

A Comprehensive R Shiny Application for Pharmacokinetic Analysis

Keyi Zhang, Dizal (Shanghai) Pharmaceutical Co., Ltd.

Abstract

This paper describes an interactive R Shiny application that streamlines PK statistical analysis across common clinical pharmacology study types. The tool supports plasma concentration profiling, non-compartmental analysis (NCA), dose proportionality assessment, food effect evaluation, drug-drug interaction (DDI) analysis, bioequivalence (BE) testing, special populations exploration, human ADME mass balance evaluation, and BE sample size estimation. For each module, the app generates summary tables of key PK parameters and produces customizable visualizations (e.g., concentration time curves, boxplots, forest plots, and dose proportionality scatter plots). Users can upload their own datasets in standard formats, select analysis parameters interactively, and download results in publication ready formats. The application is built with modular R code to facilitate maintenance and extension. By integrating multiple PK workflows into a single, user friendly interface, this tool reduces repetitive coding effort, enhances reproducibility, and supports informed decision making during drug development.

INTRODUCTION

Pharmacokinetic (PK) analysis is a cornerstone of clinical drug development, providing critical insights into drug absorption, distribution, metabolism, and elimination. A typical PK evaluation encompasses a wide range of assessments, including plasma concentration profiling, non-compartmental analysis (NCA), dose proportionality, food effect, drug-drug interactions, bioequivalence, special population studies, and human ADME mass balance. Additionally, sample size estimation for bioequivalence studies is often required. Each of these analyses generates summary statistics and visualizations that inform decision-making at various stages of development.

Despite the availability of specialized software packages (e.g., Phoenix WinNonlin, SAS, and various R libraries), conducting these analyses in a cohesive and reproducible manner remains challenging. Researchers frequently switch between multiple tools, manually compile results, and invest considerable effort in generating consistent tables and graphs. This fragmented workflow increases the risk of errors, reduces efficiency, and hampers collaboration among team members.

To address these limitations, we developed an integrated R Shiny application specifically tailored for comprehensive PK analysis. The application supports all modules within a single, user-friendly interface. Users can upload standardized datasets, select analysis parameters interactively, and obtain both summary tables and publication-ready visualizations, including concentration-time curves, boxplots, forest plots, and scatter plots. The tool is built with modular R code, ensuring maintainability and extensibility. By unifying diverse PK workflows, our application reduces repetitive programming tasks, enhances reproducibility, and expedites the interpretation of PK data. It can be deployed locally or on a server, making it accessible to statisticians, clinical pharmacologists, and regulatory writers alike. This article describes the design, implementation, and functionality of the application, and demonstrates its utility through example datasets.

APPLICATION ARCHITECTURE AND INPUT DATA

Architecture

The application is implemented in R [1] as a single-file Shiny app (app.R) [2] plus a modular R/ folder. Each analysis module follows the standard Shiny module pattern (mod_<feature>_ui / mod_<feature>_server) and shares one reactive data object produced by the Data module. Statistical engines live in dedicated utility files (pk-stats.R, nca.R, data-utils.R, be-sample-size-utils.R, cv-pooling.R, formatting.R, word-report.R); the linear mixed-effects models are fit with *lme4* [3] so that they can be unit-tested independently and reused across modules. The UI is built with *bslib* + *Bootstrap 5* (Flatly theme), *DT* [4] for tables, *ggplot2/plotly* [5] for plots, *shinycssloaders* and *waiter* for asynchronous feedback, and *officer* and *flextable* [6] for docx reports. Figure 1 depicts the layered architecture.

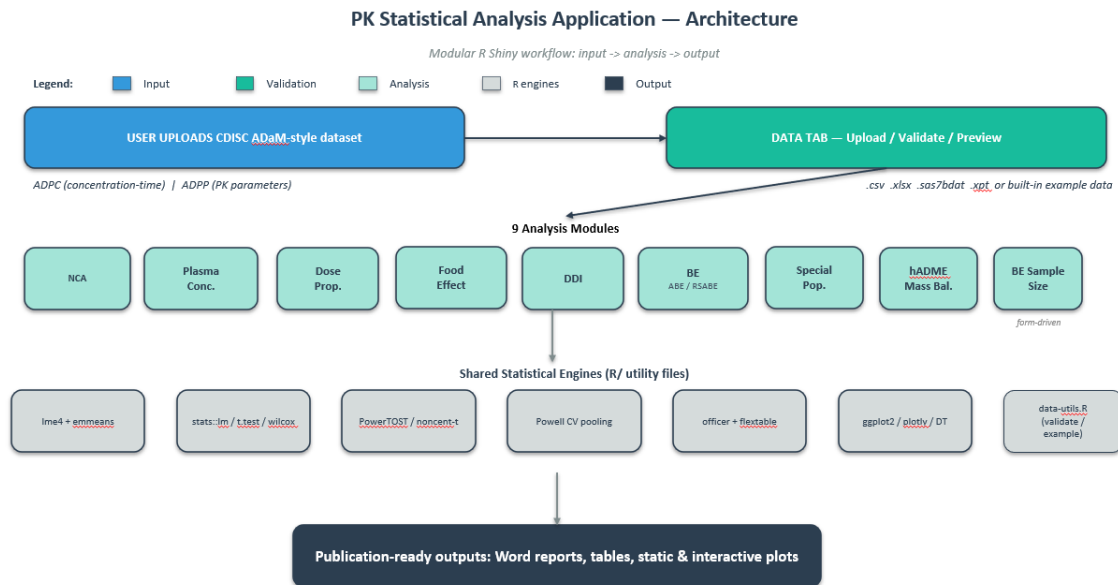


Figure 1: Layered architecture of the PK analysis application.

From Figure 1, users upload a CDISC ADaM-style ADPC or ADPP dataset, which is then routed through the Data tab for upload, validation, and preview. From there, data-driven analysis modules consume the shared reactive data object and call a common set of R statistical engines, and the final results are exported as docx reports, tables, and static or interactive plots.

