <= &NOBS1*(1/2)) THEN TREATMENT=1;
ELSE TREATMENT=2;
AB= LAG(TREATMENT);
OUTPUT;
END;
RUN;
```

This step identifies the starting UNIT and total count of subjects in the TREATMENT=2 group. It creates a new dataset GEN2 where the original TREATMENT=2 group from GEN1 is split into TREATMENT=2 and TREATMENT=3, effectively creating a design with three treatment arms (e.g., 2:1:1 ratio or any other ratio as decided by team).

```
PROC SQL NOPRINT;
SELECT COUNT(TREATMENT) INTO: NOBS2 FROM GEN1 WHERE TREATMENT=2;
SELECT UNIT INTO: X FROM GEN1 WHERE AB ^=TREATMENT=2;
QUIT;
```

```
DATA GEN2;
SET GEN1;
DO I = &X. TO (&X.+&NOBS2.);
IF TREATMENT=2 THEN DO;
IF (&X.<=UNIT <= (&X.+&NOBS2./2)) THEN TREATMENT=2;
ELSE TREATMENT=3;
END;
END;
RUN;
```

/* Randomize The Design Using Proc Plan Seed */

This code block takes an existing dataset that contains a sequential UNIT variable (and potentially other variables like a non-randomized TREATMENT assignment linked to that UNIT). It then uses PROC PLAN to randomly shuffle the order of the UNIT values (and consequently, the TREATMENT assignments associated with them) and outputs this randomized sequence into a new dataset (GEN3). This GEN3 dataset can then be merged with subject IDs to assign treatments randomly to subjects.

SEED=12345: This option sets the seed for the random number generator. Using a SEED value ensures that if you run this code multiple times with the same seed, you will get the exact same "random" sequence. This is crucial for reproducibility, especially in clinical trials or other scientific studies where the randomization scheme needs to be fixed and verifiable.

```
PROC PLAN SEED=12345;
FACTORS UNIT=&NOBS1/ NOPRINT;
      OUTPUT DATA=GEN2 OUT=GEN3;
RUN;
```

/* Merge the randomized design with the existing rand CRF data and populate scrambled treatment values as decided by team */

```
DATA TRTR (KEEP= USUBJID TRT01A NW_TRT);
MERGE AD_RAND GEN3;
BY UNIT;
IF TREATMENT = 1 THEN NW_TRT = "Study Drug C";
ELSE IF TREATMENT = 2 THEN NW_TRT = "Study Drug B";
ELSE IF TREATMENT = 3 THEN NW_TRT = " Study Drug A";
RUN;
```

```
DATA &DATAST.;
MERGE &DATAST. TRTR;
BY USUBJID;
RUN;
```

```
%MEND;
```

R CODE FOR PERFORMING TREATMENT SCRAMBLING SIMILAR TO SAS ABOVE -

#Load Libraries, Load necessary R packages. These packages provide functions for data manipulation, statistical analysis, date handling, and reading SAS files.

```
LIBRARY (ADMIRAL) # Toolkit for ADaM dataset creation (Clinical Data Standards)
LIBRARY (DPLYR) # Data manipulation functions (often loaded by tidyverse)
LIBRARY (STRINGR) # Functions for string manipulation (often loaded by tidyverse)
#LIBRARY (TIDYR) # For tidying data (e.g., pivot_longer, pivot_wider) (often loaded by tidyverse)
```

Initial processing of ADSL data for treatment arm recording

```
AD_RAND <- ADSL %>%
  SELECT (USUBJID, TRT01A) %>%
  ARRANGE (USUBJID) %>%
  )
```

Get the number of observations (subjects) from the processed ADSL data

```
NOBS1 <- NROW (ADSL1)
```

Emulating SAS 'DATA GEN1' step: Initial, unrandomized treatment assignment

```
GEN1 <- TIBBLE (UNIT = 1:NOBS1) %>% #
  MUTATE (
    TREATMENT = IF_ELSE (UNIT <= NOBS1 * (1 / 2), 1, 2),
    AB = LAG (TREATMENT)
  )
```

Prepare ADSL data for merging with randomized assignments later.

This part ensures a consistent USUBJID and generates a UNIT variable

that will be used to link to the randomized assignments.

