



Universiteit
Leiden
The Netherlands

Copulas for covariate simulation in pharmacometrics

Guo, Y.; Guo, T.; Hasselt, J.G.C. van; Zwep, L.B.

Citation

Guo, Y., Guo, T., Hasselt, J. G. C. van, & Zwep, L. B. (2026). Copulas for covariate simulation in pharmacometrics. *Cpt: Pharmacometrics & Systems Pharmacology*, 15(4). doi:10.1002/psp4.70242

Version: Publisher's Version

License: [Creative Commons CC BY-NC 4.0 license](https://creativecommons.org/licenses/by-nc/4.0/)

Downloaded from: <https://hdl.handle.net/1887/4307620>

Note: To cite this publication please use the final published version (if applicable).

TUTORIAL OPEN ACCESS

Copulas for Covariate Simulation in Pharmacometrics

Yuchen Guo  | Tingjie Guo  | J. G. Coen van Hasselt  | Laura B. Zwep 

Division of Systems Pharmacology and Pharmacy, Leiden Academic Center for Drug Research, Leiden University, Leiden, the Netherlands

Correspondence: Laura B. Zwep (l.b.zwep@lacdr.leidenuniv.nl)

Received: 5 November 2025 | **Revised:** 17 March 2026 | **Accepted:** 26 March 2026

ABSTRACT

Patient-specific covariates are commonly incorporated in pharmacometric and quantitative system pharmacology models to predict differences in pharmacokinetic or pharmacodynamic profiles between patients. When simulating new virtual populations of patients, generating realistic covariate sets that accurately reflect the correlation structures among covariates is essential to obtain reliable simulation outcomes. Copulas are joint distribution functions that characterize the dependence structures of patient covariates and enable the simulation of virtual populations. The current tutorial provides a step-by-step guide for understanding the concept of copulas and an overview of applications of copulas in pharmacometric research.

1 | Introduction

In model-informed drug development and precision dosing, patient characteristics are key elements in predicting drug exposure (i.e., pharmacokinetics, PK) and treatment response (i.e., pharmacodynamics, PD). In population PK (PopPK) or physiological-based PK (PBPK) models, patient characteristics are often included as covariates or parameters, such as age, weight and creatinine concentrations. For PD or QSP models, treatment outcomes are often modeled in relation to patient-specific disease- or physiology-related characteristics, such as disease severity and baseline values of biomarker levels. In both scenarios, these models together with specified patient-specific properties can be used to derive individualized dosing schedules or optimized trial design which takes into account these patient-specific differences. For simplicity, covariates in PopPK/PD and patient-specific parameterizations in PBPK and QSP models will be referred to as covariates.

Covariates are often correlated. Therefore, when patient-specific covariates are incorporated into pharmacometric simulations, particularly when existing models are reused for new simulation purposes, it is important to ensure realistic combinations of patient covariates are generated. Several approaches are currently used, such as bootstrapping, which requires the availability of

underlying individual-level data. However, when individual-level data are not available, alternative approaches are needed, like univariate or multivariate normal distributions (MVN) fitted on the covariate data from a relevant population [1, 2]. These methods can, however, fail to realistically capture the dependency structures among patient characteristics, and thus do not accurately reflect the real-world data [3–5].

Copulas offer a reliable way to model and simulate complex dependencies among patient characteristics, enabling more realistic and reproducible virtual populations for a variety of research and regulatory applications [5, 6]. In this tutorial, we introduce the theoretical concept of copulas in the context of pharmacometrics (Section 2). Next, we provide a step-by-step practical guide with an example of covariate simulation in R, including evaluation of copula fits, allowing you to start using copulas right away (Section 3). Lastly, we discuss a number of specific issues and guidance for applications, including cases with discrete covariates (Section 3.7).

2 | Theory

This section provides a thorough statistical understanding of copula concepts. It is however possible to first explore copula

This is an open access article under the terms of the [Creative Commons Attribution-NonCommercial](https://creativecommons.org/licenses/by-nc/4.0/) License, which permits use, distribution and reproduction in any medium, provided the original work is properly cited and is not used for commercial purposes.

© 2026 The Author(s). *CPT: Pharmacometrics & Systems Pharmacology* published by Wiley Periodicals LLC on behalf of American Society for Clinical Pharmacology and Therapeutics.

estimation yourself (Section Practical workflow for copula development) and refer back here as needed.

2.1 | Statistical Foundations

2.1.1 | Probability Density

Some covariate values are more likely to occur than other values, which is described by the probability distribution. For continuous covariates, this distribution can be expressed as a probability density function (PDF, $f(X)$). The PDF can be obtained by calculating the frequency of values occurring within certain value ranges, and is often visualized using a histogram. A PDF can also be defined by a parametric or nonparametric function from a certain distribution, such as a normal or log-normal distribution.

PDFs can describe the density of more than one covariate at once, which is called bivariate or multivariate density functions. If we consider a situation with two covariates (X_1 and X_2), their PDF ($f(X_1, X_2)$) is called the joint PDF and describes the relative likelihood of observing the combination of values for X_1 and X_2 . The covariates can be correlated or dependent on each other, which is reflected in their joint distribution, also referred to as the dependency structure. In the case of two or multiple covariates, the marginal distribution of a covariate is obtained by considering its distribution regardless of other covariates, denoted by $f_1(X_1)$. This distribution can be viewed as the projection of a joint distribution onto a single dimension. Marginal distributions and the dependency structure are fundamental aspects when characterizing a covariate dataset.

A PDF can be estimated from data by choosing a function and estimating its parameters. The function can be either a parametric distribution, for example the (multivariate) normal distribution, or a nonparametric function, which does not assume a specific parametric distribution, like a kernel density function. Parameters are usually estimated by maximizing the likelihood. For a one-dimensional density or a marginal density, this optimization is relatively straightforward. However, the complexity of estimating a joint PDF increases with dimensionality.

2.1.2 | Conditional Probability Density

The density of one covariate (e.g., X_1) conditioned on the value of another covariate (e.g., X_2) is called the conditional density ($f(X_1|X_2)$). Its relationship with joint density $f(X_1, X_2)$ and marginal density $f_2(X_2)$ is illustrated in Equation (1).

$$f(X_1|X_2) = \frac{f(X_1, X_2)}{f_2(X_2)} \quad (1)$$

Similarly, the joint density of two covariates conditioned on the third covariate is the conditional joint density ($f(X_1, X_2|X_3)$). For example, the correlation between weight and height can vary by age, with a stronger correlation at a younger age. This can be denoted as $f(\text{height, weight}|\text{age})$, the joint density of height and weight at a certain age.

2.1.3 | Probability and Cumulative Density Functions

A probability that a covariate falls within a certain range of covariate values is given by the integral of the PDF over that range. This integral is called the cumulative distribution function (CDFs, $F(x)$), which is defined as the probability of a covariate being a value less than (or equal to) a specific value of x ($P(X \leq x)$). Similarly, the CDF for two covariates ($F(X_1, X_2)$) represents the probability that two covariates are less than specific values of X_1 and X_2 . The CDF can be estimated either non-parametrically, such as through data ranks or kernel density estimation, or parametrically by fitting a chosen CDF function (e.g., normal/lognormal distribution).

2.1.4 | Uniform Distribution

The last main concept required to understand copulas is the uniform distribution. A uniform distribution is when every value on the interval has the same density. Covariates are rarely uniformly distributed, but it is possible to transform the covariates to uniform distributions with probability integral transform (PIT, Figure 1) [7]. Briefly, for any covariate X with a PDF $f(X)$ and a CDF $F(X)$, the new random variable $U = F(X)$ follows a uniform distribution on $[0, 1]$ (derivation provided in Supporting Information A1). From uniform (U) scale, the data also can be easily mapped back to the original scale with the inverse of CDF ($F^{-1}(U)$).