Datasets

The application consumes two CDISC ADaM-style datasets [7]. The ADPC dataset (concentration-time records: USUBJID, TRTP, PARAMCD, AVAL, NRRLT, ARRLT) feeds the Plasma Concentration, NCA, and human ADME Mass Balance modules. The ADPP dataset (pre-computed PK parameters: USUBJID, TRTP, PARAMCD, AVAL, DOSEA, APERIOD, TRTSEQA) feeds Dose Proportionality, Food Effect, DDI, Bioequivalence, Special Populations, and human ADME Mass Balance. Users upload their own .csv, .xlsx, .sas7bdat or .xpt files, or load a built-in example.

APP interface

Figure 2 shows the Data tab, the single entry point for every analysis. The left panel accepts the data file and lists the required and optional columns for both dataset types. The right panel reports dataset dimensions, subject counts, treatments, parameters and dose levels, and renders an interactive preview table with column filters. All data-driven modules reactively bind to this single `pk_data()` object, ensuring that the same uploaded dataset is analyzed consistently across modules.

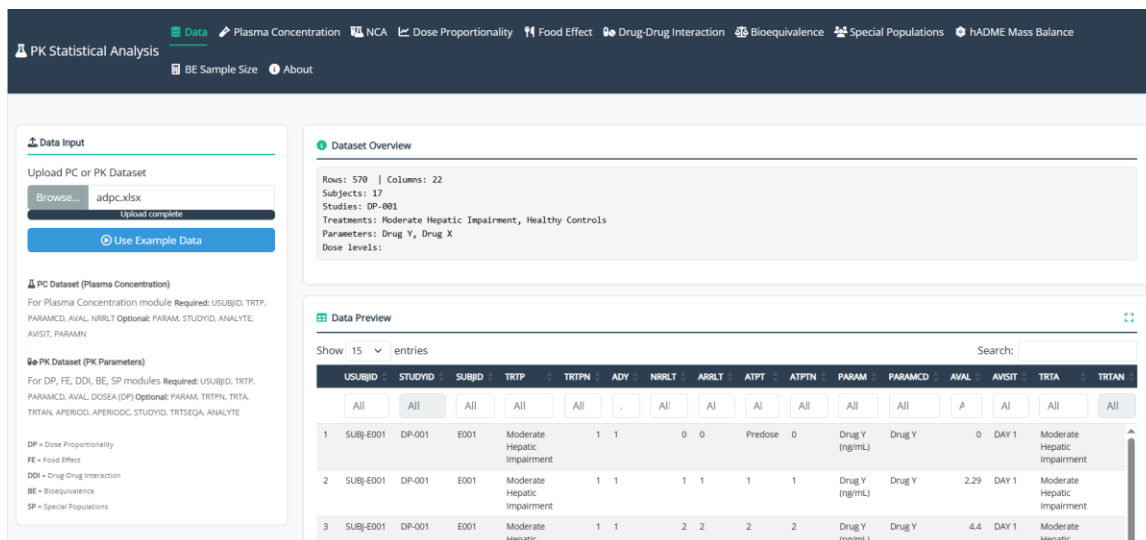


Figure 2. Data tab interface.

R Packages by Module

Module	Primary R Packages
Plasma Concentration	dplyr, ggplot2, DT, plotly, stats
NCA	dplyr, ggplot2, DT, purrr, tidyr, stats (stats::lm), htmltools, utils
Dose Proportionality	dplyr, ggplot2, DT, purrr, stats (stats::lm), emmeans
Food Effect / DDI / BE / Special Population	dplyr, ggplot2, DT, purrr, tidyr, scales (FE), lme4, emmeans, tibble (DDI + BE)
hADME Mass Balance	dplyr, ggplot2, DT, purrr, tidyr, tibble, scales, officer, flextable, stats
BE Sample Size	DT, ggplot2, purrr, tibble, shiny, stats (stats::qt + stats::pt); PowerTOST
Common (UI, tables, reports, I/O)	shiny, bslib, DT, shinycssloaders, waiter, officer, flextable

PHARMACOKINETIC ANALYSIS MODULES

Pharmacokinetic (PK) analysis sections correspond to the individual analysis modules, including *Plasma Concentration*, *NCA*, *Dose Proportionality*, *Food Effect*, *Drug-Drug Interaction*, *Bioequivalence*, *Special Populations*, *human ADME Mass Balance* and *BE Sample Size*. Every analysis module incorporates a setting component, which enables users to specify input variables and parameters, and provides visual data exploration through a summary statistics table and associated plots.

Plasma Concentration Module

This module summarizes concentration-time data with log-scale geometric statistics and renders static + interactive profiles.

Data layer. The user selects the STUDYID, ANALYTE, PARAMCD, TRTP, and AVISIT; the input is filtered accordingly. The data layer requires the NRRLT (nominal relative time) column; if it is missing or empty, the module raises an error.

Algorithm layer. The core statistical methodology consists of computing geometric means and geometric standard deviations after log-transformation, summarized by treatment group and time point. The reported analysis metrics include N (sample size), arithmetic mean, standard deviation

(SD), median, minimum, maximum, geometric mean, and geometric coefficient of variation (CV%). The summary drives a static geometric-mean $\pm$ SD profile (linear and semi-log) via *ggplot2*, and an interactive per-subject spaghetti plot via *plotly::ggplotly* for hover tooltips.

Output layer. Results are rendered as an interactive summary table, the static linear and semi-log geometric mean curves, the interactive per-subject linear and semi-log spaghetti plots, and a Word report generated via *officer* and *flextable*.

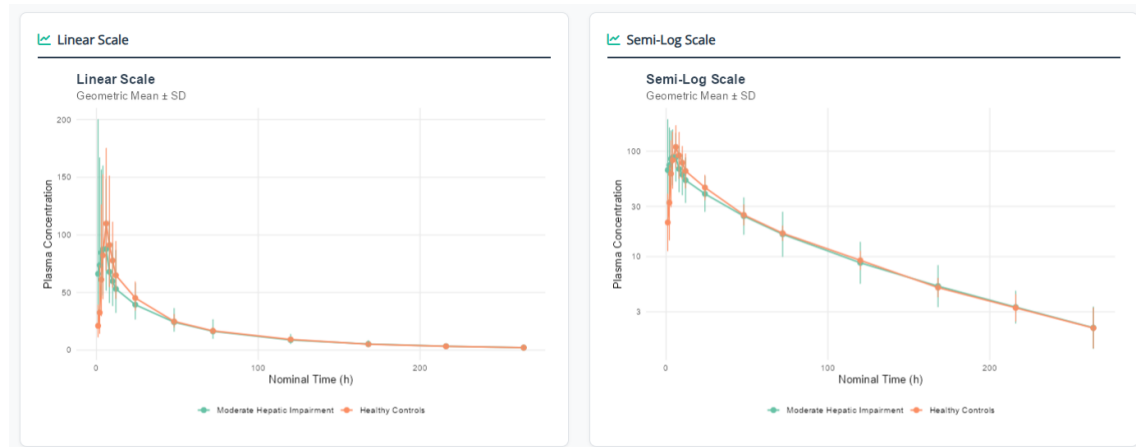


Figure 3: A representative output of Plasma Concentration module. Left: linear-scale geometric mean $\pm$ SD concentration-time profile by treatment group. Right: the same data on a semi-logarithmic Y-axis, the conventional view for visualizing the terminal elimination phase.