```
ADSL1_PROCESSED <- AD_RAND%>%
  ARRANGE (USUBJID) %>%
  MUTATE (
    USUBJID = STR_TRIM (USUBJID),
    UNIT = ROW_NUMBER () # Assigns a sequential 'UNIT' number (1 to nob1) to each subject. This
    is good practice for reproducibility if the order matters for UNIT generation.
```

Emulating SAS PROC SQL to get NOBS2 and X (Unit_X)

NOBS2: Count of subjects initially assigned TREATMENT=2 in gen1

```
NOBS2 <- GEN1 %>%
  FILTER (TREATMENT == 2) %>% # Filters for rows where TREATMENT is 2
  nrow() # Counts the number of rows
```

x_val: The first UNIT number where TREATMENT changes from 1 to 2.

In a 1:1 split, this will be (nob1/2) + 1.

```
X_VAL <- FLOOR (NOBS1 * (1 / 2)) + 1 # CALCULATES THE STARTING POINT OF
TREATMENT=2
```

Unit_X: The UNIT number where the TREATMENT value first changes to This specifically looks for the first instance where previous treatment) is not equal to the current TREATMENT, and finds the transition point

```
UNIT_X <- GEN1 %>%
  FILTER (AB != TREATMENT & TREATMENT == 2) %>%
  PULL (UNIT) # Extracts the UNIT value as a vector
CAT (PASTE ("TREATMENT 2 BEGINS", UNIT_X, "\N")) #Prints the starting unit of Treatment 2
```

```

# Subdividing TREATMENT=2 into TREATMENT=2 and TREATMENT=3
#This logic splits the original '2's into '2' and '3' to create a 3-arm design (e.g., 2:1:1 or any ratio as
decided by team)
GEN2 <- GEN1 %>%
  MUTATE (
    TREATMENT = CASE_WHEN (
      # First half of original TREATMENT=2 block remains 2
      TREATMENT == 2 & UNIT >= X_VAL & UNIT <= (X_VAL + FLOOR(NOBS2 / 2) - 1) ~ 2,
      # Second half of original TREATMENT=2 block becomes 3
      TREATMENT == 2 & UNIT >= (X_VAL + FLOOR(NOBS2 / 2))
      & UNIT <= (X_VAL + NOBS2 - 1) ~ 3,
      TRUE ~ TREATMENT # Keep other treatments (TREATMENT=1) as they are (default case)
    )
  )
SET.SEED(12345) # Sets the random seed for reproducibility (matches SAS SEED) Emulating SAS
PROC PLAN SEED=12345; Generate the randomized 'UNIT' factor.
RANDOM_UNIT_ASSIGNMENT <- SAMPLE(1:NROW(GEN2)) # GENERATES A VECTOR OF
SHUFFLED UNIT NUMBERS

# Create 'GEN3' by copying 'GEN2' and replacing its UNIT column with the randomized 'UNIT'. #
TREATMENT assignments to the new 'random_unit_assignment' order.
GEN3 <- GEN2
GEN3$UNIT <- RANDOM_UNIT_ASSIGNMENT # Replaces the sequential UNIT with the randomized
UNIT values

# Merge the randomized design to the existing subject data (adsl1_processed) This step assigns the
scrambled (randomized) treatment arm to each subject.
TRTR <- ADSL1_PROCESSED %>%
  LEFT_JOIN(GEN3, BY = "UNIT") %>%
  MUTATE (
    NW_TRT = CASE_WHEN (
      TREATMENT == 1 ~ "STUDY DRUG C",
      TREATMENT == 2 ~ "STUDY DRUG B",
      TREATMENT == 3 ~ "STUDY DRUG A",
      TRUE ~ NA_CHARACTER_ # Default for any unexpected TREATMENT values
    )
  ) %>%
  select(USUBJID, TRT01A, NW_TRT) %>% # Selects relevant columns for the final output
  ARRANGE(USUBJID) # Sorts the final dataset by USUBJID

```

REFERENCES

- Considerations for open-label clinical trials: design, conduct, and analysis Karen M. Higgins, FDA/CDER/OTS/OB/DBIII, and Gregory Levin, FDA/CDER/OTS/OB Center for Drug Evaluation and Research, Food and Drug Administration, Silver Spring, MD
- Randomization Technique, Allocation Concealment, Masking, And Susceptibility Of Trials To Selection Bias Vance W. Berger Biometry Research Group, National Cancer Institute, vb78c@nih.gov, Costas A. Christophi George Washington University

CONTACT INFORMATION

We welcome your comments and questions. Please contact the authors at:

Abhay Kumar
GlaxoSmithKline Inc
5 Crescent Drive Philadelphia, PA 19112, USA
Email: abhay.3.kumar@gsk.com

Tanya Dubinsky
GlaxoSmithKline Inc
London, UNITED KINGDOM
Email: tanya.x.dubinsky@gsk.com