2.2 | Copula

Copulas are functions describing multivariate distributions with marginal distributions that are uniform on $[0, 1]$ [8]. They are used to model the nondeterministic dependence between covariates. Therefore, derived covariates (e.g., body mass index) should not be included together with their components (e.g., height and weight). Throughout this section, we assume that all covariates are continuous and possess a joint density function. Under this assumption, Sklar's theorem implies that the joint density function of any pair of covariates ($f(X_1, X_2)$) could be expressed as the product of their marginal distributions and a copula function (Equation 2). Here, $c(U_1, U_2)$ represents the copula.

$$f(X_1, X_2) = c(U_1, U_2) \cdot f_1(X_1) \cdot f_2(X_2)$$

$$U_1 = F_1(X_1); U_2 = F_2(X_2) \quad (2)$$

The estimation of a copula can be done independently from the description of the marginal distribution. Since any marginal distribution can be transformed into a uniform distribution, copulas offer great flexibility for different marginal distributions that covariates may exhibit. Copulas include a wide variety of functions, including multivariate Gaussian (normal) distribution, but also any other (non-)parametric distributions. This flexibility allows copulas to capture diverse data structures. A bivariate copula describes the dependence structure between two covariates (Section 2.2.1), and a vine copula extends the framework to multiple covariates (Section 2.2.2). A more in-depth guide

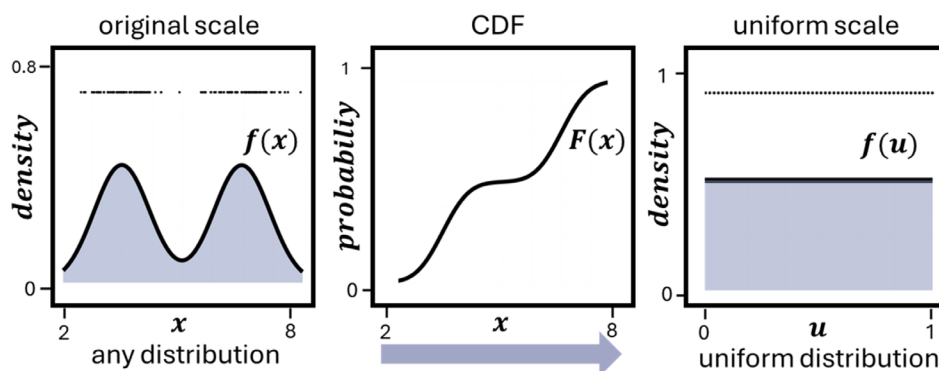


FIGURE 1 | Probability integral transform. Transforming a covariate x with an original distribution (left) to a uniform distributed covariate u (right) using the CDF (middle). The total area under the PDF, represented by the shaded region, equals 1 for both x and u . CDF, cumulative distribution function; PDF, probability density function.

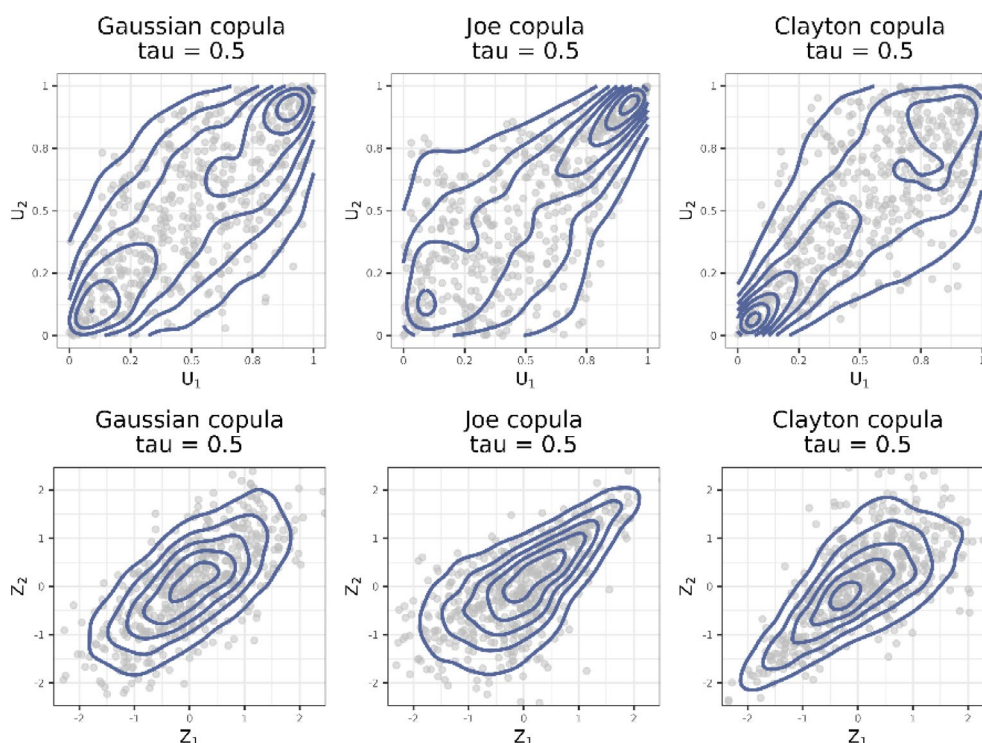


FIGURE 2 | Two-dimensional density contour plots. Example bivariate datasets that exhibit different dependency patterns but share the same correlation (Kendall's tau). U_1 and U_2 represent the uniform scale covariate data, and Z_1 , Z_2 represent the covariate data on a standard normal scale. The marginals of data were transformed from the uniform scale to the standard normal scale using standard normal quantile functions.

to the mathematical concepts of copulas is provided by Czado (2019) [9].

2.2.1 | Bivariate Copulas

Bivariate copulas, or pair copulas, are copulas for two covariates. According to Sklar's theorem, the joint distribution of two covariates X_1 and X_2 can be expressed as Equation (2), where the $c(U_1, U_2)$ is the bivariate copula. Various bivariate copula functions are available to describe different shapes of correlations, including symmetric or asymmetric, linear or nonlinear, monotonic or non-monotonic, and tail dependencies, where dependency differs over the values of the covariates [9, 10]. Bivariate copula functions include both parametric

and nonparametric, depending on whether a specific form is assumed.

A variety of copula functions are available in statistical software, and these functions allow for different dependency patterns. For instance, the Gaussian copula is a frequently used function to handle symmetric dependencies, and Joe or Clayton copulas are suited for capturing tail dependency (Figure 2). Note that the multivariate Gaussian distribution is a specific type of joint distribution with Gaussian marginal distributions and a Gaussian copula. Nonparametric approaches, such as polynomials or kernel density estimators, approximate the actual dependency structure without assuming a parametric shape [11]. Compared to parametric copulas, nonparametric copulas offer more flexibility, making them more suitable for modeling

complex dependencies, including non-monotonic dependencies. However, nonparametric estimation typically comes with higher computational costs and may encounter issues due to data sparsity.

