

INFORMAČNÍ BULLETIN



České statistické společnosti

Ročník 29, číslo 2, červen 2018

MODELÝ SMĚSÍ RIZIKOVÝCH FUNKCÍ S APLIKACÍ V ANALÝZE SPOLEHLIVOSTI MODELS OF MIXTURES OF HAZARD RATES WITH APPLICATION TO RELIABILITY ANALYSIS

Petr Volf

Adresa: ÚTIA AV ČR, Pod vodárenskou věží 4, 182 00 Praha 8

E-mail: volf@utia.cas.cz

Abstrakt: V příspěvku jsou studovány modely pro analýzu životnosti sestavené jako kombinace několika jednoduchých rizikových funkcí. Užitečnost takových modelů je ukázána na příkladech s umělými i reálnými daty.

Klíčová slova: riziková funkce, Weibullovo rozdělení, analýza spolehlivosti.

Abstract: The objective of this contribution is to present and explore some methods of construction and estimation of models based on mixtures of hazard rates. The goal is to demonstrate their usefulness and applicability with the aid of both artificial and real examples.

Keywords: hazard rate, Weibull distribution, reliability analysis.

1. Introduction, models based on mixtures

Distributions of probability constructed as mixtures of two or more simple distributions are used frequently in many fields of application, in cases when no simple model describes the data sufficiently. Such mixtures are based on the convex combination of probability densities or distribution functions, in order to ensure that the final function is also a density (or distribution function). The most popular model is composed from a set of normal densities. The simplest example is the following:

$$f(x) = p_1 \cdot f_1(x) + p_2 \cdot f_2(x). \quad (1)$$

Interpretation (used also in cluster analysis) could be such that with probability p_1 an item belongs to first group having distribution with density f_1 , with probability p_2 to the second group. It is also the way how data representing mixture (1) could be generated. Further, probability densities (and normal densities in particular) are used also as components for the construction of regression models, in the same way as regression splines. For instance a combination (now it need not be convex) of gaussians (in this context called also “radial basis functions”) is used to create a curve or surface.

In the statistical reliability analysis, when modeling the time to failure, one of often used characteristics is the hazard rate (HR) $h(t)$ or its integrated version, the cumulative hazard rate (CHR) $H(t) = \int_0^t h(z) dz$. Let us denote $F(t)$ the distribution function and $S(t) = 1 - F(t)$ the survival or reliability function, then $H(t) = -\ln(S(t))$. The use of hazard rate as one of basic characteristics suggests that the mixture models in reliability could be based as well on mixtures of hazard rates instead on mixtures of distributions.

The first step, namely the linear failure rate model $h(t) = a + bt$ was proposed already in Kodlin [5], polynomial failure rate models have then been studied in Bain [1], advanced computation methods also in Pandey et al. [7]. Nonlinear (and no polynomial) models $h(t) = a + bt^\gamma$, $a, b, \gamma > 0$ and methods of estimation have, naturally, also been examined, for instance in Salem [9]. Some systematic recent investigation is also due Tien Thanh Thach from the Technical University in Ostrava, see for instance Briš and Thach [3]. Essentially, there is no problem to extend the methods to deal with more than two components. A question arises, however, whether and when it is reasonable, improving the fit of model to data significantly. In fact, both examples used in the present paper are of such a kind, in both it is shown that two components of additive hazard rate do not suffice.

The next section introduces hazard mixture models in more details. Then, a comparison of two examples illustrates the difference among mixtures of distributions and of hazard rates. Section 4 solves an artificial example leading to the bath-tube shaped hazard rate, and, finally, Sections 5 and 6 solve a real data case, except a plain sum of hazard rate considering also an incremental model, i.e. a variant of model with change points.

2. Mixtures of hazard rates

In practical survival data analysis, rather frequently the hazard rate during the lifetime (not only of a technical device, but also of a biological object) may have a “bath-tube” shape, decreasing in the first short lifetime period, then being approximately constant for the most of lifetime, finally increasing as the object is ageing. This could be a reason for considering a mixture of three simple hazard rates (see an example in Figure 2 below). Notice that the mixture of hazard rates need not be convex, as hazard rate is not standardized (in the sense as the density function is, i.e. that the integral from density function equals one). Notice also a special property of Weibull distribution: When its hazard rate is multiplied by a constant, we obtain the hazard rate of another Weibull distribution. Recall that the Weibull cumulative hazard

rate can have two forms,

$$H(x) = a x^\beta \quad \text{or} \quad H(x) = \left(\frac{x}{\alpha}\right)^\beta.$$

Hence

$$c H(x) = c a x^\beta = a^* x^\beta \quad \text{or} \quad H(x) = \left(\frac{x}{\alpha^*}\right)^\beta, \quad \alpha^* = \alpha/c^{1/\beta}. \quad (2)$$

I shall use mostly the second notation, as it corresponds to the notation in Matlab (and in Excel, too). From above it is seen that the mixture of two Weibull hazard rates is equivalent to a simple sum $h = h_1 + h_2$ of other Weibull hazard rates, parameters of would-be combination are not identifiable.

However, other distribution types have not such property, hence it has a sense to consider a model – combination (possibly convex) of hazard rates. What is an interpretation of such a model? The sum of hazard rates $h = \sum h_j$ corresponds to the hazard rate of a serial system composed from independent items with h_j , the survival time of the system then to minimum of survival times of components. Multiplication $h^* = c \cdot h$ then corresponds to a proportional change of hazard rate (it is actually the first step to the “proportional hazard regression model” or to frailty models), then $S^* = S^c$.

Immediately a question arises when one should prefer the model based on density function or on hazard rate. Naturally, each approach has its advantages and also specific tools, model fitting methods and procedures. In general, in both cases the model estimation often uses iterative methods of optimization of a criterion based for instance on the maximal likelihood, on Bayes estimation (nowadays often connected with the Markov chain Monte Carlo method, MCMC), or on some other distance. For instance, one of traditional methods how to fit the Weibull distribution is based on the least squares and linear regression. Namely, let theoretical Weibull cumulative hazard rate be $H(t) = a \cdot t^\beta$, data T_i , $i = 1, \dots, N$, estimated CHR $\hat{H}(t)$. Then, after further logarithmization, we obtain the relation

$$\ln \hat{H}(T_i) \sim \ln(a) + \beta \cdot \ln(T_i).$$

Or, when dealing with additive hazard rate composed from two Weibull hazard rates, we can use the following:

$$\hat{H}(T_i) \sim a_1 T_i^{\beta_1} + a_2 T_i^{\beta_2},$$

which is linear at least w.r. to a_j -s. The solution is found with the aid of the least squares method, eventually weighted by the asymptotic variance

of $\hat{H}$. Other criteria can use distances of Kolmogorov-Smirnov and variants (Cramér-Von Mises, Anderson-Darling) measuring departures of parametric models from nonparametric estimates of $S(t)$ or $H(t)$.

3. Two examples of mixture models

The first example on Figure 1 shows a case of mixture of two probability densities. Such a model is applicable for instance in the analysis of unemployment duration, as, roughly, there are two types of people: one with certain qualification and willingness to find a new employment, so called “movers”, here represented by shorter distribution with density f_1 . The other type, called “stayers”, however, have problems with finding a new job, their staying time in unemployment is represented by density f_2 . Hence, considering these two groups together and their proportions, the density of common distribution of unemployment time is given by a mixture $f = p \cdot f_1 + (1 - p) \cdot f_2$. In Figure 1 f_1 corresponds to Weibull distribution with $\alpha_1 = 50$, $\beta_1 = 1.5$, f_2 to Weibull with $\alpha_2 = 150$, $\beta_2 = 3$, $p = 0.5$.

