



Universiteit
Leiden
The Netherlands

Virtual patient populations for quantitative pharmacology with its application in antibiotic treatment optimization

Guo, Y.

Citation

Guo, Y. (2026, September 11). *Virtual patient populations for quantitative pharmacology with its application in antibiotic treatment optimization*.

Retrieved from <https://hdl.handle.net/1887/4309451>

Version: Publisher's Version

License: [Licence agreement concerning inclusion of doctoral thesis in the Institutional Repository of the University of Leiden](#)

Downloaded from: <https://hdl.handle.net/1887/4309451>

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



Chapter 3

Generation of realistic virtual adult populations using a model-based copula approach

Authors

Yuchen Guo

Tingjie Guo

Catherijne A. J. Knibbe

Laura B. Zwep

J. G. Coen van Hasselt

Journal of Pharmacokinetics and Pharmacodynamics, 2024 Dec;51(6):735-46.

Abstract

Incorporating realistic sets of patient-associated covariates, i.e., virtual populations, in pharmacometric simulation workflows is essential to obtain realistic model predictions. Current covariate simulation strategies often omit or simplify dependency structures between covariates. Copula models are multivariate distribution functions suitable to capture dependency structures between covariates with improved performance compared to standard approaches. We aimed to develop and evaluate a copula model for generation of adult virtual populations for 12 patient-associated covariates commonly used in pharmacometric simulations, using the publicly available NHANES database, including sex, race-ethnicity, body weight, albumin, and several biochemical variables related to organ function. A multivariate (vine) copula was constructed from bivariate relationships in a stepwise fashion. Covariate distributions were well captured for the overall and subgroup populations. Based on the developed copula model, a web application was developed. The developed copula model and associated web application can be used to generate realistic adult virtual populations, ultimately to support model-based clinical trial design or dose optimization strategies.

3.1 Introduction

In pharmacometric modeling, patients' covariates are usually identified as a source of variability between individuals that impacts pharmacokinetics and pharmacodynamics [1]. Generation of virtual populations (VPs), i.e., realistic sets of patient characteristics or covariates, is essential to ensure that realistic responses are produced in pharmacometric simulations, eventually providing valuable information to support in silico clinical trials and optimization of dosing strategies.

Realistic VPs should reflect both the marginal distribution and dependency structure observed between covariate variables of interest. In statistics, a marginal distribution describes the probability distribution of one separate variable, and a dependence structure reveals the relationships or dependent patterns between variables in a dataset. For instance, within a real-world dataset focusing on the elderly population, age, as a covariate, may exhibit a t -distributed margin, with a certain mean and standard deviation; meanwhile, it could be negatively correlated with renal function biomarkers. Misspecification of the margins or dependency structures, i.e., in comparison with those actually observed, may impact the quality of subsequent patient responses obtained in pharmacometric simulations.

VPs can be generated using several approaches, which are either data-driven or distribution-driven. Data-driven methodologies such as the bootstrap or conditional distribution modeling [2] utilize an actual dataset of patient characteristics to sample from. Requesting such data, however, is sometimes not possible due to patient privacy regulations. Distribution-based approaches characterize the distribution of the marginals of covariates of interest but may not always capture their dependency

structure. For example, series of univariate distributions can be used to describe the marginals yet ignore interdependencies between covariates. Multivariate normal distributions [3] do consider the dependency but assume that variables are normally distributed, which may not always hold. Finally, machine learning algorithms [4, 5, 6] have been proposed, but these models are usually based on complex frameworks and often lack interpretability of underlying dependencies.

Copula modeling is a powerful tool for calculating multivariate distributions and has been widely used in various fields, such as finance [7, 8, 9], climate research [10, 11] and engineering [12, 13]. Copulas can capture the dependence structure between random variables independently from the description of the marginals [14]. Using a transformation of any marginal distribution to a uniform distribution, the dependence structure can be separated from the marginal structure. Moreover, a rich variety of copula models is available to be selected to estimate diverse dependent patterns in data [15]. An extension of the copula, the vine copula, addresses the difficulty of calculating multivariate joint distributions by using conditional dependence and bivariate building blocks [16]. Recently, copulas have been introduced to the field of pharmacometrics as a relevant key strategy for VP generation, demonstrating favorable performance in simulating realistic VPs compared to standard approaches, while their distribution-based nature facilitates sharing of covariate data within the community [17].

Here, we present a copula model for the simulation of adult virtual populations. We first developed a copula model for 12 covariates of relevance for pharmacometric models using data from adult individuals present in the NHANES database [18]. Then we evaluated the performance of the copula in simulating the overall and subgroup populations. Finally, a web application was designed for the copula model developed to facilitate generation of adult VPs.

3.2 Methods

3.2.1 Data

We used the public database from National Health and Nutrition Examination Survey (NHANES), an initiative that collects data on non-institutionalized individuals in the U.S., including laboratory measurements, physical screening, and surveys; data are released to the public every two years [18]. We combined the NHANES data for 2009~2010, 2011~2012, 2013~2014, 2015~2016, and 2017~2018 releases based on their accessibility and consistency in laboratory methods. Differences in laboratory, instruments, and methods across releases were considered by implementing the adjustment equations provided by NHANES.

We focused on the adult population aged 18-80 years, with 27,008 subjects in total. Common covariates of interest for population pharmacokinetic models were selected: sex, race-ethnicity, age, height, body weight, fat mass (Fat), serum creatinine (SCR), alanine aminotransferase (ALT), aspartate aminotransferase (AST), alkaline phosphatase (ALP), albumin and total bilirubin (BR) [19, 20, 21, 22, 23]. We

acknowledge the sensitivity regarding the use of race-ethnicity as a medical indicator. Its inclusion in our study is focused on subgroup analysis when relevant, and not intended to perpetuate stereotypes or contribute to health disparities.

Table 3.1 provides the summary statistics of covariates in the model development dataset. Of note, over 50% of fat mass data in the observed dataset were missing. Fat mass is measured via dual-energy x-ray absorptiometry (DXA) examination. Half of the missing data were due to not meeting the inclusion criteria of age (< 60 years old) during the DXA examination, while another half were due to the examination not being conducted in the 2009 2010 release. Since copulas allow to be estimated from incomplete datasets and generate complete simulation datasets, a validation analysis was conducted (**Supplementary material**) to provide more insights into the reliability of simulated fat mass for people aged ≥ 60 years. This analysis involved validating fat mass predictions after excluding measured fat mass data for individuals within specific age groups. The result showed no significant bias in the simulated fat mass for the age groups under 60 years old (Figure S3.1).

3.2.2 Vine copula model development

A vine copula was fitted to the NHANES data. First, to avoid producing covariates of negative values in VPs, biochemical measurement data were log-transformed. As copulas are joint distribution functions with uniform margins, data were then transformed into uniform distributions using the probability integral function [24] based on kernel density estimation. Candidate vine copula models consist of parametric bivariate copula functions, such as Gaussian, Clayton and Frank copulas. Each kind of bivariate copula possesses distinct strengths in depicting various dependence behaviors, and a rich variety of bivariate copulas are available to be selected to represent diverse dependence patterns in data.

The vine copula model was constructed with a tree structure, which defines the pairs of covariates and copulas to be estimated. A vine tree structure comprises a sequence of trees, with the first tree representing a group of unconditional bivariate copulas, and subsequent trees representing a group of bivariate copulas conditional on the previous trees. Each edge in a tree represents a bivariate copula between two covariates (the first tree) or two copulas (higher order trees).

The vine tree structure was sequentially selected and estimated. For the first tree, the maximum spanning tree (MST) algorithm was used to select covariate pairs in each tree by maximizing the sum of correlations over the possible pairs in each tree, then all bivariate copula functions (Gaussian, Clayton or other bivariate distribution functions) were fit for the selected pairs and parameters were estimated; the best-fitting bivariate functions were selected based on the Akaike Information Criterion (AIC). This procedure was iterated for each subsequent tree until all trees were selected and estimated. Detailed methodology was described in the literature on bivariate copulas ([25] C3), tree structures ([25] C5), and MST [26].

To incorporate the covariate 'race-ethnicity' in the copula and optimize the model, we treated race-ethnicity as an ordered categorical variable and tested copulas with all order combinations.

