

PERSPECTIVE

## The potential synergies between synthetic data and *in silico* trials in relation to generating representative virtual population cohorts

To cite this article: Puja Myles *et al* 2023 *Prog. Biomed. Eng.* **5** 013001

View the [article online](#) for updates and enhancements.

### You may also like

- [Principal component analysis for fast and model-free denoising of multi \*b\*-value diffusion-weighted MR images](#)  
Oliver J Gurney-Champion, David J Collins, Andreas Wetscherek et al.
- [The calibration of Strömberg \*uvby\*-H photometry for late-type stars: a model atmosphere approach](#)  
A Önehag
- [A synthetic data generation procedure for univariate circular data with various outliers scenarios using Python programming language](#)  
N S Zulkipli, S Z Satari and W N S Wan Yusoff

# Progress in Biomedical Engineering



## PERSPECTIVE

### The potential synergies between synthetic data and *in silico* trials in relation to generating representative virtual population cohorts

Puja Myles<sup>1,\*</sup> , Johan Ordish<sup>1</sup> and Allan Tucker<sup>2</sup>

<sup>1</sup> Medicines and Healthcare Products Regulatory Agency, London, United Kingdom

<sup>2</sup> Department of Computer Science, Brunel University London, London, United Kingdom

\* Author to whom any correspondence should be addressed.

E-mail: [puja.myles@mhra.gov.uk](mailto:puja.myles@mhra.gov.uk)

**Keywords:** *in silico* trials, synthetic data, medical devices

## Abstract

*In silico* trial methods promise to improve the path to market for both medicines and medical devices, targeting the development of products, reducing reliance on animal trials, and providing adjunct evidence to bolster regulatory submissions. *In silico* trials are only as good as the simulated data which underpins them, consequently, often the most difficult challenge when creating robust *in silico* models is the generation of simulated measurements or even virtual patients that are representative of real measurements and patients. This article digests the current state of the art for generating synthetic patient data outside the context of *in silico* trials and outlines potential synergies to unlock the potential of *in silico* trials using virtual populations, by exploiting synthetic patient data to model effects on a more diverse and representative population. Synthetic data could be defined as artificial data that mimic the properties and relationships in real data. Recent advances in synthetic data generation methodologies have allowed for the generation of high-fidelity synthetic data that are both statistically and clinically, indistinguishable from real patient data. Other experimental work has demonstrated that synthetic data generation methods can be used for selective sample boosting of underrepresented groups. This article will provide a brief outline of synthetic data generation approaches and discuss how evaluation frameworks developed to assess synthetic data fidelity and utility could be adapted to evaluate the similarity of virtual patients used for *in silico* trials, to real patients. The article will then discuss outstanding challenges and areas for further research that would advance both synthetic data generation methods and *in silico* trial methods. Finally, the article will also provide a perspective on what evidence will be required to facilitate wider acceptance of *in silico* trials for regulatory evaluation of medicines and medical devices, including implications for post marketing safety surveillance.

## 1. Introduction

*In silico* trial methods represent an opportunity to augment and streamline elements of the path to market for both medical devices and medicines. Broadly, some of the promise is that further use of such methods might reduce reliance on animal trials and bolster evidence that would otherwise be generated at risk to clinical trial participants [1]. Accordingly, so long as the evidence is robust, the more that can be mustered from modelling, the smoother the route to market will be and less risk will be borne by participants. Indeed, there is emerging consensus of *in silico* trials' potential, the trajectory being that these methods have a role in evidencing medicines and medical devices, the primary questions now being how and to what extent? However, this potential is all predicated on confidence in data quality. *In silico* models trained on poor data will themselves perform poorly; models trained on incomplete data will be incomplete.