The second example displayed in Figure 2 shows an instance of the hazard rate having “bath-tube” shape. It was constructed by mixing a decreasing HR of Weibull distribution with parameters $\alpha_1 = 1$, $\beta_1 = 0.5$, constant HR (i.e. of

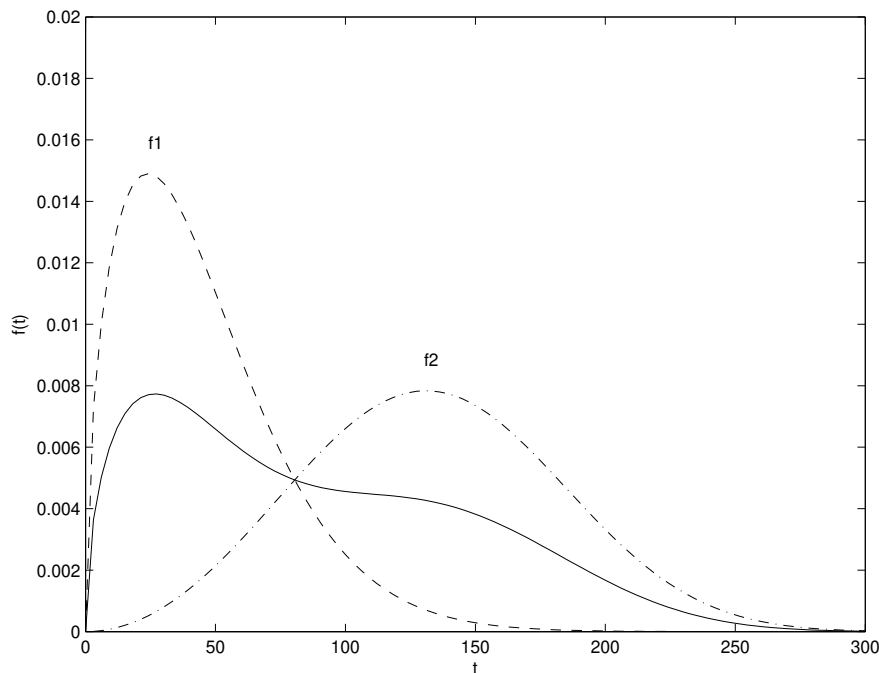


Figure 1: Example of mixture of 2 densities representing “movers”, f_1 , and “stayers”, f_2 , resulting $f = (f_1 + f_2)/2$ (thick curve).

exponential distribution, $\alpha_2 = 30, \beta_2 = 1$) and an increasing HR of Weibull distribution with $\alpha_3 = 150, \beta_3 = 5$. The final hazard rate was obtained as $h = p_1 h_1 + p_2 h_2 + p_3 h_3$ with $p_1 = 0.2, p_2 = 0.5, p_3 = 0.3$. Results are displayed in Figure 2. As it has already been said, the model is equivalent to a simple additive one $h = h_1^* + h_2^* + h_3^*$, namely with $\alpha_1^* = 25, \alpha_2^* = 60, \alpha_3^* = 190.8$ and β -s the same as above.

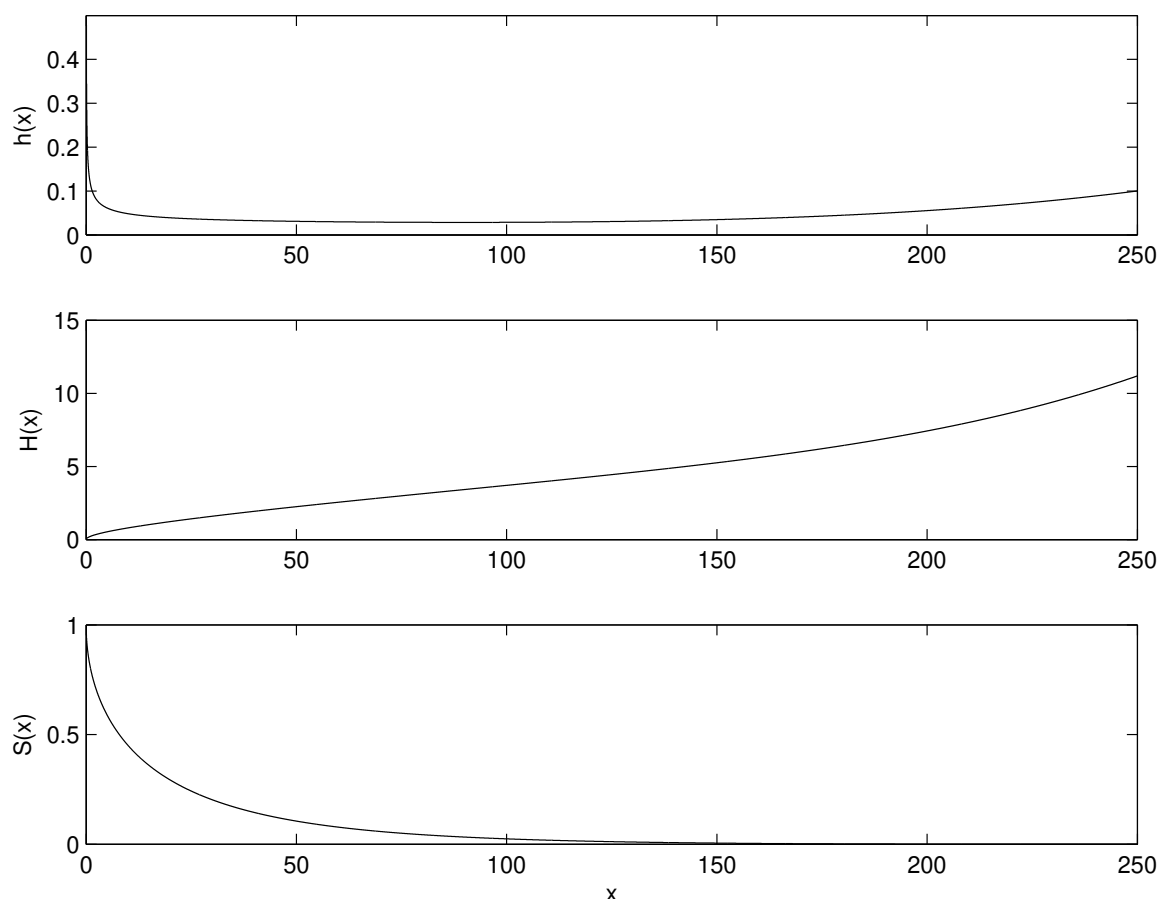


Figure 2: Example of “bath-tubed” hazard rate shape, a mixture of 3 hazard rates.

4. Artificial data example

The data $X_i, i = 1, \dots, N, N = 500$, were generated from the model displayed on Figure 2. Hence, it was expected that its HR is given by the sum of two Weibull and one exponential components. Figure 3 shows nonparametric estimates of the CHR and survival function. The task is to estimate parametric model. From several methods listed above the maximum likelihood estimate (MLE) was chosen. It is possible, for given parameters, to

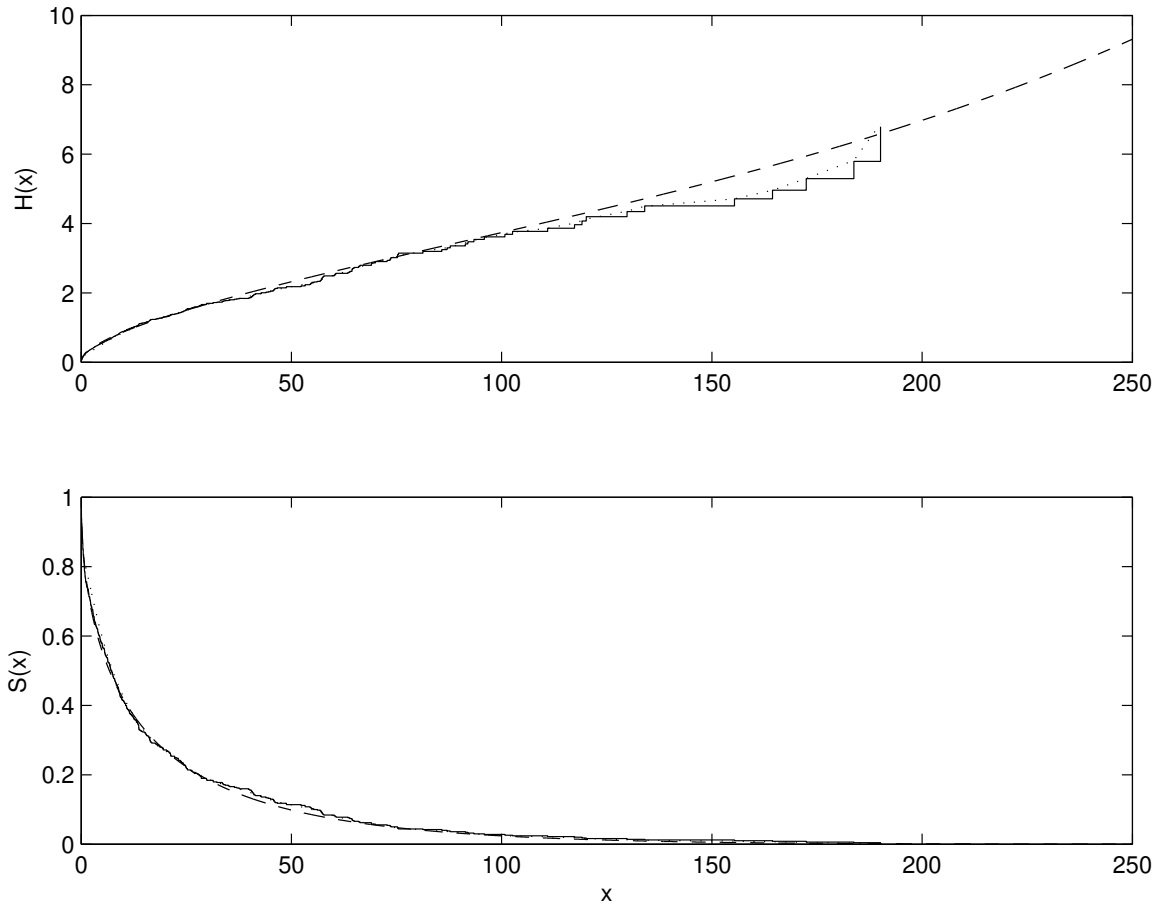


Figure 3: Stepwise: empirical CHR and survival function, dotted are their kernel-smoothed versions, dashed curves are based on the MLE.

evaluate both the log-likelihood as well as its derivatives. On the other hand, parameters maximizing the log-likelihood are found just by a random search (which is here computationally easier than for instance an iteration using the derivatives repeatedly, like the Newton-Raphson algorithm). Therefore the result is just approximate, as well the estimates of borders of 90% confidence intervals. Dashed curves in Figure 3 correspond to the model using the MLE of parameters, their values are listed in Table 1.