3.2.3 Model Evaluation

Model evaluation was conducted through a simulation-based strategy: performing 100 simulations of the original dataset and comparing the metrics between the real-world population and VPs that were back transformed to their original scales. To assess the model performance on the marginal distributions, we evaluated observed and simulated populations by comparing the frequency of each category for categorical covariates, and for continuous covariates, comparing the marginal metrics, mean, standard deviation (SD), and percentiles (5th, 50th and 95th), denoted by M , between observed and simulated data in terms of relative error (RE) (Eq.1).

$$RE = \frac{M_{sim} - M_{obs}}{M_{obs}} \quad (\text{Eq. 1})$$

where M_{sim} and M_{obs} represent the metrics for simulated population and observed population, respectively.

To assess the performance of the model on capturing the dependency structure, pairwise correlation coefficients were compared between observed and simulated datasets. Since data sharing the same correlation could display various shapes of the dependence, a two-dimensional metric was developed to quantify the overlap of the density contours in observed and simulated data. For each pair combination of covariates, 95th percentile density contours were calculated for observed and simulated populations. The overlap metric was computed as the Jaccard index [27]: the ratio between the intersection area and union area (Figure S3.2). Higher overlap indicated a better description of dependence relations. We systematically evaluated the performance of the model from the following aspects:

(1) Overall performance: the NHANES copula (full copula) was developed based on the whole set of participants of NHANES that represents a general population. Simulated populations and the real-world population were then compared.

(2) Subgroup performance: populations of interest in clinical trials and cohort studies typically comprise individuals with certain race-ethnicity or sex. To be able to create realistic VP of interest, it is important to determine whether the full copula could capture the characteristics of subgroup populations. Predictive performance of the full copula for subsets of VPs was assessed with a particular interest in the race-ethnicity and sex subgroups. For comparison, two series of subgroup copulas were also constructed using data specific of each subgroup population: 1) Hispanic copula, White copula, African American copula, Asian copula, Other race copula, 2) male copula, female copula. Virtual subgroup populations were obtained in two ways: by simulating from the full copula model and filtering out the irrelevant individuals, and by directly simulating from the subgroup copula. The performance of full copula was compared with that of subgroup copula to provide an understanding of whether the full copula was sufficient for generating subsets of VPs.

Table 3.1: Summary statistics of covariates in dataset combined from National Health and Nutrition Examination Surveys 2009–2010, 2011–2012, 2013–2014, 2015–2016 and 2017–2018. The total number of individuals was $n = 27,008$.

Variable name	Variable description	Percentage (%)	Actual N (% Missing)	Mean $\pm$ SD [range]
Sex	Gender	-	27,008 (0%)	-
	Male	48.5	13,104	-
Race-Ethnicity	Female	51.5	13,904	-
	Race-Ethnicity	-	27,008 (0%)	-
Race-Ethnicity	Hispanic*	26.6	7,176	-
	White	36.9	9,978	-
	African American	22.8	6,166	-
	Asian	10.6	2,859	-
	Other Race	3.1	829	-
	Age (year)	-	27,008 (0%)	46.05 $\pm$ 17.21 [18–79]
	Body weight (kg)	-	26,746 (1.0%)	82.02 $\pm$ 22.17 [32.3–242.6]
Weight	Standing height (cm)	-	26,757 (1.0%)	167.16 $\pm$ 10.09 [123.3–204.5]
Height	Total body fat (kg)	-	11,826 (56.2%)	27.11 $\pm$ 11.93 [4.9–102.3]
Fat	Serum creatinine (mg/dL)	-	25,313 (6.3%)	0.88 $\pm$ 0.45 [0.16–17.41]
SCR	Alanine aminotransferase (ALT, U/L)	-	25,307 (6.3%)	25.57 $\pm$ 20.42 [5–1363]
ALT	Aspartate aminotransferase (AST, U/L)	-	25,288 (6.4%)	25.93 $\pm$ 17.13 [7–882]
AST	Alkaline phosphatase (ALP, U/L)	-	25,310 (6.3%)	69.34 $\pm$ 24.59 [7–907]
ALP	Albumin (g/dL)	-	25,315 (6.3%)	4.28 $\pm$ 0.35 [2.0–5.6]
Albumin	Total bilirubin (mg/dL)	-	25,298 (6.3%)	0.62 $\pm$ 0.31 [0.0–7.3]
BR				

* Mexican American and other Hispanic in NHANES were recorded as Hispanic in our real-world dataset. Other race-ethnicity groups remained unchanged.

3.2.4 Shiny Application Development

To provide a convenient and user-friendly tool, an interactive web application that could output VPs was developed using the NHANES copula. Next to the NHANES copula, a weighted copula was estimated with the incorporation of sampling weights [28] to address the sampling bias in NHANES. The sample weights account for complex sampling design and non-response of NHANES and are associated with demographic properties of the US population. The weighted copula allows users to sample a virtual population that is representative of the actual US population.

3.2.5 Software

The analysis was performed in R 4.1.2. Processing of NHANES data was conducted with survey package. Kernel density estimation of marginal distributions was performed with *kde1d* package. Development of NHANES copula was implemented with *rvinecopulib* package. The overlap metric was calculated using *ks* and *sf* packages. R shiny application was developed using *shiny* package. Visualizations of this study were generated with *ggplot2* package. All scripts are available on https://github.com/vanhasseltlab/NHANES_copula.

3.3 Results

3.3.1 Vine copula of NHANES data

Logarithmic and uniform transformed data were fitted to estimate the underlying dependency structure with a vine copula. Instead of displaying the whole tree structure, we only showed the first tree since the first layer dependence captured the strongest correlations while trees of higher levels describe the conditional dependence, and are less influential on the overall fit than the first tree [29]. Sex was located at the center of the first tree structure, as it showed relatively strong dependence relationships with height, logBR, logALT, logAlbumin, and logSCR (Figure 3.1 A). The density contours of covariate pairs in the real-world population displayed various patterns, and the VP was found to overlap the real-world population in selected covariate pairs well (Figure 3.1 B).

Estimating the NHANES copula using the input dataset comprising 27,008 records with 12 covariates required approximately 20 minutes on a Windows computer with Intel Core i7 processor operating at 2.80 GHz. In contrast, simulations from copula are computationally efficient, with an average of 1.6 seconds per 1000 individuals simulated.

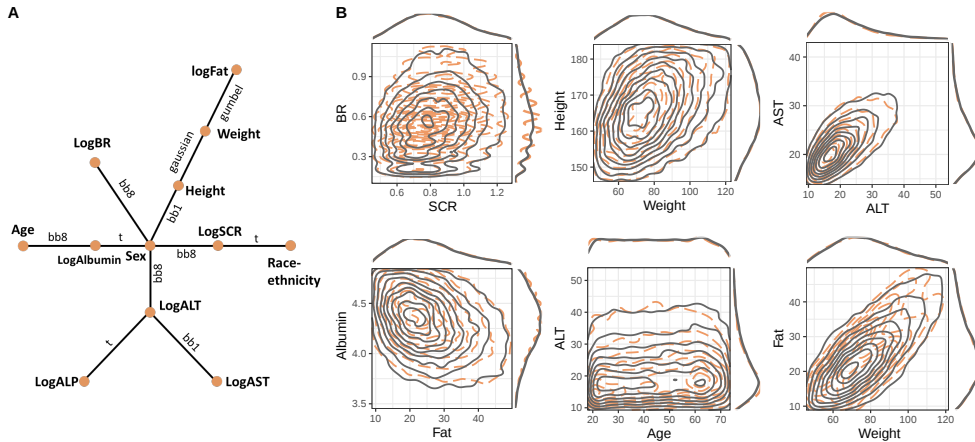


Figure 3.1: Graphical representation of dependence structure estimated by NHANES copula and bivariate densities of observed and simulated covariates. A. The first tree structure of NHANES copula with 12 covariates (nodes), and 11 copulas (edges). The texts on the edges denote the selected bivariate distribution functions: Gaussian, Gumbel, t, Clayton-Gumbel (bb1) and Joe-Frank (bb8). The tree was chosen based on the maximum spanning tree algorithm. The inter-dependent relationships between other covariate pairs were captured by subsequent trees that describe the conditional dependency. B. Density contours of different covariate pairs of the observed population (orange dashed line) and the simulated population using the NHANES copula (gray solid lines). Marginal densities were displayed on the top and right sides of each plot. Six covariate pairs were chosen to present the diverse dependent patterns as illustrative examples. Abbreviations: LogFat: log fat mass; LogSCR: log serum creatinine concentration (SCR); LogALT: log alanine aminotransferase concentration (ALT); LogAST: log aspartate aminotransferase concentration (AST); LogALP: log alkaline phosphatase concentration (ALP); LogAlbumin: log albumin concentration; LogBR: log total bilirubin concentration (BR).

