

Handling Missing Data in Analgesia Clinical Trials and Corresponding SAS Implementation

Jingbo Zhu, Jiehan Zhang, Jiangsu Hengrui Pharmaceuticals Co., Ltd.

ABSTRACT

In analgesic clinical trials, primary endpoints are mostly patient-reported pain scores, such as Numeric Rating Scale (NRS), Visual Analogue Scale (VAS), Western Ontario and McMaster Universities Osteoarthritis Index (WOMAC) etc. Missing data is prevalent, mainly caused by premature discontinuation, use of rescue therapy, adverse events, poor participants' compliance, or missed assessments. Such incompleteness increases the complexity of statistical analysis and may lead to biased estimates of clinical efficacy.

This paper addresses missing data handling in analgesia clinical trials with repeated measurements and longitudinal observations. We focus on two primary methods under the Missing at Random (MAR) assumption—Multiple Imputation (MI) and Mixed-Effect Model Repeated Measures (MMRM), also review Windowed Last Observation Carried Forward (wLOCF) as a practical imputation approach for rescue therapy. By bridging statistical methodology with practical SAS implementation, this paper helps clinical programmer to handle missing data more effectively in analgesia clinical trials.

INTRODUCTION

Analgesia clinical trials are essential for evaluating the efficacy and safety of novel analgesic agents and therapeutic interventions. These trials commonly use patient-reported outcomes (PROs) as primary endpoints, including the 0-10 Numeric Rating Scale (NRS), the Visual Analogue Scale (VAS), subscale scores from the Western Ontario and McMaster Universities Osteoarthritis Index (WOMAC), and the Patient Global Assessment of Pain Improvement (PGA). Repeated-measures designs are often employed to capture temporal changes in pain intensity over the course of treatment.

Owing to frequent data collection and extended follow-up, missing data are an almost unavoidable feature of these trials. Missingness can undermine statistical precision, introduce bias, and ultimately distort the evaluation of treatment efficacy. In practice, missing data arises from a wide array of factors, including but not limited to participant loss to follow-up, data entry errors, and technical malfunctions. The presence of missing data poses substantial challenges to statistical analysis, as distinct missing data mechanisms—namely Missing Completely at Random (MCAR), Missing at Random (MAR), and Missing Not at Random (MNAR)—demand tailored analytical strategies. To guarantee the accuracy and reproducibility of clinical trial findings, regulatory authorities including the US Food and Drug Administration (FDA), the European Medicines Agency (EMA), and the ICH E9(R1) guideline have established explicit specifications for data integrity.

Accordingly, the effective handling of missing data in analgesia trials has become an important focus in biostatistics and clinical research. This paper reviews and operationalizes a broad range of methods for missing data management in analgesia studies in order to provide a practical methodological reference for future research. Drawing on a structured literature review, we describe the major missing data mechanisms, conventional statistical approaches aligned with those mechanisms, common missing data scenarios in analgesia trials, and corresponding mitigation strategies. We further illustrate the practical implementation of these methods in SAS.

MISSING DATA PATTERNS IN ANALGESIA TRIALS

MISSING DATA MECHANISMS

Missing data can occur under three distinct mechanisms, each of which calls for a specific analytical approach:

Missing Completely at Random (MCAR): The missing values have no correlation with other values in the dataset, neither observed nor missing. Under MCAR, estimates remain consistent across a wide range of analytical methods, and complete case analysis is generally valid.

Missing at Random (MAR): the propensity of the missing value of a variable is related to the observed variables, but not related to the value of the variable itself. Under this assumption, both Multiple Imputation (MI) and Mixed Models for Repeated Measures (MMRM) can provide asymptotically unbiased estimates.

Missing Not at Random (MNAR): missingness is associated with the unobserved values themselves and may also be related to observed variables. This mechanism is analytically challenging because the reason for missingness is directly linked to the true, unrecorded outcome. For example, a participant's Numerical Rating Scale (NRS) score may be missing because pain became intolerable, a scenario that typically reflects MNAR.

Table 1 summarizes recommended statistical methods under each missing data mechanism.

Missing Mechanism	Statistical Assumption	Recommended Method(s)
MCAR (Missing Completely at Random)	Missingness is unrelated to both observed and unobserved outcome values.	Complete case analysis, simple single imputation, multiple imputation, MMRM
MAR (Missing at Random)	Missingness depends solely on observed data.	Multiple imputation, expectation-maximization (EM) algorithm, MMRM
MNAR (Missing Not at Random)	Missingness is associated with the unobserved values themselves, even after conditioning on all available observed information.	Pattern-mixture models, sensitivity analyses

Table 1. Statistical Methods by Missing Data Mechanism

A common example in analgesia trials is a participant who requests rescue therapy because of worsening pain, after which subsequent NRS assessments become unavailable. In this situation, the missingness is likely related to the underlying pain intensity, suggesting an MNAR tendency. Conservative sensitivity analyses, such as Worst Observation Carried Forward (WOCF), may therefore be warranted.

Sensitivity analyses are intended to evaluate whether conclusions are robust to departures from the assumptions used in the primary analysis. By supplementing the methods summarized in the table above with corresponding sensitivity analyses, investigators can more comprehensively evaluate the influence of missing data on analytical results and derive more reasonable interpretations.

MISSING DATA SCENARIOS IN ANALGESIA TRIALS

In repeated-measures analgesia trials, missing pain scores, such as Numerical Rating Scale (NRS) values, can be categorized according to their timing relative to the first and last valid observations:

Early missingness (pre-first valid observation): values occurring before the first valid observation are generally assumed to follow a Missing at Random (MAR) mechanism and are commonly imputed using a rule-based approach, such as imputing the median score from the same treatment group at the corresponding visit.

Intermittent missingness (between two valid observations): Missing records flanked by valid observations on both sides are presumed to satisfy the MAR assumption, linear interpolation with a time-weighted slope is routinely applied for imputation.

Late missingness (after the last valid observation):

- Missing data arising after the final assessment are generally treated under the MAR assumption and handled via Last Observation Carried Forward (LOCF).
- If discontinuation is due to adverse events, withdrawal of informed consent, or participant initiated treatment termination, the missingness is classified as Missing Not at Random (MNAR); in this

setting, Worst Observation Carried Forward (WOCF) is adopted to provide a conservative evaluation.

For participants who permanently discontinue the trial with no planned follow-up assessments, the missing data pattern is monotone missingness; under a MAR assumption, monotone regression-based Multiple Imputation (MI) is a recommended option.