5. Real data example

The data used in this part are “Aircraft windshield failure data” taken from Ruhi [8], they were analyzed also in several other papers, e.g. in Blischke et al. [2]. They concern to the “damage or delamination of the nonstructural outer ply or failure of its heating system”. These failures do not cause a severe damage to the aircraft but lead to replacement of the windshield. The

	α_1	α_2	α_3	β_1	β_3
MLE	18.0308	78.1834	197.4708	0.5047	3.6122
LCI	10.8970	42.4971	148.4805	0.4549	2.2910
UCI	25.1646	113.8697	246.4611	0.5545	4.9334

Table 1: ML estimates, LCI and UCI are the lower and upper bounds of asymptotic 90% confidence intervals.

	α_1	α_2	β_1	β_2
MLE	39.3812	3.5834	0.9571	2.9336
LCI	38.0081	2.9995	0.5560	1.9097
PCI	40.7543	4.1673	1.3582	3.9575

Table 2: ML estimates, LCI and UCI are again the lower and upper bounds of asymptotic 90% confidence intervals.

data consist of 153 observations, from them 88 are times to failures, the remaining 65 are censored times when no failure occurred during the time of observation. It means that we deal with the random right-censoring scheme. Hence the PLE (Product Limit Estimator) of Kaplan and Meier will be used as nonparametric estimator of survival function, while cumulative hazard rate will be estimated by the Nelson-Aalen Estimator (NAE), see e.g. Kalbfleisch and Prentice [4]. The unit of measurement was 1000 hours.

Both nonparametric estimates are displayed in Figure 4 (stepwise functions), together with estimates obtained from them by kernel smoothing. Ruhi [8] as well as other authors (references see in Ruhi) tried to describe the data by models mixing two plain distribution densities. Success of modelling was assessed by several criteria, as the likelihood or the Kolmogorov-Smirnov distance of model survival function from the nonparametric PLE. Naturally, these two criteria lead to different results, though both are asymptotically consistent. For comparison, I decided to use a model based on mixture of hazard rates with its parameters estimated by the maximum likelihood method. Figure 5 shows the result (dashed curve), namely the model constructed just by sum of two Weibull hazard rates. Such a model has 4 parameters. Similarly as above, the best solution was approached by a random search, then asymptotic confidence intervals were computed from the second derivative

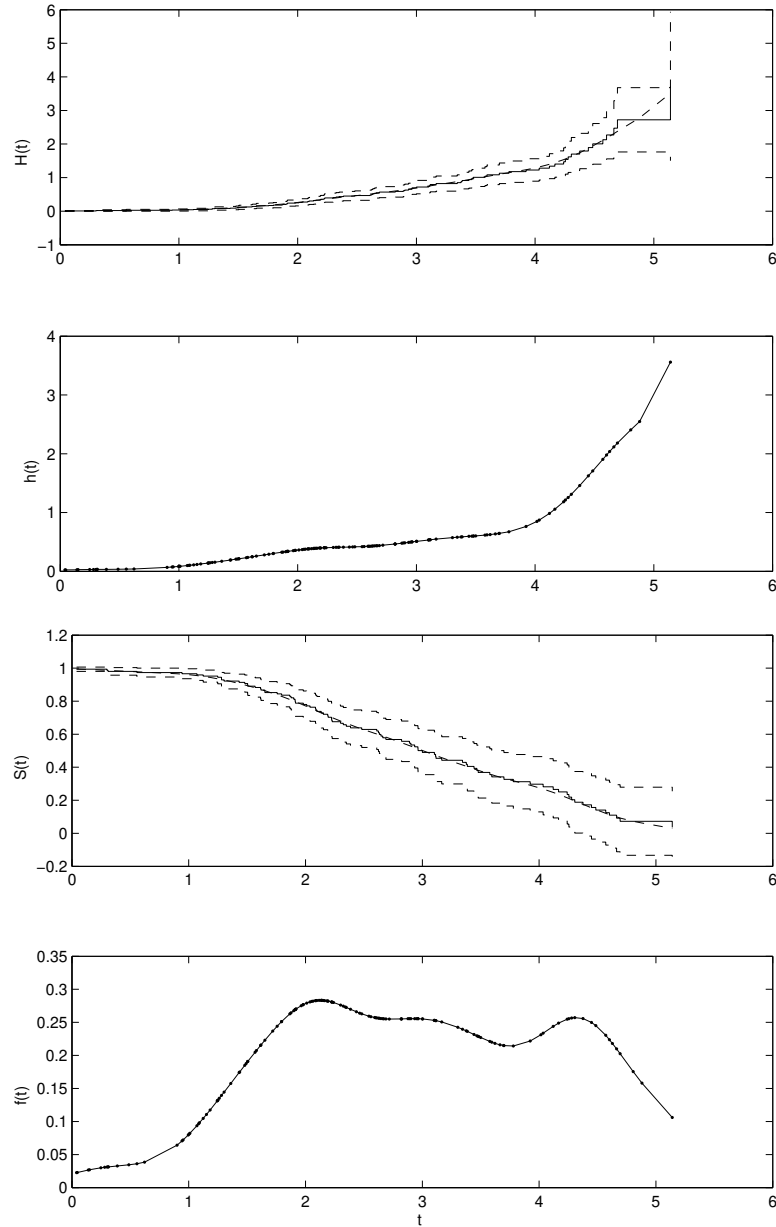


Figure 4: From above: the NAE with 95% CI-s, kernel-smoothed estimate of HR, then the PLE with 95% CI-s, kernel-smoothed estimate of density, data of Ruhi (2015).

of the log-likelihood. Figure shows that the fit should be improved yet. For comparison, achieved maximal value of the log-likelihood was -170.14 , while the best result of Ruhi [8] was -176.7 obtained by the mixture of densities of Weibull and Normal distributions having 4 parameters of components plus one of mixture.

Estimated parameters of the sum of two Weibull hazard rates are in Table 2. It is seen that the first distribution can be taken as the exponential one,

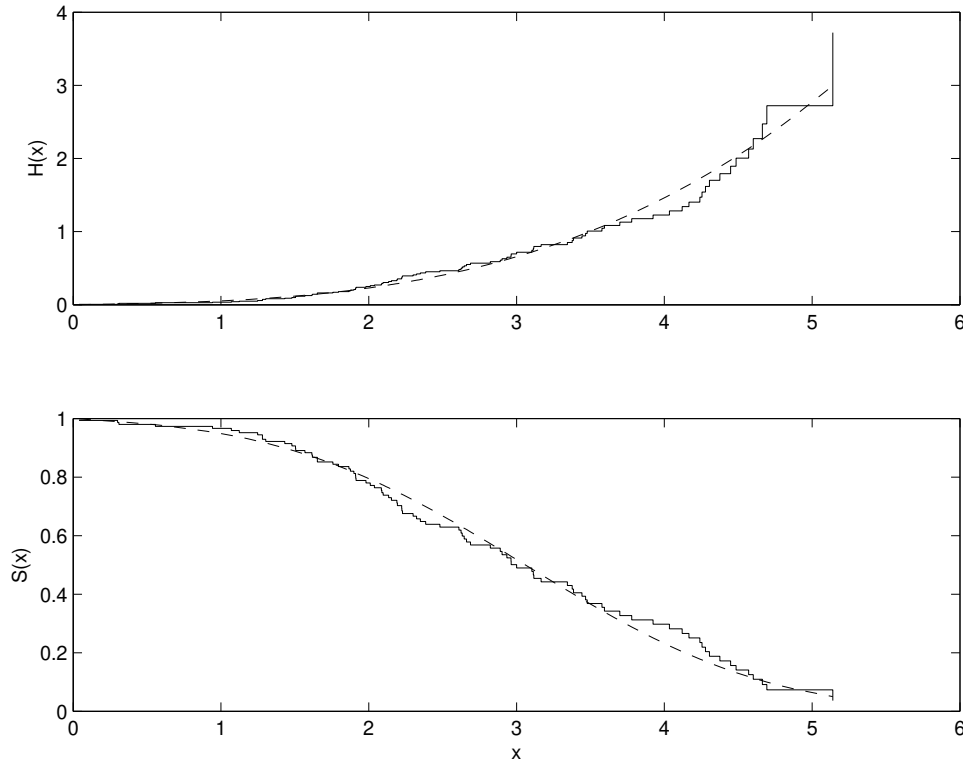


Figure 5: CHR above, with stepwise NAE, survival functions are below, with stepwise PLE. Dotted are kernel-smoothed nonparametric estimates, dashed curves correspond to mixed hazard rates model obtained by the MLE.

i.e. with $\beta_1 = 1$. Then the MLE of other parameters is comparable, as well as maximal log-likelihood. Now the model has just 3 parameters estimated as $\alpha_1 = 38.5612$, $\alpha_2 = 3.5717$, $\beta_2 = 2.9590$, with maximal log-likelihood -170.16 .