3.3.2 Overall performance

The overall simulation performance of the developed copula model was evaluated for the entire population, without specifying any subgroups. For categorical covariates, (i.e., race-ethnicity and sex), the frequency of each category in the virtual population aligned with that of the real-world population (Figure 3.2 A). For continuous covariates, density curves of each individual covariate in the simulation dataset well tracked observed ones (Figure 3.2 B); mean, standard deviation and percentiles of VP agreed with those of the observed population, with relative errors within ± 0.10 (Figure 3.2 C). For percentiles and mean metrics, coefficient of variation across simulations were all within 0.007, and those of standard deviation were within 0.09.

The simulated correlations from the copula model were very similar to observed correlations for most pair combinations of covariates, with 0.023 median error (Figure 3.3 A). Covariate pairs associated with the largest error of correlation were height-SCR (0.105) and SCR-albumin (0.102). The median overlap was 92.0% across all

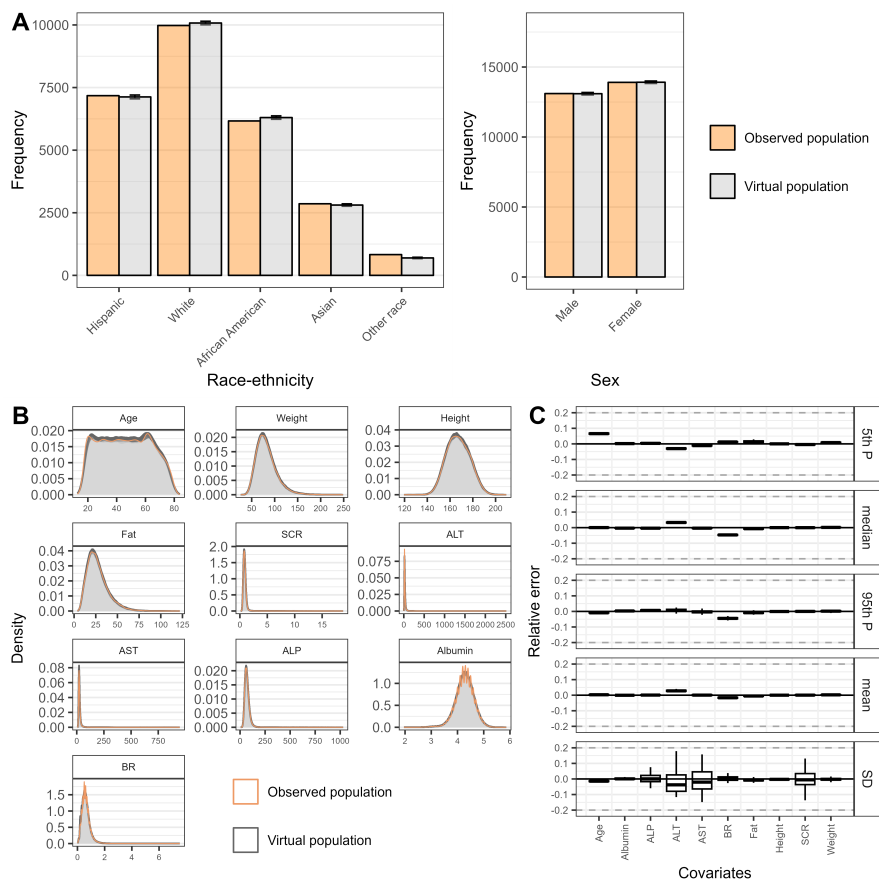


Figure 3.2: Marginal characteristics of the covariates in observed population and virtual populations simulated from NHANES copula. A. Frequency of each category in discrete covariates, race-ethnicity and sex, in the real-world population (orange columns) and virtual populations (grey columns). B. Density curves of each continuous covariate variable in the real-world population (orange line) and virtual populations (gray lines). C. Relative error of marginal metrics (percentiles, mean, and standard deviation) of continuous covariates as compared to the statistics of the real-world population. Virtual population was simulated 100 times. Error bars indicated the standard deviation of 100 simulations. Gray dashed lines indicate $\pm 20\%$ relative error.

covariate pairs and simulations, and the model achieved over 85% overlap in 96% (43/45) covariate pairs, indicating a good capture of dependency structure (Figure 3.3 B). The only two covariate pairs that did not reach 85% were ALT-BR and weight-fat, with 81.4% and 70.5% overlap percentages.

The full copula model reproduced the marginal properties as well as the dependence relations of covariates of input population data. Variability across simulations tended to be small for all metrics except for standard deviation, showing the robustness of the copula model.

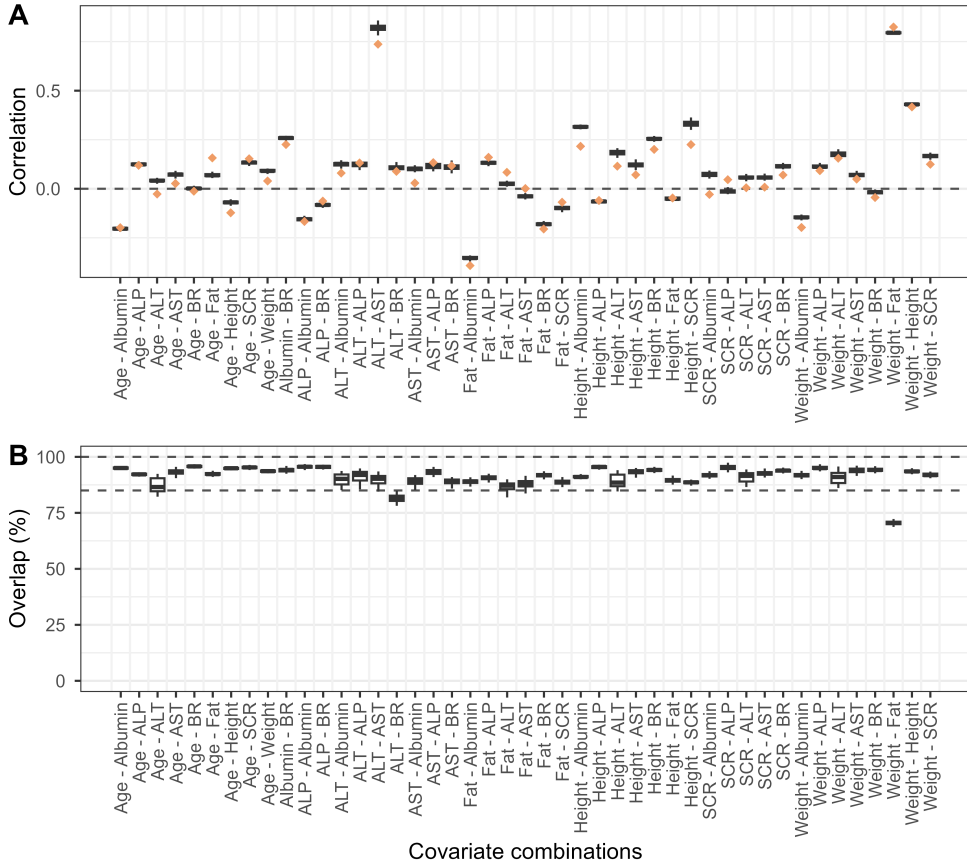


Figure 3.3: Dependency metrics of covariate pairs in the observed population and virtual populations simulated from NHANES copula. A. Correlations of each covariate pair in real-world population (orange diamond) and virtual populations (black box). Gray dashed line represents no correlation between covariate pairs. B. Overlap metric of 95th density contours of virtual population relative to observed population. Virtual population was simulated 100 times. Error bars indicated the standard deviation of 100 simulations. Gray dashed lines indicate 100% and 85% overlap percentages.

3.3.3 Subgroup performance

To gain further insights into the usefulness of the full copula for simulating subgroups of the total population, we conducted two separate investigations on the performance of full copula for VP simulation in race-ethnicity and sex subgroups.