NCA Module

Non-compartmental analysis (NCA) is a model-independent approach for estimating pharmacokinetic parameters based on statistical moment theory. It directly exploits the area under the plasma concentration–time curve and its characteristics to describe the absorption, distribution, metabolism, and excretion of a drug.

Data layer. The input is sequentially filtered by STUDYID, ANALYTE, PARAMCD, selected time variable (NRRLT or ARRLT), and TRTP.

Algorithm layer. The NCA computation is implemented as eight pure functions in *nca.R*. AUC is computed by trapezoidal integration, either linear or mixed log-linear (user-selectable), and the terminal elimination rate constant λ_z is estimated by ordinary least-squares regression of log concentration on time, using the terminal four points as the default. From λ_z , the half-life $T_{1/2}$, apparent clearance CL/F , and apparent volume of distribution V_d/F are derived. Beyond the per-subject filters, the UI exposes a Time Variable selector for Nominal versus Actual sampling time, a Dose Source panel where the user enters a fixed dose that is used to compute CL/F , or selects DOSEA from the data, and an optional Body Weight panel that supplies the body weight for normalization, with three options: fixed, read from the data, or unused. The dose and weight inputs appear conditionally via *shiny::conditionalPanel*.

Output layer. NCA is applied to each subject in parallel via a nest-map-unnest pattern, producing per-subject parameter tables, treatment-level descriptive summaries (including a log-scale geometric CV%), and C_{max} / AUC box plots, with docx report and CSV export.

Summary Statistics by Treatment

Parameter	Treatment	N	Mean	SD	CV%	Median	Min	Max	Geo Mean	Geo CV%
AUC(0-inf)	Moderate Hepatic Impairment	8	137.648	102.766	74.658	126.662	29.667	353.743	107.458	91.516
AUC(0-inf)	Healthy Controls	8	193.489	78.476	40.559	185.493	93.941	338.053	179.677	43.663
AUC(0-t)	Moderate Hepatic Impairment	8	137.648	102.766	74.658	126.662	29.667	353.743	107.458	91.516
AUC(0-t)	Healthy Controls	8	193.489	78.476	40.559	185.493	93.941	338.053	179.677	43.663
CL/F	Moderate Hepatic Impairment	8	0.608	0.502	82.563	0.397	0.141	1.685	0.465	91.516
CL/F	Healthy Controls	8	0.301	0.128	42.732	0.274	0.148	0.532	0.278	43.663
Cmax	Moderate Hepatic Impairment	8	11.234	4.184	37.243	10.2	4.9	19.3	10.537	41.021
Cmax	Healthy Controls	8	12.299	5.726	46.561	11.45	5.12	23.8	11.209	49.116
Lambda_z	Moderate Hepatic Impairment	8	0.033	0.023	70.428	0.027	0.009	0.075	0.026	88.285
Lambda_z	Healthy Controls	8	0.018	0.007	38.029	0.016	0.01	0.028	0.017	40.95
T1/2	Moderate Hepatic Impairment	8	33.703	24.507	72.715	26.095	9.194	76.163	26.441	88.285
T1/2	Healthy Controls	8	43.026	16.853	39.17	43.088	25.037	72.858	40.222	40.95
Tmax	Moderate Hepatic Impairment	8	3	1.69	56.344	3	1	6	2.539	73.175
Tmax	Healthy Controls	8	4.25	1.282	30.159	4	2	6	4.059	34.968
Vd/F	Moderate Hepatic Impairment	8	21.723	17.143	78.916	15.692	7.336	61.416	17.749	71.579
Vd/F	Healthy Controls	8	16.795	5.478	32.618	16.079	11.196	29.075	16.148	29.509

Figure 4: Representative output of the NCA module

Dose Proportionality

Dose proportionality analysis assesses whether systemic exposure (AUC and Cmax) increases proportionally with dose.

Data layer. The input is filtered by STUDYID, ANALYTE, and the selected PK parameters, the analysis visit (AVISIT) and dose levels (DOSEA).

Algorithm layer. Two complementary analyses are performed at the user-specified confidence level (default 90%, configurable via a numeric input).

- (i) A power model $\log(PK) = \alpha + \beta \log(Dose)$ is fit by linear regression; dose proportionality is concluded when the 90% confidence interval of the slope β contains 1.

```
fit_pm <- stats::lm(log(AVAL) ~ log(DOSEA), data = d)
ci_pm <- stats::confint(fit_pm, "log(DOSEA)", level = 0.90)
```

- (ii) A one-way ANOVA on log(PK) by dose level is fit, and pairwise geometric mean ratios are derived via *emmeans*; a dose-adjusted Geometric Mean Ratio (GMR) is then obtained for each pair by dividing the corresponding GMR by the test-to-reference dose ratio, and dose proportionality is concluded when its 90% confidence interval (CI) contains 100%. A descriptive dose-normalized boxplot (AVAL / DOSEA) is also produced for visual inspection.

```
# DOSE_F is factor of dose
fit_emm <- stats::lm(log(AVAL) ~ DOSE_F, data = d)
emm <- emmeans::emmeans(fit_emm, ~ DOSE_F)
contr <- emmeans::contrast(emm, method = "revpairwise")
ci_emm <- summary(contr, infer = c(TRUE, TRUE), level = 0.90)
```

Output layer. Results are rendered as a power-model summary table, a pairwise GMR table, a dose-adjusted GMR table with pass/fail flags, a log - log scatter plot with the fitted regression line, a forest plot of dose-adjusted GMRs, a dose-normalized boxplot, and a docx report generated via *officer* and *flextable*.