2.2.2 | Vine Copulas

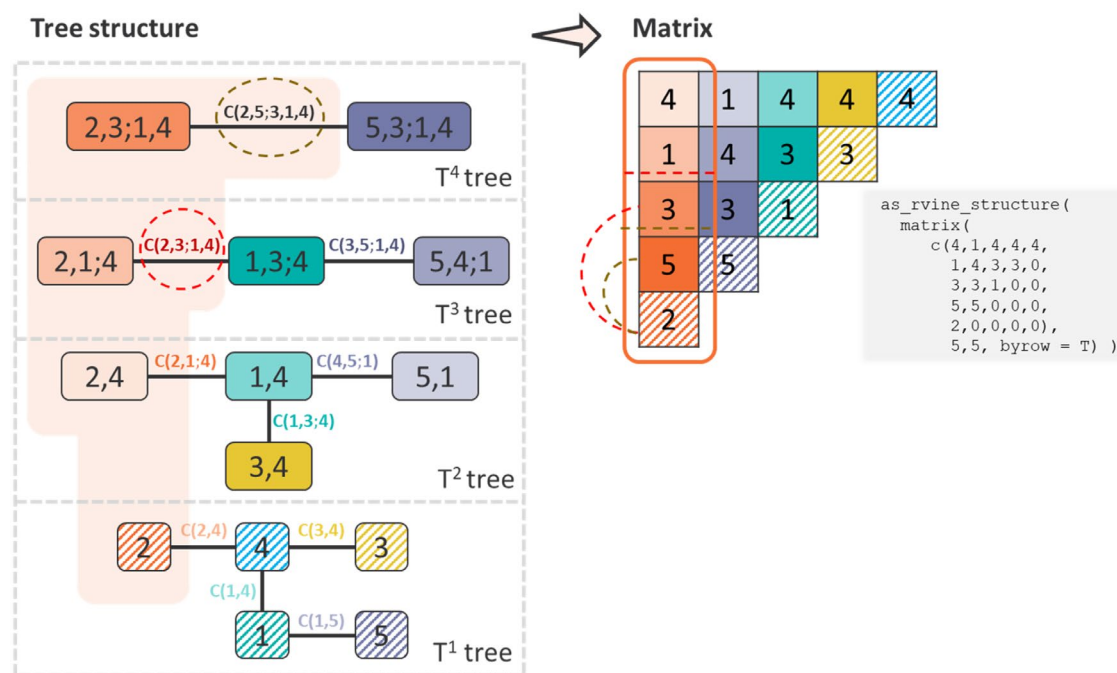
While Sklar's theorem generalizes to 3 or more dimensions, the copula function for d dimensions ($c(F_1(X_1), F_2(X_2), \dots, F_d(X_d))$) becomes exponentially more complex. Rather limited techniques are available to construct a multivariate copula [12]. To circumvent the problem, a d -dimensional copula can be constructed by decomposing the multivariate density into a set of bivariate copulas, called pair-copula construction [13]. Because the graphical representation of these copulas is similar to grape vine, copulas built via pair-copula constructions are called vine copulas.

Vine copula estimation involves an iterative process of structure learning and parameter estimation. The model structure of a vine copula, a tree structure that consists of a sequence of trees, specifies how covariates are associated or correlated with each

other. A full tree structure for a d -dimension covariate dataset consists of $d-1$ trees, and the k^{th} tree (T^k , $k = 1, 2$ to $d-1$) constitutes $d-k$ (conditional) bivariate copulas. For example, a full vine tree structure for five covariates consists of a sequence of four trees, starting with the first tree T^1 at the bottom and ending at the last tree T^4 at the top. (Box 1, left panel). T^1 tree defines the first set of pair copulas (edges), which captures the direct association between covariates (nodes). The second tree, T^2 , captures the conditional copula between the copulas from T^1 , where the T^1 copulas become the nodes and conditional copulas are the edges. This procedure continues, where the edges of each lower-level tree become the nodes in the next higher level tree, until all trees in the sequence are formed. Conditional copulas are denoted using a semicolon, e.g., $c(1, 3; 4)$ for copulas $c(1, 4)$ and $c(3, 4)$.

For each (conditional) pair copula, a bivariate copula function and its parameters are chosen and estimated independently, constituting a flexible framework for the estimation of a joint density function. Using the copula, the mathematical expression of joint density for the five covariates can be expressed as the product of their respective marginal distribution and all (conditional) bivariate copulas (Supporting Information A2).

BOX 1 | Translation of a tree structure to a vine matrix.



Example for five covariates (numbered 1 through 5). In R, a matrix is used to define the tree structure of a vine copula. To translate the tree structure into a matrix, the tree is decomposed into paths that correspond to a column in the matrix.

We start at the highest copula in the tree structure (here 2, 5; 3, 1, 4) and choose one of the conditioned covariates (2 or 5). The chosen first covariate (2) is the root covariate, and is placed on the diagonal lower left corner of the matrix. The conditioned covariate that the root covariate is coupled with is placed on the same column, descending in the tree structure and ascending in the column, so starting with the first copula (5), going down to the next (1, 3 and 4). Then we start in the second path, choosing the copula that has not been involved in the earlier path (5, 3; 1, 4) and repeat the process (choose root 5 on diagonal, descend in the tree structure: 3, 4, 1). This is repeated until arriving in the lowest tree and most right column in the matrix.

There is flexibility in choosing paths, so one tree structure can correspond to multiple matrices, but not the other way around.

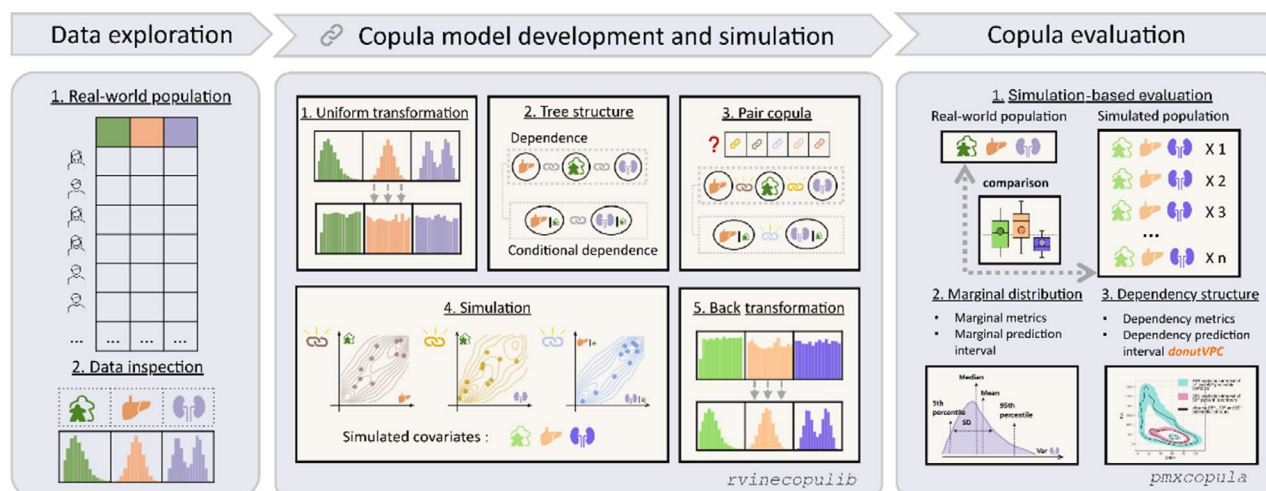


FIGURE 3 | Flow chart of copula estimation and simulation using vine copulas. For each step, the corresponding R packages used in this tutorial are noted in the boxes.

3 | Practical Workflow for Copula Development

In this section, we will first explain the procedure to estimate a copula model with a dataset of an observed (real-world) population. Next, we will demonstrate how to evaluate a copula by assessing whether the covariates simulated from a copula (virtual population) are representative of the observed covariates (Figure 3). Lastly, we will briefly showcase the application of covariates simulated from a copula in PK simulations.

3.1 | Software

The example code provided in this tutorial was developed in R 4.5.2. For updated code compatible with the latest version of R, refer to the *pmxcopula* package (<https://github.com/vanhasseltlab/pmxcopula>). Functions for kernel density integral transformation and copula estimation are from the R packages *kde1d* (v 1.1.1) and *rvinecopulib* (v 0.7.3.1.0), respectively [14, 15]. Example datasets MIMIC_5cov and mix_data, together with the evaluation tools are provided by the package *pmxcopula* (v 0.1.1). Pharmacokinetic simulation is done using *rxode2* (v 4.1.0) [16]. To navigate the package and its functionality, a summary table of key functions for copula development, simulation, and evaluation is provided in Supporting Information A3.

3.2 | Data Preparation

A copula can be estimated on data, which should represent the intended population and covariates of interest. The selection of relevant population and covariates depends on the research question. For example, the aim could be to simulate biomarkers in cancer patients, then the data need to include biomarker information of that specific population. For covariate selection it is important to consider the generalizability or specificity of the copula for future applications. By covering a relatively wide range of relevant covariates, the estimated copula can support the covariate simulation across various scenarios; however,

more dimensions in the copula require more observations (Section 3.7.2). Lastly, secondary covariates that can be calculated from primary covariates should be removed from copula estimation to avoid adding deterministic elements in the stochastic model. For instance, if body weight and height are included, BMI should be excluded.