*In silico* approaches may include one or more of the following elements: virtual participants and population, virtual exams and investigations, virtual readers (for e.g. interpretation of a virtual image generated with a virtual simulated scanner) and outcomes [2]. As access to quality data is likely to be the

foremost challenge in getting *in silico* trial methods into standard practice, it is necessary to consider methods that might both facilitate access to data and methods that might boost underrepresented subgroups within datasets. For instance, it is one question to consider to what extent *in silico* methods are appropriate for regulatory purposes for both medicinal products and medical devices. In this article we focus on one element of *in silico* approaches, namely, the generation of representative virtual participant cohorts and potential learnings from recent methodological advances in generation of synthetic patient data outside the context of *in silico* trials. We define synthetic data as well as the various approaches used to generate such data, then outline a framework to evaluate synthetic data, also considering potential synergies between synthetic data and *in silico* trial methods, and then finally consider both areas for future research and regulatory questions that require further investigation.

## 2. Defining synthetic data

Conceptually, synthetic data are artificial data that mimic the properties of and relationships in real data. The quality of synthetic data depends on the approach taken to synthetic data generation and is often described in terms of its ‘utility’ or ‘fidelity.’ A synthetic dataset that captures complex inter-relationships between various data fields and the statistical properties of real data can be referred to as a ‘high-fidelity’ synthetic dataset [3]. It would follow that a ‘high-fidelity’ synthetic dataset should also have ‘high-utility’ i.e. the capability to produce analysis results similar to the original data [4].

Using the example of patient healthcare data (the focus of this paper), a high-fidelity synthetic dataset would be able to capture complex clinical relationships and be clinically indistinguishable from real patient data. The generation of high-utility synthetic data tends to be highly resource intensive given the present state of play and depending on the application for which synthetic data are required, it may be acceptable to use low or medium utility synthetic data.

While high-fidelity synthetic data could be used as a proxy for real data (including for complex multivariable analyses involving a range of machine learning algorithms) with a high degree of confidence, medium-fidelity data would only be suitable for simple analyses like proportions, summary statistics for single variables or cross-tabulations involving two variables. Low-fidelity synthetic data on the other hand, should only be used as a sample dataset that provides an understanding of the data types, data values, data formats, data structure and table relationships in the real data that it seeks to represent. In the context of *in silico* trials, high-fidelity synthetic data would be required.

## 3. Synthetic data generation approaches

Synthetic data generation methods with respect to synthetic patient populations, that have been developed outside the context of *in silico* trials, can be broadly categorized into three groups: generating synthetic data based on statistical properties of real data; adding noise to real data; and using machine-learning techniques to generate synthetic data [5].

### 3.1. Generating synthetic data based on statistical properties of real data

This approach relies on statistical properties of real data such as population distributions—for example, mean values, standard deviation, and value ranges for data fields such as blood cholesterol measurements or known prevalence of a disease in various subgroups. Typically, in this approach one variable at a time is synthesised though it may be possible to undertake conditional generation of some variables on a limited basis (for e.g. different height distributions are inputted for the different genders). This approach is useful when the real data are difficult to access, or the distribution of events is highly imbalanced in the available real data sample. A key limitation of this approach is that, while each synthetic data field will have the statistical properties of real data at the univariate level, the complex multivariable relationship between data fields will be difficult to capture. Thus, this approach would generally yield low- or medium-fidelity synthetic data.

### 3.2. Adding noise to real data

This approach involves perturbation of some of the data fields in real data in different ways including substitution of real values with other realistic values, random shuffling of data values within a particular data field or application of a random numeric variance (for e.g.  $\pm 10\%$  applied to all data values in a field such that the data distribution is preserved). Substitution of real values can also be approached by swapping data within a data field with another sample from the same distribution [6]. These techniques can be used to generate low- or medium-fidelity data.

### 3.3. Machine learning techniques to generate synthetic data

Advanced statistical modelling and machine learning techniques such as Hidden Markov models, Bayesian networks (BNs), and deep-learning approaches such as generative adversarial networks (GANs) can be used to learn patterns and relationships between different data fields in real data. The learned patterns are then used as an input for the synthetic data generator to yield synthetic data. These methods can be used to yield medium- or high-fidelity synthetic data because they are able to capture complex multivariable relationships between various data fields.