3.3.3.1 Race-ethnicity subgroup analysis

The full copula was able to approximate the marginal characteristics of the observed population in Hispanic, White, and African American subgroups, with median relative errors of marginal metrics across covariates within $[-0.19, 0.28]$ (Figure S3.3). For Asian and Other race VP populations, median relative errors were in the ranges $[-0.21, 0.41]$ and $[-0.68, 0.20]$, respectively. For comparison, subgroup copulas for Hispanic, White, African American and Asian populations showed good performances in terms of the marginal metrics, with the median relative errors of all covariates in $[-0.10, 0.16]$. However, the relative errors were larger for Other race subgroup copula, with a range of $[-0.46, 0.09]$. Full copula and subgroup copulas showed comparable performance in capturing the marginal attributes of Hispanic, White, African American and Other race subgroups, however, subgroup copula showed superior performance in Asian population.

The full copula model achieved 84.6%, 88.6%, 87.8%, 74.0% and 80.1% median overlap percentages for Hispanic, White, African American, Asian and Other race populations, while subgroup copulas reached 89.7%, 91.0%, 89.0%, 88.7% and 85.1% (Figure 3.4 A), respectively. Subgroup copulas outperformed the full copula in simulating the dependence structure of covariates in Asian and Other race subgroups, but showed similar performance in the rest of race-ethnicity subgroups.

3.3.3.2 Sex subgroup analysis

In general, compared with subgroup copulas, the full copula model could well capture the margins and dependency structures in male and female populations. For marginal metrics, median relative errors of full copula were within the range $[-0.19, 0.09]$ and $[-0.28, 0.07]$ for male and female populations (Figure S3.4). For comparison, subgroup copulas for male and female yielded median relative errors of $[-0.33, 0.06]$ and $[-0.07, 0.08]$. The median overlap metric of full copula was calculated to be 88.5% and 88.8% for male and female populations (Figure 3.4 B), while subgroup copulas achieved 91.9% and 92.1% overlap percentages for the two populations.

3.3.4 R shiny application

The copula covariate simulator (CoCoSim) web application was developed based on the NHANES copula and made available online (<https://cocosim.lacdr.leidenuniv.nl/>, Figure 3.5). Using this application, VPs can be generated online following these steps: (1) define the population of interest by selecting race-ethnicity, sex, age, and body mass index (BMI); (2) select the covariates of interest. Secondary covariates, including BMI, lean body weight, and estimated glomerular filtration rate, can be calculated based on the covariates in NHANES dataset; (3) select the number of individuals for simulation; (4) select the weighted or unweighted NHANES copula for the virtual population simulation; (5) generate the VP and download the data.

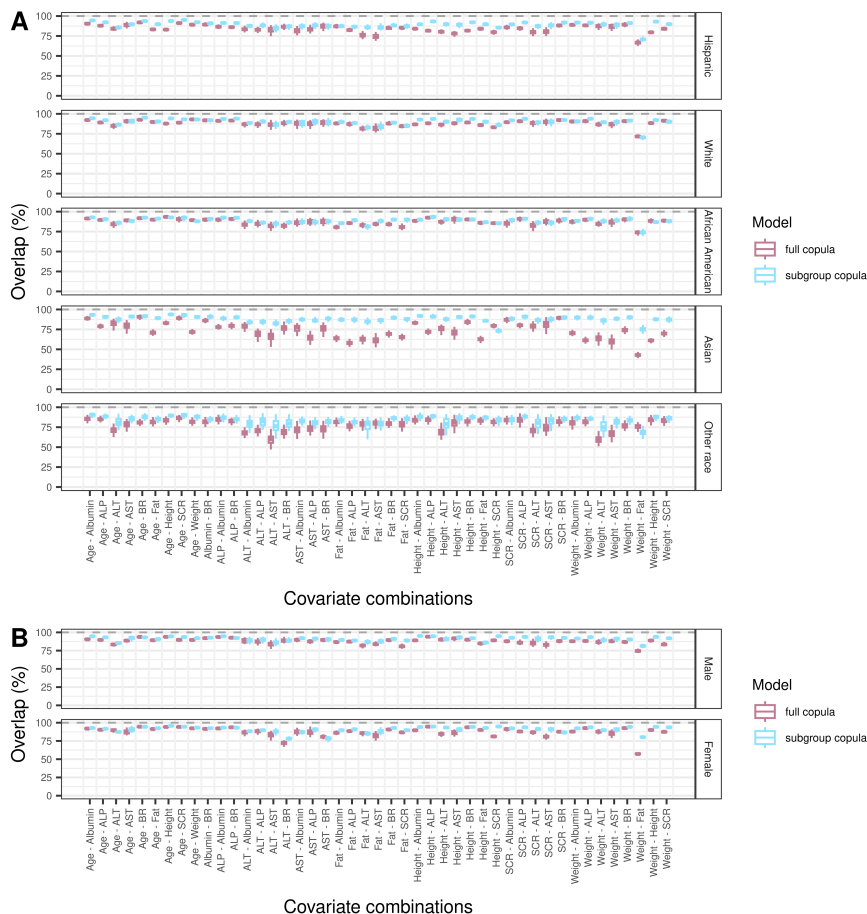


Figure 3.4: Overlap metric of each subgroup virtual population relative to corresponding real-world population. A. Overlap metrics calculated for each race-ethnicity subgroup population. B. Overlap metrics calculated for each sex subgroup population. The full copula was created utilizing the whole set of data, while subgroup copulas were developed based on each subgroup of data. Subgroup virtual populations were simulated 100 times using full copula (pink boxes) and subgroup copulas (blue boxes) for each. Error bars indicated the standard deviation of 100 simulations. Gray dashed lines indicate 100% overlap percentage.

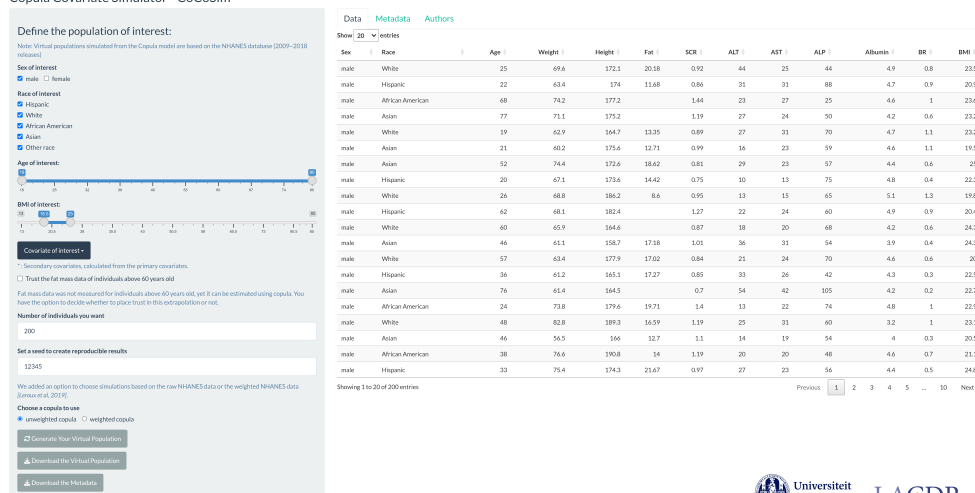


Figure 3.5: Interface of the R shiny application CoCoSim (<https://cocosim.lacdr.leidenuniv.nl/>). Virtual population can be generated according to user-defined characteristics based on NHANES copula (full copula).

With the app, users can generate virtual population with desired characteristics, including race-ethnicity, sex, age and BMI ranges. Generated virtual populations can then be used as covariate distributions for pharmacometric model-based simulations, for example as part of clinical trial simulations or dosing strategy optimization simulations.

3.4 Discussion

We developed a copula model for an adult population which adequately captured the covariate distributions as present in the NHANES database.

The tree structure of the NHANES copula revealed associations between commonly used covariates in population pharmacokinetics studies, which may help in the process of covariate model development. Identified associations were in line with the literature in which sex was found to influence height, weight, serum creatinine, and liver function biomarkers (total bilirubin and ALT) [30]. The correlation between covariates may explain the situations where sex may not be relevant as a covariate when the other covariates are included, since different PK or PD outcomes depend on underlying covariates (such as weight and serum creatinine) [31, 32].

