



What do we mean by identifiability in mixed effects models?

Marc Lavielle, Leon Aarons

► To cite this version:

Marc Lavielle, Leon Aarons. What do we mean by identifiability in mixed effects models?. 2016.
hal-01251986

HAL Id: hal-01251986

<https://hal.science/hal-01251986>

Preprint submitted on 7 Jan 2016

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

What do we mean by identifiability in mixed effects models?

Marc Lavielle · Leon Aarons

Received: date / Accepted: date

Abstract We discuss the question of model identifiability within the context of nonlinear mixed effects models. Although there has been extensive research in the area of fixed effects models, much less attention has been paid to random effects models. In this context we distinguish between theoretical identifiability, in which different parameter values lead to non-identical probability distributions, structural identifiability which concerns the algebraic properties of the structural model, and practical identifiability, whereby the model may be theoretically identifiable but the design of the experiment may make parameter estimation difficult and imprecise. We explore a number of pharmacokinetic models which are known to be non-identifiable at an individual level but can become identifiable at the population level if a number of specific assumptions on the probabilistic model hold. Essentially if the probabilistic models are different, even though the structural models are non-identifiable, then they will lead to different likelihoods. The findings are supported through simulations.

Keywords Model identifiability · Practical identifiability · Structural identifiability · Parameter estimation · Mixed effects model · Pharmacokinetics

1 Introduction

A statistical model is said to be identifiable when, given an infinite amount of data, it is possible to uniquely estimate the true values of the model parameters. The uniqueness property implies that different values of the model parameters generate different probability distributions of the observable variables. Conversely if two or more sets of parameters generate identical distributions of the observed values the model is not identifiable. However it still may be possible to uniquely identify a subset of parameters: that is, not all the parameters of an unidentifiable model are unidentifiable. In this case we say that the model is partially identifiable.

It should also be noted that even when the parameters of the model are unidentifiable, the model itself may perfectly describe the observed variables, in other words perfectly fit the data.

M. Lavielle
Inria Saclay, France
E-mail: Marc.Lavielle@inria.fr

L. Aarons
University of Manchester, UK

A distinction is made in the literature between structural identifiability and practical identifiability. Structural identifiability is related to the structure of the underlying mathematical model. It is concerned with whether the parameters of a model can be exactly identified from a given experiment with perfect input-output data [1,5,15,18,21]. Chappell *et al.* [4] compare methods for analyzing the global structural identifiability of the parameters of a nonlinear system with a specified input function. Gargash and Mital [9] study the problem of global and local structural identifiability for fixed structure compartmental deterministic model of the biological system. Frohlich *et al.* [7] investigate the effect of structural non-identifiability on the performance of frequentist methods for standard uncertainty analysis. They observe that the profile likelihood approach is the only one that properly identifies structural non-identifiability of parameters. On the other hand practical identifiability is related to the study design, i.e. the limited amount of information that can be obtained from a given experiment [11,13,15]. A link exists between practical identifiability and sensitivity analysis. Brun *et al.* [3] propose a systematic approach for tackling the parameter identifiability problem of large models based on local sensitivity analysis.

Model identifiability is closely related to model indistinguishability. The objective of indistinguishability analysis is to determine if different models are capable to fit the available input-output data [10,22]. Identifiability and distinguishability of parametric models are important properties when the parameters to be estimated have a biological meaning or when the model is to be used to reconstruct physiologically meaningful variables that cannot be measured directly [18].

The question of identifiability of pharmacokinetic (PK) and pharmacokinetic-pharmacodynamics (PKPD) models has been previously studied. Evans *et al.* [6] consider the identifiability of a parent-metabolite pharmacokinetic model for ivabradine and one of its metabolites. Shivva *et al.* [16] use a simple one compartment population pharmacokinetic model to show that identifiability of the variances of the random effects are affected by the parameterisation of the fixed effects. An approach for assessment of identifiability for fixed and mixed effects PK or PKPD models is proposed in [15]. Guedj *et al.* [11] study the identifiability of parameters in a model of HIV dynamics based on a system of non-linear Ordinary Differential Equations. Garcia *et al.* [8] discuss different types of identifiability that occur in physiologically-based pharmacokinetic (PBPk) models and give reasons why they occur.

Identifiability in mixed effects models has received much less attention. Identifiability of linear mixed effects models is considered in [19]. In the context of mixed effects models, identifiability is a fundamental prerequisite for model identification. It concerns uniqueness of the population parameters estimated from a given set of observations obtained from several individuals of a same population [2,12,20]. It was the purpose of the current investigation to examine some issues in identifiability of mixed effects models.

We define the concept of identifiability of mixed effects models in Section 2. Identifiability is usually defined for continuous data models but this concept can be easily extended to categorical, count or time to event data models. We show that a model can be identifiable even if it is not structurally identifiable. Indeed the probability distribution of the individual parameters plays an important role for characterizing the identifiability property of the model. Furthermore, a distinction is made between identifiability of the population parameters (the population parameters can be estimated successfully and unequivocally from the observed data) and identifiability of the individual parameters (the individual parameters can be estimated successfully and unequivocally from the observed data).

Sections 3 and 4 present some specific results taken from the field of pharmacokinetics. We consider first a situation where different parameterizations of a PK model are algebraically undistinguishable. The model is shown to be identifiable as soon as some hypothesis are made on the correlation structure of the PK parameters. It is usually claimed that the bioavailability F cannot be estimated using only PK measurements obtained from an oral administration: only the ratio F/V

can be properly estimated. We show that the model is identifiable under some assumptions on the probabilistic model. Then, both F and V can be simultaneously estimated.

2 Identifiability of mixed effects models

2.1 Preliminary remarks

The definition of identifiability for mixed effects models is not always very precise. As an example, the definition given in [20] reduces to: “parameters or models are called *non-identifiable* if two sets of different parameters lead to the same probability distribution”. This definitions remain quite ambiguous and needs to be clarified. In particular, we will need to distinguish the identifiability of the population parameters and the identifiability of the individual parameters.

We need also to make a clear distinction between identifiability, structural identifiability and practical identifiability, defining what these properties mean in the context of mixed effects models.