The actual choice of machine-learning algorithm is dependent on the specific requirements for synthetic data. For instance, when transparency is a key requirement, BN approaches are preferable to GAN-based approaches. Unpublished findings from the Medicines and Healthcare products Regulatory Agency's (MHRA)'s synthetic data research team suggest that GAN-based approaches may perform better than BN approaches for numerical data fields and vice versa for categorical/nominal data fields. The BN approach included latent variable modelling to deal with missing values in the real data. Hidden Markov models on the other hand, have been particularly useful for time-series data and are able to take into account missing values in real longitudinal data [7].

To summarise, recent approaches to synthetic patient data generation outside the context of *in silico* trials, tend to use advanced statistical modelling or machine learning approaches to generate synthetic patients who are statistically and clinically indistinguishable from that of the target population they intend to model. These are phenomenological data-driven models that do not describe the mechanisms that underpin the data but only the univariate and multivariable relationships. Thus, in the example of BN approaches, even though it is possible to view the relationships between data features in a graphical representation, the connections between the various nodes representing data features are not causal.

Within the *in silico* trials domain however, both phenomenological and mechanistic models can be used for generating virtual participants, simulating virtual interventions or examinations, as well as virtual outcomes. Mechanistic approaches involve fully specifying a data domain using an underlying mechanistic model (often employing differential equations). It exploits expertise and knowledge of a domain to mimic real data as closely as possible [8] and enables the fine-tuning of parameters to achieve good fit to datasets that can be too small for machine learning techniques. The resulting models enable the simulation of interacting processes under different conditions to enable the prediction of complex system behaviour. Mechanistic modelling is growing in popularity for modelling the behaviour of clinical trial outcomes [9]. In some cases, mechanistic models can be used to complement machine learning approaches [10, 11] where there is both background knowledge and sufficient data available.

## 4. Evaluation framework for synthetic data

This section outlines some of the evaluation approaches used for synthetic patient data more broadly and could be adapted for evaluating the quality of virtual participant cohorts used for *in silico* trials. Data utility measures are a good way to assess whether a synthetic dataset can justify the claims of being high-fidelity. One of the earlier papers considering evaluation of synthetic data, Snoke *et al* (2018) outlined general and specific utility measures for synthetic data [12]. They defined general utility measures as summaries of differences between real and original data as opposed to specific measures of utility that focused on results from particular analyses. They suggest that when the intended purpose of the synthetic dataset is known, specific measures of utility may be more helpful but when the intended purpose is not known, general utility measures are more appropriate.

More recently, El Emam *et al* (2020) describe three types of approaches to assessing the utility of synthetic data: workload-aware evaluations, generic data utility metrics and subjective assessments of data utility [4]. Workload-aware metrics consider which types of analyses are feasible using the synthetic data and by replicating analyses carried out in the real data using the synthetic data. Analyses can range from simple descriptive statistics to complex multivariable machine learning models. Subjective evaluations involve classification of a random mix of real and synthetic records by domain experts followed by an evaluation of the accuracy of that classification. Generic assessments include metrics like the distance between the original and transformed data; these assessments provide an assessment of fidelity with utility being inferred on this basis.

Distributions can be compared by visual examination of histograms or by using summary statistics like the Hellinger distance (a probabilistic measure between 0 and 1, where 0 indicates no difference between distributions) to measure the difference in distributions between each variable in the real and synthetic data. The median Hellinger distance across all variables should be close to 0 with very small variations, for a high-fidelity dataset. Bivariate and multivariate distance analyses typically involve correlation analyses.