For the practical example, we consider the MIMIC_5cov dataset from the *pmxcopula* package, which consists of five covariates from 1000 ICU patients, based on the MIMIC database [5]. Covariates are age, weight, blood urea nitrogen (BUN), serum creatinine and serum albumin. A heavily skewed marginal distribution can result in nonuniform transformed data with PIT. To improve uniformity, we use log-transformed BUN (logbun) and creatinine (logcrea) for copula model development.

```
library(kde1d)
library(pmxcopula)
library(dplyr)
data("MIMIC_5cov")
```

3.3 | Copula Estimation

3.3.1 | Uniform Transformation

Copula modeling typically begins with a uniform transformation of the margins, where the distribution of each covariate is rescaled to a uniform scale. Uniform distributed covariate data is obtained via PIT, which transforms any marginal distribution into a uniform distribution with the CDF of that covariate. CDF can be obtained with a fitted parametric distribution function, such as a (log)normal distribution,

```
age_unif <- pnorm(MIMIC_5cov$age, mean(MIMIC_5cov$age), sd(MIMIC_5cov$age))
```

or a nonparametric distribution function, such as kernel density estimate (KDE) of the distribution, for example using the R package *kde1d* [14].

TABLE 1 | Bivariate copula functions available for copula estimation in the package *rvinecopulib*.

Type	Name	Argument	Families
—	Independence	indep	All, nonparametric, parametric, itau
Elliptical	Gaussian	gaussian	All, parametric, onepar, elliptical, itau
Elliptical	Student t	student	All, parametric, twopar, elliptical, itau
Archimedean	Clayton	clayton	All, parametric, onepar, archimedean, itau
Archimedean	Gumbel	gumbel	All, parametric, onepar, archimedean, itau
Archimedean	Frank	frank	All, parametric, onepar, archimedean, itau
Archimedean	Joe	joe	All, parametric, onepar, archimedean, itau
Archimedean	Clayton-Gumbel (BB1)	bb1	All, parametric, twopar, archimedean, BB
Archimedean	Joe-Gumbel (BB6)	bb6	All, parametric, twopar, archimedean, BB
Archimedean	Joe-Clayton (BB7)	bb7	All, parametric, twopar, archimedean, BB
Archimedean	Joe-Frank (BB8)	bb8	All, parametric, twopar, archimedean, BB
Extreme-value	Tawn	tawn	All, parametric, threepar, ev
Nonparametric	Transformation kernel	ttl	All, nonparametric

```

crea_kde <- kde1d(MIMIC_5cov$logcrea,
mult=1)
crea_uniform <- pkde1d(MIMIC_5cov$logcrea,
crea_kde)

```

KDE-Based distributions are constructed with a smooth function and can learn the shape of density from the data without specifying a particular parametric form [17]. Here, a bandwidth multiplier of $\text{mult}=1$ was used to ensure standard smoothing. Given its flexibility, KDE can be used for each covariate in a dataset, for example using a for-loop or lapply function. The `kde1d` estimates both the CDF used for PIT ($F(x)$), and the inverse of CDF ($F^{-1}(x)$), which is essential for transforming simulated data from the uniform distribution back to the original data scale in the later section Simulation from the model developed with *vinecop*.

```

kde_list <- lapply(MIMIC_5cov, kde1d)
MIMIC_unif <- mapply(function(x, kde) pkde1d(x, kde),
MIMIC_5cov,
kde_list) %>% as.data.frame()

```

3.3.2 | Copula Fitting

A vine copula is fitted on the transformed data through pair copulas constructions, which capture the density between different pairs of covariates and copulas (see Section 2.2.2, Box 1 left panel). This process consists of two steps: structure learning and parameter learning.

3.3.2.1 | Structure Learning. The tree structure of a vine copula can be arbitrarily chosen, but the fit of a copula can be improved by placing pairs of covariates that are more correlated closer to each other in the vine.

The standard approach to determine a tree structure in *rvinecopulib* is the *maximum spanning tree* (MST) algorithm. In this algorithm, Kendall's tau correlations are first calculated for all possible covariate pairs and covariates are paired in the tree so the sum of the correlations for $d-1$ covariate pairs is maximized [18]. Next, bivariate copula functions are fitted for the selected pairs and parameters are estimated (Section 3.3.2.2). For the second tree, the MST selects the copula pairs that yield the strongest conditional correlations given the shared covariate. This procedure is iterated for each subsequent tree until all trees are selected and estimated.

There can be reasons for specifying a different tree structure, such as previous knowledge about which covariates are expected to exhibit stronger associations, or the interpretability of the copula between the covariates of interest in the first tree. This is possible within *rvinecopulib* by translating a part of, or a whole tree structure into a matrix that can be specified in R (Box 1).

3.3.2.2 | Parameter Learning. For each selected covariate pair in the vine structure (Box 1), a bivariate copula function is fitted. A candidate bivariate copula function is first specified, followed by parameter estimation. The *rvinecopulib* package offers a variety of bivariate copula functions (Table 1). Instead of specifying a particular copula function for a covariate pair, a collection or family of candidate copula functions can be defined, from which the best fitting one is selected based on criteria such as the Akaike information criterion (AIC) or Bayesian information criterion (BIC). The choice of copula functions requires balancing flexibility and computational efficiency, which can also introduce overfitting (Section 3.7.2). Including a wider range of bivariate copulas increases flexibility but may lead to higher computational costs, especially with nonparametric copulas. To increase computational speed, nonparametric options can be excluded by using predefined sets, such as the “parametric” copula family.

```
library(rvinecopulib)
MIMIC_copula_unif <- vinecop(MIMIC_unif, fam-
ily_set = "parametric")
```

Once a copula is fitted, its tree structure can be obtained as matrix form and visualization.

```
MIMIC_copula_unif$structure
plot(MIMIC_copula_unif, tree = "ALL")
```

3.3.3 | Single-Function Copula Development

While following the above steps sequentially provides flexibility and helps build intuition, it is reassuring that this is not always necessary. The *rvinecopulib* package offers the powerful `vine()` function, which combines all prior steps into one function.

```
MIMIC_copula <- vine(MIMIC_5cov,
copula_controls = list(family_set =
"parametric"),
margins_controls = list(mult = 1))
```

The function `vine()` internally uses kernel density estimation (`kde1d`) for the transformation of covariates into uniform scale and fits a copula. The arguments `copula_controls` and `margins_controls` correspond to the arguments in `vinecop()` and `kde1d()`, respectively.

For the copula model fitted with `vine()`, the tree structure is contained within the copula component and must be accessed via, for example, `MIMIC_copula$copula$structure`. Other functions can be applied following the same convention.

3.4 | Copula Simulation

Covariate simulation is often the ultimate goal of copula modeling. Once the copula has been estimated, covariates can be simulated by drawing random samples from the copula. This will produce a dataset with n rows and d covariates as columns. One point to keep in mind is that the scale of simulated outcomes differs depending on whether the model was fitted with `vinecop()` or `vine()`.

3.4.1 | Simulation From the Model Developed With Vinecop()

The covariates simulated from the model developed with `vinecop()`, that is, `MIMIC_copula_unif`, are on a uniform scale. The syntax for generating samples is similar to other random sampling functions in base R, such as `rnorm()`.

```
MIMIC_sim_unif <- rvinecop(100,
MIMIC_copula_unif)
```

To transform the uniformly distributed data back to the original scale, we use the inverse of the CDF, the quantile function $F^{-1}(u)$. In case of a parametric distributions, the (base) R functions can be used to back-transform the data using `q<dist>()` (e.g., `qnorm()`). For kernel densities, the equivalent function

`qkde1d()` can be used. Here, we need to enter the same distribution object as used in Section 3.3.1. We apply the quantile function to each covariate using `mapply()`.

```
MIMIC_sim <- mapply(function(x, kde) qkd-
eld(x, kde),
MIMIC_sim_unif, kde_list) %>% as.data.frame()
```

The outcome (`MIMIC_sim`) is a data frame with the simulated covariates on the original scale.