The concept of structural identifiability was introduced first in the area of systems and control (see [1]), where the systems are deterministic and depend on nonrandom parameters. The situation is quite different in a population approach context where the individual parameters are random variables. In order to analyse the properties of a statistical model it is necessary to take into account the algebraic properties of the structural model as well as the properties of the probabilistic model. In particular, the choice of the parameterization may have a strong impact on the probabilistic properties of the model. For instance, parameterizations of a PK model using volume V and clearance Cl , or using volume V and elimination rate constant k , where $k = Cl/V$, are absolutely equivalent from a purely algebraic point of view: we can use both interchangeably. That’s not the case if we put a distribution on these parameters.

Assume, for instance, that V and Cl are independent and log-normally distributed:

$$\begin{aligned}\log(V) &\sim \mathcal{N}(\log(V_{\text{pop}}), \omega_V^2) \\ \log(Cl) &\sim \mathcal{N}(\log(Cl_{\text{pop}}), \omega_{Cl}^2) \\ \text{corr}(\log(V), \log(Cl)) &= 0\end{aligned}$$

where $\text{corr}(\log(V), \log(Cl))$ is the linear correlation between $\log(V)$ and $\log(Cl)$.

Then, we are implicitly assuming that k is also log-normally distributed with a variance ω_k^2 larger than the variance of V and that $\log(k)$ and $\log(V)$ are negatively correlated. Indeed,

$$\log(k) = \log(Cl) - \log(V)$$

Then, since $\log(Cl)$ and $\log(V)$ are both normally distributed, $\log(k)$ is also normally distributed and

$$\log(k) \sim \mathcal{N}(\log(Cl_{\text{pop}}/V_{\text{pop}}), \omega_V^2 + \omega_{Cl}^2)$$

Furthermore, the covariance between $\log(V)$ and $\log(k)$ is given by

$$\text{cov}(\log(V), \log(k)) = -\text{var}(\log(V)) = -\omega_V^2$$

Then,

$$\text{corr}(\log(V), \log(k)) = -\omega_V/\omega_k$$

On the other hand, if we use parameters (V, k) in the model, assuming independent distributions, then, we implicitly assume that V and Cl are dependent.

2.2 Structural identifiability

Structural identifiability is related to the structure of the underlying mathematical model, for example as discussed above, the PK model.

For sake of simplicity, we will only consider models for univariate data, i.e. for a single outcome.

We need some notation. Let N be the number of individuals. Then, for $i = 1, 2, \dots, N$, let $y_i = (y_{ij}, 1 \leq j \leq n_i)$ be the n_i observations for individual i collected at times $(t_{ij}, 1 \leq j \leq n_i)$.

Let us start with a basic model for continuous data:

$$y_{ij} = f(t_{ij}; \psi_i) + \varepsilon_{ij}$$

Here, f is the *structural model*, which is a parametric function of time, ψ_i is a vector of *individual parameters*, and $(\varepsilon_{ij}, 1 \leq j \leq n_i)$ is a sequence of *residual errors*. We will assume that ε_{ij} is a sequence of random variables with mean 0 and finite variance σ^2 .

Structural identifiability of the model is directly related to the properties of the structural model f . We don't take into account possible differences between individuals.

Without any loss of generality, we will assume that f is defined for $t \geq 0$. Let $f(\cdot; \psi)$ be the function f defined for a given vector of parameters ψ (we could equivalently use the notation f_ψ). Then, we will say that the model is structurally identifiable if there exists a one-to-one mapping between the parameter ψ and the function $f(\cdot; \psi)$, i.e.

$$\psi = \psi' \Leftrightarrow f(t; \psi) = f(t; \psi') \quad \text{for any } t \geq 0$$

This definition can be easily extended to any other parametric mixed effects models:

- Continuous data model with non constant residual error model: Assume now that there exists a parametric function g such that

$$y_{ij} = f(t_{ij}; \psi_i) + g(t_{ij}; \psi_i)\varepsilon_{ij}$$

Here, the structural model is the pair (f, g) .

- Time-to-event data model: The structural model is the hazard function $h(t; \psi)$.
- Count data model: Consider for instance a Poisson model,

$$y_{ij} \sim \text{Poisson}(\lambda(t_{ij}, \psi_i))$$

The structural model is the Poisson intensity $\lambda(t; \psi)$

- Categorical data model: Assume a Bernoulli model for binary data as an example,

$$\mathbb{P}(y_{ij} = 1) = 1 - \mathbb{P}(y_{ij} = 0) = \pi(t_{ij}; \psi_i)$$

Here, the structural model is the function $\pi(t; \psi)$.

2.3 Practical identifiability

Practical non-identifiability is less clearly defined in the literature compared to structural non-identifiability.

Practical identifiability not only depends on the model structure, but is also related to the experimental conditions together with the quality and quantity of the measurements [13]. Then, a parameter that is structurally identifiable may be practically unidentifiable with a limited amount and quality of experimental data.

Deriving some relationship between identifiability and confidence interval of parameter estimates seems natural. Nevertheless, since practical identifiability is a non-asymptotic property, it is not possible to propose a rigorous definition based on asymptotic confidence intervals.

Raue *et al.* [14] propose an appealing definition of practical identifiability based on a likelihood-based confidence region instead of asymptotic confidence intervals. They suggest an approach that exploits the profile likelihood and enables the detection of both structural and practical non-identifiabilities.

Unfortunately, such an approach is cumbersome to adopt for (nonlinear) mixed effects models since it requires the estimation of the population parameters and computation of the likelihood many times.

Methods that can be used in practice for detecting some identifiability issues remain quite empirical:

- We can for instance run the estimation algorithm with different initial values. Convergence to different solutions may be suspicious.
- Even if it is an asymptotic criteria, the observed Fisher Information Matrix can also be used. Indeed, the inverse of this matrix provides an approximation of the variance-covariance matrix of the estimated parameters. A large condition number of this variance-covariance matrix (i.e. the ratio of its largest and smallest eigenvalues) reflects a strong correlation structure between estimates and may indicate some identifiability issue.