Our own experiments in synthetic data generation have used a combination of all three approaches described by Wang *et al* (2021), using generic assessments of fidelity like the univariate, bivariate and multivariate distances between variables [5]. We used the Kolmogorov–Smirnov test to determine any differences in the univariate distance between the synthetic and real datasets and nonmetric multidimensional scaling to assess multivariate distance. We also undertook a subjective evaluation whereby two independent medical assessors reviewed a sample ( $n = 100$ ) containing randomly selected records for equal number of synthetic and real patients with the aim of categorising them as synthetic or real based on the clinical characteristics. Finally, we compared the real and synthetic datasets by using stacked ensembles including six different machine learning algorithms [least absolute shrinkage and selection operator, classification and regression training, extremely randomised trees, feed-forward neural networks, non-negative least squares and random forest] to predict cardiovascular disease risk for a more rigorous test of fidelity. This approach shares a similar philosophy to the ‘all models test’ approach proposed by Tucker *et al* (2020) where all possible models are examined as it is not known *a priori* what an actual analyst would want to do with the dataset. Based on these evaluations, we posited that our approach to synthetic data generation using BNs incorporating latent variables to learn the distributions and relationships in the real data, yielded high fidelity synthetic data [7].

Such an evaluation framework could also apply to virtual patient cohorts employed in the context of *in silico* trials by providing a meaningful comparison to real patient cohorts. This would also be applicable to some degree to virtual patient cohorts that include boosted characteristics or simulated values to address missing data gaps in real data, though further work is needed in this area.

## 5. Potential synergies between synthetic data and *in silico* trials

High-fidelity synthetic patient data capture many of the complexities of real patient data. It offers the ability to infer the effects of medical interventions on a diverse population if generated using models of large national datasets. This has been possible for our approach because the CPRD database covers underlying health conditions of many different subpopulations within the UK, incorporating effects of, for example, age, ethnicity and regional disparity. Our approach to synthetic patient data generation means that we can condition our sampling of synthetic patients on evidence. For example, we may want to sample patients who suffer from a particular condition or from a specific demographic. This means that we can control for outcomes of virtual clinical trials to explore the effects more widely [13]. However, the utility of synthetic patient data can be limited by reliance on high-quality secondary data. That is, data collected for reasons other than simulating the effects of interventions [14]. This can potentially result in models that only reflect what has been measured in a population in the past and will not include effects of previously unseen interventions. One method to deal with this can be by combining the phenomenological synthetic patient data approaches outlined here with mechanistic models to simulate intervention effects to provide more realistic estimates of effectiveness in the intended target population as well as in subgroups [15].

## 6. Areas for further research

Linking high-fidelity synthetic patient data to virtual/*in silico* clinical trial data offers great potential. However, this research is still in its infancy and the identification of suitable proxies for linking the two data sources will be key to its success. There will need to be a full exploration of bias in clinical trials using appropriate metrics on sub-populations. We have begun this process on synthetic data by using boosting methods applied to certain sub-populations that are identified as under-represented based upon model performance metrics [16]. We are undertaking further experiments to determine whether such boosting is informative. Furthermore, research is required to fully understand under what circumstances *in silico* models developed on synthetic patient data would be acceptable for regulatory purposes versus real patient data and what requirements would attach to those models.

## 7. Regulatory perspective

There is growing acceptance that *in silico* modelling has a role in evidencing medicines and medical devices. Synthetic patient data that can demonstrate that it is representative of the target population of intended use has the potential to accelerate *in silico* modelling. At minimum, we suggest that *in silico* models would have to demonstrate fidelity to their real patient data counterparts or comparative performance versus models trained on real data. As described above, the methodology for evaluating synthetic patient data is still nascent and developing. In the context of regulation, until the state of synthetic data crystallises, it is likely that the use of *in silico* modelling that typically requires some simulation akin to synthetic data will be stymied. It is

therefore likely that the trajectory of further acceptance of *in silico* modelling will continue in the regulatory sphere, but further acceptance and standardisation of synthetic data will be necessary to accelerate acceptance.

## 8. Conclusion

The benefits of *in silico* modelling are plain: so long as the models are accurate, these methods provide a useful adjunct data that should increasingly make a contribution to the evidence base of medical devices and medicines. A key element of *in silico* trial methods is the use of virtual participant cohorts. Advances in synthetic patient data generation methods outside the context of *in silico* trials could present a complementary set of methods with obvious synergies to unlock and bolster virtual participant datasets that underpin *in silico* models. Consequently, if the acceptance of synthetic patient data in general is stymied, so too will the development of *in silico* models based on virtual participant cohorts. Nevertheless, as methods to assess synthetic data progress, the benefits that it provides may outweigh its risks, thereby driving its acceptance amongst regulators.

