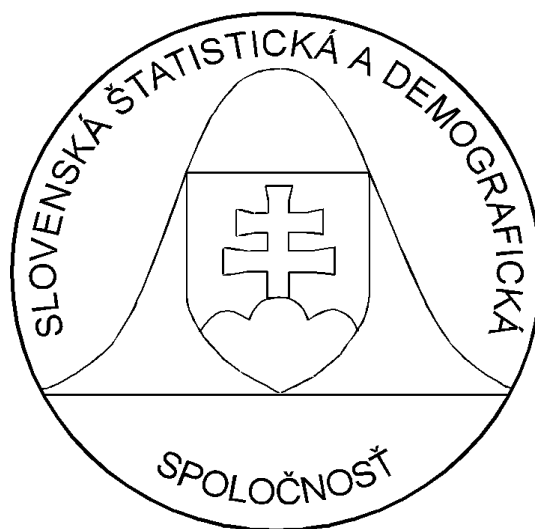


2/2010

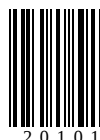
FORUM STATISTICUM SLOVACUM



ISSN 1336-7420



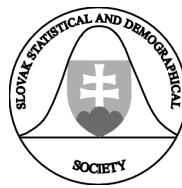
9 771336 742001



20101



**Slovenská štatistická a demografická
spoločnosť Miletičova 3, 824 67
Bratislava
www.ssds.sk**



Naše najbližšie akcie:

(pozri tiež www.ssds.sk, blok Organizované akcie)

EKOMSTAT 2010, 24. škola štatistiky

tematické zameranie: *Štatistické metódy vo vedecko-výskumnej, odbornej
a hospodárskej praxi.*

30. 5. – 4.6. 2010, Trenčianske Teplice

FernStat_CZ 2010

23. - 24. 9. 2010, Ústí nad Labem

15. SLOVENSKÁ ŠTATISTICKÁ KONFERENCIA

7. – 8. 10. 2010, Kaskády, Galanta - Únovce

19. Medzinárodný seminár VÝPOČTOVÁ ŠTATISTIKA,

2. – 3. 12. 2010, Bratislava, Infostat

Prehliadka prác mladých štatistikov a demografov

2. 12. 2010, Bratislava, Infostat

Regiónálne akcie

priebežne

Pohľady na ekonomiku Slovenska 2011

12. 4. 2011, Bratislava, Aula EU

ÚVOD

Vážené kolegyně, vážení kolegovia,

druhé číslo šiesteho ročníka vedeckého časopisu Slovenskej štatistickej a demografickej spoločnosti je zostavené z príspevkov, ktoré sú obsahovo orientované v súlade s tematikou konferencie Nitrianske štatistické dni. Táto akcia sa uskutočnila v dňoch 27. a 28. mája 2010 v Nitre. Konferenciu organizovala Slovenská štatistická a demografická spoločnosť v spolupráci s Univerzitou Konštantína Filozofa v Nitre.

Akciu z poverenia Výboru SŠDS, zorganizoval Organizačný a programový výbor. Programový výbor v zložení: prof. RNDr. Anna Tirpáková, CSc., doc. Ing. Jozef Chajdiak, CSc., STU, doc. PaedDr. Jana Kubanová, CSc., doc. RNDr. Bohdan Linda, CSc., RNDr. Ján Luha, CSc., RNDr. Jitka Poměnková, PhD., prof. RNDr. Beáta Stehliková, CSc., prof. RNDr. Ondrej Šedivý, CSc., doc. RNDr. Marta Vrábelová, CSc. Organizačný výbor v zložení PaedDr. Janka Melušová, PhD., doc. Ing. Jozef Chajdiak, CSc., RNDr. Ján Luha, CSc., doc. RNDr. Dagmar Markechová, CSc. a RNDr. Kitti Vidermanová, PhD

Na príprave a zostavení tohto čísla FORUM STATISTICUM SLOVACUM participovali: doc. Ing. Jozef Chajdiak, CSc., RNDr. Ján Luha, CSc., PaedDr. Janka Melušová, PhD., Mgr. Michal Práznovský, prof. RNDr. Beáta Stehlíková, CSc., prof. RNDr. Anna Tirpáková, CSc., RNDr. Kitti Vidermanová, PhD.

Recenziu príspevkov zabezpečili: doc. Ing. Jozef Chajdiak, CSc., doc. PaedDr. Jana Kubanová, CSc., doc. RNDr. Bohdan Linda, CSc., RNDr. Ján Luha, CSc., PhD., prof. RNDr. Beáta Stehlíková, CSc., prof. RNDr. Anna Tirpáková, CSc., RNDr. Jitka Poměnková, Ph.D.,

Výbor SŠDS

CHANGE POINT DETECTION

Jaromír Antoch, Daniela Jarušková

Abstract First part of this contribution deals with tests on the stability of statistical models. The problem is formulated in terms of testing the null hypothesis H against the alternative hypothesis A . The null hypothesis H claims that the model remains the same during the whole observational period, usually it means that the parameters of the model do not change. The alternative hypothesis A claims that, at an unknown time point, the model changes, which means that some of the parameters of the model are subject to a change. In case we reject the null hypothesis H , i.e. when we decide that there is a change in the model, we concentrate on a number of questions that arise:

- when has the model changed;
- is there just one change or are there more changes;
- what is the total number of changes etc.

The time moment when the model has changed is usually called *change point*. Aside testing for a change, our interest is to estimate change point(s) in different models. The least squares estimators are introduced. Of course, we also estimate other parameters of the model(s), show approximations to the distributions of the change point estimators and show that the estimators of the change points are usually closely related to some of the test statistics treated in the first part.

Keywords: Change point problem; LSE estimators; Bayesian, maximum-type and trimmed maximum-type statistics; location model, linear regression, abrupt and gradual changes, multiple changes; AR and ARMA processes; hypotheses testing, approximate critical values, Bonferroni inequality

1. Introduction

In meteorology, climatology, hydrology or environmental studies a certain quantity or quantities such as temperature, content of water impurity etc. are recorded sequentially in time. The observed time series exhibit random fluctuation and therefore their behavior may be described by stochastic models. While studying such series, an important question may arise whether stochastic properties of observed time series remain the same over time, i.e. the observed series are stationary, or whether at some time unknown to the observer the studied time series changed its behavior. The notion “stationary behavior of time series” covers here more situations than what the statisticians usually understand under this word. The series is considered in this paper to be stationary even if for example its mean increment remains the same over time etc. From the statistical point of view a change in series occurs if there exists a time point such that the observations follow one distribution up to that point and follow another distribution after that point.

In practice some of these changes are induced by natural variability, for example periods of warmer weather turns to periods of colder weather and vice versa. Some others are caused by human activity such as by urbanization or deforestation. The prospect of Global Change advanced by many researchers raises the possibility of yet another class of transitions of crucial climatic variables. Some other changes may be caused by an alternation in measurement process such as the relocation of a gauge or a change in the time at which the measurement is taken etc.

The field of models applied in change point analysis is very broad. The change may occur in location, variation, regression coefficients, correlation structure and so on. The stochastic models can be parametric or nonparametric. One assumes that there can be only one change in the series or one can assume that the series might change several times etc.

Statistical inference usually takes two steps. First, one decides whether the behavior of the series under the study has changed. Second, if a change was detected, the change point(s) is/are to be found. The decision whether the observed series remains stationary, i.e. follows one distribution only, or whether a change of a specific kind occurred, is usually based on hypotheses testing. To locate the change point(s) is an estimation problem.

To illustrate the change point analysis some examples are presented that show use of change point methods to ecology as well as to some related scientific branches.

Example 1

The data of interest in this example are two series of average monthly air pressure obtained from daily measurements in two Swiss meteorological stations – Saentis and Guetsch. The series of monthly averages at Saentis was homogenized and served as the reference series. The problem was to detect an inhomogeneity (if there is any) in the Guetsch series caused by changes in the measurement process. The statistical inference was based on the sequence of $n = 240$ differences between the Guetsch and Saentis series of air pressure monthly averages after the seasonality was removed. It is supposed that the alternation of measurement process might cause a sudden positive or negative shift in the mean of the considered differences. The aim of statistical analysis was to decide whether such a shift occurred and, if it happened, to estimate the change point and the magnitude of the detected shift.

If we denote by Y_i the time series under the study, then the null hypothesis H and the alternative hypothesis A of the corresponding problem of hypotheses testing may be set as follows:

$$\begin{aligned} H : Y_i &= \mu + e_i, & i &= 1, \dots, n, \\ A : \exists \mathbf{m} \in \{1, \dots, n-1\} \quad \text{such that} & & & \\ Y_i &= \mu + e_i, & i &= 1, \dots, \mathbf{m}, \\ Y_i &= \mu + \delta + e_i, & i &= \mathbf{m} + 1, \dots, n, \quad \delta \neq 0, \quad \mu \in \mathcal{R}_1. \end{aligned} \tag{1}$$

The variables $\{e_i\}$ represent some random errors.

Example 2

Earth climate seems to be threatened by global warming. For studying the effect of the global warming several artificial series have been composed to describe the behavior of the global temperature by combining temperature series measured at many meteorological sites. One of the most famous artificial temperature series is the series of global world temperature anomalies composed by Jones and his colleagues (1994). The series starts in year 1854 and consists of 140 data (each of them corresponds to one year). The goal of statistical inference is to decide whether the series may be considered to be stationary. The expected change has the form of gradual increase of the series mean. If we consider that at an unknown time point a continuous linear trend appears (it is the simplest kind of gradual change) the null hypothesis H and the alternative A are set as follows:

$$\begin{aligned}
H : Y_i &= a + e_i, & i &= 1, \dots, n, \\
A : \exists \mathbf{m} \in \{1, \dots, n-1\} \quad \text{such that} \\
Y_i &= a + e_i, & i &= 1, \dots, \mathbf{m}, \\
Y_i &= a + b \cdot \frac{i - \mathbf{m}}{n} + e_i, & i &= \mathbf{m} + 1, \dots, n,
\end{aligned} \tag{2}$$

where $a \in \mathcal{R}_1$, $b \neq 0$ and $\{e_i\}$ are random errors.

Example 3

Large scale deforestation may cause the soil to lose its capability for water retention. The researchers of the Czech Research Institute for Forest Management studied the effect of controlled deforestation on the rainfalls-runoffs relationship. The objective of statistical inference was to decide whether this relationship changed during the study. To simplify the model it was supposed that the relation between rainfalls and runoffs was linear. The problem is one of many change point problems in linear regression. One can either suppose that the change might occur in the intercept only or that it might occur in the both parameters, i.e. in the intercept and/or the slope. The second test problem has the form:

$$\begin{aligned}
H : Y_i &= a + bx_i + e_i, & i &= 1, \dots, n, \\
A : \exists \mathbf{m} \in \{2, \dots, n-2\} \quad \text{such that} \\
Y_i &= a + bx_i + e_i, & i &= 1, \dots, \mathbf{m}, \\
Y_i &= a_0 + b_0x_i + e_i, & i &= \mathbf{m} + 1, \dots, n,
\end{aligned} \tag{3}$$

where $a \neq a_0$ and/or $b \neq b_0$, and $\{e_i\}$ are random errors. The statistical inference was based on 36 paired data of the annual total amount of precipitation in the area and the corresponding water discharges of a small creek called Malá Ráztoka.

2. How to decide whether a series is stationary

As was mentioned before the decision whether a change occurred is often done by testing statistical hypotheses. The decision rule for rejecting the null hypothesis claiming that the series is stationary is based on a test statistic. Two methods are usually applied to derive test statistics, namely, the maximum likelihood method and the pseudo-Bayesian method. In the following simple example we demonstrate how both approaches can be applied.

Example 4

The same quantity is measured by two measuring devices. Since the measurement is affected by random errors the obtained values may differ. In the beginning neither of the measurements are supposed to have a systematic error so that the difference between them are supposed to vary around zero. However, it can happen due to a failure of one of the measuring devices that after an unknown time moment these differences may start to vary around an unknown non-zero constant. Notice, that this example is very similar to Example 1. The only difference is that in Example 1 the measuring gauges were not located at one site and therefore even in the beginning the mean of the calculated

differences was supposed to be an unknown constant which had to be estimated. If it is supposed, for simplicity, that the variance of the studied differences has not been changing, we test the null hypothesis H against the alternative A :

$$\begin{aligned} H : Y_i &= e_i, & i &= 1, \dots, n, \\ A : \exists \mathbf{m} \in \{0, \dots, n-1\} \text{ such that} \\ Y_i &= e_i, & i &= 1, \dots, \mathbf{m}, \\ Y_i &= \mu + e_i, & i &= \mathbf{m} + 1, \dots, n, \end{aligned} \quad (4)$$

where $\mu \neq 0$. Sometimes, the one sided alternative with $\mu > 0$ or $\mu < 0$ may be considered if the sign of the conceivable shift is supposed to be known. For in our example the random errors $\{e_i\}$ are the measurement errors and we may suppose that they are independent and distributed according to a normal distribution. For simplicity again, we suppose that their variance is known so that it may be set to one. Therefore the variables $\{e_i\}$ are iid and follow $N(0, 1)$ distribution with the density $\phi(x)$ and the distribution function $\Phi(x)$.

Likelihood ratio method

If $\mu \neq 0$, the log-likelihood ratio for testing H against A is

$$\max_{0 \leq k \leq n-1} \sup_{\mu} \log \frac{\prod_{i=1}^k \phi(Y_i) \prod_{i=k+1}^n \phi(Y_i - \mu)}{\prod_{i=1}^n \phi(Y_i)} = \max_{0 \leq k \leq n-1} \frac{1}{2(n-k)} \left(\sum_{i=k+1}^n Y_i \right)^2,$$

the test statistic usually applied for a two-sided alternative is of the form

$$\max_{0 \leq k \leq n-1} \left\{ \left| \frac{1}{\sqrt{n-k}} \sum_{i=k+1}^n Y_i \right| \right\}, \quad (5)$$

while a one-sided alternative with $\mu > 0$ is

$$\max_{0 \leq k \leq n-1} \left\{ \frac{1}{\sqrt{n-k}} \sum_{i=k+1}^n Y_i \right\}. \quad (6)$$

These statistics belong to the so-called *maximum-type statistics*.

Pseudo-Bayesian method

The method is based on the assumption that the unknown change point $\mathbf{m}$ and unknown magnitude of shift $\boldsymbol{\mu}$ are random variables such that the prior distribution of $\mathbf{m}$ is uniform, i.e. $P(\mathbf{m} = i) = 1/n$, $i = 0, \dots, n-1$, and $\boldsymbol{\mu}$ is distributed according to $N(0, \gamma^2)$. Since the density of $Y_1, \dots, Y_n$ given $\boldsymbol{\mu} = \mu$ and $\mathbf{m} = k$ is normal, i.e.,

$$f(y_1, \dots, y_n \mid \boldsymbol{\mu} = \mu, \mathbf{m} = k) = \prod_{i=1}^k \phi(y_i) \prod_{i=k+1}^n \phi(y_i - \mu),$$

the unconditional density can be expressed as

$$\begin{aligned} f(y_1, \dots, y_n) &= \sum_{k=1}^n \frac{1}{n} \int_{-\infty}^{\infty} \prod_{i=1}^k \phi(y_i) \prod_{i=k+1}^n \phi(y_i - \mu) \frac{1}{\gamma\sqrt{2\pi}} \exp\left\{-\mu^2/2\gamma^2\right\} d\mu = \\ &= \prod_{i=1}^n \phi(y_i) \left(\sum_{k=1}^n \frac{1}{n} \int_{-\infty}^{\infty} \frac{1}{\gamma\sqrt{2\pi}} \exp\left\{-\frac{1}{2} \left(-2\mu \sum_{i=k+1}^n y_i + (n-k)\mu^2 + \frac{\mu^2}{\gamma^2} \right) \right\} d\mu \right) \end{aligned}$$

and the corresponding likelihood ratio has the form

$$\frac{f_A(Y_1, \dots, Y_n)}{f_H(Y_1, \dots, Y_n)} = \sum_{k=1}^n \frac{1}{n} \sqrt{\frac{1}{1 + (n-k)\gamma^2}} \exp \left\{ \frac{\gamma^2}{2(1 + (n-k)\gamma)^2} \left(\sum_{i=k+1}^n Y_i \right)^2 \right\}.$$

Letting $\gamma \rightarrow 0$ and applying Taylor expansion, the likelihood ratio is, for the two-sided alternative, equivalent to the test statistic

$$\sum_{k=1}^n \frac{1}{n} \left(\frac{1}{\sqrt{n}} \sum_{i=k+1}^n Y_i \right)^2. \quad (7)$$

The obtained statistic belongs to the so called *sum-type* test statistics.

For the one-sided alternative with $\mu > 0$ (μ is assumed to follow a 50% truncated normal distribution) we may analogously derive the sum-type test statistic

$$\sum_{k=1}^n \frac{1}{n} \left(\frac{1}{\sqrt{n}} \sum_{i=k+1}^n Y_i \right) = \frac{1}{n} \frac{1}{\sqrt{n}} \sum_{k=1}^n (k-1) Y_k. \quad (8)$$

Critical values

For the decision whether the null hypothesis H should be rejected we need to know critical values of the suggested test statistics. This means knowing their distributions under H .

Let us start with the test statistic (5). Assume that $\{Y_i\}$ are iid and follow the $N(0, 1)$ distribution. Then all statistics

$$\frac{1}{\sqrt{n-k}} \sum_{i=k+1}^n Y_i, \quad k = 0, \dots, n-1, \quad (9)$$

have the $N(0, 1)$ distribution. To find the exact distribution of (5) means to find the distribution of the maximum of absolute values of standardized normal variables that are (unfortunately) correlated. Theoretically, it should be no problem to find the distribution of (5). However, in practice the distribution is so complex, that its quantiles (desired critical values) may be computed only for small values of n .

Sometimes, the approximate critical values may be quite satisfactory for practical use. To find approximate critical values we can use a very simple idea by applying **the Bonferroni inequality** as follows:

$$P \left(\max_{0 \leq k \leq n-1} \left\{ \left| \frac{1}{\sqrt{n-k}} \sum_{i=k+1}^n Y_i \right| \right\} > C \right) \leq \sum_{k=0}^{n-1} P \left\{ \left| \frac{1}{\sqrt{n-k}} \sum_{i=k+1}^n Y_i \right| > C \right\}.$$

Hence, the $100(1 - \alpha/(2n))\%$ -quantile of the standard normal distribution $N(0, 1)$ may serve as an upper estimate of the critical value at the significance level α for the problem (4) applying the test statistic (5). The approximate critical values obtained in this way are good enough for small samples (for small values of n), but they are too conservative for n large.

For n large, the **asymptotic behavior** of the studied test statistic (5) is of interest. It can be proved, applying the law of iterated logarithm, that

$$\max_{0 \leq k \leq n-1} \left\{ \left| \frac{1}{\sqrt{n-k}} \sum_{i=k+1}^n Y_i \right| \right\} \rightarrow \infty \quad \text{almost surely as } n \rightarrow \infty.$$

It follows that the limit distribution of (5) does not exist and that critical values increase to infinity as $n \rightarrow \infty$. The problem is caused by the behavior of the sequence $\{(n-k)^{-1/2} \sum_{i=k+1}^n Y_i, i = 1, \dots, n-1\}$ near to its end.

Therefore, some authors suggest using, instead of the statistics (6) and (5), the *trimmed maximum-type* test statistics

$$\max_{0 \leq k \leq \lfloor (1-\beta)n \rfloor} \left\{ \frac{1}{\sqrt{n-k}} \sum_{i=k+1}^n Y_i \right\}, \quad \max_{0 \leq k \leq \lfloor (1-\beta)n \rfloor} \left\{ \left| \frac{1}{\sqrt{n-k}} \sum_{i=k+1}^n Y_i \right| \right\}, \quad (10)$$

where β is a small positive constant less than one and $\lfloor x \rfloor$ denotes the integer part of x . The advantage of the statistics (10) is that they are bounded in probability. The *trimming off* a $100\beta\%$ portion of the sample (upper time points) means that one assumes that the change did not occur during this time period.

Some other authors recommend to deal with maximum of weighted statistics (9) (or their absolute values):

$$\max_{0 \leq k \leq n-1} \left\{ w(k) \frac{1}{\sqrt{n-k}} \sum_{i=k+1}^n Y_i \right\}, \quad \max_{0 \leq k \leq n-1} \left\{ w(k) \left| \frac{1}{\sqrt{n-k}} \sum_{i=k+1}^n Y_i \right| \right\}.$$

The weights are chosen to correspond to a prior knowledge of the time where the change might occur that is provided by an expert.

The asymptotic behavior of maximum of a sequence of random variables is given by the limit Gaussian process. Unfortunately, for different test statistics the limit processes differ and therefore the exceedence level probabilities differ as well. In the example introduced above the limit process is $\left\{ \frac{W(1-t)}{\sqrt{1-t}}, t \in [0, 1) \right\}$, where $\{W(t), t \geq 0\}$ denotes the Wiener process.

The approximate critical values for the over-all maximum-type test statistic (6) and (5) can be calculated from the approximations:

$$P \left(\max_{0 \leq k \leq n-1} \left\{ \frac{1}{\sqrt{n-k}} \sum_{i=k+1}^n Y_i \right\} > \frac{x + b_n}{a_n} \right) \approx 1 - \exp \left\{ -\frac{1}{2} e^{-x} \right\}, \quad x \in \mathcal{R}_1, \quad (11)$$

$$P \left(\max_{0 \leq k \leq n-1} \left\{ \left| \frac{1}{\sqrt{n-k}} \sum_{i=k+1}^n Y_i \right| \right\} > \frac{x + b_n}{a_n} \right) \approx 1 - \exp \left\{ -e^{-x} \right\}, \quad x \in \mathcal{R}_1, \quad (12)$$

where

$$a_n = \sqrt{2 \log \log n} \quad \text{and} \quad b_n = 2 \log \log n + \frac{1}{2} \log \log \log n - \frac{1}{2} \log \pi.$$

The approximate critical values for the trimmed maximum-type test statistics (10) can be calculated from the approximations:

$$P \left(\max_{0 \leq k \leq \lfloor (1-\beta)n \rfloor} \frac{1}{\sqrt{n-k}} \sum_{i=k+1}^n Y_i > x \right) \approx \frac{1}{2} x \phi(x) \log \frac{1}{\beta}, \quad x \in \mathcal{R}_1, \quad (13)$$

$$P \left(\max_{0 \leq k \leq \lfloor (1-\beta)n \rfloor} \left\{ \left| \frac{1}{\sqrt{n-k}} \sum_{i=k+1}^n Y_i \right| \right\} > x \right) \approx x \phi(x) \log \frac{1}{\beta}, \quad x \in \mathcal{R}_1. \quad (14)$$

The critical values obtained from (11) and (12) are not very good because they are too conservative. The approximations (13) and (14) gives more accurate critical values

but the choice of β affects them significantly. The critical values obtained from (13) and (14) get better as β gets larger.

For the sum-type statistic (7) obtained from the pseudo-Bayesian method the asymptotic critical values can be calculated from the convergence

$$\sum_{k=1}^n \frac{1}{n} \left(\frac{1}{\sqrt{n}} \sum_{i=k+1}^n Y_i \right)^2 \xrightarrow{\mathcal{D}} \int_0^1 W^2(t) dt. \quad (15)$$

The distribution of $\int_0^1 W^2(t) dt$ was studied by Mac Neil (1974, 1978). The approximations (13), (14) and (15) hold true even if the random errors $\{e_i\}$ are iid but not normally distributed.

The calculation of critical values of (8) derived for the one-sided alternatives is very simple as the statistic is normally distributed as $N\left(0, \frac{1}{3} - \frac{1}{2n} + \frac{1}{6n^2}\right)$.

Remark: Another possibility how to obtain critical values is via simulations. In Antoch et al. (2001) simulated critical values for all the situations considered here are available together with more detailed description of the methods.

♣ Now, after explaining the basic ideas, let us go back to our examples.

Example 1 (continuation)

If the random errors $\{e_i\}$ are iid with the normal distribution $N(0, \sigma^2)$, where σ^2 is unknown, then the maximum-type statistics have the form

$$\max_{1 \leq k \leq n-1} \left\{ \frac{1}{s_k} \left| \sqrt{\frac{k(n-k)}{n}} (\bar{Y}_k - \bar{Y}_k^o) \right| \right\}, \quad (16)$$

resp.

$$\max_{\lfloor \beta n \rfloor \leq k \leq \lfloor (1-\beta)n \rfloor} \left\{ \frac{1}{s_k} \left| \sqrt{\frac{k(n-k)}{n}} (\bar{Y}_k - \bar{Y}_k^o) \right| \right\}, \quad (17)$$

where

$$\bar{Y}_k = \frac{1}{k} \sum_{i=1}^k Y_i, \bar{Y}_k^o = \frac{1}{n-k} \sum_{i=k+1}^n Y_i, s_k^2 = \frac{1}{n-2} \left(\sum_{i=1}^k (Y_i - \bar{Y}_k)^2 + \sum_{i=k+1}^n (Y_i - \bar{Y}_k^o)^2 \right).$$

The sum-type test statistic has the form

$$\frac{1}{\hat{\sigma}^2} \sum_{k=1}^n \frac{1}{n} \left(\frac{1}{\sqrt{n}} \sum_{i=1}^k (Y_i - \bar{Y}_n) \right)^2, \text{ where } \hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n (Y_i - \bar{Y}_n)^2. \quad (18)$$

To find approximate critical values using the Bonferroni inequality we are to realize that under H all variables $\left\{ \sqrt{\frac{k(n-k)}{n}} (\bar{Y}_k - \bar{Y}_k^o) \right\}$ are distributed according to the t -distribution with $n-2$ degrees of freedom.