✓ Dose-Adjusted GMR (Dose Proportionality Assessment)

Parameter	Comparison	Dose-Adj GMR (%)	Lower (%)	Upper (%)	DP (CI includes 100%)
AUC Infinity (ng*h/mL)	100 mg / 50 mg	101.91	76.76	135.29	✓ Yes
AUC Infinity (ng*h/mL)	200 mg / 100 mg	107.16	80.72	142.26	✓ Yes
AUC Infinity (ng*h/mL)	200 mg / 50 mg	109.2	82.26	144.98	✓ Yes
AUC Last (ng*h/mL)	100 mg / 50 mg	96.65	71.67	130.34	✓ Yes
AUC Last (ng*h/mL)	200 mg / 100 mg	95.34	70.7	128.58	✓ Yes
AUC Last (ng*h/mL)	200 mg / 50 mg	92.15	68.33	124.27	✓ Yes
Cmax (ng/mL)	100 mg / 50 mg	145.97	109.79	194.07	✗ No
Cmax (ng/mL)	200 mg / 100 mg	81.9	61.6	108.89	✓ Yes
Cmax (ng/mL)	200 mg / 50 mg	119.56	89.92	158.95	✓ Yes

Figure 5: Dose-Adjusted GMR

Figure 5 shows the results of dose proportionality assessment. For Cmax and comparison of 100 mg / 50 mg, the interval does not include 100% and lies entirely above 100%, indicating a supraproportional increase in Cmax from 50mg and 100mg, therefore Dose Proportionality (DP) is No in this situation. Compared to other results in Cmax, it suggests that Cmax increased more than dose proportionally from 50 mg to 100 mg, whereas the increase from 100 mg to 200 mg was consistent with dose proportionality. The overall comparison between 200 mg and 50 mg remains compatible with proportionality.

Food Effect

The typical food effect study is a randomized, single-dose, two-period, two-sequence crossover trial in healthy volunteers. In one period, the drug is administered after an overnight fast (fasted state); in the other, it is given with a standardized high-fat, high-calorie meal (fed state). A washout interval separates the two periods to eliminate carryover effects. It is designed to evaluate the effect of food intake on drug exposure.

Data layer. The input is filtered by STUDYID, ANALYTE, the selected PK parameters, and the two treatment groups (Fed as Test, Fasted as Reference).

Algorithm layer. The GMR of fed versus fasted AUC and Cmax is estimated to have its 90% confidence interval.

- (i) For 2×2 crossover data with both sequence and period available, a linear mixed-effects model of the form $\log(PK) \sim TRTP + APERIOD + TRTSEQA + (1 | TRTSEQA:USUBJID)$ is fit with *lme4::lmer*; least-squares means and pairwise contrasts are then extracted via *emmeans* [8] with Kenward–Roger [9] degrees of freedom. When only one sequence (or a confounded period) is available, the model is reduced to $\log(PK) \sim TRTP + (1 | USUBJID)$.

```
mod <- lme4::lmer(
  log(AVAL) ~ TRTP + APERIOD + TRTSEQA + (1 | TRTSEQA:USUBJID),
  data = d
```

```
)
emm <- emmeans::emmeans(mod, specs = ~ TRTP, lmer.df = "kenward-roger")
```

Where TRT is treatment, APERIOD is analysis period, TRTSEQA is treatment sequence, and (1 | TRTSEQA:USUBJID) is random effect which captures within-subject correlation across the two periods, modelling each subject's log (AVAL) as a random effect nested within their sequence.

- (ii) Parallel designs are analyzed by a linear model of $\log(PK) \sim TRTP$.

```
mod_p <- stats::lm(log(AVAL) ~ TRTP, data = d)
emm_p <- emmeans::emmeans(mod_p, specs = ~ TRTP)
```

- (iii) Tmax and Tlag are compared by the paired Wilcoxon signed-rank test, with the median difference and its 90% CI reported via the Hodges–Lehmann estimator. A 90% CI of the GMR fully within [80%, 125%] for both AUC and Cmax indicates no clinically meaningful food effect.

```
test_vals <- d$AVAL[d$TRTP == "Fed"]
ref_vals <- d$AVAL[d$TRTP == "Fasted"]
wt <- stats::wilcox.test(test_vals, ref_vals, paired = TRUE,
                        conf.int = TRUE, conf.level = 0.90,
                        exact = FALSE, correct = TRUE)
```

Output layer. Results are rendered as treatment-level summary tables, a GMR table, a Tmax analysis table, a forest plot overlaid with the 80%–125% no-effect acceptance zone, a box plot by treatment, a clinically worded conclusion, and a docx report generated via *officer* and *flextable*.

APP interface

Geometric Mean Ratio (Fed / Fasted)								
Parameter	n (Drug 100 mg (Fed))	GLS Mean (Drug 100 mg (Fed))	n (Drug 100 mg (Fasted))	GLS Mean (Drug 100 mg (Fasted))	Ratio of GLS Means	90% CI	Intra-subject CV%	Within 80-125%
AUC Infinity (ng*h/mL)	24	618.4839	24	489.7899	126.28	(113.19, 140.88)	22.35	✗ No
AUC Last (ng*h/mL)	24	573.711	24	443.9564	129.23	(112.76, 148.09)	28.02	✗ No
Cmax (ng/mL)	24	97.0247	24	78.4717	123.64	(110.91, 137.84)	22.19	✗ No
Half-life (h)	24	7.9206	24	8.0938	97.86	(90.23, 106.13)	16.47	✓ Yes

Figure 6: GMR of Fed/Fasted with Intra-subject CV%

From Figure 6, food significantly increased both AUC and Cmax, with the upper 90% CI exceeding the 125% boundary. Therefore, a clinically relevant food effect was identified across these parameters. However, the 90% CI of Half-life falls in 80%-125% interval entirely, it shows the food effect is not clinically meaningful in Half-life.