Model-based handling of missing data: for continuous repeated measures endpoints under a MAR, Mixed Models for Repeated Measures (MMRM) are routinely used, and no separate imputation step is required.

SPECIAL MISSING SCENARIOS

Concomitant Interventions (Rescue Analgesia or Prohibited Concomitant Medications)

Define the rescue time window by setting rescue onset as the rescue start time plus one elimination half-life of the rescue medication.

- For scores within this window, apply windowed Last Observation Carried Forward (wLOCF), using the value immediately before rescue onset, or windowed Worst Observation Carried Forward (wWOCF), using the worst value before rescue onset.
- If multiple rescue medications are administered or rescue dosing is repeated, the n-th and all subsequent rescue events should be handled according to prespecified special rules. For example, if rescue medication is used cumulatively ≥ 3 times, all values recorded after the third rescue event should be replaced with the most recent valid score obtained before that third rescue event.

Item-Level Missingness of Outcome Scales (Example: WOMAC Index)

For each WOMAC subscale (Pain, Stiffness, Physical Function):

- If the number of missing items within a subscale does not exceed the prespecified threshold, the total subscale score may be imputed using the mean of the available items within that subscale.
- If the number of missing items exceeds the threshold, the total score for that subscale is set to missing, with no further imputation performed.

Missing Data Attributable to Special Reasons

Example: Missing Resting Pain Scores, for example NRS values, Due to Sleep.

- If the last resting NRS score recorded before sleep is less than 3, apply Last Observation Carried Forward (LOCF).
- If the last resting NRS score recorded before sleep is 3 or greater, assign a fixed value of 3, corresponding to the threshold for mild pain.

MISSING DATA HANDLING APPROACHES AND SAS IMPLEMENTATION

This section covers four analytic approaches: complete case analysis without imputation, single imputation, multiple imputation, and mixed models for repeated measures (MMRM).

COMPLETE CASE ANALYSIS WITHOUT IMPUTATION

Analyses are conducted directly on datasets that retain missing values.

For example, under a treatment policy strategy, all observed participant level outcome data, including post-intercurrent event measurements, are analyzed directly for between group comparisons, irrespective of whether intercurrent events occur, with no exclusion, imputation, or ad hoc adjustment of these data.

SINGLE IMPUTATION METHODS

Single imputation replaces each missing observation with a single estimated value.

Linear Interpolation

- Principle: missing values are imputed from the linear relationship between the preceding and subsequent valid observations. Mathematically, linear interpolation assumes that the outcome changes linearly between two scheduled assessment time points.
- Applicability: this method is appropriate when valid observations are available on both sides of the missing time point, the outcome changes relatively smoothly over the interval, and the scheduled time interval is well defined.
- Mathematical Formula:

$$y = k \times t + b$$

$$k = (y_2 - y_1) / (t_2 - t_1) \text{ [slope]}$$

$$b = (t_2 \times y_1 - t_1 \times y_2) / (t_2 - t_1) \text{ [intercept]}$$

Variable definitions: y denotes the imputed value, t denotes time, and the output classification is designated as "linear."

Mean, Median, or Mode Imputation

- Principle: missing values may be replaced with the mean or median for continuous variables, or the mode for categorical variables. For continuous endpoints, mean imputation is generally appropriate when the data are approximately normally distributed, whereas median imputation is preferable for markedly skewed data. For non-numeric variables, mode imputation fills missing entries with the most frequently observed category among complete observations.
- Applicability: this approach is most suitable for early missing data, or for complete absence of all scale scores in an individual participant. Example: for missing early post baseline scale scores, values may be replaced with the median Numerical Rating Scale (NRS) score among participants in the same treatment group at the corresponding scheduled visit.

LOCF, BOCF, and WOCF Imputation

- Principle: Last Observation Carried Forward (LOCF) carries the most recent valid score forward to subsequent missing time points and has been widely used in longitudinal and time series clinical studies. Baseline Observation Carried Forward (BOCF) replaces missing values with each participant's baseline measurement. Worst Observation Carried Forward (WOCF) imputes missing values using the worst observed outcome during the trial.
- Applicability: all three methods are most defensible under strong assumptions, often approximating MCAR. BOCF treats the baseline value as the final response, regardless of the reason for withdrawal or the pain score at discontinuation. WOCF is often used when participants discontinue because of lack of efficacy, thereby providing a conservative assessment of analgesic benefit.
- wLOCF / wWOCF (Windowed LOCF / WOCF): the prefix "w" generally denotes a predefined time window. These methods are commonly used in sensitivity analyses to evaluate the robustness of study conclusions. By applying a more conservative imputation strategy, investigators can assess whether the treatment effect remains evident under less favorable assumptions. In practice, the worst score within a prespecified window is typically carried forward for wWOCF; more rarely, "w" may also be interpreted as weighted. Example: when rescue analgesia is administered, the most recent pain score recorded before rescue may be used for imputation within the prespecified time window. If the observed score exceeds the substitute value, no imputation is applied.

MULTIPLE IMPUTATION

Principles and Three Stage Framework of Multiple Imputation (MI)

Multiple imputation (MI) is a model-based approach to handling missing data under the assumption that the data are at least missing at random. Rather than replacing missing values with a single substitute, MI

generates (m) completed datasets, analyzes each dataset separately, and then pools the results.

- Imputation stage: generate m complete datasets, typically with $m \geq 10$, using PROC MI based on the observed data.
- Analysis stage: fit the same statistical model independently to each imputed dataset, for example, an analysis of covariance (ANCOVA) model or PROC MIXED.
- Pooling stage: combine the m point estimates and their associated variances according to Rubin's rules using PROC MIANALYZE.

In SAS, the MI workflow is implemented by running PROC MI for the imputation stage and PROC MIANALYZE for the pooling stage. PROC MIANALYZE not only combines estimates across imputed datasets but also provides inferential statistics, including confidence intervals and P values for the pooled results.

This three-stage framework allows MI to address missing data problems efficiently while improving the accuracy and reliability of statistical inference. It also enables investigators to account for plausible missing data mechanisms more explicitly, thereby supporting more defensible interpretation of trial findings.

PROC MI Implementation

Multiple imputation generates several complete imputed datasets by performing repeated imputation on incomplete datasets with missing values.

A missing data pattern describes the structural distribution of missing values within a dataset and can be classified into two types:

- Monotone missingness: For an ordered set of variables, once a value is missing at a given variable, all subsequent variables are also missing. A typical example is permanent treatment discontinuation after study withdrawal. Under this pattern, multiple imputation (MI) can, through variable-by-variable imputation, approximate the dependence structure among imputed values; accordingly, different distributional assumptions may be specified for continuous and categorical variables, as in standard practice.