Blischke [2] as well as Ruhi [8] considered also the Akaike information criterion (AIC) for comparison of different models performance. Let us recall that $AIC = 2k - 2L$, where k is the number of model parameters and L equals achieved maximum of the log-likelihood, the model with smaller AIC is regarded as being better. Hence, in the case of our model with $\beta_1 = 1$ $AIC = 346.3$, which is less than the AIC of all models considered in Blischke or Ruhi. The model has the form of the “nonlinear failure rate model” with $h(t) = a + bt^\gamma$, where $a = 1/\alpha_1 = 0.0259$, $b = \beta_2/(\alpha_2^{\beta_2}) = 0.0684$, $\gamma = \beta_2 - 1 = 1.9590$.

As it has been said, there are also other model fit methods available. For instance the use of the MCMC methods can improve the random search results, it also offers Bayes credibility intervals for parameters. Naturally, one may consider mixtures with also other hazard rates types, not limiting

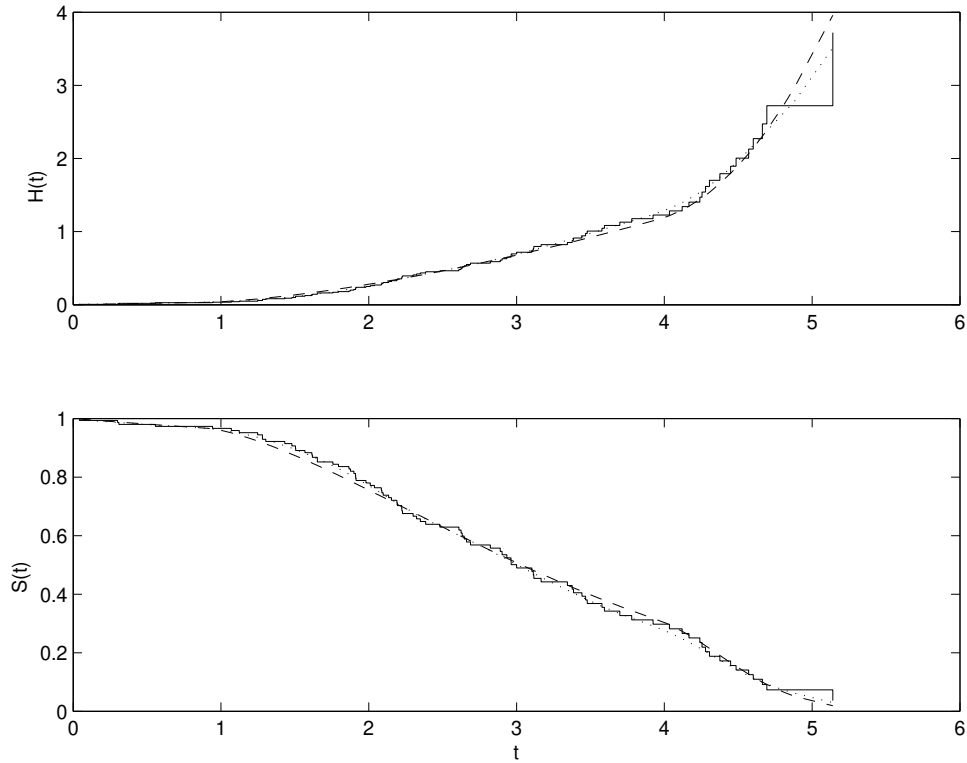


Figure 6: CHR above, with stepwise NAE, survival functions are below, with stepwise PLE. Dotted are kernel-smoothed nonparametric estimates, dashed curves corresponding to change-point model obtained by the MCMC.

himself just to Weibull, or to use more than two components. In fact, kernel estimate of density functions in Figure 4 below suggests that 3 component could be optimal. That is why the same data were analyzed also in the next section, with, it seems, a better result.

6. Change-point models for hazard rates

The task of change point detection belongs among quite well developed statistical techniques. There exist also results dealing with changes of hazard rates. While corresponding asymptotic theory connected with the MLE is rather complicated, cf. Nguyen et al. [6], practical analysis can be based on rather simple methods detecting a region where the residuals (i.e. reasonably defined deviation of data from actual model) are crossing a given border. Our problem could be viewed also as an incremental construction of a signal model, which means that in detected point of change a new component is added to current model. This is possible when the change leads to the increase of hazard.

Param.	α_1	α_2	α_3	β_1	β_2	β_3	T_{ch1}	T_{ch2}
Estim.	31.384	2.996	0.834	0.957	1.555	2.053	0.907	3.952
LCrI	20.792	2.400	0.760	0.762	1.304	2.118	0.754	3.075
UCrI	39.476	3.338	1.891	1.196	1.925	4.745	1.235	4.226
$\beta_1 = 1$	37.705	2.883	0.877	1.000	1.811	2.273	0.835	3.914

Table 3: Estimates of incremental model parameters and change-points, with LCrI and UCrI the lower and upper bounds of sample-based 90% credibility intervals. The last row – model with $\beta_1 = 1$.

The example shown in Figure 6 deals still with the data from Ruhi [8]. Figure displays again stepwise estimates (thick) and dotted smoothed estimates, the NAE of the CHR above and the PLE of survival function below. Further, dashed curves show the best “two change-points” model achieved. Namely, the model of hazard rate was constructed from one Weibull hazard rate starting from $t = 0$ and two Weibull hazard rates added at two times of change, T_{ch1} , T_{ch2} , which also had to be found optimally. Then $h(t) = h_1(t) + h_2(t - T_{ch1}) \cdot I[t > T_{ch1}] + h_3(t - T_{ch2}) \cdot I[t > T_{ch2}]$, $I[.]$ denotes the indicator function. Thus, the model had 8 parameters. The method of solution used the Metropolis MCMC algorithm generating a representation of Bayes posterior distribution of parameters, while their prior distributions were chosen to be independent uniform, in reasonable intervals. Therefore the posterior distribution was proportional to the likelihood. Together 50 000 iterations of the algorithm were performed, results were taken from last 20 000. Table 3 contains the “modes” of posterior distribution, i.e. the values maximizing the likelihood, and empirical quantiles representing the Bayes 90% credibility intervals. Achieved maximum of the log-likelihood was -165.5 .

Similarly as above it is seen that the first component is close to exponential. When we fixed $\beta_1 = 1$, the results were comparable, they are in the last row of Table 3. The max of log-likelihood was -165.6 . Model had just 7 parameters, with $AIC = 345.2$, which was even smaller than the best (regarding the AIC) result in the preceding section.

7. Concluding remarks

The problem of construction of hazard rate from additive components can be also interpreted as the problem of a regression model fitting the nonparamet-

ric estimate (e.g. smoothed from the NAE). To preserve some interpretation, the model should be constructed from a small set of given parametric function and kept non-negative.

In fact, the case studied here was simplified, just sums of Weibull hazard rates were considered. A next problem to explore is to use (convex) mixtures of hazard rates of other distributions and to study whether it is possible to estimate both parameters of distributions and coefficients of mixture and whether this task is unambiguous. Or, variantly, the flexibility of mixture models (based either on hazard rates or on probability densities) could be examined after the time is transformed (e.g. to the logarithmic scale).

Acknowledgements

The research was supported by the grant No. 18-02739S of the Grant Agency of the Czech Republic.