Tmax / Tlag Analysis (Wilcoxon Signed-Rank Test)						
Parameter	Median (Fed)	Median (Fasted)	Median Diff (H-L)	Lower CI	Upper CI	P-value
Tmax (h)	2.7422	2.5431	0.2726	-0.1982	0.9211	0.4663
Paired Wilcoxon Signed-Rank Test; Hodges-Lehmann estimator with 90% CI						

Figure 7: Tmax Analysis

Figure 7 shows the 90% CI for this difference ranged from -0.20 to 0.92 hours and included zero, and the corresponding p-value was 0.47. These results indicate no statistically significant shift in Tmax between the fed and fasted states, suggesting that food does not appreciably delay the rate of absorption as measured by Tmax.

Drug-Drug Interaction (DDI) Module

The DDI module evaluates how a concomitant medication, inhibitor or inducer affects the pharmacokinetics of a substrate drug, using a standard PK dataset.

Data layer. The input is filtered by STUDYID, ANALYTE, the selected PK parameters, and the two treatment groups (Test = drug co-administered with the perpetrator; Reference = drug alone).

Algorithm layer. The statistical pipeline mirrors the Food Effect module: a log-transformed linear mixed-effects model is fit for crossover designs, and a linear model on $\log(PK)$ is used for parallel designs, with Kenward–Roger degrees of freedom via *emmeans*. The default parameter set is AUC Infinity, AUC Last, and Cmax, and the confidence level is user-configurable (default 90%). Each interaction is then classified against the FDA DDI guidance [10] into seven mutually exclusive tiers according to the AUC ratio point estimate: ≥ 5 Strong Inhibitor, 2–5 Moderate Inhibitor, 1.25–2 Weak Inhibitor, 0.8–1.25 No DDI Effect, 0.5–0.8 Weak Inducer, 0.2–0.5 Moderate Inducer, and < 0.2 Strong Inducer. Tmax and Tlag, when included, are compared by the paired Wilcoxon signed-rank test with the Hodges–Lehmann estimator.

Output layer. Results are rendered as treatment-level summary tables, a GMR table, the seven-band DDI classification, a forest plot color-coded by category, an individual-subject scatter plot, and a docx report generated via *officer* and *flextable*.

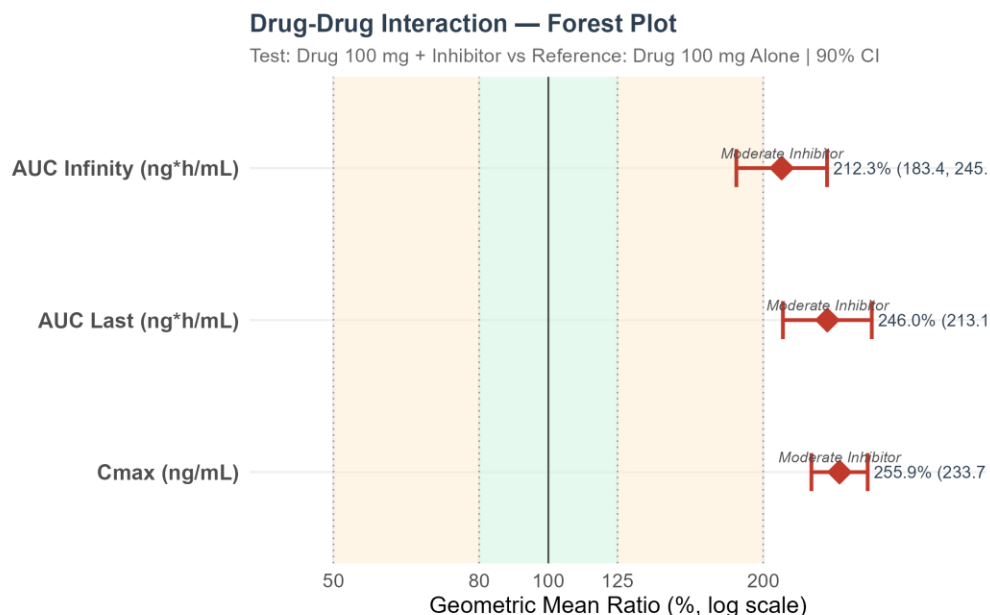


Figure 8: Forest plot of DDI

Figure 8 shows the ratio of DDI classification, according to regulatory categories, AUC Infinity, AUC Last and Cmax fall with moderate inhibition (2.0–5.0), indicating a moderate inhibitory effect on the substrate's metabolism or transport, with the seven-band classification color-coded by regulatory category.

Bioequivalence (BE) Module

The Bioequivalence (BE) module implements the FDA May 2026 guidance [11] for average, scaled-average, and narrow-therapeutic-index bioequivalence, supporting 2×2 crossover, parallel, full replicate (2×2×4), and partial replicate (2×2×3) designs.

Data layer. The input is sequentially filtered by STUDYID, ANALYTE, and the selected PK parameters and formulations (Test vs. Reference).

Algorithm layer. The implementation is dispatched by the user-selected BE method.

- (i) For standard Average BE, the same *lme4* + *emmeans* pipeline as the Food Effect module is applied: a linear mixed-effects model $\log(PK) \sim TRTP + APERIOD + TRTSEQA + (1 | TRTSEQA:USUBJID)$ is fit on crossover data with Kenward–Roger degrees of freedom, and the geometric mean ratio (GMR) and its 90% confidence interval are derived from pairwise contrasts. Replicate designs (2×2×4 and 2×2×3) use the same full formula as the 2×2 crossover. The default parameter set is AUCINF, AUCLST, and CMAX, and the BE limits are user-configurable (default 80–125%).

```
# for cross-over (2x2), replicate (2x2x4, 2x2x3) designs
mod <- lme4::lmer(
  log(AVAL) ~ TRTP + APERIOD + TRTSEQA + (1 | TRTSEQA:USUBJID),
  data = d
)
emm <- emmeans::emmeans(mod, specs = ~ TRTP, lmer.df = "kenward-roger")
```