	Y1	Y2	Y3	Y4
1	X	X	X	
2	X	X		
3	X			

Figure 1. Monotone missing pattern

- Arbitrary (non-monotone) missingness: Participants may re-enter follow-up after an earlier missed visit or become lost to follow-up again at later visits. In this setting, MI requires joint imputation of all variables, treating them as a multivariate response. A common distributional framework should be applied across both continuous and categorical data.

	Y1	Y2	Y3	Y4
1	X	X	X	X
2	X	X		
3	X		X	

Figure 2. Non-monotone missing pattern

Sample reference SAS code is provided below in accordance with official SAS documentation:

```
PROC MI <options>;
BY variables;
CLASS variables;
MONOTONE <options>;
EM <options>;
FCS <options>;
FREQ variable;
MCMC <options>;
TRANSFORM transform (variables </ options>) <...transform (variables </
options>) >;
VAR variables;
```

BY	Defines groups for which separate multiple imputation analyses are conducted.
CLASS	Identifies categorical variables among those specified in VAR; these variables may be either character or numeric.
MONOTONE	Under a monotone missing data pattern, continuous variables may be imputed using regression, predictive mean matching, or propensity-score methods. For categorical variables, logistic regression is applicable to binary, nominal, and ordinal outcomes, whereas discriminant-function methods are applicable to binary and nominal outcomes.
FCS	For arbitrary missing data patterns, FCS can be applied when a joint distribution of all variables is assumed to exist. Rather than imposing a single multivariate distribution on all variables, FCS specifies a separate conditional distribution for each variable to be imputed. Imputation generally proceeds in two stages: • an initialization stage, in which missing values are filled sequentially one variable at a time to generate starting values; • an iterative stage, in which missing values are repeatedly updated in sequence for each variable until convergence criteria are met.
MCMC	When multivariate normality is plausible, MCMC can be used to impute all missing values or to perform partial imputation that yields a monotone missing pattern. Once the data become monotone, more flexible monotone model approaches can be used (for example, monotone regression without Markov-chain simulation), and different covariate sets may be specified for different imputed variables.
EM	Assuming a multivariate normal distribution, the EM algorithm is an iterative method for obtaining maximum likelihood estimates (MLEs) in the presence of missing data.
TRANSFORM	Specifies variables to be transformed prior to imputation; imputed values are then back transformed to the original scale.
VAR	Lists numeric variables for analysis. If VAR is omitted, all numeric variables not referenced in other statements are included by default.

Table 2. PROC MI Statements

PROC MIANALYZE and Rubin's Rules

Rubin's combination formulas:

Point estimate: $\bar{Q} = (1/m) \sum \hat{Q}_j$ (average over m imputed datasets)

Within-imputation variance: $\bar{W} = (1/m) \sum \hat{U}_j$

Between-imputation variance: $B = (1/(m-1)) \sum (\hat{Q}_j - \bar{Q})^2$

Total variance: $T = \bar{W} + (1 + 1/m) \cdot B$

Degrees of freedom: Barnard–Rubin approximation (accounts for finite m)

Rubin's rules combine the mean parameter estimate across imputed datasets with both within and between imputation variability. This yields pooled estimates and standard errors that appropriately reflect uncertainty due to missing data and support more robust inference.

Sample reference SAS code is provided below in accordance with official SAS documentation:

```
PROC MIANALYZE <options>;
BY variables;
CLASS variables;
MODELEFFECTS effects;
<label:> TEST equation1 <, ..., <equationk>> </ options>;
STDERR variables;
```

BY	Defines groups for which separate analyses are conducted.
CLASS	Identifies categorical variables among those specified in MODELEFFECTS; these variables may be either character or numeric.
MODELEFFECTS	Specifies the effects (parameters) to be analyzed.
STDERR	Lists the standard errors corresponding to effects in MODELEFFECTS when parameter estimates and standard errors are stored as variables in the same DATA= data set. This statement is valid only when each MODELEFFECTS effect is represented by a single continuous variable.
TEST	Tests linear hypotheses for model parameters.

Table 3. PROC MIANALYZE Statements

Example SAS Implementation

Case Background

Variables and endpoint: Sum of Pain Intensity Differences over 0–48 hours (SPID). An analysis of covariance (ANCOVA) model was used to estimate the least squares mean (LSMean) differences between each active dose group and the placebo group.

Data description: A total of 299 participants were enrolled. Pain intensity was assessed using the 0–10 Numerical Rating Scale (NRS) at prespecified time points: before administration of the loading dose of the investigational product, immediately after completion of injection, and at 30 min, 1, 2, 3, 4, 6, 8, 10, 12, 18, 24, 36, and 48 h. The corresponding analysis variables were named NRS1 to NRS15, where NRS1 represented the baseline measurement.

Data Preparation

The original data were structured in BDS (Basic Data Structure) format under the CDISC ADaM standard. To accommodate PROC MI, the NRS data were transposed into wide format, creating a dataset where each row corresponds to a single subject and the 15 time-point NRS values are stored as separate columns (NRS1–NRS15).

	USUBJID	TRT01PN	ATPIN	AVAL	PRSCAT	NRSBL
1	01001		3	0	6 XXX	6
2	01001		3	0.17	4 XXX	6
3	01001		3	0.5	4 XXX	6
4	01001		3	1	3 XXX	6
5	01001		3	2	3 XXX	6
6	01001		3	3	3 XXX	6
7	01001		3	4	3 XXX	6
8	01001		3	6	2 XXX	6
9	01001		3	8	2 XXX	6
10	01001		3	10	2 XXX	6
11	01001		3	12	2 XXX	6
12	01001		3	18	2 XXX	6
13	01001		3	24	2 XXX	6
14	01001		3	36	XXX	6
15	01001		3	48	XXX	6

Figure 3. NRS Dataset: Example BDS Dataset

	USUBJID	TRT01PN	PRSCAT	NRSBL	NRS1	NRS2	NRS3	NRS4	NRS5	NRS6	NRS7	NRS8	NRS9	NRS10	NRS11	NRS12	NRS13	NRS14	NRS15
1	01001		3 XXX	6	6	4	4	3	3	3	3	2	2	2	2	2	2		

Figure 4. NRS_T Dataset: Example Wide Format Data

Missing Data Profile and Mechanism Assumption

PROC MI diagnostic output (NIMPUTE = 0) was used to characterize missing patterns without generating imputed datasets.

```
PROC MI data = NRS_T SIMPLE NIMPUTE = 0;
run;
```