For n large the asymptotic critical values can be found using the approximations ($x \in \mathcal{R}_1$):

$$P \left(\max_{1 \leq k \leq n-1} \left\{ \frac{1}{s_k} \left| \sqrt{\frac{n}{k(n-k)}} \sum_{i=1}^k (Y_i - \bar{Y}_n) \right| \right\} > \frac{x + b_n}{a_n} \right) \approx 1 - \exp \{ -2e^{-x} \}, \quad (19)$$

$$P\left(\max_{\lfloor \beta n \rfloor \leq k \leq \lfloor (1-\beta)n \rfloor} \left\{ \frac{1}{s_k} \left| \sqrt{\frac{n}{k(n-k)}} \sum_{i=1}^k (Y_i - \bar{Y}_n) \right| \right\} > x\right) \approx 2x\phi(x) \log \frac{1-\beta}{\beta}. \quad (20)$$

The approximate critical values for the sum-type test statistic may be obtained from the approximation

$$P\left(\frac{1}{\hat{\sigma}^2} \sum_{k=1}^n \frac{1}{n} \left(\frac{1}{\sqrt{n}} \sum_{i=1}^k (Y_i - \bar{Y}_n) \right)^2 > x\right) \approx 1 - \frac{\sqrt{2}}{\pi^{3/2} \sqrt{x}} \sum_{j=0}^{\infty} \frac{\Gamma(j+1/2)}{\Gamma(j+1)} \sqrt{2j + \frac{1}{2}} \exp\left\{-\frac{(4j+1)^2}{16x}\right\} K_{1/4}\left(\frac{(4j+1)^2}{16x}\right),$$

where $K_{1/4}(\cdot)$ denotes a modified Bessel function of the second type (see, e.g., function `besselk (nu,z)` in Matlab or function `BesselK` in Mathematica).

In the case of the Guetsch–Saentis differences the statistic (16), resp. (17), attains the value 12.9 which is highly significant regardless which approximation is used.

Example 2 (continuation)

The maximum-type statistics for the problem (2) have the form

$$\max_{1 \leq k < n} \left\{ \frac{\hat{b}_k}{s(\hat{b}_k)} \right\} \quad \text{and} \quad \max_{1 \leq k \leq \lfloor (1-\beta)n \rfloor} \left\{ \frac{\hat{b}_k}{s(\hat{b}_k)} \right\}, \quad (21)$$

where $\hat{b}_k$ is the least squares estimator of b and $s(\hat{b}_k)$ is the corresponding estimator of the standard deviation of $\hat{b}_k$ supposing that the change occurred at the moment k . For application of the Bonferroni inequality notice that if the errors are independent identically normally distributed then under H the variables

$$\frac{\hat{b}_k}{s(\hat{b}_k)} = \frac{\frac{1}{s_k} \frac{1}{\sqrt{n}} \sum_{i=k+1}^n (Y_i - \bar{Y}_n) \frac{i-k}{n}}{\sqrt{\frac{(n-k)(n-k+1)(n-k+1/2)}{3n^3} - \frac{(n-k)^2(n-k+1)^2}{4n^4}}}, \quad k = 1, \dots, n-1,$$

follow the t -distribution with $n-2$ degrees of freedom ($\{s_k, k = 1, \dots, n-1\}$ are the usual estimators of σ based on the residual sums of squares).

For n large, critical values can be obtained using the approximations ($x \in \mathcal{R}_1$):

$$P\left(\max_{1 \leq k \leq n-1} \left\{ \frac{\hat{b}_k}{s(\hat{b}_k)} \right\} > \frac{x + b_{n,2}}{a_n}\right) \approx 1 - \exp\{-e^{-x}\}, \quad (22)$$

where

$$a_n = \sqrt{2 \log \log n} \quad \text{and} \quad b_{n,2} = 2 \log \log n + \log \frac{\sqrt{3}}{4\pi},$$

$$P\left(\max_{1 \leq k \leq \lfloor (1-\beta)n \rfloor} \left\{ \frac{\hat{b}_k}{s(\hat{b}_k)} \right\} > x\right) \approx \frac{1}{\sqrt{\pi}} \phi(x) \int_0^{1-\beta} \frac{\sqrt{6t}}{(1-t)(1+3t)} dt.$$

The statistics (21) for the Jones temperature series attain the value 17.7 that is highly significant regardless which of the given approximations is applied.

Example 3 (continuation)

The maximum-type test statistics for the problem (3) are of the form

$$\max_{2 \leq k \leq n-2} \{F_k\} \quad \text{and} \quad \max_{\lfloor \beta n \rfloor \leq k \leq \lfloor (1-\beta)n \rfloor} \{F_k\}, \quad (23)$$

where

$$F_k = \frac{1}{s_k} \left(\frac{nk(\bar{Y}_k - \bar{Y}_n)^2}{n-k} + \frac{Q_{xy}^2(k)}{Q_{xx}(k)} + \frac{Q_{xy}^{o2}(k)}{Q_{xx}^o(k)} - \frac{Q_{xy}^2(n)}{Q_{xx}(n)} \right),$$

s_k being the usual estimator of the random errors variance σ^2 supposing that the change occurred at time k and

$$\begin{aligned} Q_{xx}(k) &= \sum_{i=1}^k (x_i - \bar{x}_k)(x_i - \bar{x}_k), & Q_{xy}(k) &= \sum_{i=1}^k (x_i - \bar{x}_k)(Y_i - \bar{Y}_k), \\ Q_{xx}^o(k) &= \sum_{i=k+1}^n (x_i - \bar{x}_k^o)(x_i - \bar{x}_k^o), & Q_{xy}^o(k) &= \sum_{i=k+1}^n (x_i - \bar{x}_k^o)(Y_i - \bar{Y}_k^o). \end{aligned}$$

If the random errors $\{e_i\}$ are independent identically normally distributed then under H all the variables $\{F_k, k = 2, \dots, n-2\}$ have F -distribution with 2 and $n-4$ degrees of freedom. For n large the asymptotic critical values can be calculated from the approximations ($x \in \mathcal{R}_1$):

$$P \left(\max_{2 \leq k \leq n-2} \{F_k\} > \left(\frac{x + b_{n,3}}{a_n} \right)^2 \right) \approx 1 - \exp \{ -2e^{-x} \}, \quad x \in \mathcal{R}_1, \quad (24)$$

where

$$a_n = \sqrt{2 \log \log n}, \quad b_{n,3} = 2 \log \log n + \frac{1}{2} \log \log \log n$$

and

$$P \left(\max_{\lfloor \beta n \rfloor \leq k \leq \lfloor (1-\beta)n \rfloor} \{F_k\} > x^2 \right) \approx x^2 e^{-x^2/2} \log \frac{1-\beta}{\beta}. \quad (25)$$

The statistic (23) attains the value 31.78. By applying the Bonferroni inequality the upper estimate of p -value is $8.24 \cdot 10^{-7}$ so that the null hypothesis claiming that the rainfalls-runoffs relationship is stationary is rejected.

In addition to maximum likelihood method the method of recursive residuals suggested by Brown et al. (1975) is very popular for decision whether a linear relationship between dependent and independent variables is stationary. The authors claim that their method is quite general as it is not aimed against any specific violation of stationarity.

3. Estimation of change point

The most popular method for change point estimation is the maximum likelihood method. If the change occurs in mean of normally distributed variables (as in our examples) then we deal with the problem of parameters estimation in non-linear regression model, so that the least squares method for estimation of change point may be applied. If the studied sequence is not normally distributed then non-parametric methods as well as some robust methods may be of use. In the case that the prior distribution of change point is available, the Bayesian methods are appropriate to apply.

The statisticians are usually not satisfied to find the change point estimator only but they prefer to know its distribution to be able to construct a confidence interval. The exact distribution is almost always very complex and one looks for an asymptotic distribution. Unfortunately, even if the asymptotic approach is applied the limit distribution of change point differs from one model to the other and depends heavily on the behavior of the studied sequence close to the change point, e.g. whether the change occurred suddenly or gradually.

Example 1 (continuation)

Denote by $\hat{\mu}$, $\hat{m}$ and $\hat{\delta}$ the least squares estimators of μ , m , resp. of δ , i.e. the values that minimize

$$\sum_{i=1}^k (Y_i - \mu)^2 + \sum_{i=k+1}^n (Y_i - \mu - \delta)^2, \quad (26)$$

and by $\hat{\sigma}^2$ the estimator of σ^2 , i.e.

$$\hat{\sigma}^2 = \frac{1}{n-2} \left\{ \sum_{i=1}^{\hat{m}} (Y_i - \bar{Y}_{\hat{m}})^2 + \sum_{i=\hat{m}+1}^n (Y_i - \bar{Y}_{\hat{m}}^o)^2 \right\}. \quad (27)$$

It holds

$$P \left(\frac{\hat{\delta}^2}{\hat{\sigma}^2} (\hat{m} - m) \leq x \right) \approx P(V \leq x), \quad x \in \mathcal{R}_1, \quad (28)$$

where

$$P(V \leq x) = \begin{cases} 1 + \sqrt{\frac{x}{2\pi}} e^{-x/8} - \frac{1}{2}(x+5)\Phi(-\frac{1}{2}\sqrt{x}) + \frac{3}{2}e^x\Phi(-\frac{3}{2}\sqrt{x}), & x \geq 0, \\ 1 - P(V \leq -x), & x < 0. \end{cases} \quad (29)$$

For the Guetsch–Saentis differences we get the estimators $\hat{m} = 120$, $\hat{\delta} = 0.4791$ and $\hat{\sigma}^2 = 0.082728$. By applying (28) (97.5% quantile of (29) being 11.033) we estimate that with the 95% probability the change occurred between the 116th and 124th observations.

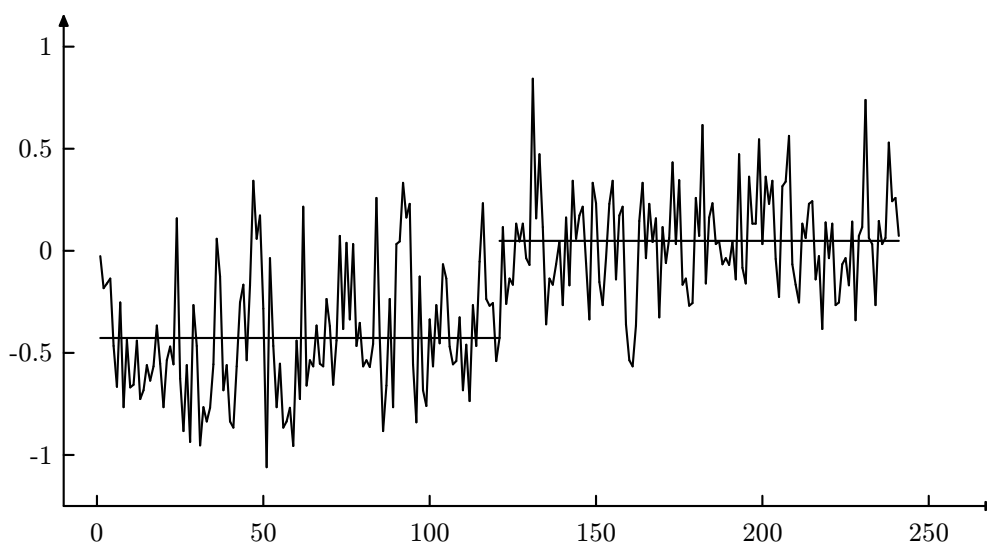


Figure 1. Guetsch-Saentis: Data and model

Example 2 (continuation)

Denote the least squares estimators of $\mathbf{m}$ by $\hat{\mathbf{m}}$, that of b by $\hat{b}$ and the estimator of σ based on residual sum of squares by $\hat{\sigma}$, then $\forall x \in \mathcal{R}_1$

$$P\left(\frac{\hat{b}\hat{\mathbf{m}} - \mathbf{m}}{\hat{\sigma}\sqrt{n}}\sqrt{\frac{(\hat{\mathbf{m}}/n)(1 - \hat{\mathbf{m}}/n)}{1 + 3\hat{\mathbf{m}}/n}} \leq x\right) \approx \Phi(x). \quad (30)$$

By applying (30) with $\hat{\mathbf{m}} = 54$ (which corresponds to the year 1907), $\hat{b} = 0.6943$ and $\hat{\sigma} = 0.117$, we estimate with the probability 95% that the linear trend appeared between the years 1906–1908. Figure 2 shows the data together with the optimal model.

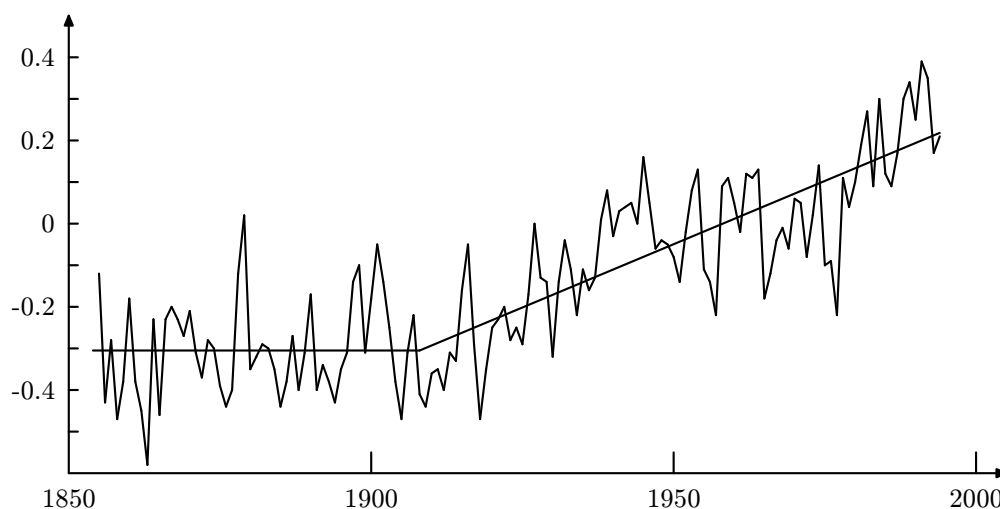


Figure 2. Global temperature anomalies: Data and model

Example 3 (continuation)

The least squares estimator of the change point $\mathbf{m}$ is equal $\hat{\mathbf{m}} = 26$. There exists asymptotic methods to find the limit distribution of $\hat{\mathbf{m}}$, see Antoch et al. (2001). Unfortunately, the number of observations $n = 36$ seems to be too small to use them.

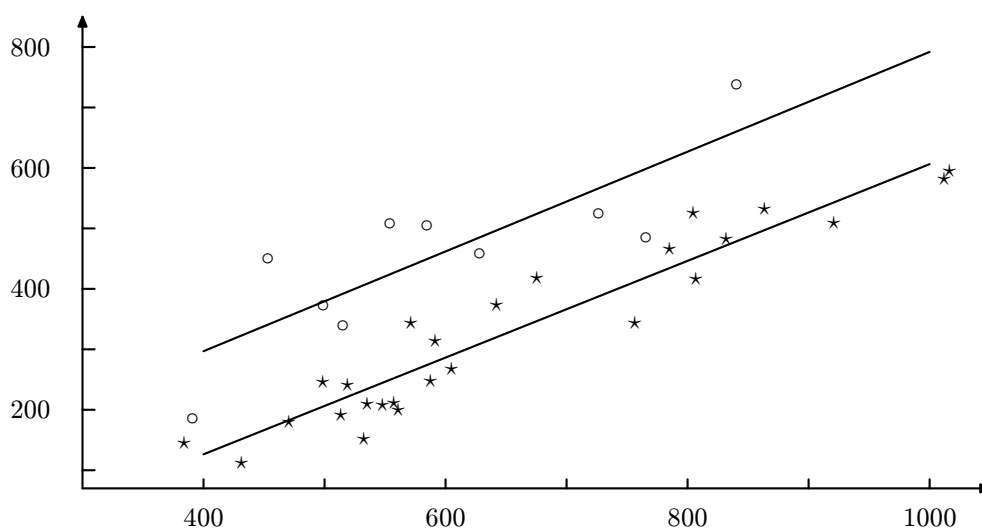


Figure 3. Malá Ráztoka: Data and model.

Figure 3 shows the linear relationship between rainfalls and run-offs in the first as well in the second time period; i.e. $y = -193.6 + 0.8x$ and $y = -33.1 + 0.82x$. By $\star$ we denote

the observations related to the first 26 years of observation, while by $\circ$ the observations from the last 10 years.

4. Problems

Dependence

One of the most typical features of the hydrological and meteorological data is dependence among neighboring observations caused by a certain persistence in the behavior of nature. The dependence has a great impact on decision whether a series is stationary as well as on the length of the confidence interval for the change point. In the models corresponding to Example 1 and 2 it is possible to show that if we suppose that the studied series form a stationary ARMA sequence, then the critical values for test statistics derived from the maximum likelihood method have to be multiplied by $\sqrt{2\pi h(0)/\gamma}$ where $h(\cdot)$ denotes a spectral density of the process $\{e_t\}$ and $\gamma = \text{var } e_t$. The limits of confidence intervals has to be adapted in the same way. Some more details can be found in Antoch et al. (1997) or Bai (1994).

5. Multiple changes

If the number of possible change points is known, we can proceed similarly as before using, e.g., the maximum likelihood method. If the number is unknown its determination is a difficult problem. The most often applied approach is to use some information criterion. If we think that the change points are not located too close to each other then the method of a moving window may be successful. For the situation described in Example 1 the method of the sequential splitting can function well.

Bibliography

- [1] Albin J. M. P. (1990). *On extremal theory for stationary processes*. Ann. Probab. **18**, 92–108.
- [2] Alexandersson H. (1986). *A homogeneity test applied to precipitation data*. J. Climatol. **6**, 661–675.
- [3] Antoch J., Hušková M. and Jarušková D. (2001). *Off-line process control*. In: Multivariate Total Quality Control. Foundation and Recent Advances. N. C. (Lauro et al eds.), Springer, Heidelberg, 1–85.
- [4] Antoch J., Hušková M. and Prášková Z. (1997). *Effect of dependence on statistics for determination of change*. J. Stat. Plan. Infer. **60**, 291–310.
- Bai J. (1994). *Least squares estimation of a shift in linear processes*. J. Time Series Analysis **15**, 453–472.
- [5] Bhattacharya P. K. (1990). *Weak convergence of the log-likelihood process in the two-phase regression problem*. Proc. of the R. C. Bose Symposium on Probability, Statistics and Design of Experiments, Wiley Eastern, New Delhi, 145–156.
- [6] Brown R. L., Durbin J. and Evans J. M. (1975). *Techniques for testing the constancy of regression relationships over time (with discussion)*. JRSS **B 37**, 149–192.
- [7] Buishand T. A. (1984). *Tests for detecting a shift in the mean of hydrological records*. J. Hydrol. **73**, 51–69.

- [8] Chernoff H. and Zacks S. (1964). *Estimating the current mean of normal distribution which is subjected to changes in time*. Ann. Math. Statist. **35**, 999–1018.
- [9] Csörgő M. and Horváth L. (1997). *Limit Theorems in Change Point Analysis*. J. Wiley, New York.
- [10] Deshayes J. and Picard D. (1986). *Off-line statistical analysis of change point models using nonparametric and likelihood methods*. In: Lecture Notes in Control and Information Sciences **77**, Basseville M. et al. (eds.), Springer Verlag, New York, 103–168.
- [11] Gardner L. A. (1969). *On detecting changes in the mean of normal variates*. Ann. Math. Statist. **40**, 116–126.
- [12] Hinkley D. V. (1969). *Inference about the intersection in two-phase regression*. Biometrika **56**, 495–504.
- [13] James B., James K. L. and Siegmund D. (1987). *Tests for change points*. Biometrika **74**, 71–84.
- [14] Jarušková D. (1998). *Change-point detection for dependent data and application to hydrology*. Istatistik. Journal of the Turkish Statistical Association **1**, 9–21.
- [15] Jarušková D. (1996). *Change point detection in meteorological measurement*. Monthly Weather Review **124**, 1535–1543.
- [16] Jarušková D. (1997). *Some problems with application of change point detection methods to environmental data*. Environmetrics **8**, 469–483.
- [17] Jones P. D., Wigley M. L. and Briffa K. R. (1994). *Global and hemispheric temperature anomalies - land and maritime instrumental records*. In Boden T. A., Kaiser D. P., Sepanski R. J. and Stoss F. W. (eds.) Trends'93: A Compendium of Data on Global Change, ORNC/CDIAC-65, Carbon Dioxide Information Analysis Center, Oak Ridge National Laboratory, Oak Ridge, Tenn., USA, 603–608.
- [18] Kander Z. and Zacks S. (1966). *Test procedures for possible changes in parameters of statistical distributions occurring at unknown time points*. Ann. Math. Statist. **37**, 1196–1210.
- [19] MacNeill I. B. (1974). *Tests for change of parameter at unknown time and distribution of some related functionals of Brownian motion*. Ann. Statist. **2**, 950–962.
- [20] MacNeill I. B. (1978). *Properties of sequences of partial sums of polynomial regression residuals with applications to tests for change of regression at unknown times*. Ann. Statist. **6**, 422–433.
- [21] Rhoadas D. A. and Salinger M. J. (1993). *Adjustment of temperature and rainfalls records for site changes*. J. Climatol. **13**, 899–913.
- [22] Vannitsem S. and Nicolis C. (1991). *Detecting climatic transitions: Statistical and dynamical aspects*. Beitr. Phys. Atmosph. **64**, 245–254.
- [23] Worsley K. J. (1983). *Testing for a two-phase multiple regression*. Technometrics **25**, 35–42.
- [24] Yao Yi-Ching and Davis R. A. (1986). *The asymptotic behavior of the likelihood ratio statistic for testing a shift in mean in a sequence of independent normal variates*. Sankhya **48**, 339–353.

Authors' addresses

Jaromír Antoch, prof. RNDr CSc.
Charles University in Prague
Faculty of Mathematics and Physics
Department of Probability and Mathematical Statistics
Sokolovská 83, CZ – 186 75 Praha 8
Czech Republic
jaromir.antoch@mff.cuni.cz
<http://www.karlin.mff.cuni.cz/~antoch>

Daniela Jarušková, prof. RNDr CSc.
Czech Technical University
Faculty of Civil Engineering
Department of Mathematics
Thákurova 7, CZ – 166 29 Prague 6
Czech Republic
jarus@mat.fsv.cvut.cz

Acknowledgements:

The work has been supported by grants GA ČR 201/09/0755 and MSM 0021620839.

Pravdepodobnosť na algebraických štruktúrach

Probability on algebraic structures

Beloslav Riečan

Cieľom tejto prednášky je vzbudenie záujmu o výskum v oblasti, v ktorej bola práve vyvinutá nová metóda ([2],[3],[4],[5]). Táto metóda bola aplikovaná na MV-algebry a čiastočne aj D-posety.

Historia Magistra Vitae

V tejto kapitole budeme hľadať inšpiráciu v Kolmogorovovej teórii pravdepodobnosti na Booleových algebrách. Základný pojem je pravdepodobnostný priestor, čo je trojica $(\Omega, \mathcal{S}, P)$,

kde Ω je neprázdna množina,

$\mathcal{S}$ je σ -algebra podmnožín množiny Ω a P je pravdepodobnosť, t.j.

$P : \mathcal{S} \rightarrow [0, 1], P(\Omega) = 1,$

P – σ -aditívna,

$$A_n \in \mathcal{S} (n = 1, 2, \dots), A_n \cap A_m = \emptyset (n \neq m) \implies P\left(\bigcup_{n=1}^{\infty} A_n\right) = \sum_{n=1}^{\infty} P(A_n)$$

V Kolmogorovovej teórii sa vyskytujú 3 základné pojmy. Okrem pravdepodobnosti ešte pojem náhodnej premennej a jej strednej hodnoty. Aj z metodického hľadiska je vhodná nasledujúca formulácia náhodnej premennej.

$$\begin{aligned} \xi : \Omega &\rightarrow R \\ \forall t \in R : \xi^{-1}((-\infty, t)) &= \{\omega \in \Omega; \xi(\omega) < t\} \in \mathcal{S} \end{aligned}$$

Náhodná premenná môže byť charakterizovaná distribučnou funkciou

$$\begin{aligned} F : R &\rightarrow [0, 1] \\ F(t) &= P(\xi^{-1}((-\infty, t))) \end{aligned}$$

Označme

$$\mathcal{K} = \{A \subset R; \xi^{-1}(A) \in \mathcal{S}\} \supset \mathcal{J},$$

kde $\mathcal{J}$ je systém všetkých intervalov typu $(-\infty, t), t \in R$. Potom $\mathcal{K}$ je σ - algebra. Preto

$$\mathcal{K} \supset \sigma(\mathcal{J}) = \mathcal{B}(R),$$

teda

$$A \in \mathcal{B}(R) \implies \xi^{-1}(A) \in \mathcal{S}.$$

Spomeňme si teraz na známy výsledok z teórie miery: Ak

$$F : R \rightarrow [0, 1]$$

F je neklesajúca, spojitá zľava,

$$\lim_{t \rightarrow \infty} F(t) = 1, \lim_{t \rightarrow -\infty} F(t) = 0,$$

tak existuje práve jedna taká pravdepodobnosť

$$\lambda_F : \mathcal{B}(R) \rightarrow [0, 1],$$

že

$$\lambda_F([a, b)) = F(b) - F(a)$$

pre všetky $a \leq b$.

Pozrime sa teraz na druhú koncepciu rozdelenia pravdepodobnosti. Opäť máme pravdepodobnosť,

$$(\Omega, \mathcal{S}, P), P : \mathcal{S} \rightarrow [0, 1]$$

a náhodnú premennú

$$\xi : \Omega \rightarrow R.$$

Definujme

$$P_\xi : \mathcal{B}(R) \rightarrow [0, 1]$$

rovnosťou

$$P_\xi(A) = P(\xi^{-1}(A)).$$

P_ξ je pravdepodobnostná miera.

Porovnajme teraz tieto 2 koncepcie: Daná je náhodná premenná

$$\xi : \Omega \rightarrow R,$$

F jej distribučná funkcia,

$$\begin{aligned} F(t) &= P(\xi^{-1}((-\infty, t)) = P_\xi((-\infty, t)), \\ \lambda_F : \mathcal{B}(R) &\rightarrow [0, 1] \\ \lambda_F([a, b)) &= F(b) - F(a). \end{aligned}$$

Potom

$$\begin{aligned} P_\xi([a, b)) &= P(\xi^{-1}([a, b))) = \\ &= P(\xi^{-1}((-\infty, b)) \setminus \xi^{-1}((-\infty, a))) = \\ &= P(\xi^{-1}((-\infty, b))) - P(\xi^{-1}((-\infty, a))) = \\ &= F(b) - F(a). \end{aligned}$$

Preto

$$\lambda_F = P_\xi.$$

Tretím základným pojmom Kolmogorovovej teórie je pojem strednej hodnoty definovaný pomocou Lebesguovho integrálu:

$$E(f \circ \xi) = \int_{\Omega} f \circ \xi dP = \int_R f dP_\xi = \int_R f(x) dF(x).$$

Práve uvedená formulácia umožňuje definovať jednoduchým spôsobom strednú hodnotu aj v algebraických štruktúrach.