To evaluate the performance of the developed copula model, we assessed whether the simulated population is realistic by comparing the marginal and dependency metrics between VPs and real-world populations. Interestingly, we observed that the pair combinations of covariates that showed the largest errors of correlation differed from those showing the lowest overlap percentages. Pearson correlation quantifies

linear association, while data sharing the same linear correlation could exhibit different dependency structures, and the overlap metric takes the shape or pattern of the dependency into account. Jaccard index is a similarity measure between two data samples [27], and the novelty of the overlap metric lies in its first application to two-dimensional densities. Pearson correlation and overlap metric collectively depicted the joint behavior at a pairwise level and addressed different perspectives, and as such should be evaluated together when assessing copulas or investigating the similarity between two populations.

The advantage of copulas is that they can model complex multivariate distributions more easily and efficiently. The local poor fit for dependency structure between fat-weight conditional on sex (Figure 3.4 B) is probably due to the parametric bivariate copulas used for the estimation of vine copulas. The fat-weight distribution of the real-world population exhibited a heart-shaped contour (Figure 3.1 B), which might be better captured by non-parametric bivariate copulas. Yet, deploying non-parametric copulas is computationally expensive and is prone to overfitting [33].

This study is focused on the adult. The pediatric population was not considered in this analysis mainly because some critical covariates for the pediatric population are lacking in NHANES database, such as birth weight, postnatal age and gestational age. Additionally, the pediatric population differs from the adult population in anatomical, physiological and biochemical characteristics [34], and the developmental changes over age may lead to drastic changes in the dependency structure between covariates. This could lead to inferior performance if the copula was estimated on both populations. We have thus chosen to focus on adults only.

In this study, we incorporated not only continuous but also categorical variables in the estimation of the NHANES copula. Currently, copula models for unordered categorical variables are not fully identifiable [24]. To include race-ethnicity (an unordered categorical variable) in copula, we estimated vine copulas by iterating through all possible orders of race-ethnicity and selected the model with the lowest AIC value. Since there were five categories in race-ethnicity, we considered 120 unique order possibilities of race-ethnicity categories, which was time-consuming and computationally expensive. Since this type of variable is common in clinical studies, such as disease classification, an algorithm that could efficiently deal with unordered categorical covariates is yet to be developed. When categorical variables are transformed to a uniform scale, each value does not have the same probability of occurring, and it is still inherently discrete. Instead of calculating correlations, we chose to perform a subgroup analysis for categorical variables.

Copula models can be useful to support model-based dosing optimization or clinical trial simulation. For such applications, a focus on subjects with specific covariate characteristics usually exists [35, 36]. To this end, it is important to confirm whether a copula model correctly reflects covariate distributions for relevant population subgroups of interest. In our analysis, compared with subgroup copulas, the full copula model showed comparable performance across different race-ethnicity and sex subgroups except for Asian and Other race subgroups, likely due to the relatively small number of individuals within the entire dataset. In particular, the Asian population

has distinct marginal distributions of weight, height and logFat, compared to other race-ethnicity populations (Figure S3.5). This led to a stronger correlation between height and weight, the relation between which predominates the underlying dependency structure (Figure S3.6). The ability to adequately simulate subgroups from a large copula is of great importance since creating copulas for each subpopulation of interest, including e.g. different age and BMI ranges creates a nearly infinite amount of possible subgroups.

The NHANES population represents the non-institutionalized population of America and cannot be classified as healthy subjects or patients, indicating that the virtual population simulated from full copula should be interpreted with care. In this dataset, a significant portion of fat mass data was missing due to the age-eligible criterion (< 60 years old) of examination. However, copulas allowed for interpolation and extrapolation of VPs, as they support the generation of fat mass data for individuals above 60 years old via conditional density functions [37]. Of note, we removed the extrapolated fat mass data during the evaluation of copula performance. Although no significant bias was revealed in the validation analysis, simulated fat mass for people above 60 years old should be used with caution.

To make the full copula more accessible to the community, a web application was developed to facilitate the simulation of VPs with user-defined properties. The application allows to generate virtual populations with specific demographic attributes such as race-ethnicity, sex and BMI. Yet, the performance of NHANES copula on special populations, has not been able to be validated in this study. Additional data related to these populations are necessary to further investigate copulas for specific types of patients, such as pediatric, obese, pregnant, and renally impaired patients. This work served as a basis for building a copula library in future, for sharing the copulas of special patient populations and supporting simulation studies. Collaborative efforts could be initiated to gather large-scale data to build copulas for various target populations.

A copula can generate virtual populations that accurately represent the input population, and allows for adjusting sampling weights in the estimation in case the input population is not representative of the real-world population. In this way, copulas facilitate the generation of virtual populations that are representative of the actual real-world populations if sampling weights are available. The marginal distribution of unweighted and weighted copula differed mainly in different race-ethnicity groups (Figure S3.7). Our web application includes both unweighted and weighted copulas. In the analysis, we focused on the unweighted copula, because the comparison between the virtual population and the input population is only possible with the unweighted copula.

3.5 Conclusion

In this study, we demonstrated the development and evaluation of a copula model using NHANES database to simulate commonly used covariates in pharmacometric

modeling, which can be used as part of clinical trial design and dose strategies optimization. A user-friendly web application was developed to facilitate the use of the developed copula model for covariate simulation.

References

1. Duffull S and Gulati A. Potential Issues With Virtual Populations When Applied to Nonlinear Quantitative Systems Pharmacology Models. *CPT: Pharmacometrics & Systems Pharmacology*. 2020; 9:613–6
2. Smania G and Jonsson EN. Conditional distribution modeling as an alternative method for covariates simulation: Comparison with joint multivariate normal and bootstrap techniques. *CPT: Pharmacometrics & Systems Pharmacology*. 2021; 10:330–9
3. Teutonico D et al. Generating Virtual Patients by Multivariate and Discrete Re-Sampling Techniques. *Pharmaceutical Research*. 2015 Oct; 32:3228–37
4. McComb M and Ramanathan M. Generalized Pharmacometric Modeling, a Novel Paradigm for Integrating Machine Learning Algorithms: A Case Study of Metabolomic Biomarkers. *Clinical Pharmacology and Therapeutics*. 2020 Jun; 107:1343–51
5. Nair R, Mohan DD, Setlur S, Govindaraju V, and Ramanathan M. Generative models for age, race/ethnicity, and disease state dependence of physiological determinants of drug dosing. *Journal of Pharmacokinetics and Pharmacodynamics*. 2023 Apr; 50:111–22
6. McComb M, Blair RH, Lysy M, and Ramanathan M. Machine learning-guided, big data-enabled, biomarker-based systems pharmacology: modeling the stochasticity of natural history and disease progression. *Journal of Pharmacokinetics and Pharmacodynamics*. 2022 Feb; 49:65–79
7. Brechmann EC, Hendrich K, and Czado C. Conditional copula simulation for systemic risk stress testing. *Insurance: Mathematics and Economics*. 2013 Nov; 53:722–32
8. De Lira Salvatierra I and Patton AJ. Dynamic copula models and high frequency data. *Journal of Empirical Finance*. 2015 Jan; 30:120–35
9. Arreola Hernandez J, Hammoudeh S, Nguyen DK, Al Janabi MAM, and Reboredo JC. Global financial crisis and dependence risk analysis of sector portfolios: a vine copula approach. *Applied Economics*. 2017 May; 49:2409–27
10. Wang W, Dong Z, Lall U, Dong N, and Yang M. Monthly Streamflow Simulation for the Headwater Catchment of the Yellow River Basin With a Hybrid Statistical-Dynamical Model. *Water Resources Research*. 2019; 55:7606–21
11. Schölzel C and Friederichs P. Multivariate non-normally distributed random variables in climate research – introduction to the copula approach. *Nonlinear Processes in Geophysics*. 2008 Oct; 15:761–72
12. Kilgore RT and Thompson DB. Estimating Joint Flow Probabilities at Stream Confluences by Using Copulas. *Transportation Research Record*. 2011 Jan; 2262:200–6
13. Kumar P. Copula Functions and Applications in Engineering. *Logistics, Supply Chain and Financial Predictive Analytics: Theory and Practices*. Ed. by Deep K, Jain M, and Salhi S. Singapore: Springer, 2019 :195–209
14. Czado C and Nagler T. Vine Copula Based Modeling. *Annual Review of Statistics and Its Application*. 2022 Mar; 9:453–77
15. Dewick PR and Liu S. Copula Modelling to Analyse Financial Data. *Journal of Risk and Financial Management*. 2022 Mar; 15:104
16. Aas K, Czado C, Frigessi A, and Bakken H. Pair-copula constructions of multiple dependence. *Insurance: Mathematics and Economics*. 2009 Apr; 44:182–98
17. Zwep LB, Guo T, Nagler T, Knibbe CA, Meulman JJ, and Hasselt JGC van. Virtual Patient Simulation Using Copula Modeling. *Clinical Pharmacology & Therapeutics*. 2024; 115:795–804
18. Department of Health and Human Services CfDcAP. Centers for Disease Control and Prevention (CDC). National Center for Health Statistics (NCHS). National Health and Nutrition Examination Survey Data.
19. Joerger M. Covariate Pharmacokinetic Model Building in Oncology and its Potential Clinical Relevance. *The AAPS Journal*. 2012 Mar; 14:119–32
20. Scarpignato C et al. A Population Pharmacokinetic Model of Vonoprazan: Evaluating the Effects of Race, Disease Status, and Other Covariates on Exposure. *The Journal of Clinical Pharmacology*. 2022; 62:801–11