Missing Data Patterns																																				
Group	Missing Data Patterns															Group Means																				
	TRT01PN	NRSBL	NRS1	NRS2	NRS3	NRS4	NRS5	NRS6	NRS7	NRS8	NRS9	NRS10	NRS11	NRS12	NRS13	NRS14	NRS15	Freq	Percent	TRT01PN	NRSBL	NRS1	NRS2	NRS3	NRS4	NRS5	NRS6	NRS7	NRS8	NRS9	NRS10	NRS11	NRS12	NRS13	NRS14	NRS15
1	X	X	X	X	X	X	X	X	X	X	X	X	X	X	X	X	X	269	99.00	2.903378	4.935911	4.935911	3.389940	3.216996	3.169940	3.209499	3.243243	3.250000	3.422287	2.945945	2.783784	2.726351	2.627703	2.716996	2.396245	1.789405
2	X	X	X	X	X	X	X	X	X	X	X	X	X	X	X	X	X	1	0.33	2.000000	7.000000	7.000000	5.000000	6.000000	5.000000	5.000000	5.000000	5.000000	5.000000	3.000000	3.000000	6.000000	3.000000	3.000000	3.000000	
3	X	X	X	X	X	X	X	X	X	X	X	X	X	X	X	X	X	2	0.67	3.000000	5.000000	5.000000	3.500000	3.500000	2.500000	2.500000	2.500000	2.500000	2.000000	2.000000	2.000000	2.000000	2.000000	2.000000	2.000000	

Figure 5. Missing Data Distribution

Missing values exclusively occurred at the end of the observation window (44h and 48h), forming a monotone missing pattern that satisfies the prerequisite for monotonic imputation methods.

Missingness in this trial was predominantly driven by premature participant discontinuation. The reasons for withdrawal (including insufficient therapeutic effect and adverse events) were correlated with baseline characteristics and early analgesic responses. All participants had complete early NRS measurements (NRS1–NRS13), which served as strong predictive variables for the imputation model. Accordingly, the Missing at Random (MAR) mechanism was assumed for the analysis.

Imputation Procedures Using PROC MI

Reference code for this case is shown below:

```
PROC MI data = NRS_T out = MI SEED = 20260723 NIMPUTE = 10
  MIN =. . 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0
  MAX =. . 10 10 10 10 10 10 10 10 10 10 10 10 10 10 10 10
  ROUND =. . 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1 ;
  CLASS trt01p prscat;
  MONOTONE reg(nrs2--nrs15);
  VAR trt01p prscat nrs1 nrs2--nrs15;
RUN;
```

SEED: fixes the random number seed to ensure analytical reproducibility.

NIMPUTE=10, 10 complete datasets with distinct imputed values are generated.

MIN / MAX: in this case, the range is set to 0–10 consistent with the valid NRS scoring range, and the number of boundary values matches the total variables listed in the VAR statement.

MONOTONE imputation method: Applies monotonic regression imputation, applicable only to monotone missing patterns. Under such a pattern, if a variable is missing, all subsequent sequential variables are also missing.

IMPUTATION	USUBJID	TRT01PN	PRSCAT	NRSBL	_NAME_	NRS1	NRS2	NRS3	NRS4	NRS5	NRS6	NRS7	NRS8	NRS9	NRS10	NRS11	NRS12	NRS13	NRS14	NRS15
1	01001	3	xxx	6	AVAL	6	4	4	3	3	3	3	2	2	2	2	2	2	2	2
2	01001	3	xxx	6	AVAL	6	4	4	3	3	3	3	2	2	2	2	2	2	3	2
3	01001	3	xxx	6	AVAL	6	4	4	3	3	3	3	2	2	2	2	2	2	2	1
4	01001	3	xxx	6	AVAL	6	4	4	3	3	3	3	2	2	2	2	2	2	1	2
5	01001	3	xxx	6	AVAL	6	4	4	3	3	3	3	2	2	2	2	2	2	2	1
6	01001	3	xxx	6	AVAL	6	4	4	3	3	3	3	2	2	2	2	2	2	3	2
7	01001	3	xxx	6	AVAL	6	4	4	3	3	3	3	2	2	2	2	2	2	3	0
8	01001	3	xxx	6	AVAL	6	4	4	3	3	3	3	2	2	2	2	2	2	3	2
9	01001	3	xxx	6	AVAL	6	4	4	3	3	3	3	2	2	2	2	2	2	1	0
10	01001	3	xxx	6	AVAL	6	4	4	3	3	3	3	2	2	2	2	2	2	5	3

Figure 6. MI Dataset: Completed Data After Imputation

For subsequent analytical procedures, the wide-format dataset is typically converted back to BDS dataset.

IMPUTATION	USUBJID	TRT01PN	PRSCAT	NRSBL	_NAME_	AVAL
1	01001	3	xxx	6	NRS15	2
2	01001	3	xxx	6	NRS15	2
3	01001	3	xxx	6	NRS15	1
4	01001	3	xxx	6	NRS15	2
5	01001	3	xxx	6	NRS15	1
6	01001	3	xxx	6	NRS15	2
7	01001	3	xxx	6	NRS15	0
8	01001	3	xxx	6	NRS15	2
9	01001	3	xxx	6	NRS15	0
10	01001	3	xxx	6	NRS15	3

Figure 7. MI_T Dataset: Final Transformed Data

Analysis stage

After obtaining the output dataset with all non-missing values, analysis results can be readily generated using either model-based or non-model-based methods. In this example, the MI imputed dataset was used to calculate SPID and perform model-based analysis, thereby generating model estimates.

In this example, the MI_T dataset was first used to compute SPID, with results output to the MI_T_SPID dataset. An analysis of covariance (ANCOVA) model was then fitted to estimate treatment effects, incorporating baseline resting NRS score (nrsbl) as a covariate and treatment group (trt01pn) and stratification factor (prscat) as fixed effects.

```
ODS OUTPUT LSMEANS = MIXED;
PROC MIXED data = MI_T_SPID;
  BY _imputation_;
  CLASS trt01pn prscat;
  MODEL MI_T_SPID = trt01pn prscat nrsbl;
  LSMEANS trt01pn /cl pdiff tdiff;
run;
```