3.4.2 | Simulation From the Model Developed With Vine()

Covariates at the original scale can be directly simulated from the copula object estimated with the `vine()` function, for example, `MIMIC_copula`. To prevent sampling out of the feasible space, it is possible to set boundaries (minimum and maximum) for marginal characterization during copula estimation.

```
MIMIC_sim <- rvine(100, MIMIC_copula)
```

In case of the situation where one specific subgroup population is relevant for further pharmacometric simulation, conditional sampling is supported by *pmxcopula*, with the option to define the population of interest with constrained conditions for both continuous and discrete variables.

```
MIMIC_VP <- rscopula(copula = MIMIC_copula,
n_sim = 100,
xmin = c(18, 50, NA, NA, NA)
xmax = c(50, 120, NA, NA, NA),
var_types = c('c', 'c', 'c', 'c', 'c'))
```

3.5 | Evaluation Methods

The evaluation methods for copulas differ from those of pharmacometric models. Copula evaluation aims to not only assess how well the model fits the data, but also to examine whether the copula fulfills its intended purpose: simulating virtual populations that accurately represent the observed population. This section will cover how to critically examine the performance of a copula model.

3.5.1 | Evaluation of Model Fit

Different approaches can be used to evaluate the overall performance of a copula. Goodness-of-fit (GOF) measures, such as loglikelihood, AIC, and BIC, can be used to select the most suitable copula. Overall model fit can be evaluated using the GOF measures, as well as the GOF tests based on the statistics of, for instance, the Kolmogorov–Smirnov statistic, Cramér–von Mises statistic, and Kullback–Leibler (KL) divergence [19–21]. These GOF tests, often summarized with bootstrapped p -values, can compare the performance of candidate copula models and facilitate model selection. However, these summary metrics do not provide detailed and diagnostic insights into how well the copula is fit for purpose.

3.5.2 | Evaluation for Covariate Simulation

For covariate simulation, the purpose of copula evaluation is to determine whether simulated populations resemble the observed population. Therefore, we propose to evaluate a copula with simulation-based diagnostics. This involves comparing distribution properties between the observed population and simulated populations, considering two aspects: (1) the marginal distribution, which is the probability distribution of each individual covariate, such as age, and (2) the dependency structure, which is the relationship between covariates [6]. Both aspects can be evaluated visually and quantitatively. The simulation-based tools in *pmxcopula* also facilitate comparison of different covariate simulation methods, including but not limited to copulas. In this section, the performance of MIMIC_copula on the overall population is assessed.

3.5.2.1 | VP Simulation for Evaluation. A strategy involving multiple simulations of the original population is generally recommended to account for and assess the sampling variability, for example, 100 repeated simulations. In function *rcopula()*, *sim_nr* is specified to indicate the number of simulation replicates. Each replicate contains the same number of individuals as the original dataset. Next to the columns representing each covariate, a column is included to indicate the simulation

run in the simulated dataset. This column is required for subsequent evaluation functions.

```
set.seed(12345)
MIMIC_sim <- rcopula(MIMIC_copula, sim_nr = 100)
```

3.5.2.2 | Visualization of Distributions. Density plots are a visual method to examine how the copula describes the observed data. The function *plot_1dist()* visualizes the density curves for each covariate from observed and multiple simulated datasets (Figure S1). To assess the dependent relations, the function *plot_mvdist()* displays two-dimensional density contours or scatter plots for all pair combinations of covariates, comparing the observed data with one of the simulated datasets (Figure S2). A density contour is a line or curve representing the region of equal probability density derived from the data. Together, these functions offer a quick overview of the model fit. Ideally, the density curves and density contours of the observed and simulated populations should closely align.

```
plot_1dist(sim_data = MIMIC_sim,
  obs_data = MIMIC_5cov,
  sim_nr = 100,
  pick_color = c("#F68E60", "#747AA9"))
plot_mvdist(sim_data = MIMIC_sim,
  obs_data = MIMIC_5cov,
  pick_color = c("#F68E60", "#747AA9"),
  plot_type = "density",
  bins=8)
```

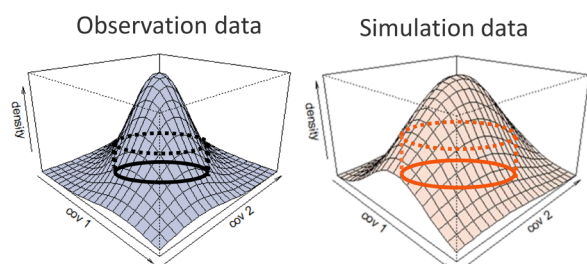
3.5.2.3 | Quantitative Evaluation Metrics. Quantitative evaluation metrics can assist copula evaluation in a quantitative manner. For this purpose, mean, median, standard deviation, 5th and 95th quantiles are included as marginal metrics in *pmxcopula* package. These metrics measure the centrality, limit behaviors, and variability of each covariate. The differences in marginal metrics between virtual and observed populations can be expressed as relative errors. Relative errors close to zero imply that the margins of virtual populations are in good agreement with those of the observed population.

```
calc_margin(sim_data = MIMIC_sim,
  obs_data = MIMIC_5cov,
  sim_nr = 100,
  var = NULL,
  aim_statistic = c("mean", "median"))
```

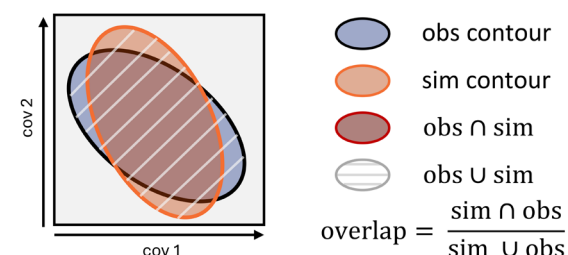
To quantify the difference between the dependency structure of observed and simulated distributions, two measures are available in the *pmxcopula* package. One is Pearson's correlation, which quantifies the linear relationship. When the Pearson correlations between observed and simulated covariates are close, it suggests that the linear correlation between covariates has been well captured. Another measure is the overlap metric, which assesses the nonlinear dependencies. Overlap metrics account for the shape of dependency pattern by calculating of overlaps between the bivariate density contours from observed covariates and those of simulated covariates (Box 2). The overlap metric quantifies the ratio of intersection area to union area

BOX 2 | Illustration of the overlap measure.

Step 1: obtain the bivariate density contour



Step 2: calculate the Jaccard index



A density value can be defined and the density contours at this level are compared using the overlap measure. The probability density surface is estimated from Kernel density estimate for the observed and simulated population (Step 1). The *p*th percentile contour is related to *p*% data point coverage, e.g., the 95th percentile contour encloses 95% of points. For a specified density level, the overlap metric is measured by the Jaccard index, the ratio of the intersection to the union of observed and simulated density contours (Step 2).