Algebraické štruktúry

Nebudeme precizovať pojem algebraickej štruktúry. Totiž práve výskum tejto témy by mal viesť k okruhu štruktúr, na ktoré je možno aplikovať metódu, ktorú tu propagujeme. Predsa však príkladom, na ktorý sa táto metóda dá použiť je MV-algebra. Klasickým príkladom MV-algebry je jednotkový interval

$$M = [0, 1]$$

s operáciami

$$\begin{aligned} a \oplus b &= (a + b) \wedge 1 \\ \chi_A \oplus \chi_B &= \chi_{A \cup B} \\ a \odot b &= (a + b - 1) \vee 0 \\ \chi_A \odot \chi_B &= \chi_{A \cap B} \end{aligned}$$

Stav je zobrazenie

$$m : M \rightarrow [0, 1],$$

pozorovateľná je zobrazenie

$$\begin{aligned} x : \mathcal{J} &\rightarrow M \\ \mathcal{J} &= \{(-\infty, t); t \in R\}. \end{aligned}$$

Zložené zobrazenie

$$m \circ x : \mathcal{J} \rightarrow [0, 1]$$

umožňuje definovať distribučnú funkciu

$$F(t) = m(x((-\infty, t)))$$

=

$$P(\xi^{-1}((-\infty, t))).$$

Pokiaľ ide o MV-algebry, v minulosti bola rozpracovaná iná koncepcia ([6]). Stav je síce zobrazenie

$$m : M \rightarrow [0, 1],$$

ale pozorovateľná je definovaná na Borelových množinách

$$x : \mathcal{B}(R) \rightarrow M.$$

Zložené zobrazenie je mierou, ale na Borelových množinách

$$m \circ x : \mathcal{B}(R) \rightarrow [0, 1]$$

$$m_x(A) = m(x(A)).$$

To odpovedá koncepcii rozdelenia pravdepodobnosti

$$P_\xi(A) = P(\xi^{-1}(A))$$

. V prípade MV-algebrií bola táto koncepcia uvedená v práci [6].

Pravdaže, v oboch koncepciách možno definovať strednú hodnotu pomocou Stieltje-
sovhovho integrálu

$$\begin{aligned} E(f \circ \xi) &= \int_R f d\lambda_F \\ E(x) &= \int_R t dF(t) (= \sum_i t_i p_i, \text{ resp. } \int_R t \varphi(t) dt), \\ D(x) &= \int_R t^2 dF(t) - E(x)^2. \end{aligned}$$

Majme teda

$$\begin{aligned} m : M &\rightarrow [0, 1] \\ x : \mathcal{J} &\rightarrow M \\ F(t) &= m(x((-\infty, t))). \end{aligned}$$

Ale čo je $x + y$?

Opäť budeme hľadať inšpiráciu u Kolmogorova. Nech

$$\xi, \eta : \Omega \rightarrow R$$

sú náhodné premenné. Potom

$$\begin{aligned} T &= (\xi, \eta) : \Omega \rightarrow R^2 \\ g : R^2 &\rightarrow R, g(u, v) = u + v \\ \xi + \eta &= g \circ T \\ (\xi + \eta)^{-1} &= T^{-1} \circ g^{-1}, T^{-1} : \mathcal{B}(R) \rightarrow \mathcal{S} \\ g^{-1}((-\infty, t)) &= \{(u, v); u + v < t\}. \\ x + y : \mathcal{J} &\rightarrow M \\ x + y &= h \circ g^{-1} \\ (x + y)((-\infty, t)) &= h(g^{-1}((-\infty, t))) = h(\Delta_t), t \in R \\ h : \{\Delta_t; t \in R\} &\rightarrow M. \end{aligned}$$

Na riešenie problému sú v danej štruktúre potrebné dve tvrdenia. Treba zostrojiť zobrazenie h tak, aby

1.

$$\begin{aligned} z((-\infty, t)) &= h(g^{-1}((-\infty, t))) \\ z : \mathcal{J} &\rightarrow M \end{aligned}$$

bola pozorovateľná

2. Ak

$$x_1, x_2, x_3, \dots : \mathcal{J} \rightarrow M$$

je postupnosť nezávislých pozorovateľných, tak

$$m \circ h(\Delta_t^n) = \lambda_{F_1} \times \dots \times \lambda_{F_n}(\Delta_t^n), t \in R$$

Tá druhá vlastnosť je potrebná na to, aby sme mohli použiť Kolmogorovovu vetu o konzistencii, totiž reprezentovať konkrétnu postupnosť pozorovateľných pomocou Kolmogorovho pravdepodobnostného prietoru

$$(R^N, \sigma(\mathcal{C}), P)$$

$$P \circ \pi_n^{-1} = \lambda_{F_1} \times \dots \times \lambda_{F_n}$$

Literatúra

- [1] Dvurečenskij, A., Pulmannová, S.: New trends in quantum structures. Kluwer Academic Publisher 2000.
- [2] Kelemenová, J., Kuková, M.: Central limit theorem on MV-algebras. In: MANYVAL'2010, Varese, 22 - 25.
- [3] Lašová, L.: Individual ergodic theorem on Kôpka D-posets. In: MANYVAL'2010, Varese, 25 - 27.

- [4] Riečan, B.: On a new approach to probability theory on MV-algebras (to appear).
- [5] Riečan, B., Lašová, L.: On the probability on the Kôpka D-posets (to appear).
- [6] Riečan, Mundici: Probability on MV-algebras. In: Handbook on Measure Theory (E.Pap ed.), Elsevier 2002, 869 - 909.

Beloslav Riečan, Dr.h.c. prof. RNDr., DrSc.
Faculty of Natural Sciences, Matej Bel University
Department of Mathematics
Tajovského 40
974 01 Banská Bystrica, Slovakia
and
Mathematical Institute of Slovak Acad. of Sciences
Štefánikova 49
SK-81473 Bratislava
e-mail: *riecan@fpv.umb.sk*

Kvantifikácia biodiverzity Quantification of biodiversity

Mária Bauerová, Ján Brindza, Beáta Stehlíková, Anna Tirpáková

Abstract: The paper presents measures of biodiversity. There are described the most common measures of biodiversity: Index species richness, Simpson index and Shannon-Wiener index, whose calculation is illustrated with an example.

Key words: Measures of biodiversity, Index species richness, Simpson index, Shannon-Wiener index

Kľúčové slová: Miery biodiverzity, Index druhového bohatstva, Simpsonov index, Shannon-Wienerov index

1. Úvod

Svetový summit OSN v roku 1992 v Rio de Janeiro definoval *biodiverzitu ako rôznorodosť všetkých živých organizmov vrátane existujúcich suchozemských, morských a ostatných vodných ekosystémov a ekologických komplexov, ktorých sú súčasťou, čo zahŕňa rôznorodosť v rámci druhov, medzi druhmi a ekosystémami*.

Komplexnosť konceptu biodiverzity sa odráža aj vo veľkom počte definícií tohto pojmu. Jutro (1993) identifikoval v odbornej literatúre až 14 definícií, ktoré sú široko používané, uvádzané a prijaté jednotlivými organizáciami a krajinami.

Wilson (1988) vidí biodiverzitu ako transformáciu biológov z pozície častíc a prvkov na holistickú pozíciu, ktorá zdôrazňuje celistvosť a pokladá celok za niečo vyššie ako súhrn častí. Podľa Harpera a Hawkswortha (1994) to boli Norse a kol. (1986) tí, ktorí zaviedli tradičné používanie troch úrovní biologickej diverzity – genetickej, druhovej, a ekologickej. Szaro a Shapiro (1990), Noss (1992) a ďalší autori zaviedli novú úroveň biodiverzity – krajinnú.

V literatúre bolo popísaných mnoho mier biodiverzity. Norton (1994) dokazuje, že jednoduchá objektívna vedecká definícia biodiverzity v zmysle definovania spôsobu ako ju merať neexistuje. Norton (1994) tvrdí, že čím viac porozumieme biodiverzite, tým menej je pravdepodobné, že budeme chcieť skonštruovať jednoduchú objektívnu mieru biodiverzity. Roebuck a Phifer (1999) argumentujú, že zachovanie biodiverzity je zakorenené primárne v etike a my nesmieme argumentovať hodnotami ukazovateľov.

Je potrebné rozlišovať dve poňatia genetickej biodiverzity – druhovú bohatosť a relatívne zastúpenie druhov. *Druhové bohatstvo* hovorí o celkovom počte rozličných genotypov v populácii. *Relatívne zastúpenie druhov* hovorí o relatívnej početnosti jednotlivých genotypov v danej populácii.

2. Miery biodiverzity

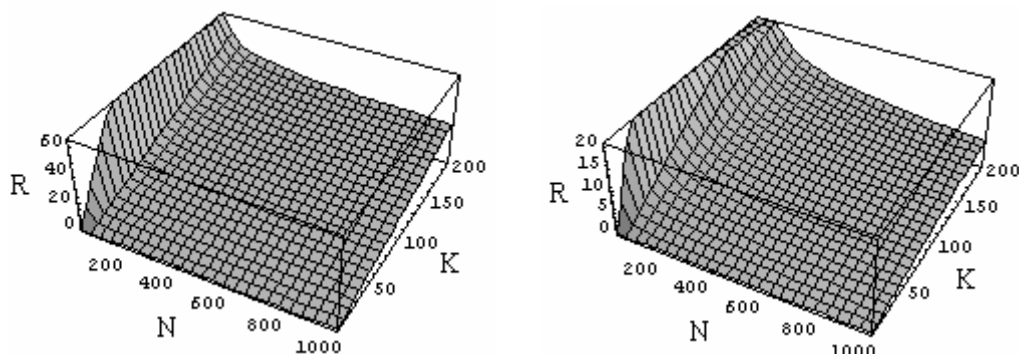
V literatúre sa najčastejšie uvádzajú nasledovné tri miery pre meranie biodiverzity: *Index druhového bohatstva*, *Simpsonov index* a *Shannon-Wienerov index*. Index druhového bohatstva je

$$R = \frac{K - 1}{\ln N},$$

kde K je počet druhov, N je počet organizmov. Známý je tiež pod menom *Margalefov index*. Index druhového bohatstva sa uvádza aj v tvare

$$R = \frac{K}{\sqrt{N}},$$

kde K je počet druhov, N je číslo organizmov. Vyššia hodnota indexu korešponduje s vyššou biodiverzitou. V literatúre je známy tiež pod menom *Menhinickov index* (Obrázok 1).



Zdroj: Vlastné zobrazenie

Obrázok 1: Margalefov index a Menhinickov index

Výberové metódy priestorovej štatistiky (Stehlíková, 2003) umožňujú dosiahnuť reprezentatívny výber. Napriek tomu, určovanie počtu organizmov je zložitá odborná záležitosť. Napríklad odnože jedinej rastliny z rodu *Iris* môžu pokryť územie niekoľkých metrov štvorcových.

Simpsonov index meria pravdepodobnosť, že dva organizmy náhodne vybrané sú toho istého druhu

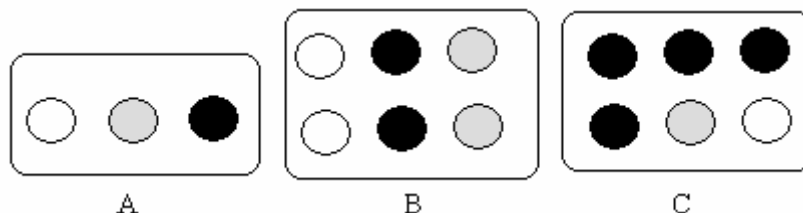
$$D = \sum_{i=1}^K \left(\frac{n_i}{N} \right)^2,$$

respektíve,

$$D = \frac{\sum_{i=1}^K n_i(n_i - 1)}{N(N - 1)},$$

kde K je počet druhov, N je počet organizmov, n_i je počet organizmov i -teho druhu ($i = 1, 2, \dots, K$). Index D nadobúda hodnotu z uzavretého intervalu $\langle 1/K, 1 \rangle$. Čím je hodnota indexu D vyššia, tým je diverzita nižšia. Hodnota $D = 1/K$ predstavuje najvyššiu diverzitu. $D = 1$ práve vtedy, keď neexistuje žiadna diverzita. Obe verzie sú v literatúre akceptované. V literatúre sa doporučuje sa, aby počet organizmov N bol väčší ako 1000.

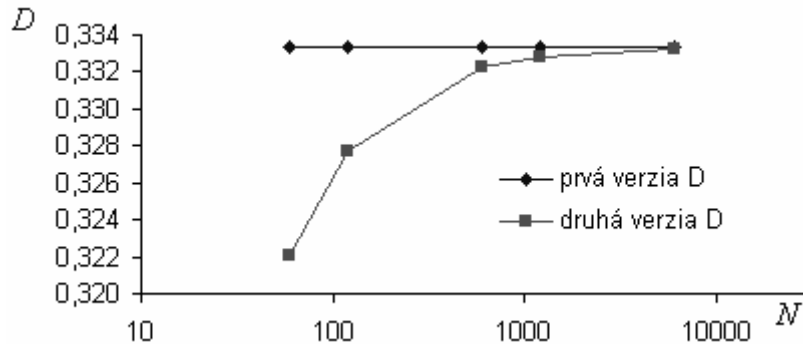
Ilustrujme výpočet jednotlivých mier pre spoločenstvá, ktoré sú znázornené na obrázku 2. Na obrázku 2 spoločenstvá A , B , C majú rovnaký počet druhov (3). Relatívne zastúpenie druhov v spoločenstve A a B je rovnaké (1/3).



Zdroj: Vlastné zobrazenie

Obrázok 2: Rôzne spoločenstvá (A , B , C) druhov (každý druh je znázornený inou farbou)

V tabuľke 1 sú hodnoty *Simpsonovho indexu* počítané podľa prvého, resp. druhého vzťahu odlišné a zvyknú sa uvádzať ako dva odlišné indexy. Nie je to úplne správne, lebo pre veľké hodnoty N sa oba indexy rovnajú, t.j. sú asymptoticky zhodné. Hodnoty oboch verzií Simpsonovho indexu pre $K = 3$, $n_i = N/3$ ($i = 1, 2, 3$) a rôzne hodnoty N (60, 120, 600, 1200, 6000) v logaritmickú škálu sú na obrázku 3.



Zdroj: Vlastné zobrazenie

Obrázok 3: Dve verzie Simpsonovho indexu

Pre kvantifikáciu biodiverzity sa používa tiež hodnota $1 - D$, tzv. *Simpsonov index diverzity* a *Simpsonov reciproký index*. Hodnota $1 - D$ je pravdepodobnosť, že dva náhodne vybrané organizmy sú rozličných druhov. Simpsonov index diverzity nadobúda hodnoty od 0 do $1 - 1/K$. Simpsonov reciproký index sa často označuje ako *Hillov index* N_2 . Nadobúda hodnoty z intervalu od 1 do K .

Shannonov-Weaverov index celkovej rozmanitosti spája v sebe aspekt druhovej pestrosti aj vyrovnanosti. Biodiverzita z hľadiska *Shannonovho indexu* nie je len o tom, koľko biologických druhov žije na skúmanom území, ale aj o tom, ako často sa s týmito druhmi môžeme stretnúť. Pri rovnakom počte druhov v dvoch výberových súboroch bude mať vyššiu hodnotu indexu ten, v ktorom sú jednotlivé druhy rovnomernejšie zastúpené. *Shannonov index* je

$$H = - \sum_{i=1}^K p_i \ln p_i,$$

kde K je počet druhov, $p_i = \frac{n_i}{N}$, pričom n_i je počet organizmov i -teho druhu ($i = 1, 2, \dots, K$) a

$N = n_1 + n_2 + \dots + n_K$. Hodnota H je maximálna ($\ln K$), ak

$$p_1 = p_2 = \dots = p_K.$$

Hodnota *Shannonovho indexu* je vo všeobecnosti od hodnoty 0 do hodnoty $\ln K$. V praktických pozorovaniach nadobúda hodnoty od 1 do 4,5. *Shannonov index* je citlivý na rozsah výberu. V prípade dlhodobých pozorovaní biodiverzity, t.j. vývoja biodiverzity v čase sa najčastejšie používa práve *Shannonov index*. Príčinou chyby pri odhade diverzity pomocou tohto indexu je nemožnosť zaradiť do výberu všetky druhy reálneho spoločenstva.

Index celkovej rozmanitosti sa uvádza aj v tvare

$$J = \frac{H}{K},$$

kde H je *Shannonov-Weaverov index celkovej rozmanitosti* a K je počet druhov. Takýto tvar indexu zohľadňuje počet druhov K v spoločenstve. Z obdobnej filozofie vychádza index

$$E = \frac{\exp(H)}{K},$$

kde H je *Shannonov-Weaverov index* celkovej rozmanitosti a K je počet druhov.

Tabuľka 1: Miery biodiverzity pre spoločenstvá z obrázku 1

Spoločenstvo	K	N	Index druhového bohatstva R	Index druhového bohatstva R	Simpsonov index D	Simpsonov index D
A	3	3	1,820	1,732	0,333	0,000
B	3	6	1,116	1,225	0,333	0,200
C	3	6	1,116	1,225	0,500	0,400

Zdroj: Vlastné prepočty

Tabuľka 1: Miery biodiverzity pre spoločenstvá z obrázku 1 (dokončenie)

Spoločenstvo	Shannonov-Weaverov index H	Maximum H	Koeficient J	Koeficient E
A	1,099	1,099	1,000	1,000
B	1,099	1,099	1,000	1,000
C	0,868	1,099	0,790	0,794

Zdroj: Vlastné prepočty

Ďalším indexom pre hodnotenie biodiverzity je *Brillouin index*, ktorý sa vypočíta podľa vzťahu

$$HB = \frac{\ln N! - \sum_{i=1}^K \ln n_i!}{N}$$

kde K je počet druhov, $p_i = \frac{n_i}{N}$, pričom n_i je počet organizmov i -teho druhu ($i = 1, 2, \dots, K$) a $N = n_1 + n_2 + \dots + n_K$. V literatúre sa odporúča použiť vtedy, ak nemôže byť garantovaný náhodný výber.

Bergerov - Parkerov index sa rovná hodnote najvyššieho podielu z podielov pre všetky druhy v danej spoločnosti

$$d = \max \{p_i; i = 1, 2, \dots, K\},$$

kde K je počet druhov, $p_i = \frac{n_i}{N}$, pričom n_i je počet organizmov i -teho druhu ($i = 1, 2, \dots, K$).

Zo štatistického hľadiska sa odporúča opakovanie menších rozsahov, čo umožňuje konštrukciu intervalov spoľahlivosti, namiesto jedného veľkého výberu.

Buzas, Hayek (1998) a Small, McCarthy (2002) uprednostňujú *SHE* analýzu (S-Species richness (druhovú bohatosť), H-information (informácia), E-evenness (relatívne zastúpenie)). *SHE* analýza je rozklad *Shannonovho-Weaverovho indexu celkovej rozmanitosti* na dve zložky, ktorý vyplýva priamo zo zlogarovania vzťahu pre výpočet koeficienta E . Platí

$$H = \ln(K) + \ln(E).$$

3. Záver

Vypočítané miery biodiverzity pre spoločenstvá sú ukážkami kvantifikácie biodiverzity. Biodiverzita je, ako už bolo povedané, komplikovaný, zložitý a mnohotvárný

fenomén. Preto je nemožné komplexne sumarizovať biodiverzitu a vyjadriť pomocou jednoduchého indexu. Viacerí autori preto pristupujú zložitejším metódam.

4. Literatúra

- [1] ALATALO, R. V. 1981. Problems in the measurement of evenness in ecology. *Oikos* 37: 199-204
- [2] Biodiversity online, Copyright © 2000 Drs. Jacqueline McLaughlin and Stam Zervanos
<https://royercenter.cwc.psu.edu/biodiversity/>
- [3] BULLA, L. 1994. An index of evenness and its associated diversity measure. *Oikos* 70: 167-171.
- [4] HARPER, J. L. - HAWKSWORTH, D. L. 1994. Biodiversity - measurement and estimation - preface. *Philosophical Transactions of the Royal Society of London. Series B: Biological Sciences* 345, 5-12.
- [5] JUTRO, P.R. 1993. Human influences on ecosystems: dealing with biodiversity, in M.J. McDonnell and S.T.A. Pickett (eds) *Humans as Components of Ecosystems*, pp. 246-256. New York: Springer. ISBN 0-387-98243-d
- [6] NORSE, E.A.-ROSENBAUM, K.L.-WILCOVE, D.S.-WILCOX, B.A.-ROMME, W.H.-JOHNSTON, D.W.-STOUT, M.L. 1986. Conserving biological diversity in our national forests. The Wilderness Society, Washington, DC.
- [7] NORTON, B. G. 1994. On what we should save: the role of cultures in determining conservation targets. (eds. P. Forey et al.) *Systematics and conservation evaluation*.
- [8] NORSE, E.A.- MCMANUS, R.E. 1980. Ecology and living resources biological diversity. In: Council on Environmental Quality: The eleventh annual report of the Council on Environmental Quality, Washington D.C.: 31-80
- [9] NOSS, R. F. 1992. The Wildlands Project: land conservation strategy. *Wild Earth* (Special Issue): 10-25.
- [10] ROEBUCK, P.-PHIFER, P. 1999. The persistence of positivism in conservation biology. *Conservation Biology* 13, 444-446.
- [11] SZARO A.- SHAPIRO, B. 1990. Conserving Our Heritage: America's Biodiversity. The Nature Conservancy, Arlington, Virginia: The Nature Conservancy. 16 p.
- [12] STEHLÍKOVÁ, B. 2003. *Priestorová štatistika*. SPU Nitra, ISBN 80-8069-046-4
- [13] WILSON, E.O. 1988. Biodiversity. New York: National Academy Press, 538 p. ISBN-10: 0-309-03739-5, ISBN-13: 978-0-309-03739-6

Adresy autorov:

Mária Bauerová, prof. RNDr., CSc.
Katedra botaniky a genetiky
FPV UKF
Nábřežie mládeže 91,
949 74 Nitra
mbauerova@ukf.sk
Ján Brindza, Doc. Ing., CSc.
Katedra genetiky a šľachtenia rastlín,
FAPZ SPU
A. Hlinku 2,
949 76 Nitra
Jan.Brindza@uniag.sk

Beáta Stehlíková, prof. RNDr., CSc.
Katedra Matematiky
FPV UKF
Tr. A. Hlinku 1
949 01 Nitra
bstehlikova@ukf.sk
Anna Tirpáková, prof. RNDr., CSc.
Katedra Matematiky
FPV UKF
Tr. A. Hlinku 1
949 01 Nitra
atirpakova@ukf.sk

Porovnávanie mier biodiverzity s využitím štatistických metód Comparing measures of biodiversity using statistical methods

Ján Brindza, Mária Bauerová, Beáta Stehlíková, Anna Tirpáková

Abstract: The contribution measures of biodiversity indices compared with the selected statistical methods: Spearman correlation coefficient and the multiplicity of empirical verification of compliance and frequency of individuals of each species calculated using the Power law model was described swallowed a χ^2 - test. There are presented methods for comparing the diversity of communities - method Rényiho diversity ranks also way to estimate the number of species: jackknife estimate and the estimate by bootstrapping.

Key words: Spearman correlation coefficient, Measures of biodiversity, Method Rényiho diversity ranks, Jackknife estimate, Estimate by bootstrapping

Kľúčové slová: Spearmanov korelačný koeficient, Miery biodiverzity, Metóda Rényiho radov diverzity, Jackknife odhad, Odhad pomocou bootstrappingu

1. Úvod

Biodiverzita je zastrešujúci pojem pre popis organizmov, ktoré sa nachádzajú na svete, t.j. pre popis počtu, rozmanitosti a variability žijúcich organizmov. Biodiverzita je rozsiahly a komplexný koncept, ktorého dôsledky hlboko zasahujú do všetkých sfér života a aktivity ľudí, píše Krishnamurthy (2004). Norton (1994, 2001) poukázal na to, že jednotlivé prístupy k definovaniu biodiverzity opisujú cesty zodpovedajúce daným presným zámerom a výber medzi týmito rozličnými modelmi biodiverzity bude vždy závisieť na tom, aké hodnoty budú dôležité pre rozhodovateľa. Tento prístup je charakterizovaný ako post-pozitivistický, pretože uznáva biodiverzitu ako nevyhnutnú hodnotovo zaťaženú kategóriu. Podľa Di Castri a Younes (1996) biodiverzita je komplex s interakciami jej hierarchických zložiek: genetickej, taxonomickej a ekologickej na rozličnej úrovni integrácie. Podľa Brindzu (2001) sa biodiverzita spravidla delí na nasledovné kategórie - ekosystémová, druhová, genetická a kultúrna, čo je v súlade s členením UNEP. Di Castri a Younes (1996) konštatujú, že biodiverzitu nemožno chápať ako jednoduchý zastrešujúci pojem nad heterogénnymi aktivitami, ale ako zloženú entitu. Identifikácia je prvý krok k determinovaniu stavu zložiek biologickej diverzity. Komplexné biologické a ekologické inventarizácie sú jednou zo základných metód získavania poznatkov využiteľných na ochranu biodiverzity a trvalo udržateľné využívanie biologických zdrojov. K tomu, aby sa mohla hodnotiť biodiverzita, je potrebné vyjadriť ju v hodnotách, ktoré sa dajú porovnávať, t.j. musí sa merať.

Pre výpočet mier diverzity existujú špeciálne softvéry, napríklad, EstimateS¹, Ecosim², Species Diversity and Richness Calculates³, Exeter Software⁴, PowerNiche⁵.

¹ <http://viceroy.eeb.uconn.edu/EstimateS>

² <http://www.garyentsminger.com/ecosim/index.htm>

³ <http://www.pisces-conservation.com/>

⁴ <http://www.exetersoftware.com/cat/ecometh/ecomethodology.html>

⁵ <http://www.entu.cas.cz/png/powerniche/index.html>

2. Porovnávanie indexov pre hodnotenie biodiverzity

Pre meranie biodiverzity existuje viacero indexov pre hodnotenie biodiverzity. Beisell a kol. (2003) skúmali závislosť a citlivosť vybraných indexov a zisťovali ich závislosť použitím Spearmanovho koeficienta korelácie. Pre simulovaných 189 spoločností sú hodnoty Spearmanových korelačných koeficientov medzi uvedenými indexmi uvedené v tabuľke 1.