References

- [1] Bain L. J.: Analysis for the linear failure-rate life-testing distribution. *Technometrics*, 16, 551–559, 1974.
- [2] Blischke W. R., Karim M. R., Murthy D. P.: *Warranty data collection and analysis*. Springer, 2011.
- [3] Briš R., Thach T. T.: Bayesian approach to estimate the mixture of failure rate model. *Proceedings of the 1st International Conference on Applied Mathematics in Engineering and Reliability*, CRC Press, 9–17, 2016.
- [4] Kalbfleisch J. D., Prentice R. L.: *The Statistical Analysis of Failure Time Data*. Wiley, 2002.
- [5] Kodlin D.: A new response time distribution. *Biometrics*, 2, 227–239, 1967.
- [6] Nguyen H. T., Rogers G. S., Walker E. A.: Estimation in change-point hazard rate models. *Biometrika* 71, 299–304, 1984.
- [7] Pandey A., Singh A., Zimmer W. J.: Bayes estimation of the linear hazard-rate model. *IEEE Transactions on Reliability*, 42, 636–640, 1993.
- [8] Ruhi S.: Application of mixture models for analyzing reliability data: A case study. *Open Access Library Journal*, 2: e1815, 2015. <http://dx.doi.org/10.4236/oalib.1101815>
- [9] Salem S. A.: Bayesian estimation of a non-linear failure rate from censored samples type II. *Microelectronics and Reliability*, 32, 1385–1388, 1992.

POUŽITÍ BOOTSTRAPOVÝCH METOD PRO SROVNÁNÍ ZPŮSOBŮ LÉČBY KRITICKÉ KONČETINOVÉ ISCHEMIE

COMPARING METHODS OF TREATMENT OF CRITICAL LIMB ISCHEMIA: AN APPLICATION OF BOOTSTRAP METHODS

Lenka Příbylová

Adresa: Trh. 17. listopadu 15, 708 33 Ostrava-Poruba

E-mail: lenka.pribylova@vsb.cz

Abstrakt: V tomto příspěvku nastiňujeme přístup, jak porovnat způsoby léčby kritické končetinové ischemie. Náš přístup je založen na aplikaci bootstrapových metod s použitím softwaru R. Porovnáváme empirické hustoty a empirické distribuční funkce bootstrapových průměrů vybraných biochemických veličin ve třech léčebných skupinách. Dále porovnáváme 95% klasické a *t*-bootstrapové konfidenční intervaly pro očekávané hodnoty a 95% klasické a *t*-bootstrapové konfidenční intervaly pro populační mediány vybraných biochemických veličin v léčebných skupinách.

Klíčová slova: bootstrapové metody, kritická končetinová ischemie, kmenové buňky.

Abstract: In this paper, we discuss a way, how to compare methods of treatment of critical limb ischemia. Our approach is based on an application of bootstrap methods using software R. We compare empirical densities and empirical cumulative distribution functions of bootstrap means of several biochemical parameters in treatment groups. Further we compare 95% classical and *t*-bootstrap confidence intervals for expected values, 95% classical and *t*-bootstrap confidence intervals for medians, of several biochemical parameters in treatment groups.

Keywords: bootstrap methods, critical limb ischemia, stem cell.

1. Bootstrapové metody

Pojem *bootstrap* zavedl Bradley Efron v roce 1979. Podstatou bootstrapu¹ je vyšetřování chování statistiky odhadující skutečné rozdělení pomocí opakovaných náhodných výběrů z vhodné aproximace tohoto rozdělení. Apro-

¹Podrobněji je tato problematika zpracována např. v [1], [2].

ximaci lze provést pomocí empirické distribuční funkce, popřípadě parametricky, s parametry odhadnutými např. metodou maximální věrohodnosti.

Názvu bootstrap předcházelo označení *resampling* (převzorkování), které vystihuje podstatu konstruování bootstrapových výběrů. Proces tvorby bootstrapových výběrů se dá stručně popsat následujícím způsobem. Z malého vstupního datového souboru je vybrán soubor stejného anebo menšího rozsahu s opakováním. Postup se B -krát opakuje. Z nově vytvořených výběrů jsou vyčísleny realizace odhadu neznámého parametru θ . Množina realizací odhadu neznámého parametru θ nyní reprezentuje datový soubor většího rozsahu, jehož rozsah je dán počtem bootstrapových výběrů. Užitím metod elementární teorie pravděpodobnosti lze dokázat, že nově vzniklý datový soubor obvykle dobře aproximuje rozdělení odhadovaného parametru původního souboru.

Bootstrapové metody jsou vhodné pro aplikaci na data malého rozsahu, protože tyto metody mají větší rychlost konvergence ke skutečnému rozdělení než při užití centrálních limitních vět. Při použití bootstrapových metod není nutné splnění předpokladu normality dat. Nevýhodou bootstrapových metod je skutečnost, že bootstrapové výběry mohou občas poskytovat zkreslený obraz reality, zvláště u dat s odlehlými pozorováními. Proto je nezbytné při aplikování bootstrapových metod zjistit, zda bootstrapový odhad není zkreslený díky odlehlému pozorování. Pro identifikaci odlehlých pozorování se v těchto případech používá metoda zvaná *jackknife*².

Bootstrapové metody mohou být aplikovány pro odhady parametrů, o nichž je známo, že splňují následující teoretické předpoklady. Odhady musí být konzistentní a asymptoticky stabilní. Tyto předpoklady splňuje většina známých statistik, jako je výběrový průměr, výběrový rozptyl, odhady parametrů lineárních a zobecněných lineárních modelů.

1.1. Parametrický bootstrap

Nechť $\mathbf{X} = (X_1, X_2, \dots, X_N)$ je náhodný výběr z rozdělení $F(\mathbf{X}, \theta)$. Předpokládáme parametrický tvar modelu. Vektor parametrů $\theta = (\theta_1, \theta_2, \dots, \theta_k)$ vhodně odhadneme vektorem statistik $\hat{\theta}$, např. metodou maximální věrohodnosti. Odhady vektoru parametrů θ jsou simulovány pomocí B bootstrapových výběrových souborů, kde $N < B \ll N^N$. Pro každý bootstrapový výběr je vypočítán odhad vektoru parametrů θ , který je označen

²Metodou jackknife je vypočten bootstrapový odhad statistiky pro každý soubor, který vznikne z původního datového souboru vynecháním jednoho pozorování. Pokud se bootstrapový odhad dané statistiky po vynechání nějakého pozorování výrazně odlišuje od ostatních, příslušné pozorování je klasifikováno jako odlehlé.

$\hat{\theta}_{N,i}^*$, $i = 1, 2, \dots, B$. Vypočtené odhady jsou použity pro aproximaci rozdělení odhadu vektoru parametrů θ .

1.2. Neparametrický bootstrap

Mějme náhodný výběr $\mathbf{X} = (X_1, X_2, \dots, X_N)$ z rozdělení s neznámou distribuční funkcí $F(\mathbf{X}, \theta)$. Pro jednoduchost popisu se nadále budeme soustředit na popis situace s jedním skalárním parametrem θ , který odhadujeme z dat. V případě neparametrického bootstrapu není známa apriorní informace o tvaru rozdělení F . Neznámá distribuční funkce F je aproximována empirickou distribuční funkcí $\hat{F}_N$.

Z náhodného výběru $\mathbf{X} = (X_1, X_2, \dots, X_N)$ je vybráno N prvků s opakováním (první bootstrapový výběr). Tento postup je opakován B -krát. Pro každý³ bootstrapový výběr $\mathbf{X}^* = (X_1^*, X_2^*, \dots, X_N^*)$ je vypočtena realizace statistiky $\hat{\theta}_{N,i}^* = \hat{\theta}_i\{X_1^*, X_2^*, \dots, X_N^*\}$, $i = 1, 2, \dots, B$. Tedy $\hat{\theta}_{N,i}^*$ označuje odhad parametru θ získaný z i -tého bootstrapového výběru. Pro B bootstrapových výběrů získáváme tedy B odhadů parametru θ . Z rozdělení odhadů $\hat{\theta}_{N,i}^*$ je aproximován odhad parametru θ původního rozdělení $F(\mathbf{X}, \theta)$.

Výše uvedené definice lze použít pro nezávislá data. V případě závislých dat je nutno použít párové bootstrapové výběry. [2]

1.3. Bootstrapové intervaly spolehlivosti

Existují tři typy bootstrapových konfidenčních intervalů: BC_α , ABC , t -bootstrapové intervaly. Vzhledem k cíli této práce se zaměříme pouze na t -bootstrapové intervaly.

Nechť $\mathbf{X}^* = (X_1^*, X_2^*, \dots, X_N^*)$ označuje bootstrapový náhodný výběr z rozdělení $F(\mathbf{X}, \theta)$, $\hat{\theta}_{N,i}^* = \hat{\theta}_i\{X_1^*, X_2^*, \dots, X_N^*\}$, $i = 1, \dots, B$.