④ _IMPUTATION_	④ EFFECT	④ TRT01PN	④ ESTIMATE	④ STDERR	④ DF	④ TVALUE	④ PROBT	④ ALPHA	④ LOWER	④ UPPER
1	TRT01PN	3	-128.28	6.0071	293	-21.36	<.0001	0.05	-140.11	-116.46
2	TRT01PN	3	-128.70	5.9919	293	-21.48	<.0001	0.05	-140.49	-116.91
3	TRT01PN	3	-128.57	5.9931	293	-21.45	<.0001	0.05	-140.37	-116.78
4	TRT01PN	3	-128.87	5.9998	293	-21.48	<.0001	0.05	-140.68	-117.06
5	TRT01PN	3	-128.72	5.9953	293	-21.47	<.0001	0.05	-140.52	-116.92
6	TRT01PN	3	-128.27	5.9949	293	-21.40	<.0001	0.05	-140.07	-116.47
7	TRT01PN	3	-128.72	5.9953	293	-21.47	<.0001	0.05	-140.52	-116.92
8	TRT01PN	3	-128.27	5.9993	293	-21.38	<.0001	0.05	-140.08	-116.46
9	TRT01PN	3	-128.74	6.0090	293	-21.43	<.0001	0.05	-140.57	-116.92
10	TRT01PN	3	-128.24	6.0017	293	-21.37	<.0001	0.05	-140.05	-116.43

Figure 8. MIXED Dataset: Model Analysis Results

Pooling Stage

Rubin's rules are applied to pool the parameter estimates obtained from model across all imputed datasets.

```
ODS OUTPUT PARAMETERESTIMATES=LSMEANS_STD;
PROC MIANALYZE data=MIXED;
  BY trt01p;
  MODELEFFECTS estimate;
  STDERR stderr;
RUN;
```

The MIANALYZE Procedure

周期01计划治疗(N)=3

Model Information	
Data Set	WORK.MIXED
Number of Imputations	10

Variance Information (10 Imputations)							
Parameter	Variance			DF	Relative Increase in Variance	Fraction Missing Information	Relative Efficiency
	Between	Within	Total				
estimate	0.060699	35.984803	36.051572	2.62E6	0.001855	0.001853	0.999815

Parameter Estimates (10 Imputations)									
Parameter	Estimate	Std Error	95% Confidence Limits		DF	Minimum	Maximum	Theta0	t for H0: Parameter=Theta0 Pr > t
estimate	-128.538781	6.004296	-140.307	-116.771	2.62E6	-128.873097	-128.238653	0	-21.41 <.0001

Figure 9. PROC MIANALYZE Results

In practice, the valid application of Rubin's rules depends on choosing appropriate statistical models and assumptions. Investigators must also pay close attention to the number and quality of imputed datasets to ensure reliable and valid results. When used appropriately, PROC MIANALYZE together with Rubin's rules provides an efficient framework for analyzing clinical trial data with missing observations.

MIXED-EFFECT MODEL REPEATED MEASURES

MMRM Principle

Mixed Model Repeated Measures (MMRM) is a widely used method for handling missing and unbalanced longitudinal data. It analyzes repeated outcomes through a linear mixed model incorporating fixed and random effects. In analgesic trials, MMRM effectively accommodates unequally timed observations and missing assessments at some visits.

Core MMRM principles are summarized below:

- Data structure: MMRM applies to longitudinal datasets with serial measurements over multiple visits. Each participant can contribute multiple observations, typically continuous and occasionally ordinal/categorical outcomes.
- Model specification: fixed effects are used to estimate treatment effects, time effects, and treatment-by-time interactions. Random effects account for between participant variability under a prespecified distribution (for example, normality).
- Missing data handling: MMRM estimates parameters using all available observed data rather than only complete cases. This supports valid inference in the presence of missingness and can improve analytical efficiency and precision.
- Applicable scenarios: MMRM is particularly suitable for datasets with incomplete observations, a common issue in trials due to withdrawal and early discontinuation.

Advantages:

- No single imputation required: REML based likelihood uses all available observations, including partial data from early dropouts.
- Flexible model structure: fixed effects (treatment, time, interaction, baseline covariates) plus within participant covariance specification for repeated measures.
- MAR assumption: when missingness depends only on observed data, MMRM provides asymptotically unbiased estimation.

This method is especially suitable for continuous repeated-measures endpoints, such as WOMAC total/subscale scores, weekly mean NRS pain scores, and changes in PGA from baseline to end of treatment, etc.

PROC MIXED Implementation

Core PROC MIXED statements for MMRM include model definition, covariance-structure specification, and key output configuration.

Sample reference SAS code is provided below in accordance with official SAS documentation:

```
PROC MIXED <options>;
BY variables;
CLASS variables;
ID variables;
MODEL dependent = <fixed-effects> </ options>;
RANDOM random-effects </ options>;
REPEATED <repeated-effect></ options>;
PARMS (value-list) ...</ options>;
PRIOR <distribution> << / options>;
CONTRAST 'label' <fixed-effect values ...><| random-effect
values ...>, ...</ options>;
ESTIMATE 'label' <fixed-effect values ...><| random-effect values ...></
options>;
LSMEANS fixed-effects </ options>;
LSMESTIMATE model-effect lsestimate-specification < / options>;
SLICE model-effect </ options>;
STORE <OUT=>item-store-name </ LABEL='label'>;
```

WEIGHT variable;

Statement	Description	Important Options
PROC MIXED	Invokes the procedure	DATA= specifies input data set, METHOD= specifies estimation method
BY	Performs multiple PROC MIXED analyses in one invocation	None
CLASS	Declares qualitative variables that create indicator variables in design matrices	None
ID	Lists additional variables to be included in predicted values tables	None
MODEL	Specifies dependent variable and fixed effects, setting up X	S requests solution for fixed-effects parameters, DDFM= specifies denominator degrees of freedom method, OUTP= outputs predicted values to a data set, INFLUENCE computes influence diagnostics
RANDOM	Specifies random effects, setting up Z and G	SUBJECT= creates block-diagonality, TYPE= specifies covariance structure, S requests solution for random-effects parameters, G displays estimated G
REPEATED	Sets up R	SUBJECT= creates block-diagonality, TYPE= specifies covariance structure, R displays estimated blocks of R , GROUP= enables between-subject heterogeneity, LOCAL adds a diagonal matrix to R
PARMS	Specifies a grid of initial values for the covariance parameters	HOLD= and NOITER hold the covariance parameters or their ratios constant, PARMSDATA= reads the initial values from a SAS data set
PRIOR	Performs a sampling-based Bayesian analysis for variance component models	NSAMPLE= specifies the sample size, SEED= specifies the starting seed
CONTRAST	Constructs custom hypothesis tests	E displays the L matrix coefficients
ESTIMATE	Constructs custom scalar estimates	CL produces confidence limits
LSMEANS	Computes least squares means for classification fixed effects	DIFF computes differences of the least squares means, ADJUST= performs multiple comparisons adjustments, AT changes covariates, OM changes weighting, CL produces confidence limits, SLICE= tests simple effects
LSMESTIMATE	Provides custom hypothesis tests among the least squares means	ADJUST= determines the method for multiple comparison adjustment of LS-mean differences, JOINT requests a joint <i>F</i> or chi-square test for the rows of the estimate
SLICE	Performs a partitioned analysis of LS-means for an interaction	ADJUST= determines the method for multiple comparison adjustment of LS-mean differences, DIFF requests differences of LS-means
STORE	Saves the context and results of the analysis	LABEL= adds a custom label
WEIGHT	Specifies a variable by which to weight R	None

Figure 10. PROC MIXED Statements

The RANDOM and REPEATED statements require special attention:

The random components are specified using the RANDOM statement, and the correlation components from repeated measurements for participant are generally specified with the REPEATED statement.