Tabuľka 1: Spearmanove korelačné koeficienty pre jednotlivé miery diverzity

Index	E_{Pielou}	$E_{Hurlbert}$	E_{Heip}	E_{I-D}	$E_{I/D}$	$E_{-ln(D)}$	$F_{2,1}$	$G_{2,1}$	O	E	E_{MI}	E'
$E_{Hurlbert}$	0,94											
E_{Heip}	0,97	0,89										
E_{I-D}	0,96	0,97	0,88									
$E_{I/D}$	0,85	0,76	0,94	0,76								
$E_{-ln(D)}$	0,98	0,95	0,95	0,97	0,87							
$F_{2,1}$	0,68	0,69	0,71	0,72	0,81	0,79						
$G_{2,1}$	0,68	0,69	0,71	0,72	0,81	0,79	1,00					
O	0,93	0,84	0,97	0,80	0,90	0,89	0,60	0,60				
E	0,97	0,91	0,96	0,89	0,83	0,93	0,59	0,59	0,97			
E_{MI}	0,97	0,97	0,92	0,99	0,81	0,99	0,76	0,76	0,84	0,91		
E'	0,89	0,78	0,97	0,77	0,96	0,87	0,66	0,66	0,97	0,90	0,81	
E_{var}	0,69	0,51	0,72	0,53	0,59	0,59	0,19	0,19	0,74	0,73	0,56	0,75

Legenda: $E_{Hurlbert}$ - Hurlbertov index, E_{Pielou} - Pielouov index, E_{Heip} - Heipov index, E_{I-D} - index bez pomenovania Smith a Wilson (1996), $E_{I/D}$ - index bez pomenovania Smith a Wilson (1996), $E_{-ln(D)}$ - index bez pomenovania Smith a Wilson (1996), $F_{2,1}$ - Hillov modifikovaný index, $G_{2,1}$ - index bez pomenovania Molinari (1989), O - index bez pomenovania Bulla (1994), E - Bullov index, E_{MI} - McIntoshov index, E' - Camargov index, E_{var} - Smithov a Wilsonov index A

Zdroj: Smith a Wilson, 1996

Všetky hodnoty Spearmanovho korelačného koeficienta sú štatisticky preukazné (pre $p < 0,01$). Smith, J.B. a Wilson, B. (1996) skúmali pomocou simulácií aj citlivosť jednotlivých indexov na výskyt zriedkavých, dominantných a priemerne sa vyskytujúcich druhov. Zistili, že na výskyt zriedkavých druhov sú citlivé indexy E_{Heip} , $E_{I/D}$ a E' , na výskyt priemerne sa vyskytujúcich druhov sú citlivé indexy $E_{Hurlbert}$ a $G_{2,1}$, pričom index $G_{2,1}$ je citlivý aj na zmenu výskytu početnosti dominantných druhov.

Jednou z metód pre porovnanie diverzity spoločenstiev je metóda Rényiho radov diverzity

$$H_{\alpha} = \frac{\ln \sum_{i=1}^K p_i^{\alpha}}{1 - \alpha},$$

kde K je počet druhov, $p_i = \frac{n_i}{N}$, pričom n_i je počet organizmov i -teho druhu ($i = 1, 2, \dots, K$).

Hodnota α nadobúda hodnoty od 0 do nekonečna. Spoločnosť s vyššou diverzitou má celú krivku nad krivkou spoločnosti s nižšou diverzitou. Ak sa krivky dvoch spoločností pretnú, tieto spoločnosti sa nedajú z hľadiska diverzity usporiadať.

Druhovú bohatosť K závisí od veľkosti územia A podľa vzťahu, známeho tiež pod názvom Arrheniusov vzťah,

$$K = aA^b.$$

Hodnota parametra a sa mení tiež v závislosti od toho, či sa jedná o ostrov, časť kontinentu, v závislosti od zemepisnej šírky (Lyons a Willig, 2002; Crawley a Haral, 2001). Závislosť

počtu druhov K od plochy územia A , z ktorého bol uskutočnený výber, sa modeluje tiež pomocou regresnej priamky tvaru

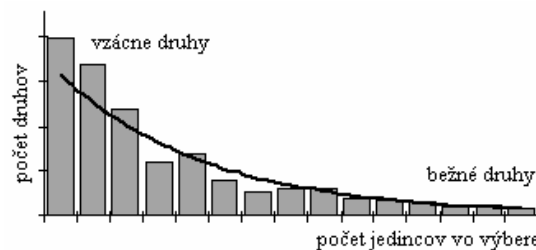
$$\log(K) = a + b \log(A)$$

alebo

$$K = a + b \log(A).$$

Je potrebné vedieť, že indexy determinácie pre pôvodné a transformované údaje sú rôzne. Vyplýva to z odlišných vlastností lineárnych a nelineárnych modelov (Lamoš, Potocký, 1989). Najdôležitejší, zatiaľ však nevyriešený, problém teoretickej ekológie sa týka *existencie pravdepodobnostného rozdelenia, ktoré popisuje relatívne rozdelenie druhov v spoločnosti vo všeobecnosti a všeobecného mechanizmu*, ak existuje, ktorý objaňuje rozdelenie druhov v spoločnosti.

Pri analýze spoločností je obvyklé, že vo výbere sa vyskytujú vzácne a bežné druhy (Obrázok 1). Vo výbere je počet jedincov jednotlivých vzácných druhov nízky. Počet jedincov bežného druhu je v jednotlivých výberoch vysoký.



Zdroj: Vlastné znázornenie

Obrázok 2: Závislosť počtu druhov a počtu jedincov

Označme ax počet druhov v celkovom výbere, ktoré sú reprezentované jediným jedincom, ax^2 počet druhov v celkovom výbere, ktoré sú reprezentované dvoma jedincami, atď. Postupnosť $ax, \frac{ax^2}{2}, \frac{ax^3}{3}, \frac{ax^4}{4}, \dots$ predstavuje geometrický rad, ktorého súčet je $a \ln(1 - x)$ a predstavuje celkový počet druhov

$$K = a \ln \left(1 + \frac{N}{a} \right),$$

kde K je počet druhov, N je počet jedincov vo výbere, a je index diverzity. Iteráciami dostaneme riešenie x rovnice

$$\frac{K}{N} = \frac{1-x}{x} (-\ln(1-x)).$$

Odhad indexu diverzity a je

$$\hat{a} = \frac{1-x}{x} N.$$

Rozptyl odhadu indexu diverzity a je

$$\text{var}(\hat{a}) = \frac{\hat{a}}{-\ln(1-x)}.$$

Vhodnosť modelu, t.j. zhodu empirických početností a početností jedincov jednotlivých druhov vypočítaných pomocou popísaného modelu, overíme pomocou štatistického χ^2 -testu dobrej zhody (Anděl, 2002).

Pre odhad počtu druhov K sa používa aj jeho *jackknife odhad* (Shao a Tu, 1995). Jeho výhodou je, že umožňuje tiež konštrukciu intervalov spoľahlivosti. Vychádza sa z výsledkov výberového skúmania v n kvadrantoch. Nech s je počet druhov, a k je počet jedinečných

druhov (t.j. druhov, ktoré sa vyskytli v jedinom kvadrante). *Jackknife odhad* počtu druhov K je

$$\hat{K} = s + \frac{n-1}{n}k.$$

Rozptyl odhadu počtu druhov K je

$$\text{var}(\hat{K}) = \frac{n-1}{n} \left[\sum_{j=1}^s j^2 f_j - \frac{k^2}{n} \right],$$

kde f_j je počet kvadrantov obsahujúcich j jedinečných druhov. Interval spoľahlivosti pre odhad počtu druhov K je

$$\left\langle \hat{K} - t_{1-\alpha/2} \sqrt{\text{var}(\hat{K})}, \hat{K} + t_{1-\alpha/2} \sqrt{\text{var}(\hat{K})} \right\rangle,$$

kde $t_{1-\alpha/2}$ je $1-\alpha/2$ kvantil Studentovho t rozdelenia s $n-1$ stupňami voľnosti.

Ďalší spôsob odhadu počtu druhov K je pomocou *bootstrappingu* (Shao a Tu, 1995). Dewdney (1997) skonštruoval dynamický systém založený na jedincoch, ktorý obsahoval niekoľko tisíc jedincov patriacich medzi niekoľko sto druhov. Jedince systému hynú, resp. sa rozmnožujú. Systém produkuje rozdelenie druhov, ktorý je štatisticky nerozlíšiteľný od rozdelení, ktoré biológovia získajú systematickým výberom (Dewdney, 2000). K tomuto záveru autor dospel porovnaním výsledkov simulácie s rozdeleniami 125 výberov, ktoré uskutočnili biológovia.

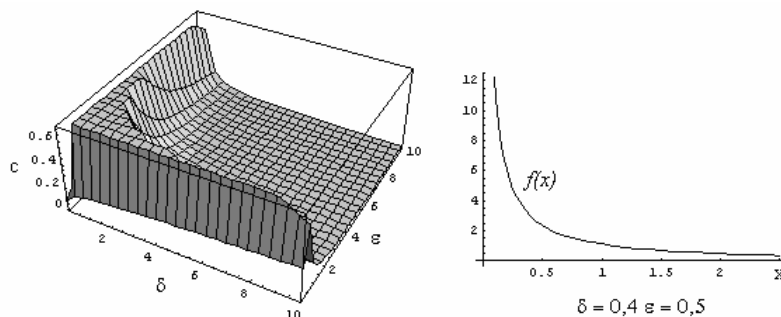
Závislosť početností a druhov v ekvilibriu dynamického systému má J rozdelenie. Funkcia hustoty J rozdelenia je

$$f(x) = c \frac{1 - \delta x'}{x'}, \quad \text{pre } 0 \leq x \leq \Delta',$$

$$0, \text{ inde.}$$

kde $x' = x + \varepsilon$, $\delta' = 1 / \Delta'$, $\Delta' = \Delta + \varepsilon$, c je funkcia závislá od ε a δ

$$c(\varepsilon, \delta) = \frac{1}{-\ln \varepsilon \delta - 1 + \varepsilon \delta}.$$



Zdroj: Vlastné výpočty

Obrázok 3: Funkcia c a funkcia hustoty J rozdelenia pre $\delta = 0,4$ a $\varepsilon = 0,5$

3. Záver

V príspevku sú porovnávané indexy miery biodiverzity pomocou vybraných štatistických metód: Spermanovho korelačného koeficienta, pričom sa určila aj citlivosť jednotlivých indexov na výskyt zriedkavých, dominantných a priemerne sa vyskytujúcich druhov. Sú tu prezentované metódy pre porovnanie diverzity spoločenstiev (metóda *Rényiho radov diverzity*) a skúmaná existencia pravdepodobnostného rozdelenia, ktoré popisuje relatívne rozdelenie druhov v spoločnosti vo všeobecnosti a všeobecného mechanizmu, ak existuje, ktorý objaňuje rozdelenie druhov v spoločnosti. Vhodnosť modelu, t.j. zhodu

empirických početností a početností jedincov jednotlivých druhov vypočítaných pomocou popísaného modelu je overená pomocou štatistického χ^2 -testu dobrej zhody. Ďalší spôsob odhadu počtu druhov K je pomocou *jackknife odhadu* a pomocou *bootstrappingu*

4. Literatúra

- [1] ANDĚL, J. 2003. Statistické metody. MATFYZPRESS, Praha, 299 s. ISBN 80-86732-08-8
- [2] BEISEL, J.N. – USSEGLIO/POLATERA, P.- BACHMANN, V. – MORETEAU, J.C.: A comparative Analysis of Evenness Index Sensitivity. Internat. Rev.Hydrobiol. 88, 2003, 1, s. 3-15
- [3] BRINDZA, J. 2001. Ochrana genofondu rastlín. Nitra: SPU, ISBN 80-7137-974-3
- [4] BROWNE, J. B.–PECK, S. B. 1996. The long-horned beetles of south Florida (Cerambycidae: Coleoptera): biogeography and relationships with the Bahama Islands and Cuba. Canadian Journal of Zoology 74:2154-2169.
- [5] BULLA, L.1994. An index of evenness and its associated diversity measure. Oikos,167-171.
- [6] CRAWLEY, M. J.- HARRAL, J. E. 2001. Scale Dependence in Plant Biodiversity. Science 2 February 2001:Vol. 291. no. 5505, pp. 864 – 868, ISSN 0036-8075
- [7] DEWDNEY, A. K. 2000. A dynamical model of communities and a new species-abundance distribution. The Biol. Bull. 198(1): 152-163.
- [8] DI CASTRI F. - YOUNÈS T. (Eds.). 1996. Biodiversity, Science and Development. Towards a new partnership. CAB International, Wallingford, UK, 646 p. Publisher: CAB International ISSN: 08519897
- [9] KRISHNAMURTHY, K.V.2003. Text book of biodiversity. Publied by Science Publishers, NH, USA, 242 p. ISBN 1-57808-325-7
- [10] LAMOŠ, F.- POTOCKÝ, R. 1989. Pravdepodobnosť a matematická štatistika. Bratislava: Alfa, ISBN 80-05-00115-0
- [11] LYONS, S. K. and WILLIG, M. R. 2002. Species richness, latitude, and scale-sensitivity. Ecology 83: 47-58
- [12] MILLER, R.G. 1974. The jackknife - a review. Biometrika 61: 1-15.
- [13] NORTON, B. G. 1994. On what we should save: the role of cultures in determining conservation targets," in (eds. P. Forey et al.) Systematics and conservation evaluation
- [14] NORTON, B. G. 2001. Conservation biology and environmental values: can there be a universal earth ethic?" in (eds. C. Potvin, et al.) *Protecting biological diversity: roles and responsibilities*. Montreal: McGill-Queen's University Press
- [15] SHAO, J. -TU, T. 1995. The Jackknife and Bootstrap, Springer, New York
- [16] SMITH, B. - WILSON, J.B. 1996. A consumer's guide to evenness indices. Oikos 76:70-82.

Adresy autorov:

Ján Brindza, Doc. Ing., CSc.
Katedra genetiky a šľachtenia rastlín, FAPZ SPU
A. Hlinku 2,
949 76 Nitra
Jan.Brindza@uniag.sk
Mária Bauerová, prof. RNDr., CSc.
Katedra botaniky a genetiky FPV UKF
Nábřežie mládeže 91,
949 74 Nitra
mbauerova@ukf.sk

Beáta Stehlíková, prof. RNDr., CSc.
Katedra Matematiky UKF
Tr. A. Hlinku 1
949 01 Nitra1
bstehlikova@ukf.sk
Anna Tirpáková, prof. RNDr., CSc.
Katedra Matematiky UKF
Tr. A. Hlinku 1
949 01 Nitra1
atirpakova@ukf.sk

Prognózovania v športe Forecasting in sport

Jaroslav Broďáni

Abstract: Forecasting (futorológia) as a science dealing with the theory and methodology of forecasting, found justification for the blasts in the sport has 60's. Sporting forecasting aims at predicting performance, individual performance, overall sports development, planning and managing the training process, modeling techniques and selection of talented youth. An important part of knowledge is also developing such views on doping, the future of the sponsors, mass media and so on.

Key words: sports, forecasting

Kľúčové slová: šport, prognózovanie,

1. Úvod

Prognostika (futorológia) ako vedný odbor zaoberajúci sa teóriou a metodikou prognózovania, našla opodstatnenie v oblastiach športu už v 60-tych rokoch. Športové modelovanie a prognózovanie sa zameriava na predpovedanie výkonov, individuálnej výkonnosti, celkového vývoja športu, plánovania a riadenia tréningového procesu, modelovania techniky a výberu talentovanej mládeže. Dôležitou časťou je taktiež poznanie vývoja názorov napr. na doping, budúcnosť sponzorov, masových médií a pod.

2. Prognózovanie v športe

Šport je fascinujúcim fenoménom našej doby. Rôznym formám športu sa venujú milióny ľudí. Športové súťaže sa konajú pred tisíckami prítomných divákov či za asistencie televíznych kamier a nechávajú masy jasat' alebo nariekať nad krásou športových stretnutí. Šport je vítaným predmetom diskusie divákov, čitateľov novín, je témou záujmu veľkého percenta populácie. Samozrejme aktívny šport je vyhradený určitej vekovej kategórii, diváci či fandenie alebo aspoň občasný záujem o šport v najrôznejších podobách zahŕňa množstvo väčší počet ľudí najrôznejšieho veku (Pupiš, 2008; Tilinger, 2004).

Vplyv spoločenských, sociálnych, ekonomických, politických, právnych, vedecko-technických, atď. zmien ovplyvňuje rozvoj jednotlivých športových odvetví. Ich rozvoj sa dá predvídať vo vzájomnom kontexte. Známa je štúdia Martina et al. (1997) z dlhodobého pôsobenia vývojových tendencií systému riadenia športového tréningu a súťaženia v štvorročných olympijských cykloch. Diegel (2000) uvažuje vo svojej práci o nadchádzajúcom vývoji svetového športu zo sociálneho, ekonomického a výkonnostného hľadiska a zmieňuje sa o možnom negatívnom vývoji. Politický rozmer versus Olympijské hry prezentuje Senn (1999), atď.

Oblasť vrcholového športu a štátnej reprezentácie kladie na prognostiku mimoriadne a pritom značne špecifické požiadavky. Vo svetovej literatúre možno nájsť detailne prepracované analýzy minulého vývoja, prognostické modely ďalšieho rozvoja výkonnosti, ale aj celého systému vrcholového športu. Deje sa tak na úrovni limitnej ľudskej výkonnosti, svetových a olympijských rekordov, ale aj na úrovni širšej výkonnostnej špičky. Prognostika tohto typu umožňuje objektívne analyzovať trendy svetovej výkonnosti a predikovať aj vplyvy zásahov ostatných spoločenských javov do tejto oblasti (Moravec, 2007; Rozim, 2009; Tilinger, 2004; Měkota - Cuberek, 2007; a iní).

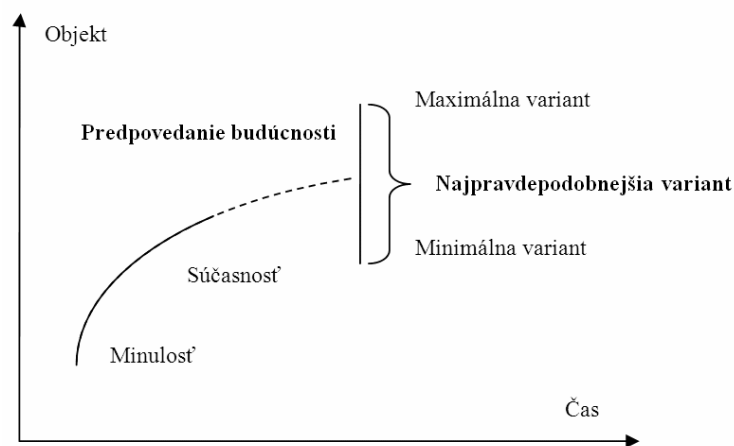
S prognózovaním sa stretávame už pri **výbere talentovanej mládeže** na šport a pri **intraindividuálnom** predpovedaní a hodnotení výkonnostného rastu pretekárov,

intenzifikácie tréningového zaťaženia v individuálnych športoch ako napr. v atletike, plávaní, behu na lyžiach, zjazdovom lyžovaní, vzpieraní, športovej gymnastike, kanoistike, veslovaní a pod.

Aby bolo možné kvalitným spôsobom **riadiť zložitý proces** športového tréningu je dôležité získať potrebné informácie o budúcnosti, na akej úrovni budú musieť budúci olympionici viesť boj o olympijské medaily za štyri alebo osem rokov.

3. Prognózy športových výkonov a výkonnosti

Prognózy delíme na operatívne (1 - 2 mesiace), krátkodobé (2 - 12 mesiacov), strednodobé (1 - 4 roky) a dlhodobé (4 - 8 rokov). Snahou je určenie najoptimálnejšieho parametra prognózy, pričom o spoľahlivosti a pravdivosti prognózy môžeme hovoriť až po realizácii a verifikácii skúmaného javu v budúcnosti. Variabilita rozptylu prognóz je pritom stanovená určením maximálneho a minimálneho tempa rastu sledovaného parametra v čase (obr. 1). Samotná športová prognóza plní poznávaciu a riadiacu funkciu, pričom popisuje rôzne varianty, pravdepodobnosti realizácie budúcej udalosti.



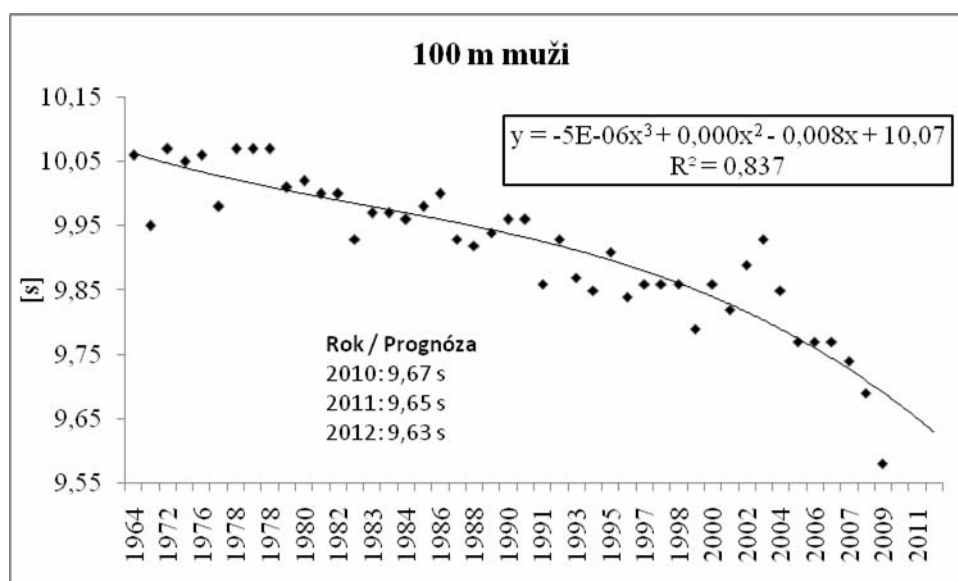
Obrázok 1 Prognostická schéma (Havlíček, 1996)

2.1 Prognózovanie v športovej výkonnosti

Základom prognózy pre predpovedanie športovej výkonnosti sa stali **expertízne metódy** (metóda brainstormingu, Delfská metóda), **metódy extrapolácie** (párová a mnohonásobná korelačná analýza, analýza časových radov), **metódy analógie** a **modelovania**. Do analýz vstupujú údajové databázy z olympijských hier, majstrovstiev sveta, najlepších sezónnych svetových výkonov až do súčasnosti. Sledované sú výkony u víťazov, výkony potrebné na kvalifikáciu, resp. výkony potrebné na účasť vo finále.

Paralelne so získanými výkonmi, sú k analýzam a prognózam využívané dostupné informácie z odbornej literatúry, praxe, štruktúre športového výkonu, kvalitatívne a kvantitatívne poznatky o športovom tréningu, informácie o vývoji materiálov, pravidiel atď. Dôležitým a rozhodujúcim zdrojom informácií sa stávajú názory odborníkov z príslušnej športovej špecializácie (tréneri, vedeckí pracovníci, športovci).

Súčasťou vyhodnotenia teoretických kriviek rastu výkonnosti bývajú údaje o veľkosti štandardnej chyby odhadu a koeficienty korelácie. Determinanty korelačných hodnôt bývajú do určitej miery vzaté ako spoľahlivosť prognózy. Výsledky sú podrobené dôkladnej analýze a poskytnuté ako súčasť dotazníka už spomenutým odborníkom. Prezentované sú v grafickej a číselnej forme (graf 1, tab. 1). Najpoužívanejšie sú čiarové grafy s preloženou trendovou čiarou a vhodne sú doplnené tabuľkou obsahujúcou prognózované a svetové výkony, výkony v posledných rokoch, regresné rovnice, spoľahlivosť prognózy.



Graf 1: Vývoj a prognóza najlepších svetových výkonov v behu na 100 m mužov (Brod'áni, 2010)

Po samotnej realizácii výkonov prichádza na rad verifikácia spoľahlivosti prognóz. Realizácia spočíva v porovnávaní skutočnej hodnoty a percentuálnej odchýlky od stanovenej prognózy, pričom prognóza predstavuje 100 %. Pri úplnej zhode prognózy so skutočným výkonom považujeme diferenciu, ktorá nepresiahla 1 % za veľmi presnú prognózu. Rozdiel do 2 % považujeme z praktického hľadiska za presnú predpoveď. Skutočné hodnoty, ktoré sa nelíšia od prognóz viac ako o 3 % hodnotíme ako „blízke očakávania“. Väčší rozdiel hodnotíme ako nezhodu.

Tabuľka 1: Prognózy vybraných najlepších svetových atletických výkonov u mužov do roku 2012 (Brod'áni, 2010)

Disciplína	Prognózované roky			Regresné parametre		Výkon v roku		
	2010	2011	2012	Rovnica	R ²	2008	2009	WR
100 m [s]	9,67	9,65	9,63	$y = -5E-06x^3 + 0,000x^2 - 0,008x + 10,07$	83,7	9,69	9,58	9,58
200 m [s]	19,37	19,30	19,22	$y = -5E-05x^3 + 0,002x^2 - 0,030x + 20,05$	39,80	19,30	19,19	19,19
400 m [s]	43,65	43,63	43,60	$y = -0,022x + 44,35$	22,10	43,75	44,06	43,18
110 m prekážok [s]	12,88	12,88	12,87	$y = -0,01x + 13,21$	50,10	12,87	13,04	12,87
400 m prekážok [s]	47,34	47,33	47,33	$y = -0,15\ln(x) + 47,89$	13,00	47,25	47,91	46,78
4x100 m [s]	37,47	37,44	37,41	$y = -0,000x^2 - 0,005x + 37,88$	18,5	37,10	37,31	37,10

WR - svetový rekord

Pri aplikovaní extrapolácie sa zvyčajne používa postup so 4 úrovňami. 1. **Určenie parametrov trendu** vychádzajúc z cieľov a zamerania prognózy. 2. **Pred výberom údajov charakterizujúcich minulé vývoj** je nutné poznať trend vývoja. 3. **Pri voľbe dĺžky extapolovaného obdobia** prihliadame na to či majú údaje dlhodobý charakter alebo údaje z minulosti neovplyvňujú vývoj budúcnosti. V prvom prípade je možné analyzovať minulosť čo najväčší časový úsek spätne a v druhom vychádzame z krátkodobých bezporuchových období. Druhý spôsob je pre športovú prognostiku vhodnejší. 4. **Voľba najvhodnejšieho tvaru krivky** je založená na „skusmom“ grafickom preložení, znižovaní rozptylu hodnôt metódou najmenších štvorcov a pomocou programov výber matematickej funkcie, ktorá by čo najlepšie vyrovnávala súbor sledovaných hodnôt.