21. Karlsson MO and Sheiner LB. The importance of modeling interoccasion variability in population pharmacokinetic analyses. *Journal of Pharmacokinetics and Biopharmaceutics*. 1993 Dec; 21:735–50
22. Morse JD, Stanescu I, Atkinson HC, and Anderson BJ. Population Pharmacokinetic Modelling of Acetaminophen and Ibuprofen: the Influence of Body Composition, Formulation and Feeding in Healthy Adult Volunteers. *European Journal of Drug Metabolism and Pharmacokinetics*. 2022 Jul; 47:497–507
23. Gupta A, Jarzab B, Capdevila J, Shumaker R, and Hussein Z. Population pharmacokinetic analysis of lenvatinib in healthy subjects and patients with cancer. *British Journal of Clinical Pharmacology*. 2016; 81:1124–33
24. Geenens G. Copula modeling for discrete random vectors. *Dependence Modeling*. 2020 Jan; 8:417–40
25. Czado, Claudia. Analyzing Dependent Data with Vine Copulas. 2019; 222
26. Czado C, Brechmann EC, and Gruber L. Selection of Vine Copulas. *Copulae in Mathematical and Quantitative Finance*. Ed. by Jaworski P, Durante F, and Härdle WK. Lecture Notes in Statistics. Berlin, Heidelberg: Springer, 2013 :17–37
27. Jaccard P. The Distribution of the Flora in the Alpine Zone.1. *New Phytologist*. 1912; 11:37–50
28. Leroux A et al. Organizing and Analyzing the Activity Data in NHANES. *Statistics in Biosciences*. 2019 Jul; 11:262–87
29. Kraus D and Czado C. Growing simplified vine copula trees: Improving Dißmann's algorithm. *arXiv preprint*. 2017 Mar
30. Jamei M, Dickinson GL, and Rostami-Hodjegan A. A framework for assessing inter-individual variability in pharmacokinetics using virtual human populations and integrating general knowledge of physical chemistry, biology, anatomy, physiology and genetics: A tale of 'bottom-up' vs 'top-down' recognition of covariates. *Drug Metabolism and Pharmacokinetics*. 2009; 24:53–75
31. Gandhi M, Aweeka F, Greenblatt RM, and Blaschke TF. Sex Differences in Pharmacokinetics and Pharmacodynamics. *Annual Review of Pharmacology and Toxicology*. 2004; 44:499–523
32. Wilson K. Sex-Related Differences in Drug Disposition in Man. *Clinical Pharmacokinetics*. 1984 Jun; 9:189–202
33. Nagler T, Schellhase C, and Czado C. Nonparametric estimation of simplified vine copula models: Comparison of methods. *Dependence Modeling*. 2017 Jan; 5:99–120
34. Fernandez E, Perez R, Hernandez A, Tejada P, Arteta M, and Ramos JT. Factors and Mechanisms for Pharmacokinetic Differences between Pediatric Population and Adults. *Pharmaceutics*. 2011 Mar; 3:53–72
35. Perez-Ruixo JJ, Piotrovskij V, Zhang S, Hayes S, Porre PD, and Zannikos P. Population pharmacokinetics of tipifarnib in healthy subjects and adult cancer patients. *British Journal of Clinical Pharmacology*. 2006; 62:81–96
36. Schaefer C, Cawello W, Waitzinger J, and Elshoff JP. Effect of Age and Sex on Lacosamide Pharmacokinetics in Healthy Adult Subjects and Adults with Focal Epilepsy. *Clinical Drug Investigation*. 2015 Apr; 35:255–65
37. Hollenbach FM, Bojinov I, Minhas S, Metternich NW, Ward MD, and Volfovsky A. Multiple Imputation Using Gaussian Copulas. *Sociological Methods & Research*. 2021 Aug; 50:1259–83

Supplementary material

Validation analysis regarding missing fat mass data

Method

In our real-world dataset, a significant portion of fat mass data was found to be missing due to the age-eligible criterion (< 60 years old) during the dual-energy x-ray absorptiometry examination. However, copulas allowed us to obtain a complete dataset for simulated VPs as it leveraged conditional density functions.

To provide insight into the reliability of simulated fat data for people above 60 years old, we applied the same missing data mechanism on the real-world dataset and investigated the performance of imputation of fat mass data for people with unmeasured fat mass data. We first extracted a complete dataset (N=11,179) from the NHANES dataset. Based on this dataset, fat mass data for individuals aged between specific cut-off ages (47, 38, and 34) and 60 years old were removed to generate incomplete datasets with 30%, 50%, and 60% missing fat mass data.

Four vine copulas were developed based on the complete and incomplete datasets. Through 100 simulations, we assessed the extrapolation ability in generating the fat data for individuals lacking fat observation. For each level of missing fat data, we compared the marginal metrics of the fat data for people aged between cut-off ages and 60 years old in the real-world population, VPs simulated from complete-data-based copula, and VPs simulated from incomplete-data-based copula.

Result

Compared with observed fat mass data, fat mass simulated from the complete-data-based and incomplete-data-based copula models had most of marginal metrics within ± 0.20 relative error (Figure S3.1), demonstrating the reliability of simulated fat mass data and the extrapolation ability of the copula model. With the increase of missing proportion (30%, 50% and 60%), incomplete-data-based simulations showed a tendency of underestimation, with -0.05 , -0.09 and -0.13 median relative errors for 5th percentile of fat mass, respectively; 50th percentile (0.03, 0.00 and -0.03) and mean (0.03, 0.00 and -0.03) exhibited a similar trend.

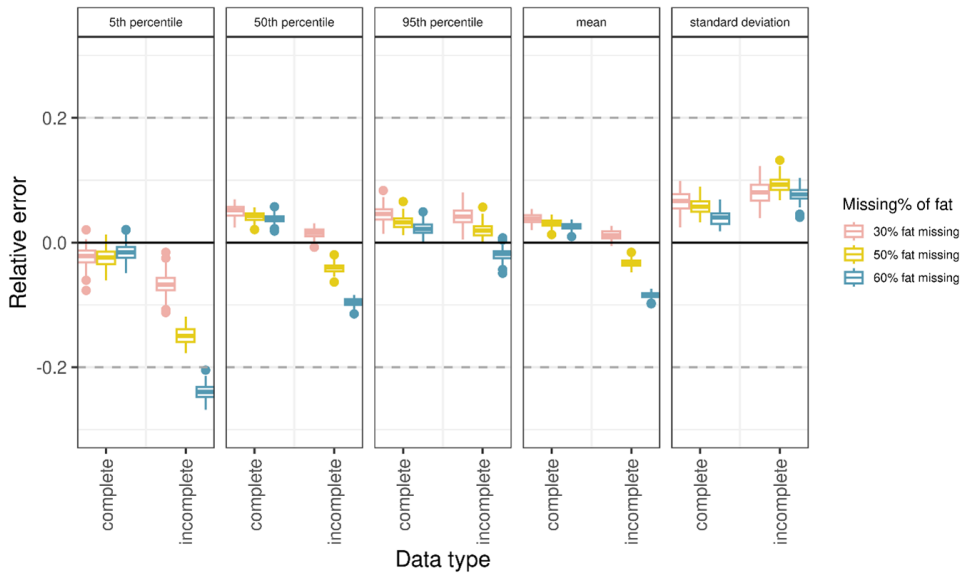


Figure S3.1: Relative error of marginal metrics of simulated fat mass data as compared to the statistics of observed fat mass data. Based on the complete dataset extracted from NHANES dataset, fat mass data for individuals aged between specific cut-off ages and 60 years old were removed to generate incomplete datasets with 30%, 50%, and 60% missing fat mass data. Each row represents the marginal metrics of fat data for people aged between a certain cut-off age and 60 years. Each subplot compared the performance of complete-data-based copula with incomplete-data-based copula. 30%, 50%, and 60% missing fat mass data corresponded to the cut-off ages of 47, 38, and 34 years old. Virtual population was simulated 100 times. Error bars indicated the standard deviation of 100 simulations.