The REPEATED effect must contain only classification variables, ensuring that the levels of the repeated effect are different for each observation within a participant.

MMRM only uses REPEATED statement, unless a random participant effect is specified.

In analgesic clinical trials, SAS PROC MIXED is a robust tool for repeated-measures analysis with missing data. It supports complex covariance structures, which are essential for accurate parameter estimation and valid statistical inference.

Example SAS Implementation

Case Background

Variables and endpoints: Change from baseline in WOMAC pain subscale score at week 4. The least squares mean difference of the change from baseline in the WOMAC pain subscale score at week 4 of treatment was estimated for each treatment group, the active control group and the placebo group, as well as the 95% CI and two-sided P value. Least squares mean differences between groups are aggregated at the population level.

Data description: A total of xx participants were included. Pain intensity was assessed using the WOMAC pain subscale score at predetermined time points: baseline, weeks 1, 2, 3, and 4. The corresponding analysis variables are named CHG (change from baseline), BASE (baseline), TRTN (treatment group, 4 presenting control group), AVISITN (visit), USUBJID (participant).

Model description: All planned visit data from baseline to 4 weeks were included in the MMRM model for analysis, with baseline WOMAC pain subscale scores as covariates, and treatment group, visit, and visit-treatment group intersection terms as fixed effects. In the model, an unstructured (UN, Unstructured) covariance matrix is used to fit the correlation between repeated measurements of participants (if the model does not converge, covariance matrix structures such as TOEPH, ARH(1) or CS can be used in order); the constrained maximum likelihood method (Restricted maximum likelihood; REML) is used to estimate the covariance matrix parameters; the Kenward-Roger method is used to calculate.

Data Preparation

Use BDS dataset for analysis. BDS (Basic Data Structure) is one of the core data structures under the CDISC ADaM (Analysis Data Model) standard, specifically designed to store longitudinal repeated measures data. Each record contains key variables such as USUBJID (participant), CHG (change from baseline), BASE (baseline), TRTN (treatment group), AVISITN (visit), etc., and can be directly imported into the MMRM model for analysis.

1. Data set for direct analysis

Used for direct fitting of MMRM model and comparisons between groups, and does not involve missing data filling. This data set is based on the original BDS data, keeping all original observation values unchanged. When estimating the model, all available data (including subjects with incomplete observations) are used for parameter estimation without deleting subjects with missing records.

2. Data set for prediction imputation

Used to fill missing data based on MMRM prediction values as a complementary strategy to sensitivity analysis. The construction of this data set needs to meet the following conditions:

- Use the same BDS dataset as the main analysis.;
- Contains complete values for all fixed effect variables (BASE, TRTN, AVISITN, TRTN*AVISITN), i.e. these variables should not be missing in the data set.

TYPE Statement Selection

This case MMRM model included data from all scheduled visits from baseline to week 4 for analysis. to determine the most suitable covariance structure for this data set, the unstructured covariance matrix (UN) and the heterogeneous Toeplitz covariance matrix (TOEPH) were fitted as examples, and the model was compared and selected through the model fitting goodness index.

```
ODS OUTPUT CONVERGENCSTATUS = CONVERGED;
```

```
PROC MIXED DATA = BDS_DATA METHOD = REML;  
  CLASS trtn(REF="4") avisitn usubjid;  
  MODEL chg = base trtn avisitn trtn* avisitn / HTYPE=3 SOLUTION DDFM=KR;  
  REPEATED avisitn / TYPE = XX SUBJECT = usubjid RCORR;  
RUN;
```

Among them, the TYPE = option is used to specify the covariance structure inside the repeated measures data block. The two structural characteristics compared in this case are as follows:

UN (Unstructured): A complete covariance matrix that does not impose any parameter constraints. The variance of each visit time point and the covariance between pairs are freely estimated. It has the highest flexibility, but there are many parameters to be estimated;

TOEPH (Heterogeneous Toeplitz): Heterogeneous Toeplitz structure, assuming that different time points have different variances, and the covariance changes with the lag order (time interval). The covariances of the same lag order are equal, there are fewer parameters to be estimated, and the structure is relatively simple.

Convergence status: Both covariance structure models converged successfully. The CONVERGED data set exported by ODS OUTPUT shows that both the UN and TOEPH models meet the preset convergence criteria, indicating that the parameter estimation process is stable and reliable, and the fitting results can be used for subsequent comparisons.

	REASON	⊕ STATUS	⊕ PDG	⊕ PDH
1	Convergence criteria met.	0	1	1

Figure 11. CONVERGED Dataset: Converged Fitting Successful

Comparison of goodness of fit: Akaike Information Criterion (AIC) and -2 Res Log Likelihood value (-2 Res Log Likelihood) are used as evaluation indicators for model selection. The smaller the values for

both, the better the model fit. The comparison results of the goodness of fit between the two covariance structures are shown in Table 4.

The covariance parameter estimation results under the two structures of UN and TOEPH are shown in Figure 12 and Figure 13 respectively.

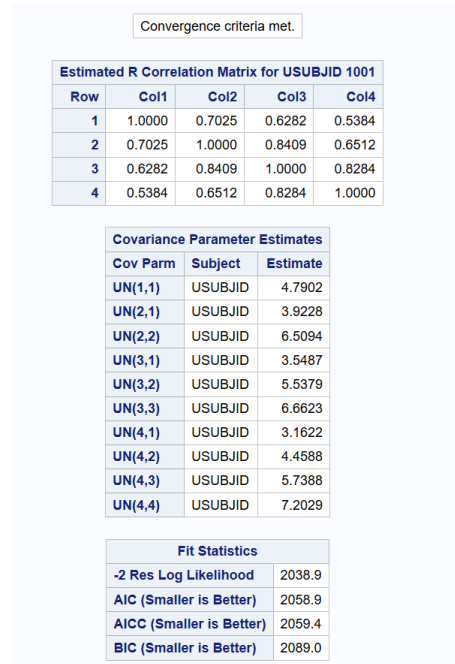


Figure 12. UN Covariance Results



Figure 13. TOEPH Covariance Results

Comparison Item	Heterogeneous Toeplitz (TOEPH)	Unstructured (UN)	Explain
Convergence Status	Convergence criteria met	Convergence criteria met	Both models reached the convergence criterion.
-2 Res Log Likelihood	2054.9	2038.9	Smaller is better, currently UN is better.
AIC	2068.9	2058.9	Smaller is better, currently UN is better.