Základnou prvkom extrapolácie v športe je analýza súvislosti medzi skúmaným objektom a faktorom času. Dynamický rad býva aproximovaný matematickým vzťahom (viď vzorec), kde „y“ je prognózovaná hodnota, „t“ je čas (roky) počas prognózovaného obdobia a $t_0, t_1, t_2 \dots t_n$ sú koeficienty rovnice:

$$y = a_0 + a_1t_1 + a_2t_2 + a_3t_3 \dots a_nt_n$$

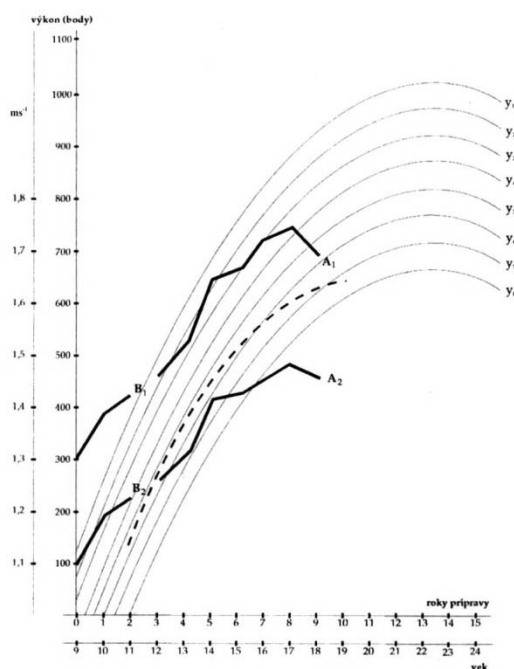
Na výpočet extrapolácie sú v praxi najčastejšie aplikované funkcie: lineárna ($y = a_0 + a_1t$), parabolická ($y = a_0 + a_1t_1 + a_2t_2$), kubicko-parabolická ($y = a_0 + a_1t_1 + a_2t_2 + a_3t_3$), stupňovitá funkcia ($y = a^{tb}$), exponenciálna funkcia ($y = a \cdot e^{bt}$), modifikovaný exponent ($y = K - a \cdot e^{bt}$), kombinovaná exponenciálno-stupňovitá funkcia ($y = e^{at^b}$), logistická krivka ($y = \frac{K}{1 + b \cdot e^{-at}}$); kvadraticko-logistická funkcia ($y = \frac{K}{1 + b \cdot e^{-at}}$) a funkcia Gompertzova ($y = K a^{bt}$). V prípade, že dochádza k interakcii viac než dvoch parametrov, využívajú sa metódy mnohonásobnej korelácie (Brodáň, 2009).

Najpoužívanejšie a najvhodnejšie trendy vývoja, ktoré športový výskum aplikuje pri extrapolácii sú lineárna krivka, polynom max 3 stupňa, exponenciálna a logaritmická krivka.

3.2. Prognóza individuálneho rastu športovej výkonnosti

Pravidelná analýza vývoja výkonnosti a objavovanie všeobecne platných zákonitostí, má dôležitý význam pre priebežné posudzovanie perspektívnosti jednotlivých športovcov vo všetkých jeho etapách športovej prípravy. Prognózovanie sa zameriava na predpoveď výkonu bezprostredne pred štartom, počas jednotlivých období športovej prípravy (v týždennom mikrocykle, v mesačných mezocykloch, v ročnom makrocykle) a predpoveď výkonu, ktorú by mohol dosiahnuť o niekoľko rokov (napr. v olympijskom mikrocykle).

Pri intraindividuálnej predikcii je veľmi dôležité vedieť, v akom štádiu rozvoja sa daný jedinec nachádza a podstatné je oddeliť nadanie a talent od tréningovej práce. V týchto prípadoch sa využívajú nomogramy (obr. 2).



Obrázok 2 Kvadratická (A_1 a A_2) a exponenciálna (B_1 a B_2) extrapolácia 200 m znak - muži (Turek, 1996; Turek - Ružbarský, 2001)

Po porovnaní priebehu dynamiky nomogramu so sklonom skutočnej dynamiky konkrétneho jedinca môžeme po dvoch až 4 rokoch zaradiť športovca do príslušných výkonnostných pásiem a pomerne dobre stanoviť roky tréningu. V prípade, že dynamika výkonnosti športovca dlhodobo zotrúva v rovnakom pásme, je potrebné zmeniť systém doterajšej prípravy. Správne postavená individuálna prognóza je nevyhnutná pre určenie potencionálnych možností u pretekárov. Pre úspešný výkonnostný rast je nevyhnutné poznať zákony dynamiky výkonnosti v závislosti od veku a faktory ovplyvňujúce športový výkon (kondičné, antropometrické, technické, taktické, psychologické, genetika ... atď.).

Skúmaná závislosť športového výkonu a času je predpokladom intraindividuálneho prognózovania. Objektívnym vyjadrením rastu športovej výkonnosti sú v prvom rade úspechy športovca, ktorých dynamiku je možné chápať ako funkciu času (vzorec 1). V takýchto prípadoch funkciu, popisujúcu štatistickú tendenciu dynamiky športovej výkonnosti vyjadrujeme kvadratickou rovnicou. Pre popis funkcie je najvhodnejší polynóm 2. stupňa (vzorec 2), kde „y“ je bodová hodnota výkonu, „a₀ a₁ a₂“ sú konštanty danej športovej disciplíny a „x“ je čas.

$$[1] \quad y = f(x)$$

$$[2] \quad y = a_0 + a_1x + a_2x^2$$

Odpoveď pri výbere tvaru priebehu výkonnostnej krivky nachádzame v komplexnom prístupe (dostatočné množstvo informácií o faktoroch ovplyvňujúcich športový výkon) a v analytickom prístupe (hypotéza vychádza z predpokladu, že športovec dosahuje stabilizované najlepšie svetové výkony počas celého roka).

Nasledovnú aplikáciu krokov intraindividuálnej extrapolácie je možné charakterizovať:

1. **Získavame množstvo údajov** o priebehu výkonnosti v danej disciplíne.
2. **Výpočet teoretických trendov výkonnosti** vychádza z aktuálnych výkonov športovca s uvádzaním koeficientov rovnice koeficientov korelácie, determinácie medzi skutočným a teoretickým priebehom.
3. **Vecnou analýzou dynamiky výkonnosti** vykonáme rozbor priebehov výkonnosti, teoretických a skutočných vzťahov a rozdelenie športovcov do skupín.
3. Posledným krokom je **výpočet dynamiky výkonnosti**, ktorého výsledkom je priemerný trend, resp. trendy doplnené o súbežné krivky vytvárajúc nomogram dynamiky výkonnosti s pásmami nízkeho a vysokého tempa výkonnosti.

Uplatnenie nomogramov je možné pri analyzovaní dynamiky výkonnosti konkrétneho športovca ktorý ukončil športovú činnosť alebo je na začiatku svojej športovej kariéry.

V **prvom prípade** je možné zistiť s využitím matematickej metódy odchýlok najmenších štvorcov a s maximálnou presnosťou koeficienty „a₀ a₁ a₂“. Koeficienty sú stanovené tak, aby sa teoretická krivka priebehu športovej výkonnosti priblížila k empirickému a aby bola miera korelácie čo najvyššia. Po dosadení známych hodnôt x₁; y₁; y_{max}; x_{max} do rovnice [2] získame po zjednodušení rovnice [3, 4, 5] pre koeficienty „a₀ a₁ a₂“ (Svoboda, 1984).

$$[3] \quad a_0 = y_1 - a_1 \frac{2x_{\max} - 1}{2x_{\max}}$$

$$[4] \quad a_1 = \frac{2x_{\max}(x_{\max} - y_1)}{(x_{\max} - 1)^2}$$

$$[5] \quad a_2 = \frac{y_1}{2x_{\max}}$$

V **druhom prípade**, keď športovec začína zo športovou kariérou, nepoznáme celý priebeh krivky. Výpočet má menšiu praktickú presnosť. Vyžívame pritom, Na základe najlepších výkonov športovca vypočítame príslušné koeficienty „a₀ a₁ a₂“ a získame priemerný rast výkonnosti.

4. Záver

Proces prognózovania bol vždy procesom od známej minulosti k neznámej budúcnosti. Na základe poznania minulosti i poznatkov minulého a aktuálneho vedeckého predvídania znižujeme mieru neistoty budúceho vývoja. Hodnota kritických prognóz prispieva k riadeniu a hľadaniu optimálnych riešení zameraných k prestavbe ľudského snaženia. Je neoddeliteľnou súčasťou ekonomizácie činnosti v širokej oblasti športu a vo všetkých jeho zložkách.

Športová prognostika a samotný šport v úzkej participácii so všetkými systémami spoločenského života napomáha k zvyšovaniu sociálnej mobility Slovenskej republiky v celosvetovom meradle.

5. Literatúra

- [1] BROŽÁNI, J. 2009. Využitie mnohonásobnej korelačnej a regresnej analýzy časových radov vo vrcholovom športe. In: Forum Statisticum Slovakum. – ISSN 1336-7420. Roč. 5., č. 1. (2009), s. 6-9.
- [2] BROŽÁNI, J. 2010. Prognóza najlepších svetových výkonov bv mužských atletických disciplínach do roku 2012. In Exercitatio Corpolis Motus Salus : Slovak journal of sports science. ISSN. 1337-7310. Roč. 2. Č. 1 (2010).
- [3] DIEGL, H. 2000. Perspectives of sports at the begining of new century. New Studia in Athletics, 2000, MON, Vol. 15, no1, p.7-15.
- [4] MARTIN, D. et al. 1997. Entwicklungstendenzen der Trainings und Wettkampfsysteme im Spitzensport mit Folgerungen für den Olympiacyklus 1996 bis 2000. Leistungssport, 27, 1997, 4.1, s.25-31.
- [5] MĚKOTA, K. - CUBEREK, R. 2007. Prognózování sportovní výkonnosti. In Pohybové dovednosti – činnosti – výkony. Olomouc: UP, 2007. S. 135-142.
- [6] MORAVEC, R. 1993. Predpoklady športovej úspešnosti v atletických skokoch. Bratislava : VSTVŠ, 1993.
- [7] MORAVEC, R. 2007. Prognózovanie a modelovanie v športe. In Teória a didaktika výkonnostného a vrcholového športu. Bratislava: ICM AGENCY, s. 216-222.
- [8] PLATONOV, V. N. et al. 1987. Prognozovanie i modelovanie v športe. In: Teoria sporta. Kijev: Višča škola, 1987, p. 350 – 371.
- [9] PUPIŠ, M. 2008. Vývoj atletickej chôdze na 20 km na vrcholných podujatiach v rokoch 1964-2008. In ATLETIKA 2008. Nitra : UKF, s. 114.
- [10] ROZIM, R. 2009. Retrospektívna analýza a vývojová predikcia výkonov v skoku do výšky mužov na Slovensku. In Atletika 2009. Banská Bystrica : Dali BB. s.97-105.
- [11] SENN, A. E. 1999. Power, Politics and the Olympic Games. Champaign: Humman Kinetics, 1999.
- [12] SVOBODA, V. 1984. Statistics. Swimming World. 25, 1984, p. 4.
- [13] TILINGER, P. 2004. Prognózování vývoje výkonnosti ve sportu. Praha : Univerzita Karlova v Praze, Karolinum, 2004,
- [14] TUREK, M. 1996. Prognózovanie v športe. Prešov: PF UPJŠ Košice, 1996, 107s.
- [15] TUREK, M. - RUŽBARSKÝ, P. 2001. Športová prognóza v praxi. Prešov: PF PU Prešov, SVSTVŠ, 2001, 83s.

Adresa autora:

Jaroslav Brožáni, doc. PaedDr. PhD.
Katedra telesnej výchovy a športu PF UKF Nitra
Tr. A. Hlinku 1
949 74 Nitra
jbrodani@ukf.sk

Vývoj nominálnych a reálnych úrokových sadzieb v SR¹ Nominal and real interest rates development in Slovak Republic

Ľudmila Fabová

Abstract: Commercial banks intermediate money transfers from creditors to debtors. The price for lent money is a interest, affected by interest rate, mainly. Significant factor, influencing interest rates, is inflation. Investors are interested not only in nominal interest rates, that are collected from or have to be paid to commercial banks but also in real interest rates determined by inflation in an economy. This article looks into the development in nominal and real interest rates of short-term deposits and loans of households in Slovak Republic during the period between 2000 through 2009.

Key words: interest, interest rate, nominal interest rate, inflation, real interest rate.

Kľúčové slová: úrok, úroková sadzba, nominálna úroková sadzba, inflácia, reálna úroková sadzba.

1. Úvod

Úlohou komerčných bánk v ekonomike je sprostredkovať pohyb peňazí, resp. kapitálu od veriteľov k dlžníkom, teda od ekonomických subjektov, ktoré majú dočasný prebytok peňazí k ekonomickým subjektom, ktoré majú dočasný nedostatok peňazí. Medzi najväčších veriteľov bánk patria domácnosti, ktoré v bankách ukladajú svoje úspory vo forme vkladov. K najväčším dlžníkom komerčných bánk zase patria firmy, ktorým banky poskytujú úvery na zabezpečenie ich činnosti, investovanie a zavádzanie inovácií, ale aj domácnosti, ktoré pomocou úverov financujú svoje potreby.

2. Nominálne a reálne úrokové sadzby

Cenou za poskytnutie kapitálu je úrok, ktorý závisí najmä od úrokovej sadzby, ktorá patrí medzi najdôležitejšie nástroje riadenia komerčných bánk. Aktívne aj pasívne operácie komerčných bánk totiž súvisia s prijímaním, resp. vyplácaním úrokov, ktorých výšku ovplyvňujú viaceré činitele. Jedným z nich je situácia na peňažnom trhu, kde banky získavajú potrebné zdroje a umiestňujú dočasne voľné finančné prostriedky. Pri prebytku peňazí v ekonomike majú úrokové sadzby na peňažnom trhu tendenciu klesať a naopak, pri nedostatku peňazí majú tendenciu stúpať. Situáciu na peňažnom trhu okrem toho ovplyvňujú aj opatrenia menovej politiky centrálnej banky, ktorá zmenou úrokových sadzieb ovplyvňuje ponuku peňazí v ekonomike.

Oficiálne úrokové sadzby centrálnej banky (do roku 2009 NBS, od roku 2009 ECB) predstavujú pre obchodné banky hraničné náklady pre získavanie zdrojov a okrem toho ich môžu používať ako referenčné sadzby pre určovanie klientskych úrokových sadzieb, preto bezprostredne ovplyvňujú ich úrokovú politiku.

Významným činiteľom, ovplyvňujúcim úrokové sadzby je inflácia. Investorov totiž nezaujímajú iba nominálne úrokové sadzby, ktoré od nich požadujú, resp. im vyplácajú komerčné banky, ale aj reálne úrokové sadzby, ovplyvnené infláciou v danej ekonomike. Na dobre fungujúcom finančnom trhu by úrokové sadzby mali odrážať aj predpokladanú mieru

¹ Príspevok bol spracovaný v rámci riešenia úlohy VEGA č. 1/0536/10 "Inovácie ako strategický základ zvyšovania konkurenčnej schopnosti SR."

inflácie a obchodné banky by ju mali zohľadňovať vo svojej úrokovej politike. Nízke, alebo dokonca záporné reálne úrokové sadzby zvýhodňujú dlžníkov, preto je v takejto situácii zvýšený záujem o úvery, ale poškodzujú veriteľov, preto klesá záujem o sporenie. A naopak, vysoké reálne úrokové sadzby priťahujú vkladateľov, ktorí zvyšujú mieru úspor, ale sú nevýhodné pre reálne investície, pretože zdražujú úvery.

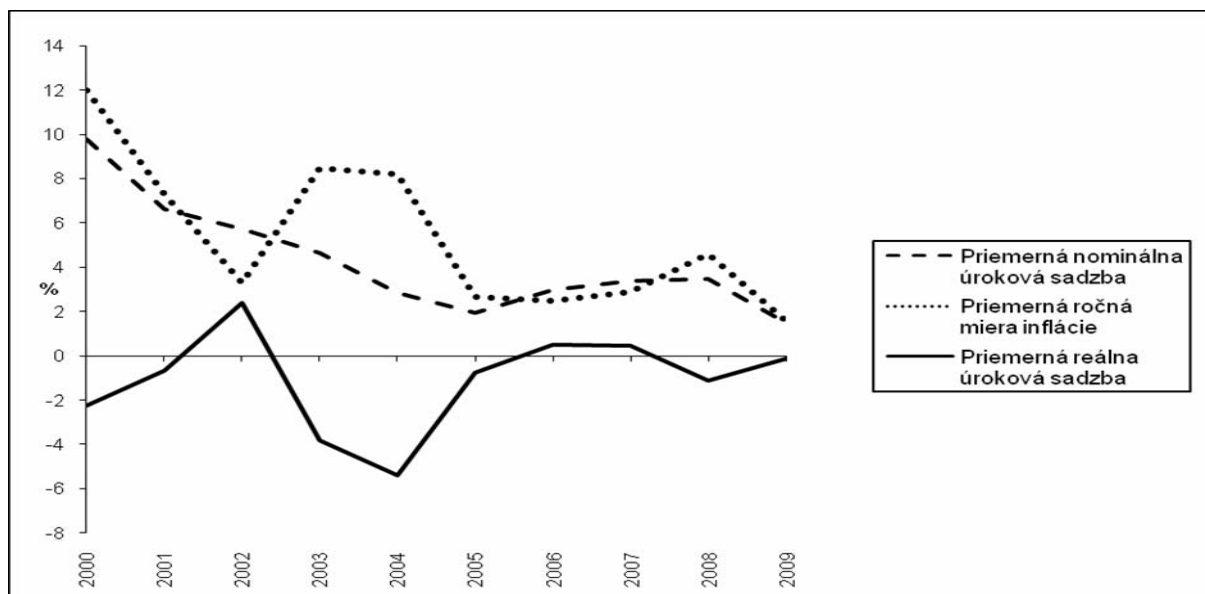
Vývoj situácie na slovenskom finančnom trhu v sektore domácností v rokoch 2000 – 2009 dokumentujú údaje v nasledujúcich tabuľkách a z nich odvodené grafy.

Tabuľka 1: Priemerné úrokové sadzby nových vkladov so splatnosťou do 1 roka v %

Ukazovateľ	2000	2001	2002	2003	2004	2005	2006	2007	2008	2009
Priem. nomin. úroková sadzba	9,76	6,62	5,70	4,65	2,82	1,94	2,99	3,37	3,47	1,49
Priem. ročná miera inflácie	12,00	7,30	3,30	8,50	8,23	2,70	2,50	2,90	4,60	1,60
Priem. reálna úroková sadzba	-2,24	-0,68	2,40	-3,85	-5,41	-0,76	0,49	0,47	-1,13	-0,11

Zdroj: NBS [1], ŠÚ SR [2] a vlastné výpočty

Vývoj priemerných nominálnych a reálnych úrokových sadzieb z nových krátkodobých vkladov domácností so splatnosťou do 1 roka v období rokov 2000 – 2009 zobrazuje tabuľka 1 a graf 1. Z ich analýzy vyplýva, že nominálne aj reálne úrokové sadzby z ročných vkladov, ktoré boli až do roku 2000 relatívne vysoké, začali postupne klesať. Vysoké úrokové sadzby boli pre vkladateľov zaujímavé a spôsobovali rast objemu vkladov v tomto období. Rast vkladov sa však v roku 2001 zastavil, čo bolo spôsobené rapidným poklesom nominálnych úrokových sadzieb (v porovnaní s rokom 1999 takmer na polovicu), pričom reálne úrokové sadzby vkladov vďaka rastúcej inflácii už od roku 2000 dosahovali dokonca záporné hodnoty.



Graf 1: Vývoj nominálnych a reálnych úrokových sadzieb krátkodobých vkladov

Najnižšiu úroveň dosiahli reálne úrokové sadzby v roku 2004, keď ich priemerná hodnota bola -5,41 %. Príčinou bola relatívne vysoká miera inflácie a pokles nominálnych úrokových sadzieb. Majitelia vkladov v bankách teda reálne strácali, čo sa prejavilo v opätovnom poklese vkladov domácností, ktoré začali hľadať iné možnosti zhodnotenia svojich úspor. V roku 2005 sa vďaka poklesu inflácie reálne úrokové sadzby zvýšili a v rokoch 2006 a 2007 aj zásluhou rastu nominálnych sadzieb opäť dosiahli plusové hodnoty. V posledných dvoch rokoch sledovaného obdobia však opäť dosiahli záporné hodnoty. Vkladateľom sa teda opäť znižuje reálna hodnota ich úspor.

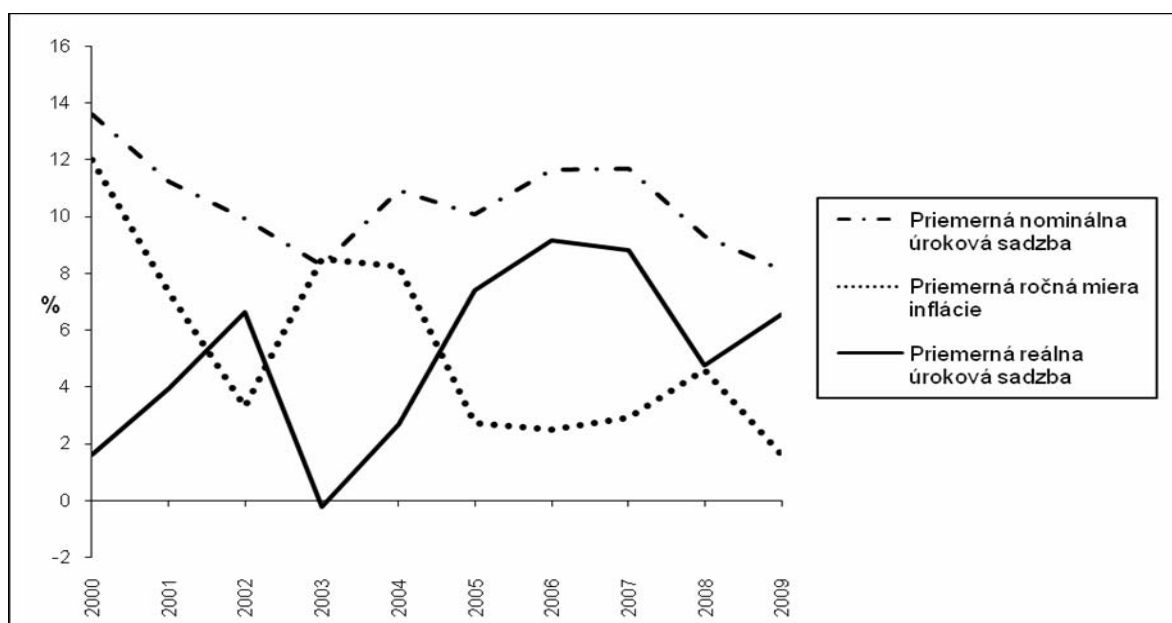
V tabuľke 2 a na grafe 2 vidno vývoj priemerných nominálnych a reálnych úrokových sadzieb z novo poskytnutých krátkodobých úverov, t. j. úverov so splatnosťou do jedného roka. Na základe ich analýzy možno konštatovať, že situácia na trhu úverov bola v období rokov 2000 – 2009 pre reálne investovanie relatívne priaznivá, pretože nominálne aj reálne úrokové sadzby z úverov v porovnaní s predchádzajúcim obdobím (1995 – 1999) vďaka rastúcej inflácii výrazne poklesli a v roku 2003 bola priemerná reálna úroková sadzba dokonca záporná.

Tabuľka 2: Priemerné úrokové sadzby nových úverov so splatnosťou do 1 roka v %

Ukazovateľ	2000	2001	2002	2003	2004	2005	2006	2007	2008	2009
Priem. nomin. úroková sadzba	13,61	11,24	9,93	8,29	10,93	10,08	11,65	11,70	9,34	8,14
Priem. ročná miera inflácie	12,00	7,30	3,30	8,50	8,23	2,70	2,50	2,90	4,60	1,60
Priem. reálna úroková sadzba	1,61	3,94	6,63	-0,21	2,70	7,38	9,15	8,80	4,74	6,54

Zdroj: NBS [1], ŠÚ SR [2] a vlastné výpočty

Pokles úrokových sadzieb úverov bol spôsobený okrem iného aj prípravou Slovenska na vstup do Európskej únie, v rámci ktorej NBS postupne znižovala svoju základnú úrokovú sadzbu a prispôbovala ju nižšej úrovni v EÚ. To sa samozrejme premietlo aj do znižovania klientskych úrokových sadzieb v komerčných bankách.



Graf 2: Vývoj nominálnych a reálnych úrokových sadzieb krátkodobých úverov

Tým sa vytvorili priaznivé podmienky pre úverovanie reálnych investícií, čo sa odrazilo aj v raste objemu poskytovaných úverov v Slovenskej republike. Od roku 2004 síce nominálne aj reálne úrokové sadzby krátkodobých úverov začali opäť mierne rásť, ale posledné dva roky sledovaného obdobia nominálne sadzby klesali a reálne mierne stúpili, čo bolo spôsobené nízkou mierou inflácie. Vďaka priaznivým úverovým podmienkam bola od roku 2006 na Slovensku zaznamenaná v sektore domácností zvýšená úverová aktivita, napriek tomu zadlženosť obyvateľstva v Slovenskej republike patrí medzi najnižšie v Európskej únii.

3. Záver

Záverom musíme skonštatovať, že situácia na Slovensku v sledovanom období nebola priaznivo naklonená domácnostiam, ktoré mali záujem sporiť prostredníctvom vkladov

v komerčných bankách. Nízke nominálne a dokonca záporné reálne úrokové sadzby vkladov nielen že im nepriniesli žiadne výnosy, ale dokonca reálne prišli o časť svojich úspor. Priaznivejšia situácia bola v úverovaní domácností, čo sa prejavilo aj v raste objemu poskytnutých úverov. Komerčné banky na Slovensku teda aj naďalej zarábajú viac na svojich veriteľoch, od ktorých získavajú podhodnotené zdroje kapitálu, ako na dlžníkoch, ktorým poskytujú úvery za relatívne nízku cenu. Situácia sa nezmenila ani po vstupe Slovenska do Európskej únie, resp. eurozóny, ako sa očakávalo, pretože slovenské banky zatiaľ nie sú vystavené väčšej konkurencii európskych bánk a preto nie sú nútené prehodnocovať svoju úrokovú politiku.