2.4 Identifiability of the population parameters

Identifiability of the population parameters θ is related to the properties of the statistical model of the observations $y = (y_i, 1 \leq i \leq N)$:

$$p(y; \theta) = \prod_{i=1}^N p(y_i; \theta) \tag{1}$$

$$= \prod_{i=1}^N \int p(y_i, \psi_i; \theta) d\psi_i \tag{2}$$

We are here in a “classical” situation where the statistical model is identifiable if the mapping between θ and the probability distribution $p(y; \theta)$ is one-to-one (we use indiscriminately $p(y; \theta)$ for the probability distribution function (pdf) and for the distribution of y).

By definition, the Maximum Likelihood (ML) estimate of θ maximizes $p(y; \theta)$. If the model is identifiable, then, the ML estimate converges to the true value of θ when N increases, under very general regularity conditions [12]. That means that it is possible to estimate θ as precisely as required as soon as the size N is large enough. On the other hand, if the model is not identifiable,

then, maximizing $\mathbf{p}(y; \theta)$ with respect to θ will not lead to a unique solution. We can find, for instance, $\hat{\theta}^{(1)}$ and $\hat{\theta}^{(2)}$ such that, for any vector of population parameters θ ,

$$\mathbf{p}(y; \theta) \leq \mathbf{p}(y; \hat{\theta}^{(1)}) = \mathbf{p}(y; \hat{\theta}^{(2)})$$

A number of particular cases have been reported. For instance, Wang *et al.* [19] study the identifiability of the covariance parameters in a linear mixed effects model. They focus on those models that are not over-parameterized and derive conditions of identifiability and study commonly used covariance structures. In an unpublished work, Nuñez and Concordet consider the identifiability problem in a nonlinear mixed effects model for continuous data, assuming a nonparametric distribution for the individual parameters. They provide several explicit conditions on f which ensure the identifiability of the model.

Since $\mathbf{p}(y; \theta) = \int \mathbf{p}(y, \psi; \theta)$, identifiability of the complete model $\mathbf{p}(y, \psi; \theta)$ is a necessary condition to ensure the identifiability of the observed model $\mathbf{p}(y; \theta)$. Unfortunately, it is not a sufficient condition.

For example consider the linear model

$$y_{ij} = (a_i + b_i)t_{ij} + \varepsilon_{ij} \quad (3)$$

where a_i and b_i are normally distributed with unknown means m_1 and m_2 . Here, the vector of population parameters θ includes m_1 and m_2 . The model $\mathbf{p}(y; \theta)$ is unidentifiable since the population parameters m_1 and m_2 cannot be estimated from a sequence of observations (y_{ij}) : only the sum $m_1 + m_2$ can be estimated since $a_i + b_i$ is normally distributed with mean $m_1 + m_2$. On the other hand, the joint model $\mathbf{p}(y_i, a_i, b_i; \theta)$ is identifiable since m_1 and m_2 can be estimated using sequences $(a_i, 1 \leq i \leq N)$ and $(b_i, 1 \leq i \leq N)$.

Analyzing the structural identifiability of the model is important and useful, but it is not sufficient for concluding if the model is identifiable or not. For instance, a statistical model may be identifiable even if it is not structurally identifiable. Indeed, when data coming from various individuals is available, and under some hypothesis, we can take advantage of the probability distribution of the individual parameters to estimate the population parameters. It is then the combination of algebraic relationships and probabilistic relationships that make the model identifiable.

Consider again the linear model (3). The structural model $f(t; a_i, b_i) = (a_i + b_i)t$ is not identifiable since only $a_i + b_i$ can be estimated. We have seen in the previous example that the statistical model of the observations is not identifiable when a_i and b_i are both normally distributed. Surprisingly, the model becomes identifiable in most cases when a_i and b_i are not both normally distributed.

Assume for instance that a_i and b_i have exponential distributions with parameters λ_1 and λ_2 . Then, assuming that $\lambda_1 \neq \lambda_2$, the sum $a_i + b_i$ is a random variable with pdf

$$p(z; \lambda_1, \lambda_2) = \frac{\lambda_1 \lambda_2}{\lambda_1 - \lambda_2} (e^{-\lambda_2 z} - e^{-\lambda_1 z})$$

There is a one-to-one mapping between (λ_1, λ_2) and this pdf: the model is identifiable. Identifiability of the model remains a theoretical property: in practice, accurate estimation of λ_1 and λ_2 will require a huge amount of data.

2.5 Identifiability of the individual parameters

We can see the problem of estimating the individual parameters (ψ_i) as an inverse problem: we aim to recover these unobserved vectors of parameters using the observations (y_i) .

Structural unidentifiability means that the problem is *ill posed*: we can find different vectors ψ_i and ψ'_i that produce the same structural predictions:

$$\psi_i \neq \psi'_i \quad \text{and} \quad f(t; \psi_i) = f(t; \psi'_i) \quad \text{for any } t \geq 0$$

And then,

$$\mathbf{p}(y_i | \psi_i) = \mathbf{p}(y_i | \psi'_i)$$

Each individual conditional model $\mathbf{p}(y_i | \psi_i)$ is therefore unidentifiable, but in a mixed effects model, we usually don't estimate each individual parameter ψ_i by maximizing this conditional distribution $\mathbf{p}(y_i | \psi_i)$. Indeed, each individual parameter ψ_i is a random vector with distribution $\mathbf{p}(\psi_i; \theta)$. Then, for a given vector of population parameter θ (given or previously estimated), we rather consider the *posterior* distribution

$$\mathbf{p}(\psi_i | y_i; \theta) = \frac{\mathbf{p}(y_i | \psi_i) \mathbf{p}(\psi_i; \theta)}{\mathbf{p}(y_i; \theta)}$$

Thus, ψ_i is identifiable as soon as $\mathbf{p}(\psi_i; \theta) \neq \mathbf{p}(\psi'_i; \theta)$.

In other words, the population distribution of the individual parameters plays the role of a *prior* distribution that now makes the problem *well posed*.