Konstrukce t -bootstrapových intervalů spolehlivosti je založena na předpokladu, že rozdělení statistiky T^*

$$T^* = \frac{\hat{\theta}^* - \hat{\theta}}{\hat{\sigma}^*} \quad (1)$$

se příliš neliší od rozdělení statistiky T

$$T = \frac{\hat{\theta} - \theta}{\hat{\sigma}}, \quad (2)$$

³Při značení bootstrapového výběru bývá obvykle pro přehlednost zápisu vynechán index.

kde $\hat{\sigma}$ označuje odhad směrodatné odchylky $\hat{\theta}$ a $\hat{\sigma}^*$ jeho bootstrapový protějšek. Tvar t -bootstrapového $100(1 - \alpha)\%$ intervalu spolehlivosti pro odhad parametru θ původního rozdělení $F(X, \theta)$ lze vyjádřit následujícím vztahem

$$P\left(\hat{\theta} - t_{1-\frac{\alpha}{2}}^* \hat{\sigma} < \theta < \hat{\theta} - t_{\frac{\alpha}{2}}^* \hat{\sigma}\right) = 1 - \alpha, \quad (3)$$

kde t_{α}^* je α -kvantil statistiky T^* . Bootstrapové testování hypotéz je vyvíjeno jako rozšíření bootstrapových konfidenčních metod.

1.4. Některé výpočetní aspekty bootstrapových metod

Z teoretického hlediska by bylo vhodné provést všech $\binom{2N-1}{N}$ různých výběrů, resp. N^N všech možných výběrů. Již pro $N = 10$ je však počet různých náhodných výběrů roven asi 92 000 a počet všech možných výběrů 10 miliard. Při použití bootstrapových metod se tedy neprovádí všechny možné výběry, ale pouze náhodný výběr z těchto výběrů. Obvykle je na bootstrapový výběr aplikována metoda Monte Carlo, kde je B -krát opakovaně simulován výběr, přičemž hodnota počtu výběrů B závisí na metodě, kterou chceme použít. Pro určení konfidenčních intervalů je např. doporučován počet výběrů $B = N^2 \log(\log N)$.

Obecně je doporučováno minimálně $B = 1000$ bootstrapových výběrů pro odhad kvantilů a distribuční funkce. V případě výběrového průměru a podobných statistik lze obdržet dobré výsledky již při 500 opakováních. Počet B bootstrapových výběrů musí být větší než rozsah N dat. [3]

2. Charakteristika lékařské studie

Hlavním cílem léčby kritické končetinové ischemie je zlepšení průtoku krve nemocnou končetinou. Tradiční přístupy jsou založeny na procesu odstranění blokády pomocí chirurgického odstranění překážky, resp. pomocí vytvoření bypassu v příslušné cévě. Nový přístup léčby kritické končetinové ischemie je založen na aplikaci preparátu vlastních kmenových buněk pacienta. [5]

Studie zabývající se léčbou kritické končetinové ischemie byla realizována ve Fakultní nemocnici Ostrava. V rámci studie bylo rozděleno 53 pacientů do čtyř léčebných skupin. Data pacientů byla zaznamenána v rámci sedmi léčebných kontrol. Aplikace kmenových buněk byla realizována cestou do svalu v blízkosti poškozené tkáně (skupina A), nitrožilně (skupina B), resp. intraarteriálně (skupina C). Skupina D byla kontrolní skupinou pacientů, u nichž nebyl aplikován preparát kmenových buněk. Tito pacienti byli léčeni některým z tradičních způsobů léčby. Rozdělení pacientů do skupin bylo provedeno náhodně užitím příslušného software. Ve všech případech byly pacien-

tovi aplikovány vlastní kmenové buňky, které byly získány odběrem kostní dřeně. U kontrolní skupiny D, ve které byli pacienti léčeni standardními léčebnými postupy, nebyl proveden odběr kostní dřeně, tudíž u této skupiny chybí hematologické parametry získané rozbořem kostní dřeně.

Do studie byli zařazeni pacienti s diagnostikovanou kritickou ischemií způsobenou PAOD a diabetickou nohou s defektem končetiny ve stadiu Rutherford IV-V/Fontaine III-IV, tedy pacienti s diagnostikovaným diabetem mellitem II. typu, kteří byli na inzulinoterapii.

2.1. Cíle studie

Hlavním cílem lékařské studie bylo charakterizovat účinnost preparátu z kmenových buněk v jednotlivých léčebných skupinách. Stanovit, který způsob aplikace preparátu z kmenových buněk je nejúčinnější. Kritériem účinnosti byla změna hladin biochemických a hematologických veličin. Předpokládalo se, že aplikace preparátu z kmenových buněk nitrožilní cestou povede k omezení potřebné dávky inzulinu, snížení hladin glukózy, cholesterolu a zlepšení parametrů imunitní odpovědi oproti aplikaci do svalu nebo intraarteriální cestou. [6]

3. Použití bootstrapových metod pro porovnání způsobů léčby kritické končetinové ischemie

Rozsah vstupního datového souboru byl malý, tudíž nebylo možno použít elementární statistické metody, u nichž se předpokládá normalita dat. Přístup, který uvádíme v tomto článku je založen na *bootstrapových metodách*.

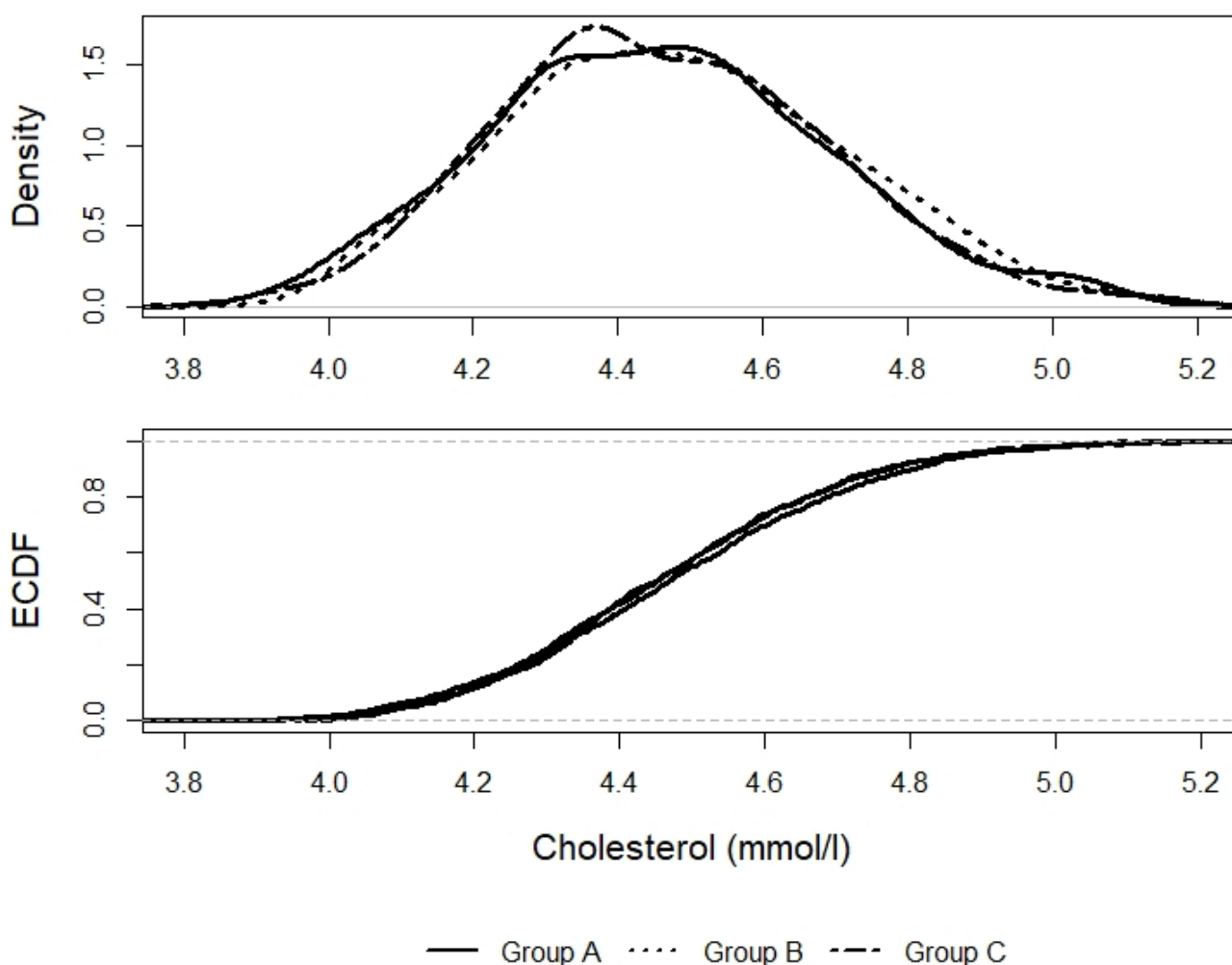
Pro porovnání účinnosti způsobu aplikace preparátu jsme použili dva přístupy založené na bootstrapových metodách. Prvním přístupem bylo srovnání empirických distribučních funkcí a empirických hustot bootstrapových průměrů změn vybraných biochemických veličin ve třech léčebných skupinách. Druhým přístupem bylo srovnání bootstrapových konfidenčních intervalů pro populační průměry a mediány změn těchto veličin. Z množství různých veličin, které byly monitorovány (krevní tlak měřený na palci dolní končetiny, LDP, tcpO, ABI, TP apod.), jsme pro ilustraci v této práci zvolili analýzu hladin cholesterolu v léčebných skupinách.