4. Literatúra

[1] <http://www.nbs.sk>

[2] <http://portal.statistics.sk>

Adresa autora:

Ľudmila Fabová, Ing., PhD.

Ústav manažmentu STU

Vazovova 5

812 43, Bratislava

ludmila.fabova@stuba.sk

Aktivita versus pasivita vybraných centrálních bank **The activity vs. passivity of selected central banks**

Radka Hájková

Abstract: This paper discusses the relationship between monetary policy implementation by central banks and the real values of macroeconomic indicators (especially inflation and economic growth). The aim of this paper is also to clarify the characteristics, their activity or passivity. Active central banks are identified as central banks, which are often changing character of monetary policy, passive conversaly. Through regression analysis, coefficient of variation and changes in the frequency calculations, this contribution seeks to determinate which of selected central banks are more active or passive, and how this relates to the real macroeconomic indicators.

Key words: Macroeconomics and Monetary economics, Taylor rule

Kľúčové slová: Makroekonómie a monetárni ekonómie, Taylorovo pravidlo

1. Úvod

Jednou z hlavních funkcí centrálních bank v tržní ekonomice je implementace monetární politiky. Prostřednictvím svých nástrojů působí na nabídku a poptávku po penězích, které v konečném důsledku ovlivní vývoj ekonomických proměnných, primárně vývoj cenové hladiny a nepřímo další ekonomické ukazatele, např. ekonomický růst, zaměstnanost a další. V současné době je tento proces implementace monetární politiky založen především na tom, že centrální banky regulují úrokové sazby. Monetární politiku lze tedy označit jako významný pilíř hospodářské politiky každého státu, jelikož díky monetární politice a jejímu udržování cenové stability mohou být efektivní nástroje fiskální politiky státu, které působí na reálné ukazatele v ekonomice. Mnoho autorů (ECB, 2006) se zabývá interakcí monetární a fiskální politiky a využívá při tom hodnocení aktivní a pasivní centrální banky, proto je tento příspěvek zaměřen na charakteristiku aktivní a pasivní centrální banky, a to, jak tento fakt souvisí s reálným makroekonomickými proměnnými. Makroekonomickými proměnnými, které jsou zde sledovány, jsou inflace a hospodářský růst, které jsou analyzovány na základě Taylorova pravidla. Spojitost charakteristiky aktivní a pasivní centrální banky využili například ve svých pracích Carvalho a Moura (Carvalho, Moura, 2008).

Cílem tohoto příspěvku je určit vztah mezi charakterem monetární politiky, zde označované jako aktivita či pasivita centrálních bank, na makroekonomické ukazatele.

2. Metodika

V analýze je sledováno využívání monetárních nástrojů centrálních bank a vztah mezi těmito nástroji a vývojem reálných hospodářských ukazatelů, jako jsou např. inflace a HDP. Mezi analyzovanými centrálními bankami je ČNB, ECB a FED.

Aby bylo možné rozhodnout, která z vybraných centrálních bank je spíše aktivnější a která spíše pasivnější a tyto výsledky pak porovnat s hlavními makroekonomickými ukazateli, je součástí práce empirická část. Tato empirická část zahrnuje výpočty variačních koeficientů pro měnový agregát a úrokové sazby, na základě kterých bude možné učinit jeden z dílčích závěrů, která centrální banka je aktivnější či pasivnější. Variační koeficient je počítán dle vzorce:

$$v_x = \frac{S_x}{\bar{x}}, \text{ pro } \bar{x} \neq 0. \quad (1)$$

Další součástí analýzy je vícerozměrná regresní analýza, která by měla odhalit závislost mezi proměnnými na základě Taylorova pravidla. Taylorovo pravidlo je použito ve tvaru:

$$i = i^* + a(\pi - \pi^*) + b(y - y^*), \quad (2)$$

kde i je výsledná požadovaná úroková sazba, i^* rovnovážná úroková sazba (používají se dlouhodobé vládní cenné papíry), π míra inflace, π^* inflační cíl, y růst HDP a y^* potenciální růst. Část rovnice $(y - y^*)$ tedy vyjadřuje mezeru růstu. Parametry a , b vyjadřují váhy inflace a mezery růstu, které specifikují charakter centrální banky. Pokud $b=0$, lze hovořit o centrální bance silně konzervativní, která sleduje pouze inflační cíl, nikoliv hospodářský růst a naopak.

Na základě regresní analýzy pomocí odhadu parametrů bude možné učinit další závěr ohledně aktivity či pasivity centrálních bank. Poslední součástí této empirické analýzy je výpočet četností, který by měl také přispět dílčím závěrem tím, že určí, která ze zvolených centrálních bank je aktivnější a které jsou spíše pasivnější.

3. Empirická část

První provedenou analýzou empirické části je analýza měnových agregátů. Za monetaristického předpokladu, že změna peněžní zásoby v dlouhém období přispěje ke zvýšení cenové hladiny, můžeme z analýzy měnových agregátů, též formulovat závěr o tom, která ze zvolených centrálních bank je aktivnější či pasivnější. Nástroje centrálních bank neovlivňují množství peněz přímo, protože pokud uvažujeme s endogenitou peněz je množství peněz ovlivněno multiplikačním efektem peněz ale i přesto nástroje centrálních bank nepřímo množství peněz ovlivňují. Základem této analýzy je výpočet variačního koeficientu pro měnové agregáty. Hodnoty časových řad pro výpočet variačního koeficientu jsou hodnoty měnového agregátu M3 pro období 1999q1-2009q4, vzhledem k nedostupnosti dat pro měnový agregát M3 v USA (zveřejněné hodnoty jsou pouze v časovém období do února 2006) jsou pro USA použity hodnoty měnového agregátu M2.

Tabulka 1: Variační koeficient pro M3 (M2 pro USA)

v %	ČNB	ECB	FED
variační koeficient	0,24	0,25	0,19

Zdroj: Vlastní výpočet

Z tabulky 1 vyplývá, že nejvíce mění měnový agregát ECB, nejméně FED. Z tohoto můžeme tedy formulovat závěr, že ECB je z hlediska dodávání likvidity nejaktivnější centrální bankou, již méně aktivní je ČNB a nejméně aktivní tedy spíše pasivní centrální bankou je FED. Výpočet variačního koeficientu může být ale ovlivněn tím, že v případě FEDu jsou použity hodnoty pro měnový agregát M2. Výpočty může také ovlivnit kratší časová řada v případě ČR, kde jsou použity z důvodu omezené dostupnosti těchto dat hodnoty z časového období 2001q1-2009q4.

Další provedenou analýzou je analýza krátkodobých úrokových sazeb na mezibankovním trhu. Analýza tohoto operačního cíle centrálních bank nám poskytne objektivnější závěr o změnách charakteru monetární politiky vybraných centrálních bank, jelikož tento operační cíl mají centrální banky pod přímou kontrolou. Na základě kolísání krátkodobých úrokových sazeb na mezibankovním trhu pak tedy bude možné učinit závěr, která centrální banka více využívá operačního cíle, tedy je aktivnější a které jej využívá méně a je tedy pasivnější. Aby bylo možné učinit závěr o kolísání úrokových sazeb na mezibankovním trhu je jako v předchozí podkapitole proveden výpočet variačního koeficientu. Pro výpočet variačního koeficientu jsou použity hodnoty 3-měsíční úrokové

sazby za období 1999q1-2009q4 z důvodu, že pouze tato úroková sazba je dostupná pro ČR, Eurozónu i USA. Výpočty variačního koeficientu jsou uvedeny v tabulce 2.

Tabulka 2: Variační koeficient pro 3-měsíční úrokové sazby

v %	ČNB	ECB	FED
variační koeficient	0,45	0,95	0,56

Zdroj: Vlastní výpočet

Z tabulky lze vidět, že nejvíce kolísají úrokové sazby na mezibankovním trhu v Eurozóně, dále pak v USA a nejméně kolísají v ČR. Lze tedy říci, že ECB mění charakter monetární politiky nejčastěji, tedy je nejaktivnější centrální bankou. Méně aktivní centrální bankou je FED, který mění charakter monetární politiky méně často v porovnání s ECB a je proto méně aktivní centrální banka. ČNB je podle této analýzy pasivní centrální bankou.

Aktivita a pasivita centrálních bank byla zatím sledována především z pohledu změn charakteru monetární politiky a následným ovlivňováním operačních a zprostředkujících cílů centrálních bank, které následně ovlivní konečné cíle. ČNB a ECB mají jasně stanovený primární konečný cíl ve formě cenové stability. FED se zaměřuje na plnění několika cílů, přičemž jedním z nich je také cenová stabilita. Pro potřeby regresní analýzy budeme vycházet z Taylorova pravidla tedy z toho, že všechny tři vybrané centrální banky sledují dva konečné cíle a to cenovou stabilitu a hospodářský růst, kterého chtějí dosáhnout prostřednictvím stabilní cenové hladiny. Cílem regresní analýzy tedy bude identifikovat vztah mezi charakterem monetární politiky, které zde bude představovat krátkodobá úroková sazba na mezibankovním trhu, cenovou stabilitou, měřenou harmonizovaným indexem spotřebitelských cen a hospodářským růstem. Regresní analýza by měla také dopomoci k určení závěru, jestli centrální banky sledují spíše cíl v podobě cenové stability či hospodářského růstu a jestli mezi charakterem monetární politiky a těmito proměnnými existuje vůbec nějaká významná závislost. Regresní analýza by nám také mimo jiné měla potvrdit závěr, že tyto dvě proměnné jsou při tvorbě monetární politiky významné proměnné.

V této vícerozměrné regresní analýze budou využity časové řady meziročních změn krátkodobých úrokových sazeb na mezibankovním trhu v ČR, Eurozóně a USA (IR_t), meziroční změny harmonizovaného indexu cen v ČR, Eurozóně a USA (P_t) a meziroční změny hospodářského růstu pro ČR, Eurozónu a USA (Y_t) v období 1999q1-2009q4. Tato regresní analýza využívá vzhledem k tomu, že pomocí absolutních hodnot by nebylo možné porovnání parametrů, využity relativní hodnoty všech parametrů přepočítané na meziroční změny. V případě Eurozóny autorka vycházela pouze z údajů pro původní státy Eurozóny (tedy EA 12) z důvodu z důvodu zamezení ovlivnění dat vstupem dalších států do Eurozóny. Výsledky této analýzy pro ČR jsou uvedeny v tabulce 3.

Tabulka 3: Regresní analýza pro ČR

	Odhad parametrů	SE	T-statistika	F-test	P-hodnota	DW	n
Model				31,06	0,000	0,338	40
Konst.	-0,434	0,056	-7,77431		0,000		
P_t	0,088	0,015	5,800		0,000		
Y_t	0,043	0,009	4,453		0,001		

Zdroj: Vlastní výpočet

Regresní analýza identifikovala významnou závislost na 1% hladině významnosti mezi inflací, hospodářským růstem a charakterem monetární politiky, který představuje krátkodobá úroková sazba na mezibankovním trhu. Analýza též potvrdila teoretická východiska a tedy, že ECB se zaměřuje při implementaci své monetární politiky spíše na cenovou hladinu. Hodnoty Durbin-Watsonova testu jsou velice nízké. To může být způsobeno tím, že model obsahuje jen proměnné dle Taylorova pravidla a mohl by obsahovat i jiné proměnné. Cílem článku však není modelování reakční funkce centrální banky, ale pouze porovnání stability a proto model nebyl doplněn o další proměnné. Důvodem ale také může být fakt, že nízké hodnoty Durbin-Watsonova testu vedou spíše k autokorelaci reziduí. Důvodem autokorelace může být to, že ČNB pracuje při tvorbě monetární politiky s predikovanými hodnotami, které jsou zde nahrazeny reálnými hodnotami. Vzhledem k nedostupnosti predikovaných hodnot za celé období 1999q1-2009q4 byly použity reálné hodnoty těchto proměnných. Výsledky vícerozměrné regresní analýzy pro Eurozónu jsou uvedeny v tabulce 4.

Tabulka 4: Regresní analýza pro Eurozónu

	Odhad parametru	SE	T-statistika	F-test	P-hodnota	DW	n
Model				54,5	0,000	0,336	40
Konst.	-0,364	0,099	-3,680		0,000		
P_{t+1}	0,087	0,045	1,922		0,063		
Y_t	0,144	0,015	9,322		0,001		

Zdroj: Vlastní výpočet

Z této regresní analýzy vyplývá také významná závislost mezi charakterem monetární politiky, inflací a hospodářským růstem ovšem v tomto případě na rozdíl od ČR pouze na 10% hladině významnosti. V případě Eurozóny uvažujeme v modelu dopředné zpoždění hodnot inflace (P_{t+1}). Zajímavým závěrem této regresní analýzy je, že na rozdíl od ČNB se ECB zaměřuje více na hospodářský růst než na cenovou hladinu, i přesto, že konečný primární cíl je definován pouze v podobě cenové stability. Hodnoty Durbin-Watsonova testu jsou i zde nízké ze stejného důvodu jako v případě ČR, tedy použitím reálných hodnot proměnných nebo nezařazením dalších vysvětlujících proměnných. Regresní analýza pro USA je uvedena v tabulce 5.

Tabulka 5: Regresní analýza pro USA

	Odhad parametru	SE	T-statistika	F-test	P-hodnota	DW	n
Model				12,93	0,001	0,260	40
Konst.	-0,544	0,141	-3,871		0,004		
P_{t+1}	0,098	0,047	2,094		0,043		
Y_t	0,156	0,042	3,714		0,007		

Zdroj: Vlastní výpočet

Z tabulky 5 vyplývá, že i v tomto případě existuje významná závislost mezi krátkodobou úrokovou sazbou na mezibankovním trhu, inflací a hospodářským růstem na 5% hladině významnosti. I v tomto případě uvažujeme s dopředným zpožděním hodnot inflace (P_{t+1}). Z regresní analýzy vyplývá, že se FED také obdobně jako ECB zaměřuje spíše na hospodářský růst. Z odhadu parametru vyplývá, že FED je oproti ČNB i ECB aktivnější.

Hodnoty Durbin-Watsonova testu jsou ve srovnání s předchozími testy pro ČNB i ECB ještě nižší, což odpovídá tomu, že FED má definováno více konečných cílů. S největší pravděpodobností budou takto nízké hodnoty způsobeny absencí další vysvětlujících proměnných v modelu. Toto však může být způsobeno jako i u ČR a Eurozóny tím, že v modelu nejsou využity predikované hodnoty inflace i hospodářského růstu.

Poslední součástí empirické analýzy je určení četnosti změn v úrokových sazbách centrálních bank. V této analýze jsou použity hodnoty oficiálních úrokových sazeb z úvěrů za časové období 1999m1-2009m2. Výpočet četností změn je uveden v tabulce 6.

Tabulka 6: Četnost změn oficiálních úrokových sazeb centrálních bank

	ČNB	ECB	FED
Četnost změn	33	30	37

Zdroj: Vlastní výpočet

Z tabulky je patrné, že úrokové sazby nejvíce mění FED a to i přesto, že časová řada použitá pro výpočet četnosti je v případě FEDu pouze 1999m1-2007m9. Méně již mění své úrokové sazby ČNB a nejméně pak ECB. Z těchto výpočtů vyplývá, že nejaktivnější centrální bankou je FED, méně aktivní je pak ČNB a za pasivní centrální banku můžeme označit ECB.

Aby bylo možné formulovat závěr ze všech provedených analýz o tom, která centrální banka je aktivnější či pasivnější a jak tento fakt, ovlivňuje reálné makroekonomické ukazatele, jsou výsledky všech provedených analýz shrnuty v následující tabulce, která je doplněna i o informace o HDP a inflaci:

Tabulka 7: Souhrn výpočtů

Ukazatel	ČNB	ECB	FED
Var.koeficient pro M3	0,24	0,25	0,19
Var.koeficient pro ur.sazby	0,45	0,95	0,56
Cetnost zmen	33	30	37
Regresni analyza (zaměřena spíše na)	cenová stabilita	hospodářský růst	hospodářský růst
Aktivní/Pasivní	spíše pasivnější	méně aktivní	nejvíce aktivní
Průměrný ek.růst	3,19	1,49	2,16
Rozptyl ek.růstu	9,22	4,52	4,72
Průměrná hodnota inflace	2,52	2,02	2,43
Rozptyl inflace	3,72	0,65	3,12

Zdroj: Vlastní výpočet

4. Závěr

Cílem tohoto příspěvku bylo určit vztah mezi charakterem monetární politiky, zde označované jako aktivita či pasivita centrálních bank, na již uvedené makroekonomické ukazatele. Ze závěrečného souhrnu vyplývá, že nejaktivnější centrální bankou je ECB, dle výpočtů variačních koeficientů pro M3 i pro úrokové sazby. Ovšem závěr z výpočtů pro úrokové sazby je z hlediska aktivity či pasivity přesnější, protože v úrokových sazbách se odráží veškerá aktivita centrálních bank. Z regresní analýzy vyplývá, že nejaktivnější centrální bankou je FED. Z regresní analýzy je též možno formulovat závěr, že obě dvě tyto centrální banky (ECB i FED) sledují spíše cíl v podobě hospodářského růstu. Výpočty variačního koeficientu pro M3 mohou být ale ovlivněny faktem, že pro USA byl z důvodu nedostupnosti dat použit měnový agregát M2, proto nemusí být tyto výpočty zcela přesné. Zkresleny mohou

být i výpočty regresní analýzy, protože neberou v úvahu predikované hodnoty inflace a hospodářského růstu, ale pouze reálné hodnoty. Nejrelevantnější výsledek tedy podává výpočet variačního koeficientu pro krátkodobé úrokové sazby na mezibankovním trhu, který určil jako nejaktivnější centrální banku ECB, dále pak méně aktivní ČNB a nejméně aktivní, tedy spíše pasivní FED. ECB tedy jako nejaktivnější centrální banka dosahuje nejlepších výsledků v podobě cenové stability. Co se týče hospodářského růstu největšího průměrného hospodářského růstu dosahuje ČR, která je všemi analýzami označena za méně aktivní, ovšem rozptyl hospodářského růstu je v porovnání s Eurozónou a USA největší, což nenaznačuje pozitivní vývoj. Lepších výsledků dosahuje FED, který jako nejméně centrální banka dosahuje druhého nejvyššího průměrného hospodářského růstu, a i výpočet rozptylu naznačuje stabilní pozitivní vývoj tohoto ukazatele.

Závěrem lze tedy na základě všech provedených analýz konstatovat, že aktivní centrální banky dosahují lepších výsledků v porovnání cenové stability. Méně aktivní centrální banky dosahují lepších hodnot ve srovnání hospodářského růstu.

Předkládaný příspěvek vznikl za podpory interního grantového projektu 71/2010 "Stabilizační funkce monetární politiky v souvislostech hospodářského cyklu České republiky".

5. Literatura

- [1]CARVALHO, A., MOURA, M. 2008. What Can Taylor Rules Say About Monetary Policy in Latin America? [online]. [citováno 2. března 2010]. Dostupné z: <http://bugarin.ibmecsp.edu.br/ISMP2008/Papers/P4_Carvalho_Moura.pdf>
- [2]ČESKÁ NÁRODNÍ BANKA. ARAD systém časových řad [online]. [citováno 30.dubna 2010]. Dostupné z:<http://www.cnb.cz/cnb/STAT.ARADY_PKG.STROM_DRILL?p_strid=AAAD&p_lang=CS>
- [3]EUROPEAN CENTRAL BANK. Monetary aggregates [online]. [citováno 30.dubna 2010]. Dostupné z:<https://stats.ecb.europa.eu/stats/download/bsi_ma_historical_sa_dp/bsi_ma_historical_sa_dp/bsi_hist_sa_u2_2.pdf>
- [4]EUROSTAT. Statistics [online]. [citováno 22.dubna 2010]. Dostupné z: <<http://epp.eurostat.ec.europa.eu/portal/page/portal/statistics/themes>>
- [5]FEDERAL RESERVE SYSTEM. Money stock measures [online]. [citováno 4.května 2010]. Dostupné z: <<http://www.federalreserve.gov/releases/h6/hist/h6hist1.txt>>
- [6]EUROPEAN CENTRAL BANK. Monetary aggregates [online]. [citováno 30.dubna 2010]. Dostupné z:<https://stats.ecb.europa.eu/stats/download/bsi_ma_historical_sa_dp/bsi_ma_historical_sa_dp/bsi_hist_sa_u2_2.pdf>
- [7]ECB. Monetary and fiscal interactions in a new Keynesian model with capital accumulation and non-Ricardian consumers [online]. [citováno 2. března 2010]. Dostupné z: <http://www.ecb.int/pub/pdf/scpwps/ecbwp649.pdf>

Adresa autora:

Radka Hájková, Bc.

Ústav financí

PEF, MENDELU v Brně

Zemědělská 1

613 00 Brno

pomenka@mendelu.cz

Stupeň dôležitosti zdrojov informácií pri inovačných aktivitách Degree of Importance Sources of Information for Innovation Activities

Jozef Chajdiak

Abstract:

Článok obsahuje opis postupu určenia stupňa dôležitosti zdrojov informácií pri inovačných aktivitách.

Paper includes description technique specification degrees of importance source of information for innovation activities.

Key words: Štatistické zisťovanie o inováciách za rok 2006 Inov 1-99, zdroje informácií, podsystém Kontingenčná tabuľka, priemer, usporiadanie

Kľúčové slová: The Community Innovation Survey 2006 (CIS 2006), Information source, subsystem PivotTable, average, sorting

1. Úvod

Výkaz Inov 1-99 Štatistické zisťovanie o inováciách za rok 2006 obsahuje aj Modul 652 Zdroje informácií a spolupráca pri inovačných aktivitách. Anonymizovaná verzia výkazu za rok 2006 obsahuje údaje za 2678 vykazujúcich jednotiek.

V Module 652 je sformulovaná otázka a pokyn:

„Aký dôležitý bol z hľadiska inovačných aktivít Vášho podniku počas 2004-2006 každý z týchto zdrojov informácií? Uveďte, prosím, zdroje, ktoré poskytli informácie pre nové inovačné projekty alebo prispeli k dokončeniu existujúcich inovačných projektov.

Označte „**nepoužitý**“, ak ste z daného zdroja nezískali žiadne informácie.“

Vo výkaze sa stupeň dôležitosti uvádza ako „**vysoký**“ (1), „**stredný**“ (2), „**nízky**“ (3) a „**nepoužitý**“ (4). Anonymizovaný súbor z Eurostatu používa stupnicu z hodnotami „**vysoký**“ (3), „**stredný**“ (2), „**nízky**“ (1) a „**nepoužitý**“ (0) pri inovačne zahrnutých firmách a chýbajúcu hodnotu (blank) pri inovačne nezahrnutých firmách.

Zdroje informácií sa členia na:

Vnútročné zdroje

1. V rámci Vášho podniku alebo skupiny podnikov

Trhové zdroje

2. Dodávatelia zariadení, materiálov, komponentov alebo softvéru

3. Klienti alebo zákazníci

4. Konkurenti alebo iné podniky vo Vašom sektore

5. Konzultanti, komerčné laboratória alebo súkromné VV inštitúcie

Inštitucionálne zdroje

6. Univerzity alebo iné inštitúcie vyššieho vzdelania

7. Vládne alebo verejné výskumné inštitúcie

Iné zdroje

8. Konferencie, veľtrhy, výstavy

9. Vedecké časopisy a obchodné/technické publikácie

10. Odborné a odvetvové združenia

V ďalšom texte sú jednotlivé zdroje kódované ako Z1, Z2, ..., Z10.

Úlohou je jednotlivé zdroje informácií štatisticky analyzovať a v druhej úlohe usporiadať podľa dôležitosti.

2. Štatistická analýza súboru údajov o zdrojoch informácií

V rámci štatistickej analýzy sme súbor údajov vytriedili podľa vykázaných hodnôt určenia stupňa dôležitosti pre jednotlivé zdroje informácií. Spravodajské jednotky možno rozčleniť do dvoch skupín. Prvú väčšiu skupinu predstavujú spravodajské jednotky, ktoré nemali inovačné aktivity (v čiastkových frekvenčných tabuľkách Tabuľky 1 sú ich počty uvedené v riadku „blank“). Spravodajské jednotky, ktoré mali inovačnú aktivitu, stupeň dôležitosti príslušného zdroja informácií v anonymizovanej verzii hodnotili na stupnici:

- 3 - vysoký,
- 2 - stredný,
- 1 - nízky,
- 0 - nepoužitý.

Zistili sme absolútne triedne početnosti výskytu spravodajských jednotiek pri jednotlivých stupňoch dôležitosti, vypočítali percentuálny relatívne početnosti výskytu (percentuálny podiel zo všetkých spravodajských jednotiek) a ako dopĺňujúcu charakteristiku sme pre jednotlivé triedy špecifikované stupňom dôležitosti sme zistili absolútny tržieb za rok 2006 a vypočítali ich percentuálny podiel na celkovom objeme tržieb. K vlastnému triedeniu a výpočtom sa použil podsystém PivotTable (Kontingenčná tabuľka) systému Excel 2007.

Vytriedené a vypočítané údaje pre jednotlivé, vo výkaze Inov 1-99 špecifikované, zdroje informácií sú uvedené v Tabuľke 1.

3. Usporiadanie zdrojov informácií podľa dôležitosti

K usporiadaniu jednotlivých zdrojov informácií podľa stupňa dôležitosti sme použili priemerné hodnoty stupňa dôležitosti SD. Vyšli sme z predpokladu, že poradovú stupnicu vysoký, stredný, nízky a nepoužitý (v kódovanej verzii 3, 2, 1 a 0), môžeme transformovať na numerickú stupnicu s diskretnými hodnotami 3, 2, 1 a 0 vyjadrujúcimi mieru dôležitosti príslušného zdroja informácií.

Keďže súbor obsahuje aj chýbajúce hodnoty (blank), k celkovému agregovanému hodnoteniu stupňa dôležitosti sme použili verziu výpočtu priemernej hodnoty stupňa dôležitosti SD1. V danej verzii SD1 rozsah súboru (počet spravodajských jednotiek) – menovateľ priemeru – sa rovná počtu jednotiek, ktoré vykázali konkrétnu hodnotu stupňa dôležitosti (súčet počtov jednotiek vykazujúcich 0, 1, 2 alebo 3 resp. celkový počet jednotiek znížený o počet neodpovedajúcich (blank)). Vzorec pre výpočet SD1 je nasledujúci:

$$\overline{SD1} = \frac{0 * n_0 + 1 * n_1 + 2 * n_2 + 3 * n_3}{n_0 + n_1 + n_2 + n_3} = \frac{n_1 + 2 * n_2 + 3 * n_3}{n_0 + n_1 + n_2 + n_3} = \frac{\sum_{i=0}^3 x_i n_i}{\sum_{i=0}^3 n_i}$$

Kde n_0 , n_1 , n_2 a n_3 sú absolútne počty vykazujúcich jednotiek hodnoty príslušného stupňa dôležitosti.