When only one sequence is available, the model reduces to $\log(PK) \sim TRTP + (1 | USUBJID)$ if period is confounded, or $\log(PK) \sim TRTP + APERIOD + (1 | USUBJID)$ if period is not confounded. When only one period is available, the model is $\log(PK) \sim TRTP + (1 | USUBJID)$.

```
# 1 sequence + period confounded
mod <- lme4::lmer(
  log(AVAL) ~ TRTP + (1 | USUBJID),
  data = d
)
emm <- emmeans::emmeans(mod, specs = ~ TRTP, lmer.df = "kenward-roger")

# 1 sequence + period NOT confounded
mod <- lme4::lmer(
  log(AVAL) ~ TRTP + APERIOD + (1 | USUBJID),
  data = d
)
emm <- emmeans::emmeans(mod, ~ TRTP, lmer.df = "kenward-roger")

# 1 period, NO sequence
mod <- lme4::lmer(
  log(AVAL) ~ TRTP + (1 | USUBJID),
  data = d
)
emm <- emmeans::emmeans(mod, ~ TRTP, lmer.df = "kenward-roger")
```

Parallel designs use a linear model on $\log(PK)$ followed by *emmeans* contrasts.

```
# parallel design
mod <- stats::lm(log(AVAL) ~ TRTP, data = d)
emm <- emmeans::emmeans(mod, ~ TRTP)
```

- (ii) For highly-variable drugs (HVD) reference-scaled ABE, subject-level contrasts $lij = \text{mean}(\log PK) T - \text{mean}(\log PK) R$ and within-reference replicates Dij are constructed, the within-reference intra-subject CV (sWR) is estimated, and a 95% upper confidence bound (UCB) on the linearized criterion is computed; mixed scaling is applied such that standard ABE is used when $sWR < 0.294$ and RSABE otherwise. BE is declared when the $95\% \text{ UCB} \leq 0$ and the point estimate of GMR lies within [80%, 125%].
- (iii) For narrow therapeutic index (NTI) drugs, all four sub-criteria must pass: the $95\% \text{ UCB} \leq 0$, the GMR point estimate lies within [80%, 125%], the 90% CI lies within [80%, 125%], and the 90% CI upper limit of the σ_{WT}/σ_{WR} ratio does not exceed 2.5.

Output layer. Results are rendered as treatment-level summary tables, the GMR and LS-means tables, the BE assessment (overall pass/fail with sub-criteria), a customized forest plot with the 80% - 125% acceptance zone, an individual-subject scatter plot, a clinically worded conclusion, and a docx report generated via *officer* and *flextable*.

APP interface

Bioequivalence Assessment							
Parameter	n (Drug 100 mg (Test))	GLS Mean (Drug 100 mg (Test))	n (Drug 100 mg (Reference))	GLS Mean (Drug 100 mg (Reference))	Point Estimate (%)	90% CI	Intra-subject CV% BE Met
AUC Infinity (ng*h/mL)	36	492.1002	36	492.6695	99.88	(90.17, 110.64)	26.09 Yes
AUC Last (ng*h/mL)	36	460.9079	36	446.9535	103.12	(95.66, 111.16)	19 Yes
Cmax (ng/mL)	36	81.8485	36	80.5041	101.67	(94.07, 109.88)	19.68 Yes

Figure 9: Bioequivalence assessment

Figure 9 shows all three parameters showed point estimates close to 100% (99.88%–103.12%) with narrow 90% confidence intervals fully contained within the 80.00%–125.00% boundary, demonstrating that there are no meaningful differences in the rate or extent of absorption between the test and reference formulations. Therefore, the test formulation is considered bioequivalent to the reference formulation.



Figure 10: Forest plot and individual subject plot of bioequivalence

Figure 10 show point estimates of all three parameters near 100% and 90% CIs fully within 80% - 125%, the forest plot with all confidence intervals comfortably inside the green acceptance region, supporting bioequivalence. The paired subject plot shows points symmetrically clustered around the line of identity, indicating strong within-subject agreement between the test and reference formulations.

Special Population

The special populations module evaluates the effect of hepatic or renal impairment on drug pharmacokinetics using a standard PK dataset and provides corresponding dose adjustment recommendations.

Data layer. The input is sequentially filtered by STUDYID, ANALYTE, the selected PK parameters, and the treatments; multiple impaired groups (e.g., mild, moderate, severe hepatic or renal impairment) are compared against a single normal/healthy reference.

Algorithm layer. The analysis reuses the same *lm* + *emmeans* parallel-design pipeline as the Food Effect module: for each PK parameter, a linear model $\log(PK) \sim TRTP$ is fit on the log-transformed values with no random effects, since the parallel-group design has no within-subject replication. LS-means are extracted via *emmeans*, and pairwise contrasts yield the GMR of each impaired group versus the reference, together with its 90% CI. Tmax, when included, is compared by the unpaired Wilcoxon rank-sum test with the Hodges–Lehmann estimator.

```
fit <- stats::lm(log(AVAL) ~ TRTP, data = d)
emm <- emmeans::emmeans(fit, specs = ~ TRTP)
```

Output layer. Results are rendered as treatment-level summary statistics, a GMR table with one row per impaired group per PK parameter, a forest plot by treatment group, a box plot, a clinical dose-adjustment recommendation following the FDA convention (GMR < 1.25: no adjustment; 1.25 - 2: monitor closely; 2 - 5: dose reduction; ≥ 5: not recommended), and a docx report generated via *officer* and *flextable*.

APP interface

➡ Geometric Mean Ratio vs. Reference

Parameter	Group	n	GLS Mean	Ratio of GLS Means	90% CI of the Ratio	Fold Change	Action
AUC Infinity (ng*h/mL)	Normal	8	634.4876				
AUC Infinity (ng*h/mL)	Mild HI	8	709.3036	111.79	(81.97, 152.46)	1.12	✅ No adjustment
AUC Infinity (ng*h/mL)	Moderate HI	8	984.6969	155.2	(121.04, 198.98)	1.55	⚠️ Monitor closely
AUC Infinity (ng*h/mL)	Severe HI	8	1202.8399	189.58	(149.78, 239.94)	1.9	⚠️ Monitor closely
AUC Last (ng*h/mL)	Normal	8	434.0596				
AUC Last (ng*h/mL)	Mild HI	8	585.0651	134.79	(100.46, 180.85)	1.35	⚠️ Monitor closely
AUC Last (ng*h/mL)	Moderate HI	8	804.748	185.4	(145.02, 237.02)	1.85	⚠️ Monitor closely
AUC Last (ng*h/mL)	Severe HI	8	1318.3822	303.73	(230.15, 400.84)	3.04	⚠️ Dose reduction
Cmax (ng/mL)	Normal	8	93.5236				
Cmax (ng/mL)	Mild HI	8	102.3523	109.44	(72.19, 165.91)	1.09	✅ No adjustment
Cmax (ng/mL)	Moderate HI	8	120.2926	128.62	(87.94, 188.13)	1.29	⚠️ Monitor closely
Cmax (ng/mL)	Severe HI	8	196.6864	210.31	(151.99, 291.01)	2.1	⚠️ Dose reduction

Figure 11: Geometric Mean Ratio of impaired group vs. normal group

Figure 11 shows the statistics metrics, such as population group, number of subjects of each group, Geometric Least Square (GLS) Mean, Ratio of GLS Means and its 90% CI and fold change in the table. According to GMR, recommendation of dose modification is presented in the *Action* column.