Table 4. UN vs TOEPH TYPE Statement Selection Comparison

As can be seen from Table 4, the UN structure is better than TOEPH in both -2 Res Log Likelihood and AIC (differences are 16.0 and 10.0 respectively), indicating that the UN structure has a better fitting effect on this data. Although the UN structure requires estimating more parameters, its characterization of data covariance characteristics is more precise, and the improvement in the goodness of fit is enough to compensate for the cost of model complexity caused by the increase in parameters.

Based on the convergence status and goodness-of-fit index, this case finally selected the unstructured covariance matrix (UN) as the covariance structure of the REPEATED statement for subsequent efficacy estimation and inter-group comparison.

MMRM - PROC MIXED

1. Direct use of MMRM for intergroup comparison and efficacy estimation

This case uses repeated measures mixed-effects model (MMRM) to estimate and test the difference in baseline change values between treatment groups. Model fitting and inter-group comparisons are directly calculated based on the original data and do not involve the filling of missing data.

```

ODS OUTPUT DIFFS = DIFFS;
PROC MIXED DATA = BDS_DATA METHOD = REML;
  CLASS TRTN(REF="4") AVISITN USUBJID;
  MODEL CHG = BASE TRTN AVISITN TRTN*AVISITN / HTYPE=3 SOLUTION DDFM=KR;
  REPEATED AVISITN / TYPE=UN SUBJECT=USUBJID RCORR;
  LSMEANS TRTN*AVISITN / DIFF CL ALPHA=0.05;
RUN;

```

METHOD: REML, Restrictive maximum likelihood estimation method, which is the default method for variance component estimation and can provide unbiased and effective variance estimation.

DDFM: KR (Kenward-Roger), recommended for small samples

DIFF: LAST, contrast vs last scheduled visit

HTYPE: Type 3 hypothesis tests for fixed effects, partial (marginal) tests that account for other effects in the model, suitable for unbalanced designs.

SOLUTION: Produces parameter estimates for fixed effects, including regression coefficients and class-level effect estimates with standard errors.

CL: Requests confidence limits for LS-means (default 95% when ALPHA=0.05 is specified).

The core efficacy indicator in this case is the least squares mean difference (LS-mean difference) of baseline change values between treatment groups. Taking the comparison of treatment group 1 (TRTN=1) and reference group 4 (TRTN=4) at week 4 as an example, the estimated value of the difference between the groups and the statistical inference results can be obtained directly from the DIFFS data set exported by ODS OUTPUT.

Where | 查询生成器 | 任务

TRTN=4 and AVISITN=1004

	EFFECT	TRTN	AVISITN	_TRTN	_AVISITN	ESTIMATE	STDERR	DF	TVALUE	PROBT	ALPHA	LOWER	UPPER
1	TRTN*AVISITN	1	1004	4	1001	-2.2432	0.5904	212	-3.80	0.0002	0.05	-3.4069	-1.0795
2	TRTN*AVISITN	1	1004	4	1002	-1.0257	0.6386	201	-1.61	0.1098	0.05	-2.2850	0.2335
3	TRTN*AVISITN	1	1004	4	1003	-0.3809	0.6509	168	-0.59	0.5592	0.05	-1.6659	0.9041
4	TRTN*AVISITN	1	1004	4	1004	0.4217	0.6748	130	0.62	0.5332	0.05	-0.9134	1.7568
5	TRTN*AVISITN	2	1004	4	1001	-2.8770	0.5970	206	-4.82	<.0001	0.05	-4.0541	-1.6999
6	TRTN*AVISITN	2	1004	4	1002	-1.6596	0.6448	196	-2.57	0.0108	0.05	-2.9312	-0.3879
7	TRTN*AVISITN	2	1004	4	1003	-1.0148	0.6570	164	-1.54	0.1244	0.05	-2.3120	0.2824
8	TRTN*AVISITN	2	1004	4	1004	-0.2122	0.6807	127	-0.31	0.7557	0.05	-1.5592	1.1348
9	TRTN*AVISITN	3	1004	4	1001	-2.7398	0.6131	207	-4.47	<.0001	0.05	-3.9484	-1.5312
10	TRTN*AVISITN	3	1004	4	1002	-1.5223	0.6597	198	-2.31	0.0220	0.05	-2.8232	-0.2215
11	TRTN*AVISITN	3	1004	4	1003	-0.8776	0.6715	166	-1.31	0.1931	0.05	-2.2034	0.4483
12	TRTN*AVISITN	3	1004	4	1004	-0.07497	0.6948	128	-0.11	0.9142	0.05	-1.4496	1.2997

Figure 15. DIFFS dataset: Results Between MMRM groups

2. Missing data filling based on predicted value

In the PROC MIXED process, the marginal prediction results are output through the OUTP= option of the MODEL statement, and the data set is named MMRM_PRED. Marginal Prediction is calculated only based on the fixed effect part of the model and does not include random effect components. For the REPEATED structure used in this case, it is equivalent to a prediction estimate that does not introduce random errors at the individual level.

```

MODEL CHG = BASE TRTN AVISITN TRTN*AVISITN / HTYPE=3 SOLUTION DDFM=KR
OUTP=MMRM_PRED;

```

The MMRM_PRED dataset contains the following key variables:

PRED: Fixed effect prediction value, that is, marginal prediction value;

RESID: Marginal residual, that is, the observed value minus the fixed effect predicted value;

STDERRPRED: Standard Error of the Predicted Mean;

LOWER / UPPER: 95% lower and upper confidence limits for the predicted mean (based on fixed effects);

and all variables in the original data set.