Před použitím bootstrapových metod bylo třeba zjistit, zda v souboru nejsou odlehlá pozorování. Pro identifikaci odlehlých pozorování jsme použili metodu jackknife. Pro dále sledovaný parametr cholesterol byla ve skupině A nalezena dvě odlehlá pozorování a ve skupině B jedno odlehlé pozorování.

3.1. Porovnání způsobů léčby pomocí empirické distribuční funkce bootstrapových průměrů

Užitím softwaru R jsme vytvořili 1000 bootstrapových výběrů zvolených biochemických veličin ve třech léčebných skupinách. Sestrojili jsme empirické hustoty a empirické distribuční funkce průměrů bootstrapových výběrů těchto veličin a porovnali jejich grafy. Pro ilustraci jsme na obrázku 1 uvedli grafy empirických hustot a empirických distribučních funkcí bootstrapových průměrů hladin cholesterolu.

Nebyl nalezen významný rozdíl mezi tvary empirických distribučních funkcí a empirických hustot bootstrapových průměrů hladin cholesterolu v léčebných skupinách. (Posouzení bylo provedeno vizuálně, na základě explorační analýzy.)



Obrázek 1: Empirické hustoty a ECDF bootstrapových průměrů

Pozn. Pokud by při opakovaných simulacích empirických distribučních funkcí a empirických hustot bootstrapových průměrů vycházely významně odlišné výsledky, ukazovala by tato skutečnost na nevhodnost použití bootstrapových metod pro daná data, resp. na nereprezentativnost dat.

3.2. Porovnání léčebných skupin pomocí bootstrapových konfidenčních intervalů

Pro analýzu srovnání účinnosti tří preparátů z kmenových buněk jsme v této kapitole použili dva přístupy. V 1. části jsme sestrojili 95% klasické a *t*-bootstrapové konfidenční intervaly pro očekávané hodnoty vybraných biochemických veličin v léčebných skupinách. V 2. části klasické a *t*-bootstrapové 95% konfidenční intervaly populačních mediánů hladiny těchto biochemických veličin.

V tabulce 1 jsou uvedeny 95% konfidenční intervaly pro populační průměry hladin cholesterolu. Před touto analýzou byla ze souboru odstraněna odlehlá pozorování identifikovaná metodou jackknife, neboť významně zkreslovala výsledky.

V druhé části analýzy byla v souboru pro srovnání ponechána odlehlá pozorování (jako obdoba neparametrického testu). V tabulce 2 jsou uvedeny 95% konfidenční intervaly pro populační mediány hladin cholesterolu.

Tabulka 1: 95% konf. int. pro očekávanou hladinu cholesterolu (mmol/l)

Skupina	Průměr	<i>t</i> -konf. int.	b-průměr	<i>t</i> -boot. int.
A	4,12	(3,88; 4,36)	4,12	(3,91; 4,33)
B	4,27	(3,62; 4,92)	4,32	(3,76; 4,88)
C	4,47	(4,04; 4,90)	4,44	(4,04; 4,84)

Tabulka 2: 95% konf. int. pro populační mediány cholesterolu (mmol/l)

Skupina	Medián	Medián int.	b-medián	<i>t</i> -boot. int.
A	4,27	(3,96; 5,15)	4,24	(3,93; 4,67)
B	4,52	(3,71; 5,86)	4,47	(3,62; 5,45)
C	4,27	(4,07; 5,16)	4,23	(3,82; 4,64)

Nebyla prokázána existence statisticky významného rozdílu mezi populačními průměry, resp. mediány, vybraných biochemických veličin v léčebných skupinách (průniky bootstrapových konfidenčních intervalů populačních průměrů, resp. mediánů, jsou neprázdné). Populační průměry a mediány vybraných biochemických veličin nejsou závislé na léčebné metodě. Závěr odpovídá testování na 5% hladině významnosti.

3.3. Srovnání klasických a bootstrapových konfidenčních intervalů

Jestliže jsou splněny výše uvedené předpoklady, tak by se bootstrapový bodový odhad parametru neměl příliš lišit od klasického bodového odhadu daného parametru. Příliš velký rozdíl mezi klasickým a bootstrapovým bodovým odhadem daného parametru ukazuje na nevhodnost použití bootstrapových metod pro daná data, resp. nutnost zpracovávat analýzu bootstrapovými metodami až po vyjmutí odlehlých hodnot.

Co se týče intervalových odhadů, tak lze zjednodušeně konstatovat následující. Pokud by nastala teoretická situace $B = N$, tak by bootstrapové intervaly byly širší než klasické konfidenční intervaly. Důvodem je větší rozptyl bootstrapových parametrů. Jelikož však v praxi vždy volíme $B \gg N$, nastává opačná situace. Bootstrapové konfidenční intervaly jsou užší než klasické konfidenční intervaly. Důvodem je větší počet pozorování použitých pro tvorbu konfidenčních intervalů u bootstrapových intervalů. Bootstrapové metody dávají tedy při dodržení výše uvedených předpokladů přesnější intervalový odhad zvoleného parametru. Nevýhodou je velký počet výpočetních operací.

4. Závěr

V příspěvku jsme demonstrovali použití bootstrapových metod pro srovnání účinností tří způsobů aplikace preparátu z kmenových buněk (do svalu, žíly nebo tepny) při léčbě kritické končetinové ischemie. Použili jsme srovnání grafů empirických distribučních funkcí a empirických hustot vybraných biochemických parametrů v léčebných skupinách. Dále jsme porovnali populační průměry a mediány vybraných biochemických parametrů v léčebných skupinách porovnáním 95% bootstrapových konfidenčních intervalů.

Nebyl nalezen prokazatelný rozdíl mezi způsoby aplikace preparátu z kmenových buněk. Nepotvrdil se tudíž lékařský předpoklad, že aplikace preparátu BMAC intravenózní cestou povede ke snížení hodnot vybraných biochemických veličin, oproti aplikaci intramuskulární nebo intraarteriální cestou.

Ovšem je třeba zdůraznit, že výše uvedený výsledek analýz neukazuje na to, že by léčba pomocí kmenových buněk nebyla účinná. Dřívější studie prokázala existenci statisticky významného vlivu aplikace preparátu BMAC na hodnoty nejen biochemických veličin pacienta oproti tradiční chirurgické metodě. Výsledky těchto analýz je možno nalézt v publikační činnosti MUDr. Václava Procházky, Ph.D., MSc., z FN Ostrava. [4]

Poděkování

Autorka děkuje České statistické společnosti za poskytnutí finančního příspěvku k pokrytí vložného na konferenci Prastan 2017.

V článku byl použit software R (R Core Team (2016). R: A language and environment for statistical computing. R Foundation for Statistical Computing, Vienna, Austria. URL: <https://www.R-project.org/>).

Reference

- [1] Davison A. C., Hinkley D. V.: *Bootstrap Methods and Their Applications*, Cambridge, Cambridge University Press, 2003. ISBN 0-521-57471-4.
- [2] Efron B., Tibshirani R. J.: *An Introduction to the Bootstrap*, New York, Chapman & Hall, 1993. 436 pp. ISBN 0-412-04231-2.
- [3] Prášková Z.: Metoda bootstrap, *Sborník prací 13. letní školy JČMF Robust 2004*, 299–314. ISBN 80-7015-972-3.
- [4] Procházka V., Gumulec J., Jalůvka F., Šalounová D., Jonszta T., Czerný D., Krajča J., Urbanec R., Klement P., Martínek J., Klement G.: Cell Therapy, a New Standard in Management of Chronic Critical Limb Ischemia and Foot Ulcer, *Cell Transplantation*, 19, 1413–1424, Cognizant Comm. Corp., USA, 2010, eISSN 1555-3892.
- [5] Teraa M., Conte M. S., Moll F. L., Verhaar M. C.: Critical Limb Ischemia – Current Trends and Future Directions, *Journal of the American Heart Association*, USA, Dallas, 2016, ISSN 2047-9980.
- [6] Procházka V., Vítková K., Robenková M., Gumulec J., Jalůvka F., Ocelka T., Pavliska L., Šalounová D.: Randomizovaná studie bezpečnosti a efektivity autologního koncentrátu aspirátu kostní dřeně k léčbě kritické končetinové ischemie u pacientů s diabetem mellitem II. typu, *Protokol klinického hodnocení, 2012-T2DM-CLI (Dialog)*, FNO Ostrava, 2012.