Example 1: Consider a bi-exponential model:

$$f(t; A_i, B_i, \alpha_i, \beta_i) = A_i e^{-\alpha_i t} + B_i e^{-\beta_i t}$$

This structural model is locally identifiable but globally unidentifiable since (A_i, α_i) and (B_i, β_i) are interchangeable. The model becomes identifiable if we introduce some information about α_i and β_i , assuming for instance that $\mathbb{P}(\alpha_i > \beta_i) > 0.5$. The use of this prior leads to select the solution where $\alpha_i > \beta_i$ and discard the other solution where $\alpha_i < \beta_i$.

Example 2: Consider again model (3) where $\varepsilon_{ij} \sim_{\text{i.i.d.}} (0, \sigma^2)$.

This structural model $f(t; a_i, b_i) = (a_i + b_i)t$ is clearly locally unidentifiable since only the sum $a_i + b_i$ can be estimated maximizing the conditional pdf of y_i . Let $c_i = a_i + b_i$ and define $\hat{c}_i$ as

$$\begin{aligned} \hat{c}_i &= \underset{c_i}{\text{Argmax}} \mathbf{p}(y_i | c_i) \\ &= \underset{c_i}{\text{Argmin}} \sum_j (y_{ij} - c_i t_{ij})^2 \end{aligned}$$

Thus, $\hat{a}_i + \hat{b}_i = \hat{c}_i = \sum_j y_{ij} t_{ij} / \sum_j t_{ij}^2$, but it is impossible to decompose this sum and compute $\hat{a}_i$ and $\hat{b}_i$ without any additional information.

Assume for instance that

$$\begin{pmatrix} a_i \\ b_i \end{pmatrix} \sim \mathcal{N} \left(\begin{pmatrix} m_1 \\ m_2 \end{pmatrix}, \Omega \right)$$

Let T_i be the $n_i \times 2$ matrix

$$T_i = \begin{pmatrix} t_{i,1} & t_{i,2} & \dots & t_{i,n_i} \\ t_{i,1} & t_{i,2} & \dots & t_{i,n_i} \end{pmatrix}'$$

The posterior distribution of a_i and b_i is now a well defined normal distribution:

$$\begin{pmatrix} a_i \\ b_i \end{pmatrix} | y_i \sim \mathcal{N} \left(\begin{pmatrix} \mu_{i,1} \\ \mu_{i,2} \end{pmatrix}, \Gamma_i \right)$$

where

$$\Gamma_i = \left(\frac{T_i' T_i}{\sigma^2} + \Omega^{-1} \right)^{-1}$$

and

$$\begin{pmatrix} \mu_{i,1} \\ \mu_{i,2} \end{pmatrix} = \Gamma_i \left(\frac{T_i' y_i}{\sigma^2} + \Omega^{-1} \begin{pmatrix} m_1 \\ m_2 \end{pmatrix} \right)$$

The maximum a posteriori (MAP) estimates of a_i and b_i are, respectively, $\mu_{i,1}$ and $\mu_{i,2}$. They are well defined and unique: the individual parameters of the model are now identifiable.

Remark: introduction of a Gaussian prior information for estimating the parameters of a linear ill-posed problem is equivalent to introduce a Tikhonov regularization term [17].

3 Illustration: the flip-flop phenomenon

In this example, the structural model is not identifiable without some physiological constraint (we can find two different sets of PK parameters that produce identical PK profiles). Nevertheless, identifiability is ensured under some assumptions on the probabilistic model.

For a sake of simplicity, we will consider a single individual and omit the subscript i in the notation.

3.1 The structural PK model

Consider a basic PK model for a single oral administration at time 0,

$$f(t; k_a, V, k) = \frac{Dk_a}{V(k_a - k)} (e^{-k t} - e^{-k_a t}). \quad (4)$$

It is easy to see that, for any $t \geq 0$,

$$f(t; k_a, V, k) = f(t; k'_a, V', k')$$

where $k'_a = k$, $k' = k_a$ and $V' = (k/k_a)V$.

Without any assumptions on the parameter values, the two solutions are indistinguishable. Some assumptions on the parameter space may make the model identifiable (e.g. $k_a > k$).

3.2 The probabilistic model

Without any assumption on the parameter space the structural model is not identifiable but the probabilistic model may be identifiable under some assumptions.

If we assume log-normal distributions for (k_a, V, k) , then (k'_a, V', k') are also log-normally distributed since $\log(V') = \log(V) + \log(k) - \log(k_a)$.

Furthermore, if we assume that $\log(k_a)$, $\log(V)$ and $\log(k)$ are uncorrelated, then $\log(V')$ and $\log(k'_a)$ are implicitly correlated, as well as $\log(V')$ and $\log(k')$.

Thus, the two following models are distinguishable:

$$\begin{aligned}\mathcal{M}_1 : \quad & \text{corr}(k_a, V) = \text{corr}(k_a, k) = \text{corr}(k, V) = 0 \\ \mathcal{M}_2 : \quad & \text{corr}(k'_a, V') = \text{corr}(k'_a, k') = \text{corr}(k', V') = 0\end{aligned}$$

In other words, if we assume that the variance-covariance matrix Ω of $(\log(k_a), \log(V), \log(k))$ is diagonal, then the model is identifiable since we cannot have simultaneously both Ω and Ω' diagonal, where Ω' is the variance-covariance matrix of $(\log(k'_a), \log(V'), \log(k'))$.

On the other hand, if we don't make any assumption on Ω , then the model is not identifiable: we can find two sets of fixed parameters $(ka_{\text{pop}}, V_{\text{pop}}, k_{\text{pop}})$ and $(ka'_{\text{pop}}, V'_{\text{pop}}, k'_{\text{pop}})$ and two covariance matrices Ω and Ω' such that

$$p(y, \psi; ka_{\text{pop}}, V_{\text{pop}}, k_{\text{pop}}, \Omega) = p(y, \psi; ka'_{\text{pop}}, V'_{\text{pop}}, k'_{\text{pop}}, \Omega')$$

where $y = (y_j, 1 \leq j \leq n)$ is a vector of observed concentrations and $\psi = (k_a, V, k)$ is the vector of individual PK parameters.