	AVISIT	AVISITN	PARAM	PARAMN	AVAL	AVALC	BASE	CHG	TRTN	TRT	USUBJID	PRED	STDERRPRED	DF	ALPHA	LOWER	UPPER	RESID
1	第1周	1001	疼痛量表评分	1	9	9	10	-1	1	低剂量治疗组	1003	-1.766364764	0.3902229073	236.04686313	0.05	-2.535129193	-0.997600336	0.7663647644
2	第2周	1002	疼痛量表评分	1			10		1	低剂量治疗组	1003	-2.615364466	1.5609127694	427.38191619	0.05	-5.683385616	0.4526566837	
3	第3周	1003	疼痛量表评分	1			10		1	低剂量治疗组	1003	-3.223691558	1.9818671047	502.09548177	0.05	-7.117465744	0.6700826273	
4	第4周	1004	疼痛量表评分	1			10		1	低剂量治疗组	1003	-3.980551352	2.2018966271	479.32731869	0.05	-8.307114095	0.3460113917	

Figure 14. MMRM_PRED Dataset: Predicted Value

For observations where the original endpoint data CHG is missing (that is, RESID is missing), the marginal prediction value PRED is used as the filling value; for observations where the original data already exists, the original value CHG is retained unchanged.

```
DATA ANA_PRED;
  SET MMRM_PRED;
  IF RESID=. THEN pred_chg=pred;
  ELSE IF RESID^=. THEN pred_chg= chg;
RUN;
```

After the above filling process, the generated new data set ANA_PRED will be used for subsequent statistical analysis.

CONCLUSION

Single imputation methods assign one fixed value to each missing observation before analysis, with LOCF and BOCF as representative examples. LOCF yields unbiased or conservative estimates only under limited and specific conditions.

Multiple Imputation (PROC MI plus PROC MIANALYZE) is a widely recognized standard under the MAR assumption. It appropriately quantifies imputation-related uncertainty using Rubin's rules and is recommended as a primary analytical strategy for trials with pain-score endpoints such as NRS.

MMRM (PROC MIXED) directly models repeated measures outcomes without single-value imputation, is asymptotically unbiased under MAR, and is highly suitable for continuous repeated endpoints such as WOMAC, weekly mean NRS, and PGA change.

Recommended practice is to use MI or MMRM for the primary analysis, with rule-based methods prespecified for sensitivity analyses. The underlying assumptions for each imputation method should be explicitly stated and fully documented in the clinical study report.

REFERENCES

- Bellamy, N., W. W. Buchanan, C. H. Goldsmith, J. Campbell, and L. W. Stitt. 1988. "Validation study of WOMAC: a health status instrument for measuring clinically important patient relevant outcomes to antirheumatic drug therapy in patients with osteoarthritis of the hip or knee." *Journal of Rheumatology*, 15(12):1833–1840.
- Gewandter, J. S., M. P. McDermott, A. McKeown, S. M. Smith, M. R. Williams, M. Hunsinger, J. Farrar, D. C. Turk, and R. H. Dworkin. 2014. "Reporting of missing data and methods used to accommodate them in recent analgesic clinical trials: ACTTION systematic review and recommendations." *Pain*, 155:1871–1877.
- Little, R. J., R. D'Agostino, M. L. Cohen, K. Dickersin, S. S. Emerson, J. T. Farrar, C. Frangakis, J. W. Hogan, G. Molenberghs, S. A. Murphy, J. D. Neaton, A. Rotnitzky, D. Scharfstein, W. J. Shih, J. P. Siegel, and H. Stern. 2012. "The prevention and treatment of missing data in clinical trials." *New England Journal of Medicine*, 367:1355–1360.
- Liu-Seifert, H., S. Zhang, D. D'Souza, and V. Skljarevski. 2010. "A closer look at the baseline-observation-carried-forward (BOCF)." *Patient Preference and Adherence*, 4:11–16.

- Lu, F., L. Gu, X. Tian, C. Song, and L. Zhou. 2022. "Federated Learning Based on Extremely Sparse Series Clinic Monitoring Data." *ZTE Communications*, 20(3):27–34.
- Rubin, D. B. 1976. "Inference and Missing Data." *Biometrika*, 63:581–592.
- Rubin, D. B. 1987. *Multiple Imputation for Nonresponse in Surveys*. New York: John Wiley & Sons.
- SAS Institute Inc. 2022a. *SAS/STAT® 15.3 User's Guide, The MI Procedure and MIANALYZE Procedure*. Cary, NC: SAS Institute Inc.
- SAS Institute Inc. 2022b. *SAS/STAT® 15.3 User's Guide, The MIXED Procedure*. Cary, NC: SAS Institute Inc.
- Siddiqui, O. 2011. "MMRM versus MI for Analyzing Incomplete Longitudinal Clinical Trial Data." *Journal of Biopharmaceutical Statistics*, 21(4):651–668.
- Tanaka, K., Y. Isaji, K. Suzuki, et al. 2025. "Imbalances in the Content of Sleep and Pain Assessments in Patients with Chronic Pain: A Scoping Review [version 3; peer review: 2 approved]." *F1000Research*, 14:605.
- von Majewski, K., O. Kraus, C. Rhein, M. Lieb, Y. Erim, and N. Rohleder. 2023. "Acute stress responses of autonomous nervous system, HPA axis, and inflammatory system in posttraumatic stress disorder." *Translational Psychiatry*, 13(1):36.
- Wang, S. and X. Luo. 2021. "Exploring the latest Pantheon SN Ia dataset by using three kinds of statistics techniques." *Communications in Theoretical Physics*, 73(4):96–102.
- Wu, T., S. Chen, Y. Tian, and P. Wu. 2020. "A Feature Optimized Deep Learning Model for Clinical Data Mining." *Chinese Journal of Electronics*, 29(3):476–481.
- Wu, W.-T., Y.-J. Li, A.-Z. Feng, L. Li, T. Huang, A.-D. Xu, and J. Lv. 2021. "Data mining in clinical big data: the frequently used databases, steps, and methodological models." *Military Medical Research*, 8(4):552–563.
- Zhao, Z., R. Liu, J. I. Groner, H. Xiang, and P. Zhang. 2023. "Data Imputation for Clinical Trial Emulation: A Case Study on Impact of Intracranial Pressure Monitoring for Traumatic Brain Injury." *AMIA Joint Summits on Translational Science Proceedings*, 2023:612–621.

ACKNOWLEDGMENTS

We thank our manager, Jiehan Zhang, for statistical and programming support and expertise.

CONTACT INFORMATION

Your comments and questions are valued and encouraged. Contact the author at:

Jingbo Zhu

jingbo.zhu.jz76@hengrui.com

Any brand and product names are trademarks of their respective companies.

SAS and all other SAS Institute Inc. product or service names are registered trademarks or trademarks of SAS Institute Inc. in the USA and other countries. ® indicates USA registration.