Figure 12: Forest plot ordered by impairment severity and boxplot by group

From forest plot of Figure 12, all point estimates lie above 100% and shift rightward with increasing impairment, moderate and severe intervals fall entirely above 125%, with AUC Last in severe impairment showing the widest CI. Through boxplot, medians and spread increase with impairment severity, most markedly in the severe group and especially for AUC Last.

hADME Mass Balance Module

The human ADME (hADME) Mass Balance module summarizes drug concentrations across biological matrices (plasma, urine, feces) using both concentration (ADPC) and PK parameter (ADPP) datasets in standard format. It supports the analysis of excretion pathways and cumulative recovery from radiolabeled or non-radiolabeled studies.

Data layer. The input is sequentially filtered by STUDYID, ANALYTE, and MATRIX (e.g., plasma, urine, feces). The two views then apply different filter logic: the **time-pivoted view** requires a valid time variable (not missing analysis time point, *ATPTN*), while the **parameter-pivoted view** is applied to all rows.

Algorithm layer. Two views are produced on the filtered dataset, with 10 user-selectable descriptive statistics (N, n (BLQ), Mean, SD, CV%, Median, Min, Max, Geometric Mean, Geometric CV%) appended as additional rows below the per-subject rows.

- (i) **Time-pivoted view.** For a user-selected parameter (*PARAMCD*), the data are filtered to rows with a valid time point (not missing *ATPTN*) and pivoted into a wide table with subjects as rows and time points (with *ATPT* text labels when available, e.g., "0.5 h" or "0 to 24 h") as columns. Each cell carries the formatted value, with missing values rendered as "ND" and zeros rendered as "BLQ". A geometric mean profile with $\pm$ SD band is rendered on linear and semi-log scales, using $\exp(\text{mean}(\log(AVAL)))$ for the central value and $\exp(\text{mean}(\log(AVAL)) \pm \text{sd}(\log(AVAL)))$ for the band, computed on positive values only.
- (ii) **Parameter-pivoted view.** For one or more user-selected PK parameters (*PARAMCD*, multi-select), the data are pivoted into a wide table with subjects as rows and parameters as columns. For subjects with multiple records of the same parameter, the per-subject mean is computed first. Missing values are rendered as "ND". Column

labels use the human-readable *PARAM* text where available. Columns are ordered by *PARAMN* if available, otherwise alphabetically.

Output layer. Results are rendered as both wide tables, the linear and semi-log geometric mean concentration-time plots, and a docx report generated via *officer* and *flextable*.



Figure 13: Mean ^{14}C -labelled total radioactivity concentration in plasma over time.

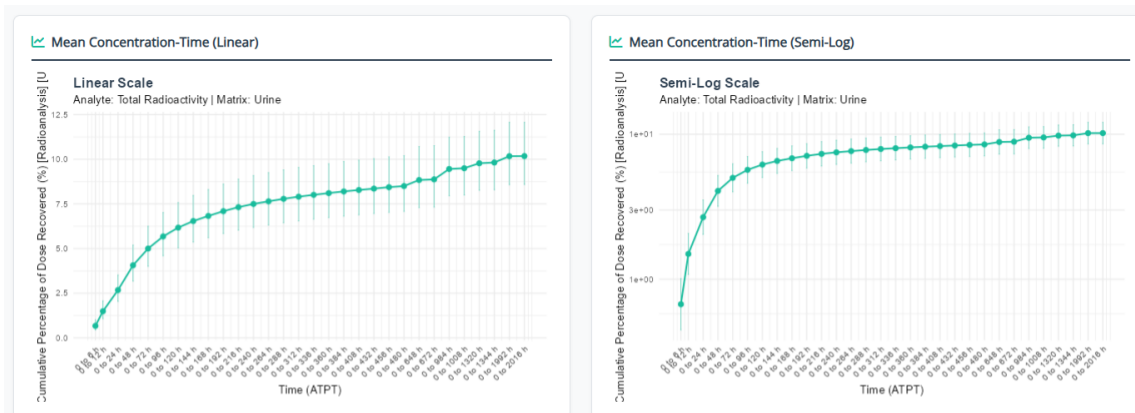


Figure 14: Mean cumulative percentage of ^{14}C -labelled total radioactivity dose recovered in urine

Figure 13 and 14 illustrate the time-pivoted view applied to two complementary record types from this module. Figure 13 is generated from ADPC and is a typical plasma concentration-time profile. Figure 14, in contrast, is generated from ADPP and is a cumulative-excretion curve. Both plots are produced by the same algorithm.

Bioequivalence (BE) Sample Size Module

BE sample size calculation is performed to ensure the study has enough statistical power to demonstrate that a generic drug and a reference (brand-name) drug are equivalent in terms of the rate and extent of drug absorption.

Data layer. The user specifies the study design (2×2 crossover, parallel, full replicate 2×2×4, or partial replicate 2×2×3), the intra-subject CV, the expected GMR, the one-sided significance level α (default 0.05), and the BE limits (default 80%–125%); a minimum-N floor of 12 (PK BE) or 24 (replicate designs) is enforced per the FDA May 2026 guidance [11].

Algorithm layer. Sample size (or achieved power) is computed by iterative search via the Two One-Sided Tests (TOST) procedure. Two interchangeable algorithms are provided: a built-in noncentral t-distribution method (Julious) [12] and the *PowerTOST* [13] package as a cross-check; if *PowerTOST* is unavailable the application automatically falls back to the noncentral-t engine. An optional CV-pooling utility combines historical intra-subject CVs from multiple prior studies via the Powell et al. (2015) method [14], with *PowerTOST::CVpooled()* available as a regulator-grade alternative.