3.3 Statistical implications

Assume now that we have some PK data and we want to use the PK model defined in (4) to fit this data. Then, it is expected to obtain different values of the likelihood under $\mathcal{M}_1$ and $\mathcal{M}_2$ since the two probabilistic models are different (we cannot have simultaneously $\text{corr}(k, V) = 0$ and $\text{corr}(k', V) = 0$).

Simulation:

Data with $N = 100$ subjects and $n = 12$ measurements per subject were simulated with the following parameter values: $ka_{\text{pop}} = 1$, $V_{\text{pop}} = 10$, $k_{\text{pop}} = 0.1$, $\omega_{k_a} = 0.25$, $\omega_V = 0.3$, $\omega_k = 0.15$.

We can then try to estimate the population parameters of the model, using any software tool such as MONOLIX or NONMEM. According to the initial value, the SAEM algorithm [12] in MONOLIX converges to two solutions: $(\hat{k}_a^{(1)}, \hat{V}^{(1)}, \hat{k}^{(1)}) = (1, 10.4, 0.099)$ and $(\hat{k}_a^{(2)}, \hat{V}^{(2)}, \hat{k}^{(2)}) = (0.097, 0.99, 1.04)$. The estimated log-likelihood for these two solutions are, respectively, -1126 and -1145 . We see that, in this example, the first solution, which is the "right" solution (i.e. close to the "true" values used for the simulation), correspond to the global maximum of the likelihood, while the other one correspond to a local maximum.

In other words, the model is identifiable if we assume a diagonal matrix Ω , and maximizing the likelihood allows to select the "right" solution.

Here, it is not the values of the population PK parameters that allows one to select a model, but the (strong) hypothesis that we make concerning the covariance structure of the random effects. Indeed, if we simulate uncorrelated PK parameters using the other parameterization, then, the likelihood criteria will select this solution.

On the other hand, if we don't make any assumption about the covariance matrix, i.e. if we estimate a full variance-covariance matrix, then the model is not identifiable and the likelihood criteria cannot select a model. Indeed, the log-likelihood still exhibits two maxima, $(1, 10.4, 0.1)$ and $(0.098, 1, 1.03)$, but with very close values: -1125.7 and -1126.2 , respectively.

Remark 1: As expected, even if independent random effects were simulated using the first parameterization, estimated variance-covariance matrix associated to the second solution is not diagonal: estimated correlations between $\log(k'_a)$ and $\log(V')$ and between $\log(k')$ and $\log(V')$ are, respectively, 0.41 and -0.62 .

Remark 2: Similar results are obtained with NONMEM and FOCE: $\hat{k}_a^{(1)} = 1.01$, $\hat{V}^{(1)} = 10.4$, $\hat{k}^{(1)} = 0.099$ using the first set of initial estimates and $\hat{k}_a^{(2)} = 0.098$, $\hat{V}^{(2)} = 0.995$, $\hat{k}^{(2)} = 1.04$ using the second one. We have simulated rich data in this example. Then, up to the exchangeability issue between k_a and k , the individual PK parameters can be estimated accurately from the individual PK data. FOCE works very well here because this algorithm is precisely based on the estimation of the individual parameters.

4 Illustration: Identifiability of the bioavailability

We show with this example that, even if the structural model only allows the identification of the ratio V/F , the probabilistic model, under some assumptions, makes the model identifiable and allows the estimation of both V and F .

We will consider again a single individual and omit the subscript i in the notation.

4.1 The model and its properties

- The PK model is a one compartment model for oral administration with parameters (F, k_a, V, k) where F is the bioavailability, i.e. the fraction of administered dose which is absorbed. Let $f(t; F, k_a, V, k)$ be the predicted concentration given by this model at time t :

$$f(t; F, k_a, V, k) = \frac{D F k_a}{V(k_a - k)} (e^{-k t} - e^{-k_a t}) \quad (5)$$

- The residual error model is an exponential model, i.e. observed concentrations are log-normally distributed:

$$\log(y_j) = \log(f(t_j; F, k_a, V, k)) + \varepsilon_j$$

- k_a , V and k are log-normally distributed while F has a logit-normal distribution:

$$\text{logit}(F) \sim \mathcal{N}(\text{logit}(F_{\text{pop}}), \omega_F^2)$$

where $\text{logit}(x) = \log(x/(1-x))$ for $0 < x < 1$.

First of all, it is easy to see that the structural model is not identifiable. Indeed, let (F, k_a, V, k) and (F', k'_a, V', k') be two set of individual parameters such that $k'_a = k_a$, $k' = k$, $V'/F' = V/F$, then $f(t; F, k_a, V, k) = f(t; F', k'_a, V', k')$ for any $t > 0$. The structural model is therefore *partially identifiable* since only k_a , k and $R = V/F$ are identifiable.

Even if the structural model is not identifiable, the model itself is identifiable. We will use a “two step procedure” to demonstrate that we can derive a consistent estimator for all the population parameters of the model. Here, consistency means that this estimator converges to the true values of the population parameters when both the number of individuals and the number of observations per individual tend to infinity.

1. The structural model is partially identifiable. Then, for each individual $i = 1, 2, \dots, N$, the set of identifiable individual parameters k_{a_i} , k_i and $R_i = V_i/F_i$ can be perfectly recovered when the number of measurements n_i for individual i tends to infinity, by maximizing the conditional distribution $\mathbf{p}(y_i | k_{a_i}, k_i, R_i)$.

2. The ratio $R = V/F$ is a random variable defined as the ratio of a logit-normal and a log-normal variable. This distribution depends on parameters F_{pop} , V_{pop} , ω_F and ω_V and there exists a one-to-one mapping between these four parameters and the distribution of R . Thus, the maximum likelihood estimator of these four parameters, derived from a N -sample $R_1, R_2, \dots, R_N$ of R , is consistent: it converges to the true values of these parameters when the number of individual N tends to infinity. ML estimators of $k_{a,\text{pop}}$ and ω_{k_a} (resp. k_{pop} and ω_k) obtained from $k_{a_1}, \dots, k_{a_N}$ (resp. $k_1, \dots, k_N$) are also consistent.