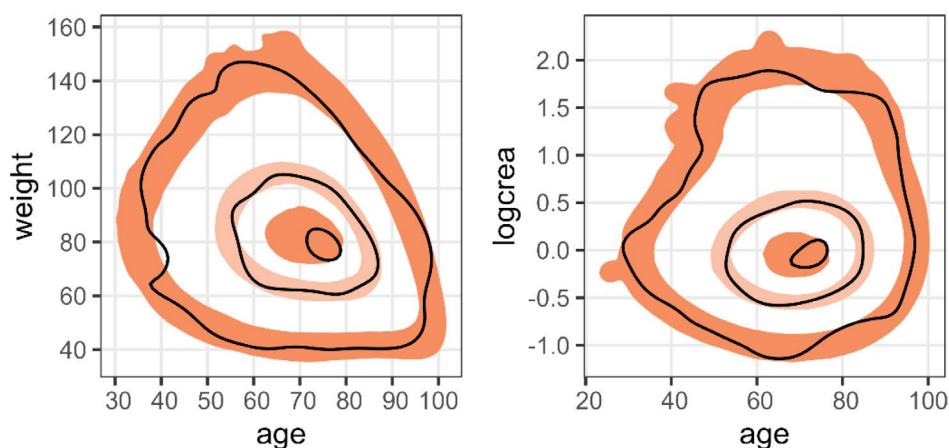


FIGURE 4 | Donut VPC, visual predictive check for dependency structures. The shaded areas represent the 95% confidence bands of the 5th, 50th and 95th percentile bivariate density contours of multiple simulations, and the black solid lines represent the density contours of the observed population.

of two density contours at the same percentile, also known as the Jaccard index. A high overlap relates to a good estimation of the dependency structure.

```
calc_dependency(sim_data = MIMIC_sim,
  obs_data = MIMIC_5cov,
  pairs_matrix = NULL,
  percentile = 95,
  sim_nr = 100,
  cores=4)
```

GOF metrics are implemented in *VineCopula*, *gofCopula*, and *kldest* packages [20–22]. In *VineCopula* and *gofCopula*, GOF tests using Kolmogorov–Smirnov and Cramér–von Mises statistics could be used to test the difference between the copula and observed population using bootstrap-based p -values, where rejecting the null hypothesis implies a bad fit of the copula. In *kldest*, KL divergence measures the distance between the distributions of the observed and virtual population, with lower KL values indicating the copula more closely resembles the observed population.

```
library(VineCopula)
copula_model <- RVineStructureSelect(MIMIC_
  unif, family = c(0:4))
MIMIC_GOFtest <- RVineGofTest(MIMIC_unif,
  copula_model,
  statistic = "CvM",
  B=100)
MIMIC_GOFtest[["p.value"]]
library(kldest)
kld_est_nn(MIMIC_sim %>% select(-simulation_
  nr), MIMIC_5cov)
```

3.5.2.4 | Visual Predictive Checks. Visual predictive checks (VPCs), commonly used for the evaluation of pharmacometrics models, are also applicable to copulas. For marginal distributions, 95% prediction intervals across all quantiles (from 1st to 99th) from multiple simulations are plotted versus that of the observed population in a quantile-quantile plot (Q-Q plot). A good fit corresponds to the ribbon tightly aligning the diagonal line (Figure S3).

```
plot_qq(sim_data = MIMIC_sim,
  obs_data = MIMIC_5cov,
  sim_nr = 100,
  conf_band = 95,
  var = c("age", "weight"),
  type = "ribbon")
```

A 2-dimensional VPC, named donutVPC, was developed to evaluate the model fitness for the pair-wise relations (Figure 4). Similar to the traditional pharmacometric VPC [23], donutVPCs are constructed with multiple simulations but focus on density contours. A density contour at the p^{th} percentile is estimated using kernel density estimation (KDE) at the region containing $p\%$ of the simulated individuals. The donutVPC is constructed by simulating new samples from the copula and estimating the KDE for each sample. Then the contour points are extracted for this p^{th} percentile. Around those contour points a new contour is estimated to get an approximation of the specified confidence band around the contours (e.g., 95% confidence, Supporting Information A4). The bands for the percentile contours are then plotted versus the observed data contours, producing a donut shaped plot. Traditional VPCs would include data points, however a scatter plot is not directly interpretable compared to density contours. A close overlap between the prediction bands and observed contours suggests that the copula accurately represents the observed population.

```
donutVPC(sim_data = MIMIC_sim,
  obs_data = MIMIC_5cov,
  percentiles = c(5, 50, 95),
  sim_nr = 100,
  pairs_matrix = matrix(c("age", "age",
    "weight", "logcrea"), 2, 2)
  conf_band=95,
  colors_bands = c("#F68E60", "#FAC1A8"),
  cores=4)
```

3.5.3 | Improve Copula Fit

Evaluation results could detect discrepancies between virtual populations and the observed population. It is possible that

copula models fail to characterize subsets of covariates or subgroup populations. This section outlines three aspects relevant to improving a copula model:

Lack of fit may arise when the transformed data deviate from a uniform distribution. A proper choice of uniform transformation could improve the copula. In the case of a heavily skewed original marginal distribution, log transformation before PIT reduces skewness and improves the uniformity of the transformed covariates.

Inadequate model performance may result from the misspecification of the dependency structure. When data is insufficient to fit a variety of bivariate copula functions or estimate the full dependency structure, restricting bivariate copula options during vine estimation is recommended. However, this requires careful consideration, as the choice of bivariate copulas can have a significant impact on copula goodness-of-fit. For example, Gaussian vine copulas use only Gaussian bivariate copulas for all covariate pairs, which overlook asymmetric or tail dependencies. This limitation can potentially lead to incorrect predictions, which contributed to the underestimation of risk in the 2008 financial crisis [24]. In most cases, a family of parametric bivariate copulas is sufficient for modeling covariate relationships. However, if covariates exhibit unusual or complex dependency patterns that parametric copulas cannot capture, nonparametric bivariate copulas, such as kernel-based copulas, may provide a better fit.

Poor fits can also arise from outlier subpopulations. Such subpopulations can be difficult to detect without subgroup evaluation, since the overall fit of the copula model may still appear satisfactory. If the poor fits are identified within certain subpopulations, e.g., male Asian population, this likely indicates that the covariate dependency structure for this group differs from that of the overall population. In such cases, one can either exclude this subgroup and estimate a separate copula model, or proceed with caution using the existing copula, acknowledging its limitations if there is insufficient data to support a separate analysis.

3.6 | Applications in Pharmacometrics

Once a copula has been estimated and evaluated, it can be used to generate virtual populations with realistic covariate distributions. These simulated populations can be easily integrated into pharmacometric simulations, such as PopPK, PD, PBPK, and QSP model-based simulations. By ensuring that the covariates or physiological parameters reflect realistic variability and correlations, copulas enhance the credibility and predictive performance of pharmacometric analyses. This approach enables the evaluation of dosing regimens in an untested population and supports individualized treatment strategies across diverse patient populations. Example code is included in the Supporting Information A5. In addition to virtual population generation, copulas have a wide set of potential uses in pharmacometrics. One of these uses is to learn parameter uncertainty from bootstrap or sampling importance resampling results [25], which can then be used to simulate parameters, incorporating the parameter uncertainty.

3.7 | Considerations for Copula Application

3.7.1 | Discrete Data

The first two sections focus on theory and practical examples within a continuous framework; however, copulas generalize to discrete covariates [26]. A common approach for modeling discrete covariates in the copula framework is to render discrete covariates continuous by assigning an order and adding small noise. The jittered data become continuous, and the original algorithm applies [27]. In *rvinecopulib*, this method is automatically applied if the covariate is specified as discrete ("d").

As ordering is one of the intrinsic natures of continuous data, an order for the discrete covariate is required. For ordered categorical covariates, such as disease stage (stage I, II, III, IV) and pain scale (0 to 5), ordering can be set as the inherent order using the base R data type ordered. Then, the discrete covariate is converted into a factor variable with an order, and the copula can be fitted with the specification of covariate data type.

To illustrate copula estimation with both discrete and continuous covariates, we consider a mixed dataset `mix_data` from *pmxcopula*. The dataset includes sex, age, weight, and albumin level. Although age is often treated as continuous, its integer nature allows it to be modeled as discrete. In this example, age is treated as discrete to demonstrate copula modeling with mixed data, with sex and age handled as ordered categorical covariates.

```
data("mix_data")
mix_data$Sex <- ordered(mix_data$Sex,
  levels = c("male", "female"),
  labels = c("1", "2"))
mix_data$Age <- ordered(mix_data$Age)
vine_discrete <- vine(dat = mix_data,
  copula_controls = list(
    family_set = "parametric",
    var_types = c("d", "d", "c", "c")))
```

In the case of dichotomous covariates, the order is arbitrary; for ordinal covariates, the order is known. There are, however, some limitations when dealing with unordered categorical covariates, also known as nominal covariates. If the order is unknown, for example for ethnicity and disease classification, the copula is not always able to provide a good fit, due to non-monotonicity. With a small number of unordered discrete covariates and categories, a good fit can be found by brute forcing, going through all order combinations and choosing the one with the best performance. An alternative approach is one-hot encoding, that is, converting nominal variables into binary indicators. For example, a three-category variable can be encoded as [1, 0, 0], [0, 1, 0] and [0, 0, 1] before fitting a copula. However, a key limitation is that a copula can simulate values that do not correspond to exactly one category (e.g., [1, 1, 0]). Although vine copulas provide useability for categorical covariates [15], more mathematical limitations have been discussed in literature as well [26, 28].