VZPOMÍNKY NA ING. MILOŠE JÍLKA, CSC. MEMORIES OF ING. MILOŠ JÍLEK, CSC.

Jan Klaschka

E-mail: klaschka@cs.cas.cz

Koncem letošního dubna mě zastihla smutná zpráva, že 22. dubna 2018 zemřel ve věku 86 let Ing. Miloš Jílek, CSc., statistik z generace Ing. Machka a Ing. Rotha.

Seznámil jsem se s ním před 40 lety jako jeho diplomant. Pracoval v Mikrobiologickém ústavu ČSAV v Praze-Krči, který posléze zůstal jeho působištěm až do důchodu. Diplomová práce se týkala simulace jistého modelu imunitní reakce. Imunologickým modelováním se Ing. Jílek zabýval od počátku 70. let a publikoval na toto téma řadu prací – zpočátku společně s imunologem prof. Šterzlem, později převážně sám nebo se svými doktorandy.

Druhou oblastí, kde Ing. Jílek zanechal trvalou stopu, jsou toleranční intervaly. Jeho zájem o ně se datoval už od konce 50. let, kdy pracoval v Centrálním výzkumném ústavu potravinářského průmyslu, a věnoval se jim i ve své kandidátské práci z roku 1962. Kromě publikování běžných úzeji zaměřených prací s obrovskou pílí a systematičností shromáždil a prostudoval množství literatury o tolerančních mezích a sepsal bibliografii [1, 2] a knihu [3], která, pokud vím, i po 30 letech zůstává nejdůkladnějším dílem v češtině o daném tématu.

Nejsou to ovšem ani tak vědecké zásluhy, jako hlavně lidské kvality, co mě k Ing. Jílkovi ještě dlouhé roky po studiích poutalo. Hned při prvním setkání jsem si nemohl nevšimnout jeho handicapu: Následkem obrny, kterou prodělal v kojeneckém věku, obtížně chodil. Všechny těžkosti života s postižením ale snášel s obdivuhodnou statečností, beze stopy zatrpklosti – naopak, nebylo snad laskavějšího člověka nad něj. Na pohřbu zaznělo, že litoval ne sebe, ale rodiče, jimž jeho nemoc musela způsobit velkou bolest.

Na rozdíl od některých svých generačních soupeřů, kteří pracovali v oboru do vpravdě kmetských let, Ing. Jílek odešel poměrně brzy do důchodu, porozdával literaturu (vzpomínám, jak jsme z jeho dejvického bytu odváželi s Jaromírem Antochem na fakultu krabice sebraných článků) a statistiku „pověsil na hřebík“. Chtěl se více věnovat své paní, která byla také „obrnářka“ a ještě hůře postižená než on. A rád také získal více času pro svou velkou zálibu vycházející z jeho hluboké křesťanské víry – studium a sběratelství Bible v nejrůznějších jazycích, překladech a vydáních.

Naposledy mi Ing. Jílek psal před loňskými Vánoci. Nestěžoval si vůbec na svoje trápení (poslední léta už ani nemohl opustit byt), ale zato se snažil,

protože jsem se mu právě svěřil se svými zdravotními problémy, povzbudit mě.

Ing. Jílek byl inspirativním živým příkladem, jak lze i nespravedlivě těžký život prožít s vlídnou tváří a úsměvem pro všechny kolem sebe. Slovem krásný člověk.

Reference

- [1] Jílek M.: A bibliography of statistical tolerance regions. *Math. Operationsforsch. Statist.*, Ser. Statistics, 12, 441–456, 1981.
- [2] Jílek M., Ackermann H.: A bibliography of statistical tolerance regions, II., *Math. Operationsforsch. Statist.*, Ser. Statistics, 20, 165–172, 1989.
- [3] Jílek M.: *Statistické toleranční meze*. SNTL, Praha, 1988.

INFORMACE Z KONFERENCE ROBUST 2018 ROBUST 2018: IMPRESSIONS

Marta Žambochová

E-mail: marta.zambochova@ujep.cz

Ve dnech od 21. do 26. ledna 2018 proběhl ve středisku Rybník na rozhraní Šumavy a Českého lesa jubilejní 20. ročník statistické konference ROBUST 2018. Konference je od roku 1980 tradičně pořádána každé dva roky a na její organizaci se podílejí Jednota českých matematiků a fyziků a Česká statistická společnost. Letošního jubilejního ročníku se zúčastnilo 86 účastníků celkem z pěti zemí. Mezi účastníky nemohl chybět Jaromír Antoch, který byl zatím přítomen na každém z pořádaných ročníků, nebo například Marie Hušková či Petr Volf, kteří se konference účastní „až“ od roku 1984.

Konference byla zahájena zajímavou zvanou přednáškou Milana Hladíka z MFF UK v Praze, a to na téma intervalové robustnosti v lineárním programování. Mezi pozvanými přednášejícími byli i ekonomové a zástupci státní statistiky. Bývalý guvernér České národní banky, nyní ředitel institucionálních vztahů a hlavní ekonom společnosti Generali CEE Holding Miroslav Singer přednesl přednášku o rozhodování za přítomnosti nejistoty. Téma bylo zaměřeno na českou měnu a kupní sílu Čechů. Ředitel Sekce makroekonomických statistik Českého statistického úřadu Jaroslav Sixta měl přednášku nazvanou Odhadování HDP: Predikce versus realita. Na obě přednášky navázala velmi široká diskuze.

Součástí konference byla i sekce mladých statistiků, v rámci níž vystoupilo 25 studentů magisterského a doktorského studia. Nejlepší příspěvky byly oceněny jednak finančně a jednak věcně. Ceny sponzorovaly společnosti RSJ Securities, a. s., a Nadace RSJ.

Celá konference proběhla jako tradičně ve velmi přátelském tónu a vše fungovalo na výbornou. Za to patří velké poděkování trojici organizátorů Jaromíru Antochovi, Gejzovi Dohnalovi a Danielu Hlubinkovi. Všichni účastníci se už nyní těší, na jaké další odlehlé místo je pořadatelé pozvou.



Logo letošního ročníku vzešlo ze soutěže v rámci ROBUST 2016.

Obsah

Vědecké a odborné články

Petr Volf

Modely směsí rizikových funkcí s aplikací v analýze spolehlivosti 1

Lenka Přibyllová

Použití bootstrapových metod pro srovnání způsobů léčby
kritické končetinové ischemie 13

Zprávy a informace

Jan Klaschka

Vzpomínky na Ing. Miloše Jílka, CSc. 22

Marta Žambochová

Informace z konference ROBUST 2018 23

Informační bulletin České statistické společnosti vychází čtyřikrát do roka v českém vydání. Příležitostně i mimořádné české a anglické číslo. Vydavatelem je Česká statistická společnost, IČ 00550795, adresa společnosti je Na padesátém 81, 100 82 Praha 10. Evidenční číslo registrace vedené Ministerstvem kultury ČR dle zákona č. 46/2000 Sb. je E 21214.

The Information Bulletin of the Czech Statistical Society is published quarterly.
The contributions in the journal are published in English, Czech and Slovak languages.

Předseda společnosti: RNDr. Marek MALÝ, CSc., Státní zdravotní ústav, Šrobárova 48, Praha 10, 100 42, e-mail: mmaly@szu.cz.

Redakce: prof. RNDr. Gejza DOHNAL, CSc. (šéfredaktor), prof. RNDr. Jaromír ANTOCH, CSc., doc. Ing. Jozef CHAJDIK, CSc., doc. RNDr. Zdeněk KARPÍŠEK, CSc., RNDr. Marek MALÝ, CSc., doc. RNDr. Jiří MICHÁLEK, CSc., prof. Ing. Jiří MILITKÝ, CSc., doc. Ing. Iveta STANKOVIČOVÁ, PhD., doc. Ing. Josef TVRDÍK, CSc., Mgr. Ondřej VENCÁLEK, Ph.D.

Redaktor časopisu: Mgr. Ondřej VENCÁLEK, Ph.D., ondrej.vencalek@upol.cz.
Informace pro autory jsou na stránkách společnosti, <http://www.statspol.cz/>.

DOI: 10.5300/IB, <http://dx.doi.org/10.5300/IB>
ISSN 1210–8022 (Print), ISSN 1804–8617 (Online)

Toto číslo bylo vytištěno s laskavou podporou Českého statistického úřadu.