Remark 1: The model is not identifiable if both V and F are log-normally distributed. Indeed, the log-ratio $\log(R)$ is normally distributed:

$$\log(R) \sim \mathcal{N}(\log(V_{\text{pop}}) - \log(F_{\text{pop}}), \omega_V^2 + \omega_F^2)$$

and there exists an infinity of possible decompositions leading to the same probability distribution.

Remark 2: This interesting result remains an asymptotical result. In practice, it means that we can expect to estimate all the population parameters of the model with a desired precision, if we have enough data for that. When the number of individuals and measurements is finite, the properties of the ML estimator cannot be derived analytically. A Monte-Carlo study can be used to evaluate these properties for a given design.

4.2 Simulation study

We simulate PK data from this model for $N = 5000$ individuals. A single dose $D = 100$ is administrated at time 0 and $n = 23$ measurements are collected at times 0.5, 1, 3, 5, $\dots$, 21, 23.

The standard deviation of the residual errors (ε_{ij}) is $\sigma = 0.1$.

Values of the population PK parameters are $k_{a,\text{pop}} = 1$, $V_{\text{pop}} = 10$, $k_{\text{pop}} = 0.1$, $\omega_{k_a} = 0.25$, $\omega_V = 0.3$ and $\omega_k = 0.15$. We will consider two logit distributions for F .

Model A: $\text{logit}(F) \sim \mathcal{N}(\text{logit}(0.9), 1)$

The pdf of F is displayed in Figure 1. We see that this distribution is very different from a log-normal distribution. We can then expect to be able to estimate the population parameters.

We used the SAEM algorithm implemented in Monolix for computing the ML estimate of the population parameters and their standard errors. Table 1 shows that population parameters are indeed very well estimated in this example.

Even if the population parameters are “almost” perfectly estimated, the individual parameters cannot be estimated very precisely. Figure 2 compares the simulated individual parameters, considered here as the “true” values, with the Maximum a Posteriori (MAP) estimates, i.e. the modes of the conditional distributions $\mathbf{p}(V_i|y_i, \hat{\theta})$ and $\mathbf{p}(F_i|y_i, \hat{\theta})$ for $i = 1, 2, \dots, N$. On the other hand, the ratio F_i/V_i is estimated very accurately.

In this example, the likelihood has a maximum which is very well defined and SAEM converges easily even with poor initial guesses. Figure 3 displays the convergence of 5 runs of SAEM obtained with different initial values. Thus, thanks to the design (i.e. a very large number of subjects and a large number of measurement per subject) and thanks to the probability distribution of the individual parameters, the model can be considered as “practically identifiable”.

Estimations obtained with NONMEM FOCE are $\hat{F} = 0.764$, $\hat{k}_a = 1.00$, $\hat{V} = 8.88$, $\hat{k} = 0.100$. We can see that FOCE introduces some bias in the estimation of F and V . Indeed, it is not possible to estimate correctly the individual parameters since the model is not structurally identifiable. Then,

any method based on the estimation of the individual parameters cannot work as well as maximum likelihood estimation.

Model B: $\text{logit}(F) \sim \mathcal{N}(\text{logit}(0.4), 0.2^2)$

Things will change with this distribution for F . Indeed, we can see Figure 4 that this distribution is now very close to a log-normal one. Even if the model remains identifiable in theory, we cannot expect anymore a good estimation of the population parameters.

Table 2 displays the results obtained with a single run of SAEM. We see that population parameters F_{pop} and V_{pop} are poorly estimated with this run. The ratio $F_{\text{pop}}/V_{\text{pop}}$ remains very well estimated (0.401 instead of 0.4) as well as the total variance $\omega_F^2 + \omega_V^2$ (0.139 instead of 0.13). We also notice a clear degradation of the results obtained with FOCE ($\hat{F}_{\text{pop}} = 0.830$, $\hat{V}_{\text{pop}} = 20.8$, $\hat{\omega}_F = 0.56$, $\hat{\omega}_V = 0.3$).

This poor estimation of the population parameters leads to a misspecified population distribution and a bias in the estimation of the individual parameters. Figure 5 shows that the V_i and the F_i are underestimated. Nevertheless, ratios $R_i = V_i/F_i$ are correctly estimated.

In this example and because of the lack of identifiability of some parameters, the likelihood does not exhibit a unique isolated maximum. Figure 6 shows that the convergence of SAEM strongly depends on the initial value. The very high correlation (0.9989) between the estimates of F_{pop} and V_{pop} confirms that we should not rely on the estimated values of these parameters.

When such lack of practical identifiability is revealed, a solution with this example consists in fixing either F_{pop} or V_{pop} . A less radical solution may consist in introducing a *prior* information on F_{pop} or V_{pop} . If we introduce, for instance, a logit-normal distribution for F_{pop} , with mean $\text{logit}(0.4)$ and standard deviation 0.2, then the population parameters are correctly estimated ($\hat{F}_{\text{pop}} = 0.41$, $\hat{V}_{\text{pop}} = 10.1$) as well as the individual parameters.

We should notice that removing the inter-individual variability of F_i is not a solution since that makes the model non identifiable. Indeed, if $F_i = F_{\text{pop}}$, then V_i/F_i follows a log-normal distribution with mean $\log(V_{\text{pop}}/F_{\text{pop}})$ (in the log domain) and standard deviation ω_F . Then, only the ratio $V_{\text{pop}}/F_{\text{pop}}$ is identifiable in this model.

5 Conclusions

We have shown that models that are non-identifiable at an individual level may become identifiable at the population level under conditions in which the probabilistic models differ between alternate models. This requires strong assumptions about the probabilistic models which may be difficult to validate in practice. Similarly, even if the models are identifiable at the individual level it may prove difficult to estimate the parameters of the model unless supported by good experimental design. From a pharmacokinetic point of view it means that the differences between individuals can break the non-identifiability seen at the population level and this may allow better mechanistic understanding of the interindividual differences in pharmacokinetics.