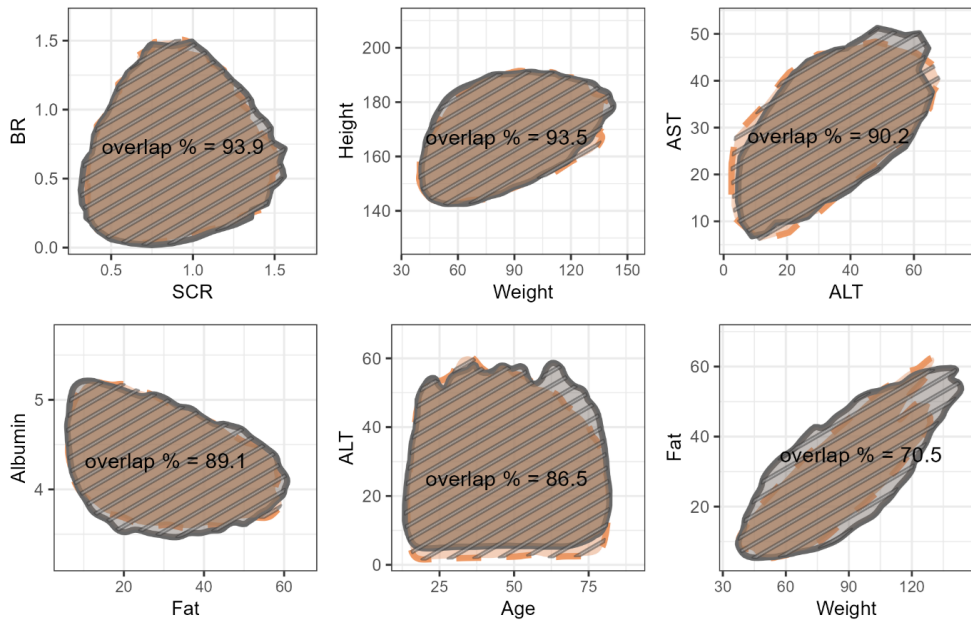


Figure S3.2: Illustration of overlap metric of 95th contours in the selected covariate pairs. Each subplot presents the 95th contour of the observation population (orange dashed line) and simulated population (gray solid line) from the NHANES copula. Brown area and striped area represent the intersection area and union area of these contours, respectively. Overlap percentage is calculated as the ratio of the intersection area over the union area. Six covariate pairs were chosen to present the diverse dependent patterns as illustrative examples. Abbreviations: Fat: fat mass; SCR: serum creatinine concentration (SCR); ALT: alanine aminotransferase concentration (ALT); AST: aspartate aminotransferase concentration (AST); ALP: alkaline phosphatase concentration (ALP); Albumin: albumin concentration; BR: total bilirubin concentration (BR).

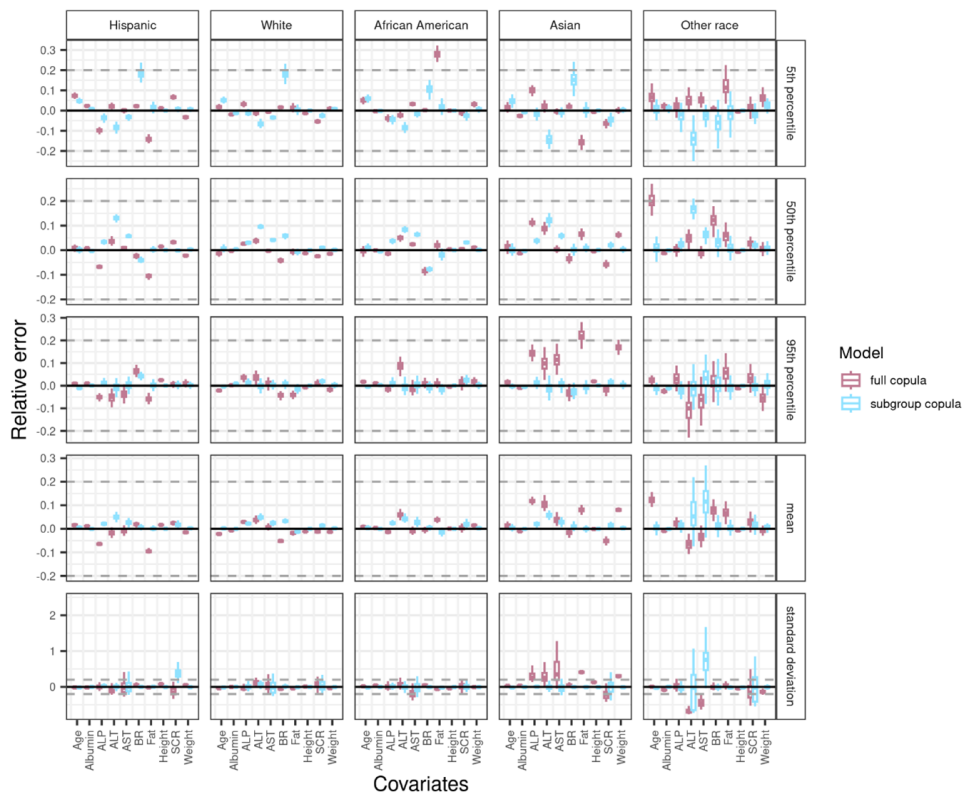


Figure S3.3: Relative error of marginal metrics in each race-ethnicity virtual subgroup population as compared to the statistics of corresponding real-world population. The full copula was created utilizing the whole set of data, while subgroup copulas were developed based on each subgroup of data. Virtual subgroup populations were simulated 100 times using full copula (pink boxes) and subgroup copulas (blue boxes) for each. Error bars indicated the standard deviation of 100 simulations. Gray dashed lines indicate $\pm 20\%$ relative error.

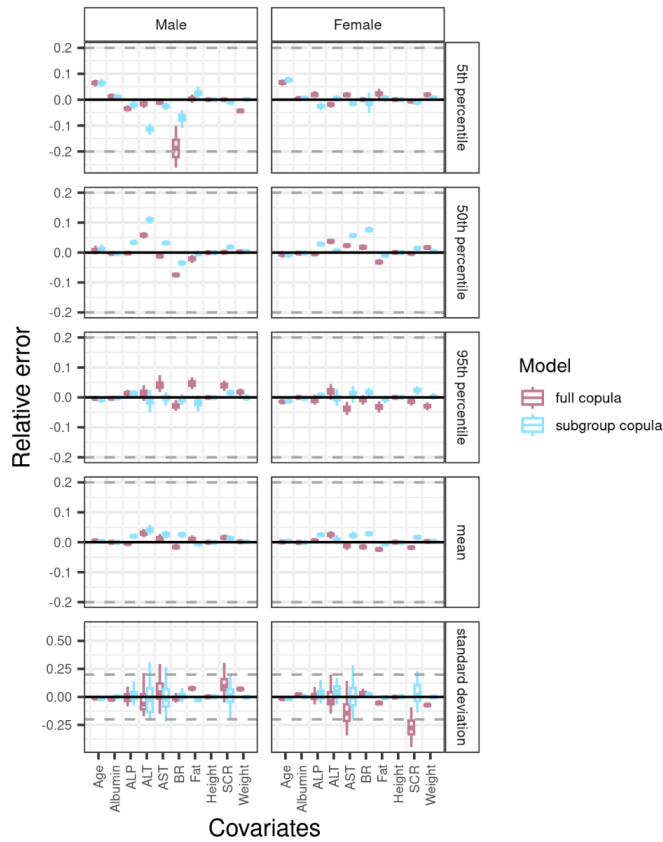


Figure S3.4: Relative error of marginal metrics in each sex virtual subgroup population as compared to the statistics of corresponding real-world population. Full copula was created utilizing the whole set of data, while subgroup copulas were developed based on each subgroup of data. Virtual subgroup populations were simulated 100 times using full copula (pink boxes) and subgroup copulas (blue boxes) for each. Error bars indicated the standard deviation of 100 simulations. Gray dashed lines indicate $\pm 20\%$ relative error.