To evaluate the performance of copulas involving discrete covariates, the marginal distributions can be evaluated by comparing the number of simulated individuals in the category of

interest (e.g., 0–10years, 10–20years) with the counts in the observed population. For both marginal and dependency structure, subgroup analysis is an option to detect any misfit within subgroup populations [6].

3.7.2 | Data Size

Copula estimation generally requires a large number of observations to obtain a reliable and robust density estimation. The required number of observations grows with the number of dimensions or covariates. For a dataset with d dimensions, the full vine copula structure includes $d(d-1)/2$ pair copulas. If using parametric bivariate copula functions, each bivariate copula requires the estimation of one or two parameters. For a 20-dimensional dataset, there are 190 bivariate copulas and 190–380 parameters to be estimated.

A rule of thumb is that an adequate dataset should have observations (n) more than the dimension (d) raised to the power of 4 (d^4) [29]. However, this can be different depending on the data and how much prior information or assumptions can be made. For instance, weak, medium, or strong dependence relations result in different probabilities of successful identification of the copula model [30]. If prior knowledge, such as the (conditional) independence relationships for certain covariate pairs, is available, vine copula structure can be simplified to improve robustness of the estimation. This is particularly important for small datasets, where model simplification helps reduce the risks of overfitting and computational burden due to the large number of possible model parameters. The risk of overfitting can also be reduced by assuming a parametric copula instead of a nonparametric copula.

Copula truncation and pruning are often used to reduce model complexity by limiting the vine structure. Specifically, truncation sets bivariate copulas to be independent in the trees above a certain level, and pruning assumes independence for covariate pairs that have weak correlations. For datasets with a large number of covariates, it is possible to assume bivariate copulas in higher trees are independent [18]. To determine an optimal truncation level, criteria such as the BIC, modified BIC and Vuong test were proposed [31]. However, when all covariates are correlated, the correlation in higher level trees may remain significant. In such cases, truncating the entire tree might result in the loss of information. Instead, only pruning the covariate pairs that have weak correlations is a more effective approach to simplify the vine structure [32]. To this end, a threshold need to be defined in order to set a bivariate copula to independence. In *rvinecopulib* package, truncation level and pruning threshold can both be either selected by modified BIC or specified empirically [29].

3.7.3 | Missing Data

Missing data is common and can affect copula estimation and statistical inference. Missing data typically occurs in three mechanisms: [33] (1) missing completely at random (MCAR): the probability of the observed missing data is independent of any covariate in the dataset; for example, covariates are lost because of random technical issue; (2) missing at random (MAR):

the probability of observing the missing data is independent of the corresponding covariate but related to other covariates; for instance, in a dataset of male and female subjects, creatinine is only recorded for females, thus, the creatinine measurements for males are MAR; (3) missing not at random (MNAR): the probability of observing the missing data depends on the values of the corresponding covariate; for example, MNAR occurs when clinical variables are below the lower limit of quantification (BLQ) and recorded as missing.

The *rvinecopulib* package allows the estimation of vine copulas from incomplete datasets by using complete pairwise observations. However, this can lead to bias, depending on the missingness mechanism. A strategy to address this issue is to impute missing data with plausible values before estimation [34]. Many imputation methods are available, but the choice of imputation method requires careful consideration of the assumptions made about the missing data mechanism. Different imputation methods have been previously described in the literature [35].

The impact of missing data relies on the modeling purpose. If the aim is to approximate the underlying true joint distribution, expert knowledge on causal relationships in the data generation process is crucial [36]. In contrast, if the aim is to generate a VP representative of the observed population, then the priority is to ensure simulated datasets resemble the observed datasets as closely as possible. Under MCAR, the missing data is a random subset of the real-world data, and standard copula estimation and simulation remain unbiased. When data are MAR or MNAR, the observed sample is not representative of the true distribution and can lead to bias in the generated virtual populations. For example, in the case of MAR, the unknown creatinine values for males can be inferred by assuming the same conditional distribution given males as observed in females. If the assumptions can be justified, simulations will not introduce bias from the missingness. If the assumption cannot be justified, simulated data outside the observed space should be removed from further simulations or treated with care during application. In the case of MNAR, such as when values below BLQ or above the upper quantification limit are missing, discarding out-of-range data and setting boundaries for marginal distribution during copula estimation could ensure that the generated virtual population remains within realistic and observed ranges.

The impact of missing data on copula-based covariate simulation also depends on the amount of missingness. When the proportion of missing data is relatively small, the effect on the estimated dependency structure is often negligible. As the proportion increases, the risk of bias correspondingly increases [37].

3.7.4 | Compositional Data

Compositional data represents a unique type of data in dependency modeling. A notable example is gut microbiome data, which has been increasingly recognized as contributors to variability in PK and clinical outcomes [38]. Microbiome data are typically reported as

compositional data, with species abundances expressed as proportions (e.g., 0%–100%) that sum to one. A key challenge in copula modeling of compositional data is the risk of generating virtual patients that fall outside bounds or violate the constant-sum constraint. Several approaches can be used to address this issue. If the absolute value is available (e.g., counts), compositional data can be mapped back to its original scale to eliminate the bounds, and copulas can be fitted on this scale, and relative abundances can be computed afterward from the simulated data. A second approach is to fit copulas directly on the compositional data and scale the simulated values by dividing by their sum to ensure the sum-to-one constraint. Alternative approaches include specifying bounds in margin characterization, or applying transformations (e.g., log-ratio transformation) [39] while modeling $d-1$ variables instead of fitting the copula on all (d) variables. However, these latter approaches still carry a risk of generating virtual patients that violate the constant-sum requirement.

3.7.5 | Time-Varying Copulas

Covariates could vary with time, and the dependency structure between covariates may also be time dependent. Time-varying copulas remain an active research area to address the dynamic relationship between variables [40–42]. Several strategies have been proposed, such as building copulas based on parameters inferred from separate mixed-effects models of the covariates [5], or incorporating time-related parameters within copulas [43]. Although time-varying copulas are frequently used in econometrics, currently, no standard has been established in pharmacometrics yet.

4 | Conclusion

This tutorial has provided an overview of the estimation, evaluation, and application of copulas for covariate simulation in R. The copula is a versatile tool to describe the dependency between covariates. A robust copula could not only help reveal the complex relationships underlying the data but also facilitate the generation of realistic virtual populations for model-informed precision dosing and in silico clinical trials. Once a valid copula is established, it serves as a vehicle to share sensitive patient data, enabling broader access to patient characteristics.

Acknowledgments

We want to thank Bart van Lieshout for his code review and valuable insights for improving the readability of the manuscript, as well as the reviewers for their thorough and thoughtful feedback.

Funding

This research was supported by the Dutch Research Council (NWO) under the Open Science (OS) Fund, grant number OSF23.2.092.

Conflicts of Interest

The authors declare no conflicts of interest.