Vypočítané hodnoty usporiadané v klesajúcom poriadku sú uvedené v Tabuľke 2.

Tabuľka 1 Rozdelenie stupňa dôležitosti jednotlivých zdrojov informácií

Z1 SENTG	1. V rámci Vášho podniku alebo skupiny podnikov			
	Počet firiem	Podiel na firmách	Tržby 2006	Podiel na tržbách
0	70	2,61%	661 016 301	1,08%
1	69	2,58%	796 873 119	1,30%
2	268	10,01%	4 459 332 484	7,27%
3	404	5,09%	34 230 753 125	55,79%
(blank)	1867	69,72%	21 207 909 479	34,57%
Spolu	2678	100,00%	61 355 884 508	100,00%
Z2 SSUP	2. Dodávateľia zariadení, materiálov, komponentov alebo softvéru			
	Počet firiem	Podiel na firmách	Tržby 2006	Podiel na tržbách
0	96	3,58%	3 049 628 021	4,97%
1	127	4,74%	6 039 174 428	9,84%
2	387	14,45%	14 111 453 120	23,00%
3	201	7,51%	16 947 719 460	27,62%
(blank)	1867	69,72%	21 207 909 479	34,57%
Spolu	2678	100,00%	61 355 884 508	100,00%
Z3 SCLI	3. Klienti alebo zákazníci			
	Počet firiem	Podiel na firmách	Tržby 2006	Podiel na tržbách
0	142	5,30%	2 668 331 809	4,35%
1	133	4,97%	9 493 039 231	15,47%
2	292	10,90%	6 374 709 962	10,39%
3	244	9,11%	21 611 894 027	35,22%
(blank)	1867	69,72%	21 207 909 479	34,57%
Spolu	2678	100,00%	61 355 884 508	100,00%
Z4 SCOM	4. Konkurenti alebo iné podniky vo Vašom sektore			
	Počet firiem	Podiel na firmách	Tržby 2006	Podiel na tržbách
0	212	7,92%	5 328 470 233	8,68%
1	211	7,88%	14 556 888 552	23,73%
2	278	10,38%	14 928 363 708	24,33%
3	110	4,11%	5 334 252 536	8,69%
(blank)	1867	69,72%	21 207 909 479	34,57%
Spolu	2678	100,00%	61 355 884 508	100,00%
Z5 SINS	5. Konzultanti, komerčné laboratória alebo súkromné VV inštitúcie			
	Počet firiem	Podiel na firmách	Tržby 2006	Podiel na tržbách
0	427	15,94%	9 260 687 094	15,09%
1	190	7,09%	6 837 049 673	11,14%
2	146	5,45%	17 680 974 909	28,82%
3	48	1,79%	6 369 263 353	10,38%
(blank)	1867	69,72%	21 207 909 479	34,57%
Spolu	2678	100,00%	61 355 884 508	100,00%
Z6 SUNI	6. Univerzity alebo iné inštitúcie vyššieho vzdelania			
	Počet firiem	Podiel na firmách	Tržby 2006	Podiel na tržbách
0	535	19,98%	16 627 010 649	27,10%
1	149	5,56%	12 208 759 251	19,90%
2	106	3,96%	7 176 220 869	11,70%
3	21	0,78%	4 135 984 260	6,74%
(blank)	1867	69,72%	21 207 909 479	34,57%
Spolu	2678	100,00%	61 355 884 508	100,00%

Tabuľka 1 Rozdelenie stupňa dôležitosti jednotlivých zdrojov informácií - pokračovanie

Z7 SGMT	7. Vládne alebo verejné výskumné inštitúcie			
	Počet firiem	Podiel na firmách	Tržby 2006	Podiel na tržbách
0	628	23,45%	22 154 124 295	36,11%
1	120	4,48%	11 771 588 184	19,19%
2	50	1,87%	5 402 264 547	8,80%
3	13	0,49%	819 998 003	1,34%
(blank)	1867	69,72%	21 207 909 479	34,57%
Spolu	2678	100,00%	61 355 884 508	100,00%
Z8 SCON	8. Konferencie, veľtrhy, výstavy			
	Počet firiem	Podiel na firmách	Tržby 2006	Podiel na tržbách
0	162	6,05%	5 279 102 946	8,60%
1	226	8,44%	10 305 263 373	16,80%
2	315	11,76%	21 164 126 068	34,49%
3	108	4,03%	3 399 482 642	5,54%
(blank)	1867	69,72%	21 207 909 479	34,57%
Spolu	2678	100,00%	61 355 884 508	100,00%
Z9 SJOU	9. Vedecké časopisy a obchodné/technické publikácie			
	Počet firiem	Podiel na firmách	Tržby 2006	Podiel na tržbách
0	192	7,17%	5 829 705 235	9,50%
1	249	9,30%	11 047 626 037	18,01%
2	302	11,28%	17 054 749 235	27,80%
3	68	2,54%	6 215 894 522	10,13%
(blank)	1867	69,72%	21 207 909 479	34,57%
Spolu	2678	100,00%	61 355 884 508	100,00%
Z10 SPRO	10. Odborné a odvetvové združenia			
	Počet firiem	Podiel na firmách	Tržby 2006	Podiel na tržbách
0	428	15,98%	16 925 969 392	27,59%
1	218	8,14%	17 247 874 763	28,11%
2	148	5,53%	5 514 226 182	8,99%
3	17	0,63%	459 904 692	0,75%
(blank)	1867	69,72%	21 207 909 479	34,57%
Spolu	2678	100,00%	61 355 884 508	100,00%

Prameň: Vlastný výpočet zo súboru anonymizovaných údajov ŠÚ SR (CIS2006 SK)

Tabuľka 2 Usporiadanie zdrojov informácií podľa stupňa dôležitosti

Kód	SD1	Zdroj informácií	Podiel 3 na tržbách
Z1	2,24	1. V rámci Vášho podniku alebo skupiny podnikov	55,79%
Z2	1,85	2. Dodavatelia zariadení, materiálov, komponentov alebo softvéru	27,62%
Z3	1,79	3. Klienti alebo zákazníci	35,22%
Z8	1,45	8. Konferencie, veľtrhy, výstavy	5,54%
Z4	1,35	4. Konkurenti alebo iné podniky vo Vašom sektore	8,69%
Z9	1,30	9. Vedecké časopisy a obchodné/technické publikácie	10,13%
Z5	0,77	5. Konzultanti, komerčné laboratória alebo súkromné VV inštitúcie	10,38%
Z10	0,70	10. Odborné a odvetvové združenia	0,75%
Z6	0,52	6. Univerzity alebo iné inštitúcie vyššieho vzdelania	6,74%
Z7	0,32	7. Vládne alebo verejné výskumné inštitúcie	1,34%

Prameň: Vlastný výpočet z údajov v Tabuľke 1

Usporiadané poradie podľa zdrojov informácií podľa prímeru SD1 je uvedené v Tabuľke 2. Tabuľka 2 v poslednom stĺpci obsahuje aj podiel stupňa dôležitosti 3 „vysoko“ na tržbách spolu. Vlastnú dôležitosť zdrojov informácií budeme hodnotiť podľa charakteristiky SD1. Podľa hodnôt SD1 jednotlivé zdroje informácií môžeme zoskupiť do štyroch skupín. Najvyšší stupeň dôležitosti má zdroj informácií „1. V rámci Vášho podniku alebo skupiny podnikov“. Hodnota prímeru SD1 rovná 2,24 a podiel odpovedí „nepoužitý“ je len 2,6 %. Druhú skupinu trochu nižšieho stupňa dôležitosti (1,85 a 1,79) tvoria zdroje informácií „2. Dodávatelia zariadení, materiálov, komponentov alebo softvéru“ a „3. Klienti alebo zákazníci“. Podiel ich odpovedí „nepoužitý“ je 3,6 % a 5,3%. Tretiu skupinu tvoria zdroje informácií „8. Konferencie, veľtrhy, výstavy“, „4. Konkurenti alebo iné podniky vo Vašom sektore“ a „9. Vedecké časopisy a obchodné/technické publikácie“ (priemerný stupeň dôležitosti SD1 1,45; 1,35 a 1,30 a podiel odpovedí „nepoužitý“ je 6,05%, 7,9% a 7,2%). Najnižší stupeň dôležitosti zdrojov informácií majú „5. Konzultanti, komerčné laboratória alebo súkromné VV inštitúcie“, „10. Odborné a odvetvové združenia“, „6. Univerzity alebo iné inštitúcie vyššieho vzdelania“ a „7. Vládne alebo verejné výskumné inštitúcie“ (SD1: 0,77; 0,70; 0,52 a 0,32; a podiel odpovedí „nepoužitý“ je 15,9%, 16%, 20% a 23,45%).

4. Záver

Pri pohľade na usporiadanú postupnosť zdrojov informácií pri inovačných aktivitách podľa stupňa dôležitosti autor cíti potrebu uviesť výrok jedného z klasikov ľudského myslenia - potreby priemyslu žnú vedecko-technický pokrok viac dopredu, ako desať univerzít.

5. Literatúra

- [1] Výkaz Inov 1-99 Štatistické zisťovanie o inováciách za rok 2006
- [2] The Fourth Community Innovation Survey (CIS IV) The Harmonised Survey Questionnaire
- [3] CHAJDIK, J. 2003. Štatistika jednoducho. Bratislava: Statis 2003. 194 s. ISBN 80-85659-28-X.
- [4] CHAJDIK, J.: 2009. Štatistika v Exceli 2007. Bratislava: Statis 2009. 312 s. ISBN 978-80-85659-49-8.

Adresa autora:

Jozef Chajdiak, Doc., Ing., CSc.
Ústav manažmentu STU
810 00 Bratislava
chajdiak@statis.biz

Vypracované v rámci riešenia úlohy VEGA č.1/0536/10 „Inovácia ako strategický základ zvyšovania konkurenčnej schopnosti SR“.

Spôsobilosť procesu úpravy úžitkovej vody vzhľadom k vybranému faktoru Process capability of supply water modification with respect to chosen factor

Ivan Janiga, Karol Žák, Andrej Valent

Abstract: In the paper is analyzed supply water with respect to the one of quality characteristic which is conductivity. Paper deals with process capability of supply water.

Key words: current conductivity of water, supply water, process capability.

Kľúčové slová: vodivosť vody, úžitková voda, spôsobilosť procesu.

1. Úvod

V priemyselnom závode sa upravuje úžitková voda, pretože surová voda má nevhodný rozsah pH, obsahuje nečistoty a má vysoký obsah mangánu a železa. Úprava vody sa robí alkalickým čírením vody s dávkovaním vápna, čím sa znižuje tvrdosť vody a zároveň sa odstraňuje vysoká prítomnosť železa a mangánu. Táto voda má vhodne upravené pH, výrazne znížený obsah železa a mangánu a je opticky číra.

Vodivosť je jeden zo znakov kvality vody. Meria sa v [$\mu\text{S}/\text{cm}$] pri 25°C. Voda je vodivá a jej nameraná hodnota vodivosti vypovedá o obsahu rozpustných látok vo vode. Čím je nižšia, tým voda obsahuje menej solí. Táto voda sa používa na priemyselné účely, preto je lepšie keď obsahuje čo najmenej rozpustných látok, teda je „čistejšia“. Voda v priemysle je pred samotným technickým použitím často ešte ďalej upravovaná. Napr. v kondenzačnej tepelnej elektrárni sa v parných okruhoch pre turbíny vyžaduje hodnota vodivosti blízka nule.

2. Spôsobilosť procesu vzhľadom na vodivosť – namerané dáta

Spôsobilosť procesu vyžaduje, aby merané hodnoty boli výberom z normálneho rozdelenia. V tabuľke 1 sú uvedené namerané hodnoty vodivosti.

Tabuľka 1: Namerané hodnoty Vodivosti

Vodivosť (25°C, $\mu\text{S}/\text{cm}$)													
1	405	8	374	15	409	22	396	29	372	36	412	43	410
2	397	9	420	16	415	23	200	30	366	37	414	44	412
3	402	10	413	17	387	24	401	31	408	38	364	45	412
4	405	11	419	18	403	25	427	32	405	39	380	46	418
5	402	12	407	19	419	26	345	33	408	40	387	47	422
6	415	13	399	20	420	27	381	34	409	41	407	48	420
7	413	14	424	21	426	28	390	35	403	42	411	49	417

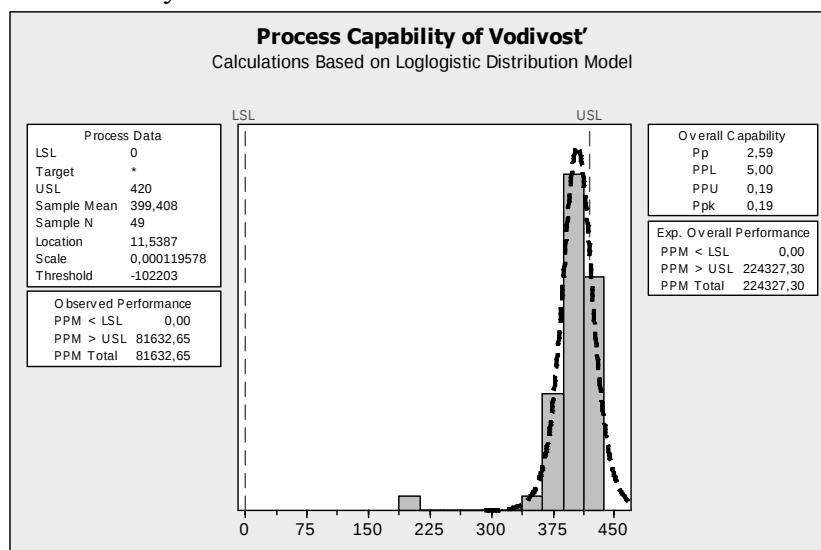
Prvou z analýz je overenie typu rozdelenia, z ktorého pochádzajú namerané hodnoty vodivosti. Použili sme pritom testy zhody pre spojité rozdelenia. Z výsledkov testov sme zistili, že dáta nepochádzajú z normálneho rozdelenia, pretože P-hodnota testu zhody je rovná 0,00872, ako možno čítať v tabuľke 2. V tejto tabuľke sú uvedené ďalšie tri rozdelenia: logistické, loglogistické a loglogistické 3 parametrické, ktoré „vykazujú“ dobrú zhodu z dátami. Z nich sme vybrali loglogistické 3 parametrické rozdelenie, ktoré má najväčšiu P-hodnotu.

Tabuľka 2: P-hodnoty testov zhody

Rozdelenia	Normálne	Logistické	Loglogistické	Loglogistické 3 parametrické
P-hodnoty	0,0087189	0,272661	0,13128	0,278331

Na výpočet spôsobilosti sa používajú tzv. indexy spôsobilosti: C_p , CPU – horný index spôsobilosti, CPL – dolný index spôsobilosti a $C_{pk} = \min(CPU, CPL)$. Na výpočet sú potrebné špecifikačné medze.

Pre znak kvality procesu Vodivosť je horná špecifikačná medza $USL = 420$ a dolná špecifikačná medza $LSL = 0$. Pomocou štatistického softvéru sme vypočítali hodnoty bodových odhadov uvedených indexov.

**Obrázok 1: Spôsobilosť procesu pre Vodivosť**

Na obrázku vidieť, že rozdelenie znaku kvality Vodivosť je posunuté k hornej špecifikačnej medzi, teda spôsobilosť procesu je určená horným indexom spôsobilosti. Hodnoty odhadu indexu spôsobilosti pre loglogistické 3 parametrické rozdelenie meraného znaku kvality Vodivosť sa vypočítajú pomocou vzťahov:

$$C_{PK} = \min(CPU; CPL) = \min\left(\frac{USL - X_{0,5}}{X_{0,99865} - X_{0,5}}; \frac{X_{0,5} - LSL}{X_{0,5} - X_{0,00135}}\right) =$$

$$= \min(0,19; 5,00) = 0,19$$

Hodnota X_p vo výpočte je odhad p-quantilu rozdelenia vypočítaný z nameraných hodnôt vodivosti (tabuľka 1)

Pre porovnanie uvádzame odhady indexov spôsobilosti troch „vyhovujúcich“ rozdelení.

Tabuľka 3: P-hodnoty testov zhody

	C_p	CPL	CPU	C_{pk}
Loglogistické 3 parametrické	2,59	5,00	0,19	0,19
Logistické	2,59	4,99	0,19	0,19
Loglogistické	2,30	4,98	0,15	0,15

Z tabuľky vidieť, že najlepšiu spôsobilosť vykazuje loglogistické 3 parametrické rozdelenie.

Na záver tejto analýzy možno konštatovať, že proces úpravy úžitkovej vody nie je spôsobilý.

3. Spôsobilosť procesu vzhľadom na vodivosť – transformované dáta

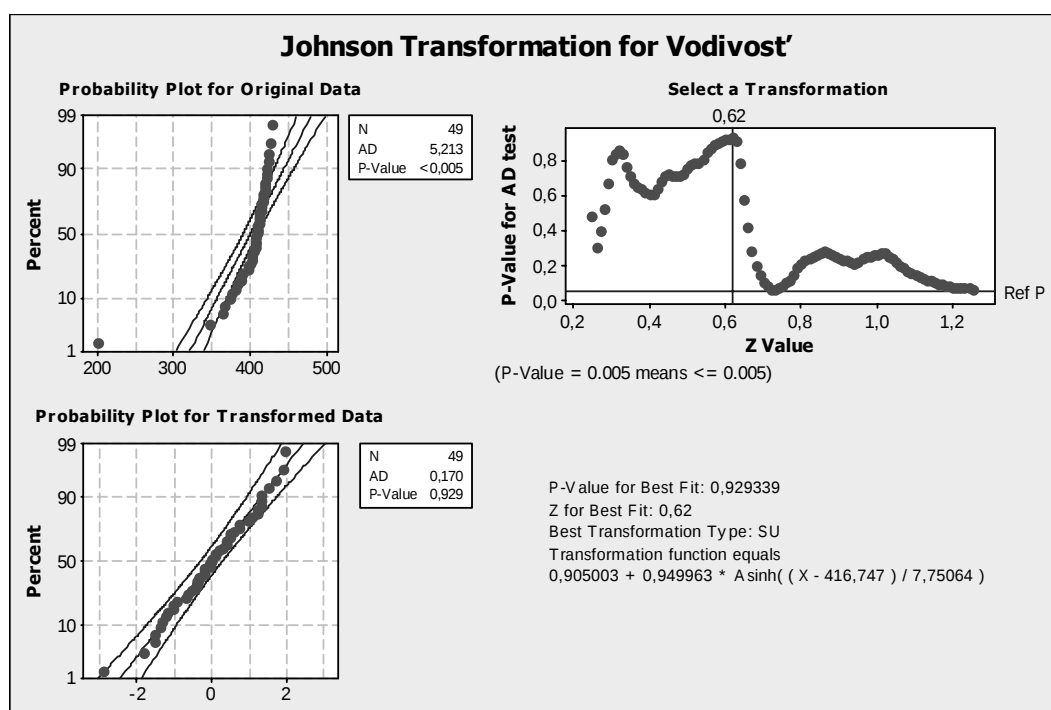
V tejto časti sme vykonali transformáciu nameraných hodnôt vodivosti tak, aby transformované dáta boli výberom z normálneho rozdelenia. Z transformácií, ktoré sme testovali, bola najlepšia Johnsonova transformácia.

Tabuľka 4: Transformované hodnoty Vodivosti

Vodivosť (25°C,μS/cm)													
1	-0,23815	8	-1,38325	15	0,06809	22	-0,72033	29	-1,42602	36	0,35456	43	0,15763
2	-0,67651	9	1,2929	16	0,69271	23	-2,91805	30	-1,544	37	0,57505	44	0,35456
3	-0,4242	10	0,46203	17	-1,04683	24	-0,47976	31	-0,01587	38	-1,58031	45	0,35456
4	-0,23815	11	1,17745	18	-0,36562	25	1,94269	32	-0,23815	39	-1,24225	46	1,05797
5	-0,4242	12	-0,09462	19	1,17745	26	-1,87023	33	-0,01587	40	-1,04683	47	1,50765
6	0,69271	13	-0,58279	20	1,2929	27	-1,21662	34	0,06809	41	-0,09462	48	1,2929
7	0,46203	14	1,69849	21	1,86642	28	-0,94946	35	-0,36562	42	0,25307	49	0,93606

V tabuľke 4 sú transformované dáta pomocou Johnsonovej transformácie. Použitá transformácia sa nachádza na obrázku 2 a má tvar:

$$0,905003 + 0,949963 \times A \sinh \frac{(X - 416,747)}{7,75054}$$



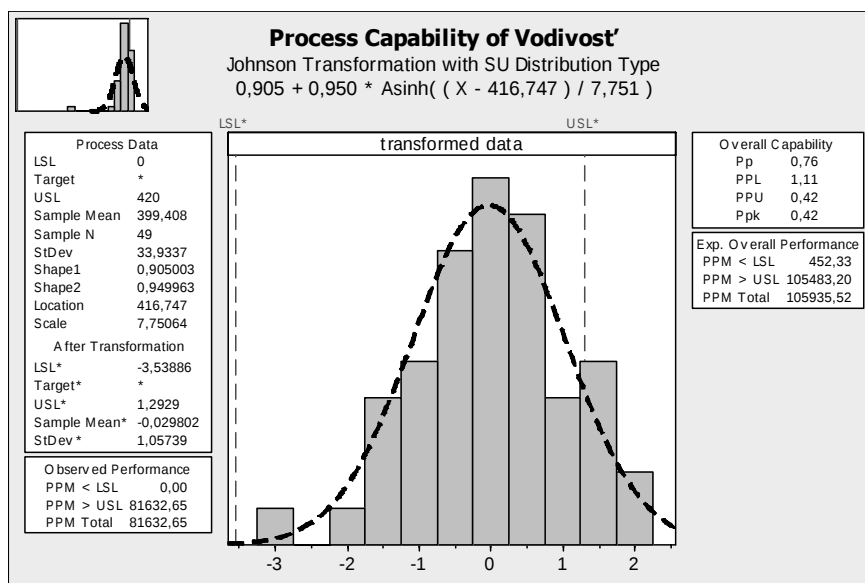
Obrázok 2: Johnsonová transformácia procesu pre Vodivosť

Z obrázku 2 vidieť, že transformované dáta pochádzajú z normálneho rozdelenia, pretože P-hodnota je 0,929. Môžeme teda korektne vypočítať indexy spôsobilosti procesu.

Tabuľka 5:

Indexy spôsobilosti pre normálne rozdelenie	Cp	CPL	CPU	Cpk
	0,76	1,11	0,42	0,42

Analýza spôsobilosti procesu uvedená na obrázku 3 ukazuje, že proces úpravy úžitkovej vody je nespôsobilý.



Obrázok 3: Spôsobilosť procesu pre Vodivost' – transformované dáta

4. Porovnanie spôsobilosti nameraných dát s transformovanými dátami

Namerané dáta, ktoré pochádzajú z loglogistického 3 parametrického rozdelenia a transformované dáta, ktoré predstavujú výber z normálneho rozdelenia majú odlišné indexy spôsobilosti (pozri tabuľku 6).

Tabuľka 6: Indexy spôsobilosti pre Loglogistické a normálne rozdelenie

	Cp	CPL	CPU	Cpk
Loglogistické 3 parametrové	2,59	5,00	0,19	0,19
Normálne rozdelenie	0,76	1,11	0,42	0,42

Indexy spôsobilosti sú výrazne pod hodnotou 1, čo znamená, že proces je nespôsobilý. Oba procesy spôsobilosti sú posunuté k hornej hranici USL. Pričom transformované dáta majú približne dvojnásobne väčšiu hodnotu Cpk.

5. Záver

Analyzovali sme namerané dáta upravenej úžitkovej vody a zistili sme, že dáta sú výberom z loglogistického 3 parametrického rozdelenia. Vypočítaný odhad indexu spôsobilosti $Cpk = 0,19$ je nízky, čo znamená, že rozdelenie je posunuté k hornej špecifikačnej medzi. Teda proces úpravy úžitkovej vody je nespôsobilý.

Transformovali sme namerané dáta na dáta, ktoré sú výberom z normálneho rozdelenia. Odhadnutý index spôsobilosti $Cpk = 0,42$ je väčší ako predchádzajúci index, čo znamená, že transformáciou došlo k zlepšeniu spôsobilosti. Obidva indexy spôsobilosti sú menšie ako 1, čo znamená, že proces je nespôsobilý.

6. Literatúra

- [1]TEREK, M. – HRNČIAROVÁ, Ľ. 2004. Štatistické riadenie kvality. IURA EDITION. 2004, Bratislava: 234 s. ISBN 80-89047-97-1.
- [2]TEREK, M. – HRNČIAROVÁ, Ľ. 2009 Metodológia šesť sigma – tri generácie implementácie. In: FORUM STATISTICUM SLOVACUM. ISSN 1336-7420, roč. V, č. 4, 2009, s. 99 – 104.

[3]WÜNSCH, J. a kol. 1981 Technická příručka pro pracovníky odboru úprav vody. Účelový náklad ČKD DUKLA. n.p. Praha, SNTL 1981 Praha: 778s. SIP-41063/03843.

Podakovanie

Tento príspevok vznikol s podporou grantových projektov VEGA č. 1/0543/10: Nové prístupy pri uplatnení metódy šesť sigma pri zlepšovaní kvality produkcie strojárskych a automobilových produktov VEGA č. 1/1247/08: Kvantitatívne metódy v stratégii šesť sigma.

Adresa autorov:

Ivan Janiga, doc. RNDr. PhD.
Karol Žák, Bc.
Andrej Valent, Bc.

Slovenská technická univerzita
Strojnícka fakulta
Námestie slobody 17
812 31 Bratislava
ivan.janiga@stuba.sk

Validation of production planning optimization model in agriculture

Validace optimalizačního modelu pro plánování produkce v zemědělství

Jitka Janová

Abstract: The design of valid applicable optimization models for production planning in agriculture is currently in focus of operations researchers worldwide. A number of models covering variety of constraints were designed aiming to be applied in the real world problems. Validating this models play a key role in the process of transferring the decision support systems from theory to practice. Nevertheless, the validation procedure of agriculture optimization models developed so far is often missing in the research articles. This stems mainly from the fact, that no unified validation process has been designed up to now for the particular use in agriculture optimization. Hence, the validation should be designed together with model developed according to the model characteristics and data availability. This paper aims at developing the particular validation procedure for the case of production planning optimization in agriculture when the crop succession requirements are considered. The validation concept is described in detail and the results of validation are visualized graphically and discussed.