We have mainly considered here the most theoretical aspects of the identifiability of a model. For the numerical examples, we have been using an EM-like algorithm, assuming that the maximum likelihood estimate of the population parameters could be computed. Our first partial results suggest that linearization methods (FO, FOCE) are more sensitive to a lack of identifiability than maximum likelihood methods with no approximation on the model. A detailed discussion around the impact of the estimation method - and its implementation in a software tool - on the results is beyond the scope of this paper. Such discussion as well as practical suggestions on how to proceed when the model shows signs of un-identifiability clearly deserve to be the subject of further work.

6 Acknowledgments

The research leading to these results received support from the Innovative Medicines Initiative Joint Undertaking under grant agreement 115156, resources of which are composed of financial contributions from the European Union's Seventh Framework Programme (FP7/2007-2013) and EFPIA companies' in kind contribution. The DDMoRe project is also supported by financial contribution from Academic and SME partners.

References

1. Bellman, R., Åström, K.J.: On structural identifiability. *Mathematical Biosciences* **7**(3), 329–339 (1970)
2. Bonate, P.L.: *Pharmacokinetic-pharmacodynamic modeling and simulation*. Springer (2011)
3. Brun, R., Reichert, P., Künsch, H.R.: Practical identifiability analysis of large environmental simulation models. *Water Resources Research* **37**(4), 1015–1030 (2001)
4. Chappell, M.J., Godfrey, K.R., Vajda, S.: Global identifiability of the parameters of nonlinear systems with specified inputs: a comparison of methods. *Mathematical Biosciences* **102**(1), 41–73 (1990)
5. Cobelli, C., Distefano 3rd, J.J.: Parameter and structural identifiability concepts and ambiguities: a critical review and analysis. *American Journal of Physiology-Regulatory, Integrative and Comparative Physiology* **239**(1), R7–R24 (1980)
6. Evans, N.D., Godfrey, K.R., Chapman, M.J., Chappell, M.J., Aarons, L., Duffull, S.B.: An identifiability analysis of a parent-metabolite pharmacokinetic model for ivabradine. *Journal of pharmacokinetics and pharmacodynamics* **28**(1), 93–105 (2001)
7. Fröhlich, F., Theis, F.J., Hasenauer, J.: Uncertainty analysis for non-identifiable dynamical systems: Profile likelihoods, bootstrapping and more. In: *Computational Methods in Systems Biology*, pp. 61–72. Springer (2014)
8. Garcia, R.I., Ibrahim, J.G., Wambaugh, J.F., Kenyon, E.M., Setzer, R.W.: Identifiability of PBPK models with applications to dimethylarsinic acid exposure. *Journal of pharmacokinetics and pharmacodynamics* **42**(6), 591–609 (2015)
9. Gargash, B., Mital, D.: A necessary and sufficient condition of global structural identifiability of compartmental models. *Computers in biology and medicine* **10**(4), 237–242 (1980)
10. Godfrey, K.R., Chapman, M.J., Vajda, S.: Identifiability and indistinguishability of nonlinear pharmacokinetic models. *Journal of pharmacokinetics and biopharmaceutics* **22**(3), 229–251 (1994)
11. Guedj, J., Thiébaud, R., Commenges, D.: Practical identifiability of hiv dynamics models. *Bulletin of mathematical biology* **69**(8), 2493–2513 (2007)
12. Lavielle, M.: *Mixed Effects Models for the Population Approach: Models, Tasks, Methods and Tools*. Chapman and Hall/CRC (2014)
13. Petersen, B., Gernaey, K., Vanrolleghem, P.A.: Practical identifiability of model parameters by combined respirometric-titrimetric measurements. *Water Science and Technology* **43**(7), 347–356 (2001)
14. Raue, A., Kreutz, C., Maiwald, T., Bachmann, J., Schilling, M., Klingmüller, U., Timmer, J.: Structural and practical identifiability analysis of partially observed dynamical models by exploiting the profile likelihood. *Bioinformatics* **25**(15), 1923–1929 (2009)
15. Shivva, V., Korell, J., Tucker, I., Duffull, S.: An approach for identifiability of population pharmacokinetic-pharmacodynamic models. *CPT: Pharmacometrics & Systems Pharmacology* **2**(6), e49 (2013)
16. Shivva, V., Korell, J., Tucker, I.G., Duffull, S.B.: Parameterisation affects identifiability of population models. *Journal of pharmacokinetics and pharmacodynamics* **41**(1), 81–86 (2014)
17. Tikhonov A., N., Goncharsky, A., Stepanov V., V., Yagola A., G.: *Numerical methods for the solution of ill-posed problems*. Springer Science & Business Media (1995)
18. Walter, E., Pronzato, L.: On the identifiability and distinguishability of nonlinear parametric models. *Mathematics and Computers in Simulation* **42**(2), 125–134 (1996)
19. Wang, W., et al.: Identifiability of linear mixed effects models. *Electronic Journal of Statistics* **7**, 244–263 (2013)
20. Wu, L.: *Mixed effects models for complex data*. CRC Press (2010)
21. Xia, X., Moog, C.H.: Identifiability of nonlinear systems with application to hiv/aids models. *Automatic Control, IEEE Transactions on* **48**(2), 330–336 (2003)
22. Yates, J., Jones, R., Walker, M., Cheung, S.: Structural identifiability and indistinguishability of compartmental models. *Expert opinion on drug metabolism & toxicology* **5**(3), 295–302 (2009)

parameter	true value	initial value	estimation	relative s.e. (%)
F_{pop}	0.9	0.4	0.909	0.8
$k_{a,\text{pop}}$	1	0.5	1.006	0.4
V_{pop}	10	5	10.188	0.9
k_{pop}	0.1	0.3	0.100	0.2
ω_F	1	3	1.007	7.0
ω_{k_a}	0.25	1	0.249	1.3
ω_V	0.30	1	0.306	1.4
ω_k	0.15	1	0.148	1.1
σ	0.10	1	0.100	0.3

Table 1 Model A: ML estimates of the population parameters.

parameter	true value	initial value	estimation	relative s.e. (%)
F_{pop}	0.4	0.2	0.245	9.9
$k_{a,\text{pop}}$	1	2	0.999	0.40
V_{pop}	10	5	6.11	9.8
k_{pop}	0.1	0.05	0.10	0.2
ω_F	0.2	1	0.289	24.0
ω_{k_a}	0.25	1	0.249	1.3
ω_V	0.30	1	0.236	20.5
ω_k	0.15	1	0.148	1.1
σ	0.10	1	0.01	0.3

Table 2 Model B: ML estimates of the population parameters.