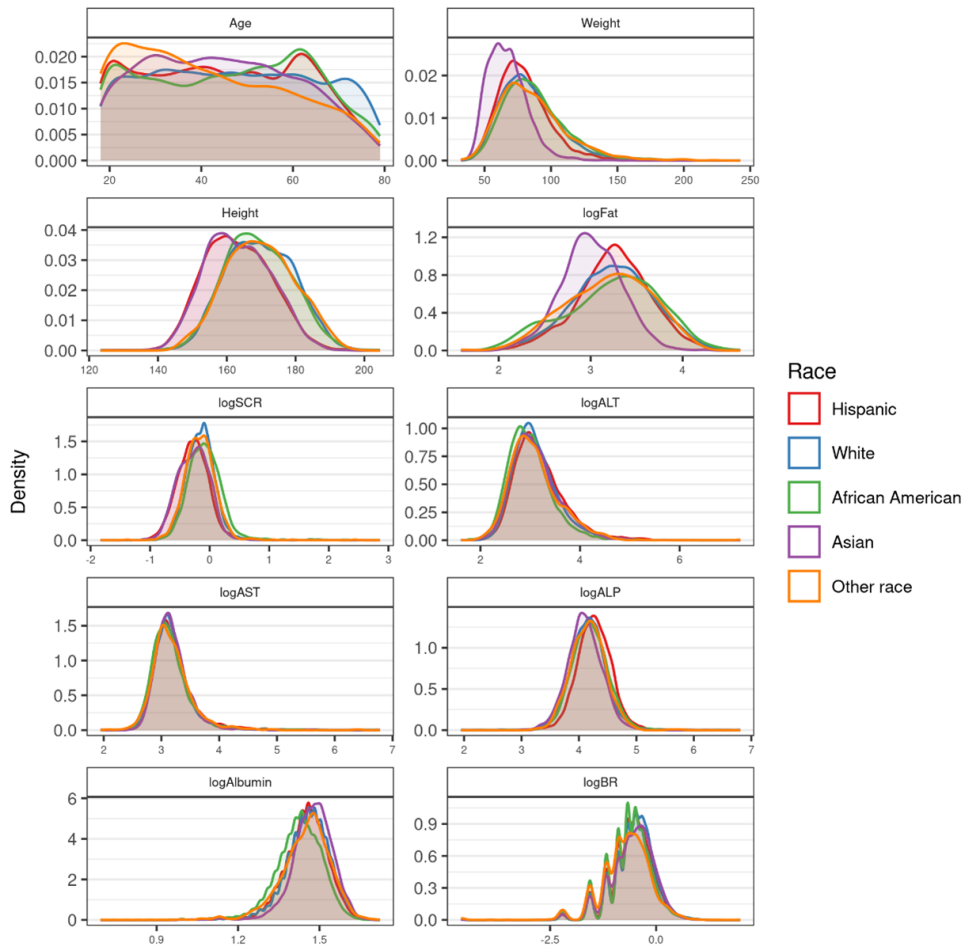


Figure S3.5: Marginal distributions of covariates in different race-ethnicity groups. Abbreviations: LogFat: log fat mass; LogSCR: log serum creatinine concentration (SCR); LogALT: log alanine aminotransferase concentration (ALT); LogAST: log aspartate aminotransferase concentration (AST); LogALP: log alkaline phosphatase concentration (ALP); LogAlbumin: log albumin concentration ; LogBR: log total bilirubin concentration (BR).

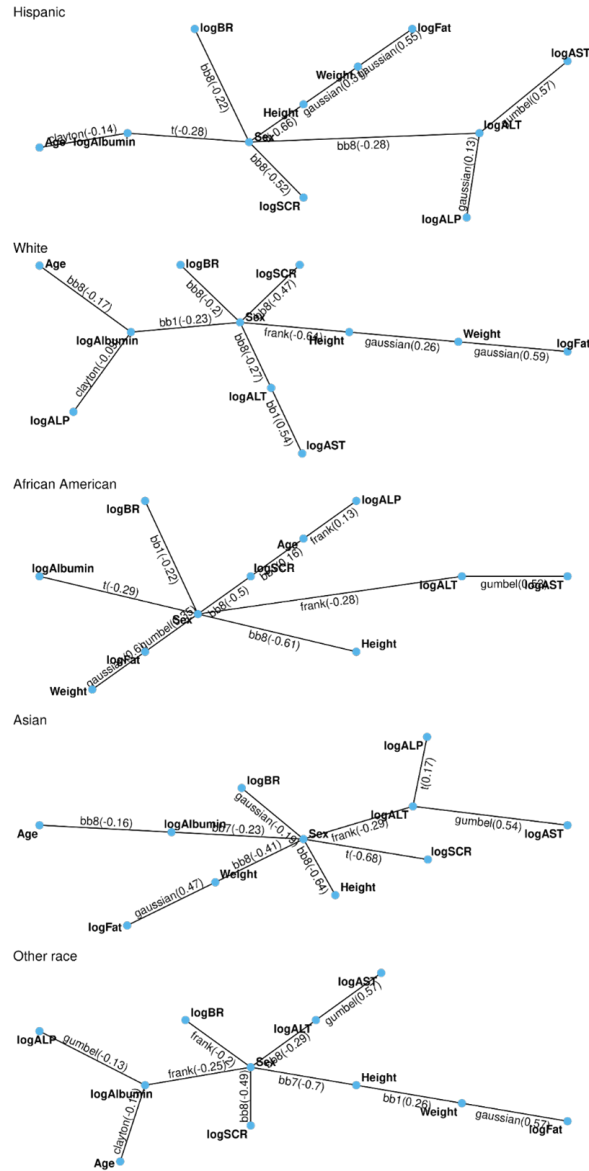


Figure S3.6: The first tree structure of race-ethnicity subgroup copulas with 11 covariates (nodes), and 10 copulas (edges). The edge texts denote the selected bivariate distribution functions: Gaussian, Gumbel, t , Clayton, Frank, Clayton-Gumbel (bb1), Joe-Clayton (bb7) and Joe-Frank (bb8). The tree was chosen based on the maximum spanning tree algorithm. The inter-dependent relationships between other covariate pairs were captured by subsequent trees that describe the conditional dependency. Abbreviations: LogFat: log fat mass; LogSCR: log serum creatinine concentration (SCR); LogALT: log alanine aminotransferase concentration (ALT); LogAST: log aspartate aminotransferase concentration (AST); LogALP: log alkaline phosphatase concentration (ALP); LogAlbumin: log albumin concentration; LogBR: log total bilirubin concentration (BR).

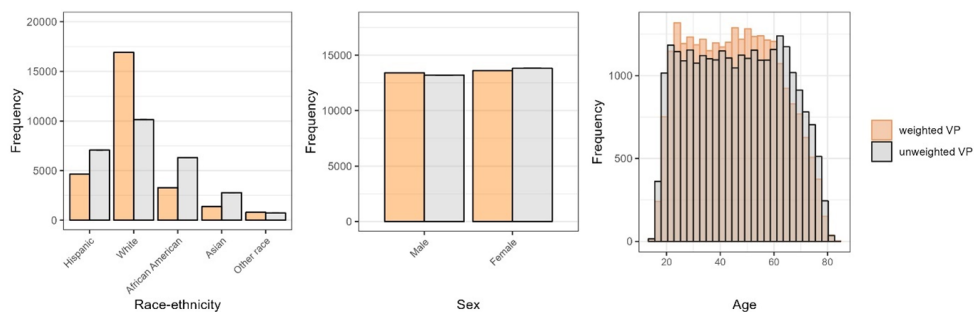


Figure S3.7: Comparison of demographic covariates between virtual populations generated from the weighted (orange) and the unweighted (grey) copulas. The same number ($N = 27,008$) of virtual individuals were simulated from each copula. The difference between these copulas is a result of the incorporation of the sampling weights provided in the NHANES database, which are calculated to correct for the over and under sampling of individuals with different demographic characteristics, such as race-ethnicity, sex and age.