Output layer. Results are rendered as a primary result card (total N and achieved power), a sensitivity table over a grid of CVs and GMRs, two power curves (power vs. N and power vs. CV), and a docx report generated via *officer* and *flextable*.

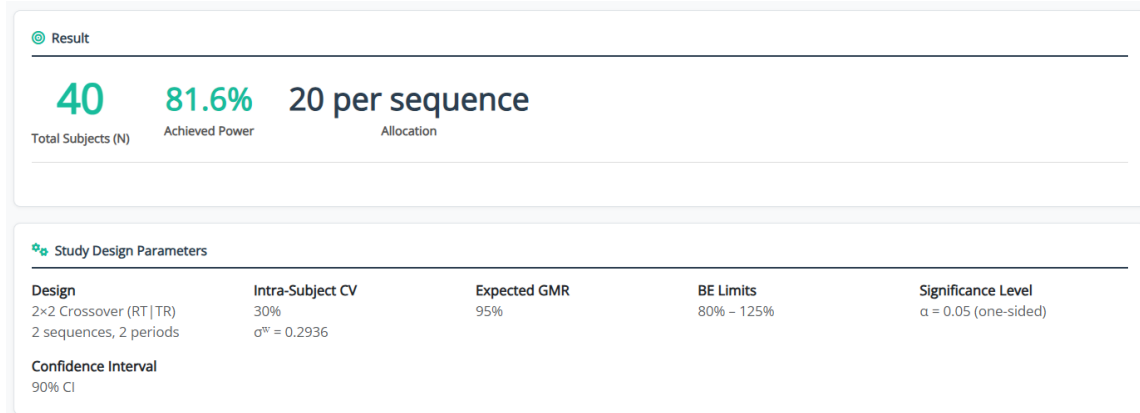


Figure 15: Result of BE sample size

Figure 15 is a representative output of the BE Sample Size module after specifying the calculation mode (Sample Size from Power), study design (2x2 crossover), CV (30%), expected GMR (95%), BE limits (80%–125%), one-sided α (0.05), and *PowerTOST* as the algorithm. The result card (top) reports the required total sample size (**N = 40**, with 20 per sequence) and the achieved power (**81.6%**); the assumptions panel (bottom) echoes the user-specified study design parameters along with the derived log-scale standard deviation ($\sigma_W = 0.2936$) and the two-sided 90% confidence interval corresponding to the one-sided $\alpha = 0.05$.

CONCLUSION

This paper describes an integrated R Shiny application that provides the complete workflow, methodology, analyses implemented in R, and interactive visual reports for a comprehensive range of pharmacokinetic evaluations in early-phase clinical pharmacology. The modules cover plasma concentration profiling, non-compartmental analysis (NCA), dose proportionality, food effect, drug-drug interaction (DDI), bioequivalence (BE), special populations, human ADME mass balance, and BE sample size estimation. The analytical approaches and decision criteria refer to current regulatory guidelines, scientific literature, and industry best practices.

REFERENCES

- [1] R Core Team. R: A Language and Environment for Statistical Computing. R Foundation for Statistical Computing, Vienna, Austria. 2024. <https://www.R-project.org/>
- [2] Chang W, Cheng J, Allaire JJ, et al. shiny: Web Application Framework for R. R package version 1.8.0. 2023. <https://CRAN.R-project.org/package=shiny>

- [3] Bates D, Mächler M, Bolker B, Walker S. Fitting Linear Mixed-Effects Models Using lme4. J Stat Softw. 2015;67(1):1-48.
- [4] Xie Y, Cheng J, Tan X. DT: A Wrapper of the JavaScript Library "DataTables". R package version 0.32. 2024. <https://CRAN.R-project.org/package=DT>
- [5] Wickham H. ggplot2: Elegant Graphics for Data Analysis. Springer-Verlag New York; 2016. (R package version 3.5.0)
- [6] Gohel D, Skintzos P. officer: Manipulation of Microsoft Word and PowerPoint Documents. R package version 0.6.5. 2024. <https://CRAN.R-project.org/package=officer>; Gohel D. flextable: Functions for Tabular Reporting. R package version 2.0.0. 2024. <https://CRAN.R-project.org/package=flextable>
- [7] CDISC Analysis Data Model (ADaM) Implementation Guide, Version 1.3. Clinical Data Interchange Standards Consortium, Inc. 2021. <https://www.cdisc.org/standards/foundational/adam>
- [8] Lenth RV. emmeans: Estimated Marginal Means, aka Least-Squares Means. R package version 1.10.0. 2024. <https://CRAN.R-project.org/package=emmeans>
- [9] Kenward MG, Roger JH. Small sample inference for fixed effects from restricted maximum likelihood. Biometrics. 1997;53(3):983-997.
- [10] U.S. Food and Drug Administration. Guidance for Industry: In Vitro Drug Interaction Studies — Cytochrome P450 Enzyme- and Transporter-Mediated Drug Interactions. Center for Drug Evaluation and Research; January 2020. <https://www.fda.gov/media/134582/download>
- [11] U.S. Food and Drug Administration. Guidance for Industry: Statistical Approaches to Establishing Bioequivalence. Center for Drug Evaluation and Research; 2001 (with 2026 draft revision). <https://www.fda.gov/media/70963/download>
- [12] Julious SA. Sample sizes for clinical trials with normal data. Stat Med. 2004;23(12):1921-1986.
- [13] Labes D, Schütz H, Lang B. PowerTOST: Power and Sample Size for Bioequivalence Studies. R package version 1.5.5. 2024. <https://CRAN.R-project.org/package=PowerTOST>
- [14] Powell A, Lawrence S, Loewen S, Tannenbaum S. A probabilistic method for combining intra-subject CV estimates from bioequivalence studies. (CV pooling method, as implemented in the app)

ACKNOWLEDGMENTS

We sincerely appreciate the collaboration from colleagues across Dizal, especially those who contributed suggestions and user feedback to this project.

CONTACT INFORMATION

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

Keyi Zhang

Dizal (Shanghai) Pharmaceutical Co., Ltd.

keyi.zhang@dizalpharma.com