## Data availability statement

No new data were created or analysed in this study.

## ORCID iD

Puja Myles  <https://orcid.org/0000-0002-8976-890X>

## References

- [1] Badano A 2021 In silico imaging clinical trials: cheaper, faster, better, safer, and more scalable *Trials* **22** 64
- [2] Abadi E, Segars W P, Tsui B M W, Kinahan P E, Bottenus N, Frangi A F, Maidment A, Lo J and Samei E 2020 Virtual clinical trials in medical imaging: a review *J. Med. Imaging* **7** 042805
- [3] Myles P, Ordish J and Branson R 2021 Synthetic data and the innovation, assessment, and regulation of AI medical devices *RF Q.* **1** 48–53 (available at: [www.raps.org/news-and-articles/rf-quarterly/2021-q2](http://www.raps.org/news-and-articles/rf-quarterly/2021-q2))
- [4] El Emam K, Mosquera L and Hoptroff R 2020 *Practical Synthetic Data Generation: Balancing Privacy and the Broad Availability of Data* 1st edn (Sebastopol, CA: O'Reilly Media)
- [5] Wang Z, Myles P and Tucker A 2021 Generating and evaluating cross-sectional synthetic electronic healthcare data: preserving data utility and patient privacy *Comput. Intell.* **37** 819–51
- [6] Sarathy R and Muralidhar K 2012 Perturbation methods for protecting numerical data: evolution and evaluation *Handb. Stat.* **28** 513–31
- [7] Tucker A, Wang Z, Rotalinti Y and Myles P 2020 Generating high-fidelity synthetic patient data for assessing machine learning healthcare software *NPJ Digit. Med.* **3** 1–3
- [8] Musumba F T et al 2021 Scientific and regulatory evaluation of mechanistic in silico drug and disease models in drug development: building model credibility *CPT Pharmacomet. Syst. Pharmacol.* **10** 804–25
- [9] Pitcher M J, Dobson S J, Kelsey T W, Chaplain M A J, Sloan D J, Gillespie S H and Bowness R 2020 How mechanistic in silico modelling can improve our understanding of TB disease and treatment *Int. J. Tuberc. Lung Dis.* **24** 1145–50
- [10] An G, Dollinger M and Li-Jessen N Y K 2022 Editorial: Integration of machine learning and computer simulation in solving complex physiological and medical questions *Front. Physiol.* **13** 949771
- [11] Antontsev V, Jagarapu A, Bunday Y, Hou H, Khotimchenko M, Walsh J and Varshney J 2021 A hybrid modelling approach for assessing mechanistic models of small molecule partitioning *in vivo* using a machine learning-integrated platform *Sci. Rep.* **11** 11143
- [12] Snoke J, Raab G M, Nowok B, Dibben C and Slavkovic A 2018 General and specific utility measures for synthetic data *J. R. Stat. Soc. A* **181** 663–88
- [13] Samei E, Kinahan P, Nishikawa R M and Maidment A 2020 Virtual clinical trials: why and what (special section guest editorial) *J. Med. Imaging* **7** 1
- [14] Boslaugh S 2007 *Secondary Data Sources for Public Health: A Practical Guide* (Cambridge: Cambridge University Press) (<https://doi.org/10.1017/CBO9780511618802>)
- [15] Sarrami-Foroushani A, Lassila T, MacRaid M, Asquith J, Roes K C, Byrne J V and Frangi A F 2021 In-silico trial of intracranial flow diverters replicates and expands insights from conventional clinical trials *Nat. Commun.* **12** 1–2
- [16] Draghi B, Wang Z, Myles P and Tucker A 2021 Bayesboost: identifying and handling bias using synthetic data generators *3rd Int. Workshop on Learning with Imbalanced Domains: Theory and Applications ECML-PKDD (Bilbao, Basque Country, Spain, 17 September 2021)* vol 154 pp 49–62