References

1. S. J. Tannenbaum, N. H. Holford, H. Lee, C. C. Peck, and D. R. Mould, "Simulation of Correlated Continuous and Categorical Variables Using a Single Multivariate Distribution," *Journal of Pharmacokinetics and Pharmacodynamics* 33 (2006): 773–794.
2. D. Teutonico, F. Musuamba, H. J. Maas, et al., "Generating Virtual Patients by Multivariate and Discrete Re-Sampling Techniques," *Pharmaceutical Research* 32 (2015): 3228–3237.
3. N. Hartung and A. Khatova, "Information-Theoretic Evaluation of Covariate Distributions Models," *Journal of Pharmacokinetics and Pharmacodynamics* 52 (2025): 21.
4. G. Smania and E. N. Jonsson, "Conditional Distribution Modeling as an Alternative Method for Covariates Simulation: Comparison With Joint Multivariate Normal and Bootstrap Techniques," *CPT: Pharmacometrics & Systems Pharmacology* 10 (2021): 330–339.
5. L. B. Zwep, T. Guo, T. Nagler, C. A. J. Knibbe, J. J. Meulman, and J. G. C. van Hasselt, "Virtual Patient Simulation Using Copula Modeling," *Clinical Pharmacology & Therapeutics* 115 (2024): 795–804.
6. Y. Guo, T. Guo, C. A. J. Knibbe, L. B. Zwep, and J. G. C. van Hasselt, "Generation of Realistic Virtual Adult Populations Using a Model-Based Copula Approach," *Journal of Pharmacokinetics and Pharmacodynamics* 51 (2024): 735–746.
7. C. Genest and L.-P. Rivest, "On the Multivariate Probability Integral Transformation," *Statistics & Probability Letters* 53 (2001): 391–399.
8. A. Sklar, "Random Variables, Joint Distribution Functions, and Copulas," *Kybernetika* 09 (1973): 449–460.
9. C. Czado, "Analyzing Dependent Data With Vine Copulas," in *Lecture Notes in Statistics*, vol. 222 (Springer, 2019).
10. M. M. Marchione and F. Baione, "Non-Monotone Dependence Modeling With Copulas: An Application to the Volume-Return Relationship," (2024), <https://doi.org/10.48550/arXiv.2403.15862>.
11. S. B. Provost and Y. Zang, "Nonparametric Copula Density Estimation Methodologies," *Mathematics* 12 (2024): 398.
12. O. Okhrin, A. Ristig, and Y. F. Xu, "Copulae in High Dimensions: An Introduction," *Applied Quantitative Finance* (2017): 247–277.
13. A. Panagiotelis, C. Czado, and H. Joe, "Pair Copula Constructions for Multivariate Discrete Data," *Journal of the American Statistical Association* 107 (2012): 1063–1072.
14. T. Nagler and T. Vatter, "kde1d: Univariate Kernel Density Estimation," (2020), <https://cran.r-project.org/package=kde1d>.
15. T. Nagler and T. Vatter, "rvinecopulib: High Performance Algorithms for Vine Copula Modeling," (2023), <https://CRAN.R-project.org/package=rvinecopulib>.
16. M. L. Fidler, M. Hallow, J. Wilkins, and W. Wang, "RxCODE: Facilities for Simulating from ODE-Based Models," (2025), <https://CRAN.R-project.org/package=RxCODE>.
17. Y. C. Chen, "A Tutorial on Kernel Density Estimation and Recent Advances," *Biostatistics & Epidemiology* 1 (2017): 161–187.
18. J. Dißmann, E. C. Brechmann, C. Czado, and D. Kurowicka, "Selecting and Estimating Regular Vine Copulae and Application to Financial Returns," *Computational Statistics & Data Analysis* 59 (2013): 52–69.
19. C. Czado, *Analyzing Dependent Data With Vine Copulas*, vol. 222 (Springer International Publishing, 2019).
20. N. Hartung, "kldest: Sample-Based Estimation of Kullback-Leibler Divergence," (2024), <https://niklhart.github.io/kldest/>.
21. O. Okhrin, S. Trimborn, and M. Waltz, "gofCopula: Goodness-Of-Fit Tests for Copulae," *R Journal* 13 (2021): 467.

22. T. Nagler, U. Schepsmeier, J. Stoeber, et al., "VineCopula: Statistical Inference of Vine Copulas," (2023), <https://CRAN.R-project.org/package=VineCopula>.
23. O. Mats and A. Karlsson, "Tutorial on Visual Predictive Checks," (2008), <https://www.page-meeting.org/default.asp?abstract=1434>.
24. D. MacKenzie and T. Spears, "'The Formula That Killed Wall Street': The Gaussian Copula and Modelling Practices in Investment Banking," *Social Studies of Science* 44 (2014): 393–417.
25. A. G. Dosne, M. Bergstrand, K. Harling, and M. O. Karlsson, "Improving the Estimation of Parameter Uncertainty Distributions in Nonlinear Mixed Effects Models Using Sampling Importance Resampling," *Journal of Pharmacokinetics and Pharmacodynamics* 43 (2016): 583–596.
26. G. Geenens, "Copula Modeling for Discrete Random Vectors," *Dependence Modeling* 8 (2020): 417–440.
27. T. Nagler, "Asymptotic Analysis of the Jittering Kernel Density Estimator," *Mathematical Methods of Statistics* 27 (2018): 32–46.
28. O. P. Faugeras, "Inference for Copula Modeling of Discrete Data: A Cautionary Tale and Some Facts," *Dependence Modeling* 5 (2017): 121–132.
29. T. Nagler, C. Bumann, and C. Czado, "Model Selection in Sparse High-Dimensional Vine Copula Models With an Application to Portfolio Risk," *Journal of Multivariate Analysis* 172 (2019): 180–192.
30. D. Huard, G. Évin, and A.-C. Favre, "Bayesian Copula Selection," *Computational Statistics & Data Analysis* 51 (2006): 809–822.
31. E. C. Brechmann, C. Czado, and K. Aas, "Truncated Regular Vines in High Dimensions With Application to Financial Data," *Canadian Journal of Statistics* 40 (2012): 68–85.
32. J. Wan and S. Li, "Modeling and Application of Industrial Process Fault Detection Based on Pruning Vine Copula," *Chemometrics and Intelligent Laboratory Systems* 184 (2019): 1–13.
33. Y. Zhang, R. Zhang, and B. Zhao, "A Systematic Review of Generative Adversarial Imputation Network in Missing Data Imputation," *Neural Computing and Applications* 35 (2023): 19685–19705.
34. E. Liebscher, "Fitting Copulas in the Case of Missing Data," *Statistical Papers* 65 (2024): 3681–3711.
35. H. Kang, "The Prevention and Handling of the Missing Data," *Korean Journal of Anesthesiology* 64 (2013): 402–406.
36. M. Kertel and M. Pauly, "Estimating Gaussian Copulas With Missing Data With and Without Expert Knowledge," *Entropy* 24 (2022): 1849.
37. H. Wang, F. Fazayeli, S. Chatterjee, and A. Banerjee, "Gaussian Copula Precision Estimation With Missing Values," in *Artificial Intelligence and Statistics* (PMLR, 2014), 978–986.
38. A. Saqr, S. Cheng, M. Al-Kofahi, C. Staley, and P. A. Jacobson, "Microbiome-Informed Dosing: Exploring Gut Microbial Communities Impact on Mycophenolate Enterohepatic Circulation and Therapeutic Target Achievement," *Clinical Pharmacology and Therapeutics* 118 (2025): 1477–1488.
39. E. Calderín-Ojeda, G. Lopez-Campos, and E. Gómez-Déniz, "A Copula Type-Model for Examining the Role of Microbiome as a Potential Tool in Diagnosis," *Mathematical Problems in Engineering* 2022 (2022): 8033806.
40. E. C. Brechmann and C. Czado, "COPAR—Multivariate Time Series Modeling Using the Copula Autoregressive Model," *Applied Stochastic Models in Business and Industry* 31 (2015): 495–514.
41. N. T. Hung, "Green Bonds and Asset Classes: New Evidence From Time-Varying Copula and Transfer Entropy Models," *Global Business Review* 26 (2025): 1405–1424.
42. H. Liu, X. Zhong, Z. Jiang, and S. Lai, "Dynamic Correlation Between Hog Futures and Industry Chain: An Empirical Study Based on

Time-Varying Copula Model," *International Journal of Systems Assurance Engineering and Management* 15 (2024): 5356–5366.

43. H. Manner and O. Reznikova, "A Survey on Time-Varying Copulas: Specification, Simulations, and Application," *Econometric Reviews* 31 (2012): 654–687.

Supporting Information

Additional supporting information can be found online in the Supporting Information section. **Data S1:** psp470242-sup-0001-DataS1.docx.