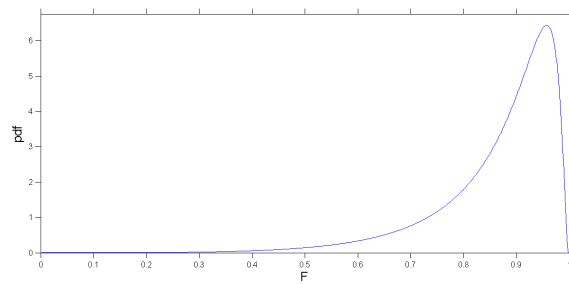


Fig. 1 Model A: pdf of the logit-normal distribution with parameters $(\text{logit}(0.9), 1)$.

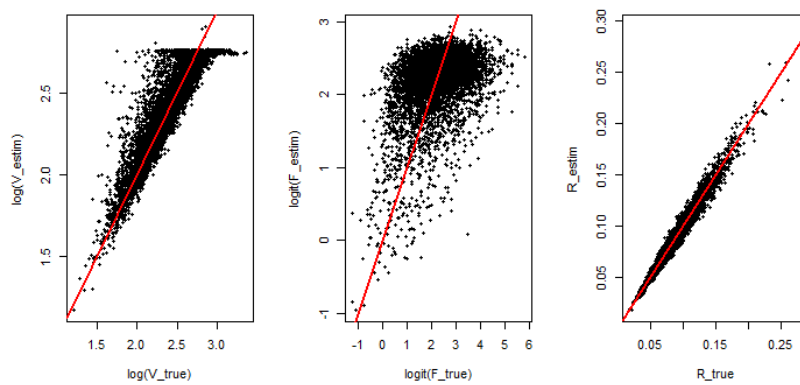


Fig. 2 Model A: estimated values versus true values of $\log(V_i)$, $\log(F_i)$ and $R_i = V_i/F_i$.

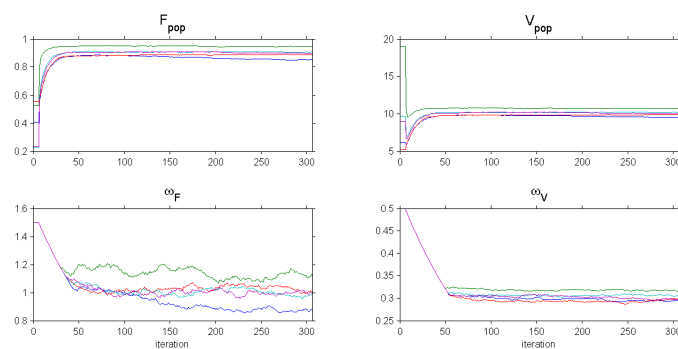


Fig. 3 Model A: convergence of 5 runs of SAEM obtained with different initial values.

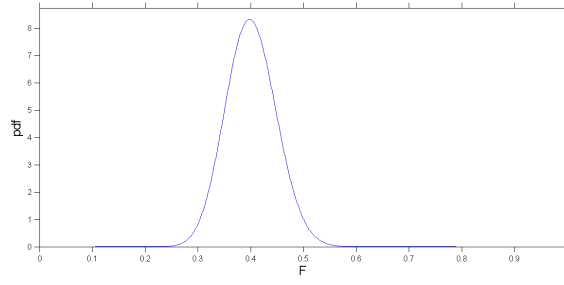


Fig. 4 Model B: pdf of the logit-normal distribution with parameters $(\text{logit}(0.4), 0.2^2)$.

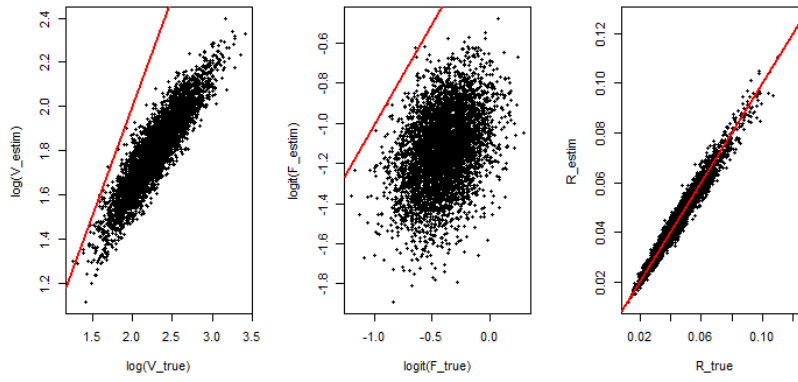


Fig. 5 Model B: estimated values versus true values of $\log(V_i)$, $\text{logit}(F_i)$ and $R_i = V_i/F_i$.

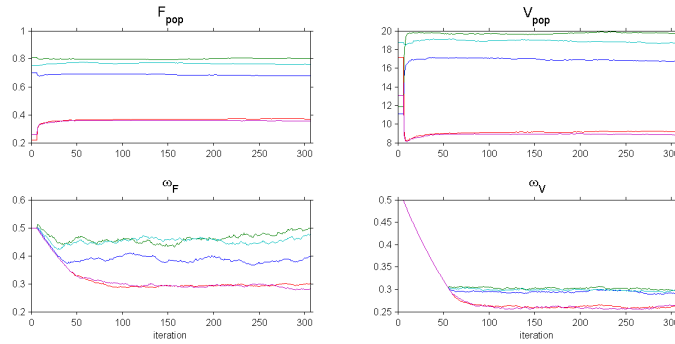


Fig. 6 Model B: convergence of 5 runs of SAEM obtained with different initial values.