Key words: crop rotation, agriculture production planning, stochastic programming, resowing constraint, validation

Klíčová slova: následnost plodin, osevní vazba, plánování produkce v zemědělství, stochastické programování, optimalizace v zemědělství, validace

1. Introduction

Currently, there is a running research on the inclusion of the crop rotation restrictions into mathematical programming models of production planning in agriculture. In [10] the method of writing the crop succession requirements as linear programming constraints in LP-based model for agricultural production planning is presented and in [5] the dynamical programming approach to crop rotation modeling is developed. Various approaches based on nonlinear, stochastic or dynamic programming are used for developing the local agricultural decision support models considering crop rotation (see [2-3] [5-13]). Some software tools supporting the agricultural decision making concerning crop rotations were developed (see [7], [2]).

The decision making and optimization problems in environmental sciences are more difficult to solve by common optimization techniques than similar problems in other areas due to the complexity of decision making that have to take into account the environmental aspects. Typically the problems deals with uncertainty and therefore the formulated mathematical programming models should be treated by the methods of stochastic programming (for introduction see [4]), which are still developing from the mathematical point of view.

The aim of this paper is to develop the particular validation procedure for the case of production planning optimization in agriculture for the case of agricultural cooperative in Czech Republic when the crop succession requirements are considered. The validation concept is described in detail and the results of validation are visualized graphically and discussed.

2. The stochastic programming model of production planning with resowing constraint

We formulate the decision problem of the typical farmer in Czech Republic as follows. The enterprise goal of the farmer or agricultural cooperative is to maximize the profit from the harvests of crops (denote $i, 1 \leq i \leq n$), when considering restrictions on:

- total area of arable land available (X),
- capital (N),
- maximal area fertilized by manures (M),
- minimum volume of feed crops (Q_i),
- minimal resp. maximal area cropped by particular crop (A_i resp. B_i),
- re-sowing of the crops.

Note, that also many others constraints may be taken into account according to particular farmer, the restrictions listed above represent the set of restrictions typical for the agriculture business in Czech Republic. We assume, that there is only one soil type in the whole area of arable land of particular farmer. Hence, all the crops may be planted on any part of the arable land. The farmer uses his own capital, which must cover the total costs of planting all the crops. The total costs C_i of planting crop i consists of the costs of seeds, labor and machine time used for planting the particular crop. The constraints on maximal area fertilized by manure and thresholds for areas cropped by particular crops stem from the livestock breeding potential and needs at the particular farm. Finally, the crop re-sowing constraints stem from the fact, that succession of certain crops on the same piece of land is not allowed or not advisable, otherwise soil fertility will decrease or the crops may become susceptible to diseases (e.g. corn may be re-sowed on the same area as soon as next period, while the oilseed rape after five or more periods, see [14]).

The agricultural cooperative maximizes its profit from the production under restrictions mentioned above. This optimization task can be formulated as a mathematical programming problem:

$$z^* = \max \sum_{i=1}^n c_i x_i \quad (1)$$

$$\sum_{i=1}^n n_i x_i \leq N \quad (2)$$

$$\sum_{i=1}^n m_i x_i \leq M \quad (3)$$

$$q_i x_i \geq Q_i \quad (4)$$

$$x_i \geq A_i \quad (5)$$

$$x_i \leq B_i \quad (6)$$

$$p=1 \text{ to } n: \sum_{s=1}^p x_{i_s} \leq X - \tilde{x}_{i_1 \dots i_p} - \tilde{y}_{i_1 \dots i_p} \quad (7)$$

$$x_i \geq 0. \quad (8)$$

where decision variables x_i stand for areas of arable land planted with crop i . In the goal function (1) the parameters c_i are the random variables of *total profit* per 1 ha of planted crop i defined as follows:

$$c_i = p_i q_i - n_i, \quad (9)$$

q_i being the random variable *volume of harvest per 1 ha* of respective crop-plant i . We consider n_i (*total costs per 1 ha of arable land planted by crop i*) and p_i (*selling price for 1 ton of crop i*), $1 \leq i \leq n$, constants of known values. The restrictions on the total available capital and manure are contained in (2) and (3), respectively, where n_i and m_i are unit costs and volume of manure needed, respectively. The restriction (4) ensures that the expected yield of feed crops will be at least at the minimum acceptable level and conditions (5) and (6) ensure maximal and minimal areas cropped by particular crops and (7) are the crop resowing constraints developed elsewhere (see [15]).

The model seems to represent a linear programming model, but as the harvests q_i must be treated as random variables the problem appears to be of stochastic nature. Note, that due to (9) the profits c_i are random as well. The prices per ton p_i change from one harvest to another, but for the purpose of this model we suppose, that the prices for next period are already given (their respective values were obtained by the expert estimation of the agricultural cooperative management).

Supposing that the harvests q_i and coefficients c_i are normally distributed, i.e.

$$q_i : N(\mu_i, \sigma_i^2), c_i : N(\gamma_i, \tilde{\sigma}_i^2)$$

we can transform the stochastic programming problem to a quadratic programming problem adopting the criterion suggested in [8].

$$z^* = \min \left(\frac{a}{2} x \Sigma x^T - \gamma x^T \right) \quad (10)$$

$$\sum_{i=1}^n n_i x_i \leq N \quad (11)$$

$$\sum_{i=1}^n m_i x_i \leq M \quad (12)$$

$$(\mu_i - F^{-1}(P)\sigma_i) \cdot x_i \geq Q_i \quad (13)$$

$$x_i \geq A_i \quad (14)$$

$$x_i \leq B_i \quad (16)$$

$$p=1 \text{ to } n: \sum_{s=1}^p x_{i_s} \leq X - \tilde{x}_{i_1 \dots i_p} - \tilde{y}_{i_1 \dots i_p} \quad (17)$$

$$x_i \geq 0. \quad (18)$$

Denoting Σ the covariance matrix of the random vector $(c_1, \dots, c_n)$ and x the decision variables vector, the maximization of the random profit (1) may be replaced by goal function (10) representing the minimization of the difference between the term representing the variability and the term representing the mean value of total profit. The parameter a denotes the risk aversion coefficient.

The constraint (4) where the random parameters q_i appear is replaced by the constraint (13), which ensures satisfaction of the respective restrictions with probability P (F is the cumulative distribution function).

The model were described in detail elsewhere (see [15]). Let us perform the results of the model for the particular case of South Moravian agriculture cooperative. In table 1 we can see the comparison of the model solution and the real decision of agriculture cooperative for the risk aversion coefficient $a=0,000001$ and probabability $P=0,75$.

Table 1: The results in of optimization for two choices of parameters a and P

a		0,000001	real decision
P		0,75	[ha]
x_1	winter food wheat	181	188
x_2	winter feed wheat	69	78
x_3	spring food wheat	0	0
x_4	spring feed wheat	19	13
x_5	winter food barley	2	0
x_6	winter feed barley	20	17
x_7	spring food barley	254	260
x_8	spring feed barley	76	74
x_9	triticale	127	0
x_{10}	corn	20	29
x_{11}	corn silage	314	230
x_{12}	oilseed rape	58	230
x_{13}	potatoe	0	0
x_{14}	grass	100	79
x_{15}	grass seed	85	64
Mean total profit [Kč]		3142000	3427000

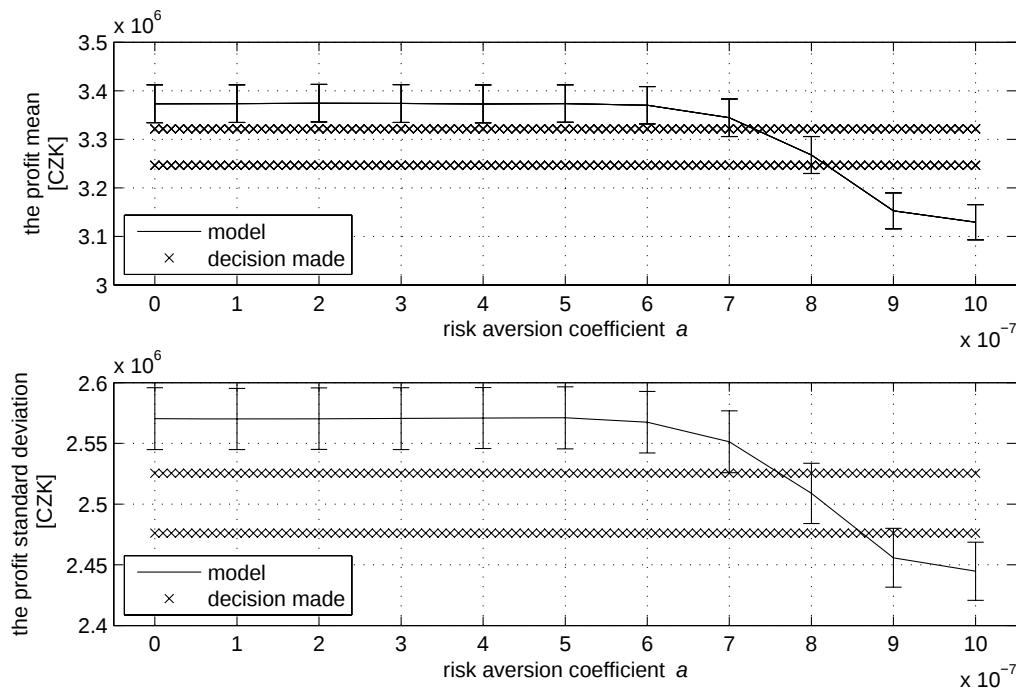
3. The validation procedure for the production planning model

The model is validated by checking whether the optimal sowing plan generates sufficient profit. The mathematical program developed ensures, that for the given goal function and constraints the solution presents the sowing plan with the extreme value of the goal function. As we have seen above, the meaning of resulting goal function differs from single profit due to the stochasticity of parameters involved. Hence, we should evaluate the profitability of the optimal sowing plan suggested by the model.

The validation is performed using the Monte Carlo simulation and comparing model solution profit characteristics with those resulting from the real sowing plan at the farm for 2009 (see table 1 for the results, when particular choice of risk aversion coefficient were made). For the purpose of the harvests simulation the probability distribution of natural yields q_i was described using the Beta distribution, which can be used in the same situations as the triangular distribution but here the underlying distribution is smoothened to reduce the peakiness of standard triangular distribution, i.e. the Beta distribution fits better the natural phenomena (see [16]). The correlation between crop yields was considered using the Spearman rank correlation. Using the harvests simulation the profit characteristics (the 95% confidence interval for mean and standard deviation) were enumerated and the results are visualized in Fig. 1.

For each risk aversion coefficient there are particular confidence intervals for mean profit and profit standard deviation for the optimal sowing plan suggested by the model. Apart from this model information, there are visualized the same profit characteristics also for the real decision of the farm. We can see, that according to the decision maker's risk attitude the profit mean and standard deviation confidence intervals resulting from the modeled optimal sowing plan can be above or below those based on the real decision. Nevertheless, there are two important results stemming from the graph:

Fig 1: Visualization of validation results for several values of risk aversion coefficients



- The profit characteristics of the model optimal solution are comparable to those obtained by the real sowing plan, which represents the current farmer's solution. This means that the model considering the randomness of the related processes provides the solution generating the profit, which is comparable to the one performed so far. The aim of using the decision support system is of course find more optimal sowing plan than the one used so far by the farmer and in this way increase the profit of the farmer. Although the current real sowing plan seems to generate very good profit characteristics, one must take into account, that this sowing plan breaks the re-sowing constraint, because for the oil seed rape there is only 131 ha of available land, while the farmer sowed by oilseed rape 230 ha of arable land. The farmer assumes, that breaking the re-sowing constraint will not affect the harvests, which is not true from the long run perspective. From the model's point of view such solution is infeasible, but if desired, the model is able to cover such exceptions.
- The model reflects the risk as expected: the higher aversion to risk the lower profit mean but also lower profit variance.

4. Conclusion

The model presented considers the randomness of the crop yield and covers the restriction on total costs, available land, available manure, minimal feasible harvests of feed crops and resowing restrictions using the resowing constraints suggested. The stochastic programming model is treated by the approach suggested by Freund [8]. The results obtained from the model represent the areas of arable land cropped by particular crop plants and these areas fulfill the requirements on crop resowing.

The validation concept presented enables the decision maker to validate the model using the data commonly available at farms. The graphical visualization makes the results of validation more clear.

Acknowledgement: The research is supported by the Grant MSM6215648904 of the Ministry of Education, Youth and Sports of Czech Republic.

References

- [1] Janová, J. Ambrožová, P: Optimization of production planning in Czech agricultural co-operative via linear programming (in czech. Avta univ. Agric. Et silvic. Mendel. Brun., 2009, LVII, No.6, 99-104.
- [2] Bachinger, J, Zander, P: ROTOR, a tool for generating and evaluating crop rotations for organic farming systems, *EUROP J AGRONOMY*, 26, (2007), 130-143.
- [3] Benjamin, LR , Milne, AE , Parsons, DJ , Cussans, J , Lutman, PJW : Using stochastic dynamic programming to support weed management decisions over a rotation *WEED RES*, 49 (2), (2009), 207-216.
- [4] Birge JR, Louveaux F: Introduction to Stochastic Programming. Corrected edition. Springer, 2002.
- [5] Castellazzi, MS et al.: A systematic representation of crop rotations. *AGR SYST*, 97, (2008), 26-33.
- [6] Detlefsen, NK, Jensen, AL ,: Modelling optimal crop sequences using network flows, *AGR SYST*, 94 (2), (2007), 566-572.
- [7] Dogliotti, S; Rossing, WAH; van Ittersum, MK: ROTAT, a tool for systematically generating crop rotations *EUROPEAN JOURNAL OF AGRONOMY*, 19 (2), (2003), 239-250.
- [8] Freund RJ: The introduction of risk into a programming model. *ECONOMETRICA* 24 (1956), 253-263.
- [9] Jatoe, JBD , Yiridoe, EK , Weersink, A , Clark, JS ,: Economic and environmental impacts of introducing land use policies and rotations on Prince Edward Island potato farms, *LAND USE POLICY*, 25 (3), (2008), 309-319.
- [10] Klein Hanveled, WK, Stegeman, AW: Crop succession requirements in agricultural production planning. *EUR J OPER RES*, 166, (2005), 406-429
- [11] Myers, P , McIntosh, CS , Patterson, PE , Taylor, RG , Hopkins, BG : Optimal crop rotation of idaho potatoes *AM J POTATO RES*, 85 (3), (2008), 183-197.
- [12] Parsons, DJ , Benjamin, LR , Clarke, J , Ginsburg, D , Mayes, A , Milne, AE , Wilkinson, DJ : Weed Manager-A model-based decision support system for weed management in arable crops, *COMPUT ELECTRON AGRIC*, 65 (2), (2009), 155-167.
- [13] Seppelt, R: Regionalised optimum control problems for agro ecosystem management, *ECOL MODEL*, 131 (2-3), (2000), 121-132.
- [14] Vach, M, Javurek, M: Rostlinná produkce s ohledem na agroekologická hlediska, Výzkumný ústav rostlinné výroby, v.v.i. 2008, ISBN 978-80-87011-58-4.
- [15] Janová, J. Production planning in agriculture: the construction of resowing constraint, In *Firma a konkurenční prostředí*, Brno, 2010
- [16] Spicka J., Boudny, J., Janotova J., The role of subsidies in managing the operating risk of agricultural enterprises, *Agric. Econ.-Czech*, 55(4), (2009), 169-179.

Corresponding adress:

Janová Jitka, Mgr. Ph.D,
PEF, MZLU v Brně
Zemědělská 1
61300 Brno
janova@mendelu.cz

Prognózovanie vývoja nových objednávok v priemyselnej výrobe SR Forecasting new orders in manufacturing in the SR

Jana Juriová

Abstract: The paper presents results of modelling for index of new orders in manufacturing in the SR that is regarded as an indicator of expected development. The results of business tendency surveys (BTS) are used as explanatory variables in the error correction models. One model is based on industrial confidence indicator (composed of balances) and the second uses information from microdata (time series of percentages of answers for questions in BTS) transformed into latent variables by means of the principal component analysis. Explanatory power of both models is compared ex-post.

Key words: Business tendency surveys, Error correction model, Latent variable, Principal component analysis

Kľúčové slová: konjunkturálne prieskumy, model s korekčným členom, latentná premenná, metóda hlavných komponentov

1. Úvod

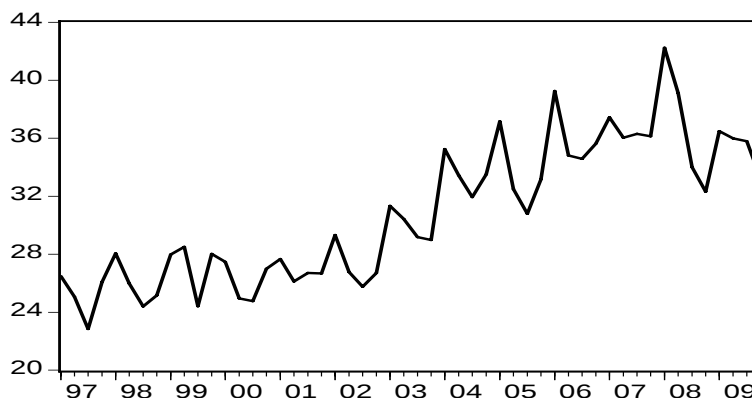
V súčasnosti, v čase globálnej hospodárskej krízy a s tým súvisiaceho nestabilného ekonomického prostredia, má čoraz väčšie opodstatnenie získavanie najskorších informácií o očakávanom vývoji v ekonomike. Veľká časť ekonomickej aktivity Slovenskej republiky (SR) sa realizuje v odvetví priemyslu, ktorý tvorí približne tretinu z celkovej hrubej pridanej hodnoty v SR a tento podiel sa v posledných rokoch ešte zvyšuje (vývoj podielu priemyslu na celkovej hrubej pridanej hodnote charakterizuje obrázok 1). Cieľom príspevku je preto poukázať na možnosti ako získať skoré informácie o očakávanom ekonomickom vývoji práve v odvetví priemyslu.

Ako jeden z hlavných kvantitatívnych ukazovateľov, ktorý charakterizuje vývoj v odvetví priemyslu sa ponúka napr. index priemyselnej produkcie. Ale vzhľadom na to, že cieľom analýzy je získanie skorých informácií, ako predmet prognózovania je zvolený ukazovateľ nových objednávok. V porovnaní s indexom priemyselnej produkcie má výhodu v tom, že poskytuje skoršiu informáciu o očakávanom vývoji v odvetví. Index priemyselnej produkcie sa z hľadiska analýzy konjunkturálnych cyklov vo všeobecnosti nepovažuje za nositeľa skorej informácie o bodoch obratu v ekonomike. A práve v súčasnosti, v čase globálnej finančnej krízy, môže index nových objednávok poskytnúť informáciu o možnom cyklickom vývoji v ekonomike.

Pre prognózovanie vývoja nových objednávok nie sú k dispozícii žiadne skoršie dostupné kvantitatívne údaje. Z toho dôvodu je buď možné sústrediť sa na modelovanie samotného časového radu indexu nových objednávok alebo využiť dostupné kvalitatívne informácie. Jedným zo zdrojov kvalitatívnych informácií sú konjunkturálne prieskumy realizované Štatistickým úradom SR (ŠÚ SR), ktoré sú zverejňované už na konci referenčného mesiaca.

Pri modelovaní indexu nových objednávok v priemyselnej výrobe bol aplikovaný ekonometrický model s korekčným členom založený na výsledkoch z konjunkturálnych prieskumov. V skutočnosti sa jedná o dve verzie tohto modelu, ktoré sa líšia použitou vysvetľujúcou premennou. V 1. verzii modelu je vysvetľujúcou premennou konjunkturálne saldo zverejňované ŠÚ SR, v 2. verzii je to latentná premenná zostrojená z mikroúdajov zisťovaných v konjunkturálnych prieskumoch, pričom pri jej konštrukcii bola využitá metóda

hlavných komponentov. Na záver bola ex-post porovnaná prognostická schopnosť oboch modelov.



Obrázok 1: Podiel priemyslu na celkovej hrubej pridanej hodnote v SR (%)

2. Metodológia

Pri modelovaní vývoja nových objednávok bol zvolený ekonometrický model s korekčným členom (Error Correction Model – ECM). Výhodou modelu ECM v porovnaní s modelmi časových radov založenými len na vlastnostiach samotného modelovaného časového radu je skutočnosť, že jeho konštrukcia je založená na ekonomickej teórii a teda na vývoji inej premennej, pre ktorú sú údaje skôr k dispozícii. To znamená, že ako vysvetľujúca premenná je vybraný taký ukazovateľ, o ktorom je odôvodnené predpokladať, že jeho vývoj súvisí s vývojom vysvetľovanej premennej. Ďalším z dôvodov použitia ECM prístupu je, že väčšina makroekonomických časových radov je nestacionárna (obsahuje trend) a keď sa použijú pri konštrukcii regresných modelov, vedie to k tzv. “zdanlivej” regresii (spurious regression) [6]. Z toho vyplýva, že v regresných modeloch by nemali vystupovať nestacionárne časové rady. Za určitých podmienok je to však možné a to vtedy, keď lineárna kombinácia nestacionárnych časových radov X_t a Y_t :

$$u_t = Y_t - \beta_0 - \beta_1 X_t \quad (1)$$

je stacionárna, t.j. premenná u_t je integrovaná $I(0)$. V takýchto prípadoch hovoríme, že časové rady X_t a Y_t sú kointegrované. Kointegrácia vyjadruje existenciu vzťahu dlhodobej rovnováhy definovanej nasledovne:

$$Y_t = \beta_0 + \beta_1 X_t + u_t, \quad (3)$$

pričom u_t je náhodná zložka, pre ktorú požadujeme, aby spĺňala podmienky bieleho šumu. Vzťah kointegrácie sa zisťuje testovaním, či chyby u_t (ktoré vyjadrujú odklon od dlhodobej rovnováhy) sú stacionárne. Keďže náhodné poruchy u_t nie sú priamo pozorovateľné, používajú sa pri testovaní stacionarity (napr. Dickey-Fullerovým testom jednotkového koreňa) ich odhady (rezíduá) získané po odhade parametrov dlhodobého rovnovážneho vzťahu (3) metódou najmenších štvorcov (OLS). Stacionárny časový rad rezíduí, ktorý sa získa z dlhodobého rovnovážneho vzťahu, sa následne využije na prepojenie krátkodobej dynamiky vývoja premennej Y_t s dlhodobým rovnovážnym vzťahom, ktoré vyúsťuje do špecifikácie modelu v tvare ECM:

$$\Delta Y_t = \alpha_0 + \alpha_1 \Delta X_t - \pi u_{t-1} + \varepsilon_t \quad (4)$$

Vzťah (4) potom v sebe zahŕňa súčasne krátkodobú aj dlhodobú elasticitu vplyvu nezávisle premennej na závisle premennú. Parameter α_1 vyjadruje krátkodobý vplyv zmeny premennej X_t na zmenu premennej Y_t . Parameter π alebo parameter tzv. korekčného člena predstavuje časť nerovnováhy z predchádzajúceho obdobia (v prípade mesačných časových radov

mesiacu), o ktorú je korigovaný vývoj Y_t v nasledujúcom období. Premenná ε_t je náhodná zložka s vlastnosťami bieleho šumu.

EC model vývoja indexu nových objednávok v priemyselnej výrobe je založený na kvalitatívnych informáciách (údajoch) vyplývajúcich z konjunkturálneho prieskumu v priemysle, ktoré sú známe už 2 dni pred koncom referenčného mesiaca. Jedná sa o informácie dostupné v tvare konjunkturálnych sáld, ktoré sú definované ako rozdiel pozitívnych a negatívnych hodnotení respondentov vyjadrený v percentách.

Konstruktia modelového vzťahu v tvare ECM je založená na východiskovej hypotéze, podľa ktorej index nových objednávok v priemyselnej výrobe (INO) rastie v čase v zásade konštantným tempom a výkyvy okolo trendu sa menia v závislosti od vývoja indikátora dôvery v priemysle (IDP) [5]. Indikátor dôvery v priemysle vzniká ako aritmetický priemer 3 konjunkturálnych sáld: očakávanej produkcie, dopytu a zásob hotových výrobkov (s opačným znamienkom). Dlhodobý vzťah v tvare ECM má potom nasledovný tvar:

$$INO = a * e^{b*time+c*IDP} \quad \text{alebo} \quad \log(INO) = \log a + b * time + c * IDP \quad (5)$$

Okrem konjunkturálnych sáld, ktoré sú zverejňované ŠÚ SR, je však možné získať aj priamo mikroudaje z konjunkturálnych prieskumov, t.j. časové rady percent pozitívnych, negatívnych a neutrálnych odpovedí na jednotlivé otázky. Jednotlivé časové rady percent odpovedí môžu byť tiež použité ako vysvetľujúce premenné v regresných modeloch [2], avšak samotné nemajú takú vypovedaciu schopnosť ako napr. zložený indikátor dôvery. Ďalšou možnosťou je použiť ich v modeloch viaceru, avšak jednotlivé percentá odpovedí na každú otázku sú medzi sebou vzájomne relatívne vysoko korelované a je preto vhodné pokúsiť sa nájsť vyhovujúcu latentnú premennú a znížiť tak počet vysvetľujúcich premenných, napr. prostredníctvom niektorej z metód viackriteriálnej analýzy.

Pri hľadaní najlepšej lineárnej kombinácie jednotlivých percent odpovedí je možné využiť *metódu hlavných komponentov* (vychádza z korelačnej matice pôvodných premenných). Jej základnou myšlienkou je redukcia nadbytočných informácií obsiahnutých vo viacerých korelovaných premenných. Metóda hlavných komponentov je výhodná aj z toho dôvodu, lebo vysvetľuje všetku variabilitu medzi pôvodnými premennými (na rozdiel napr. od faktorovej analýzy), teda snaží sa zachovať čo najviac informácií obsiahnutých v pôvodných premenných. Týmto spôsobom boli nájdené najvhodnejšie latentné premenné a následne použité v EC modeli v tvare (5) namiesto vysvetľujúcej premennej IDP.

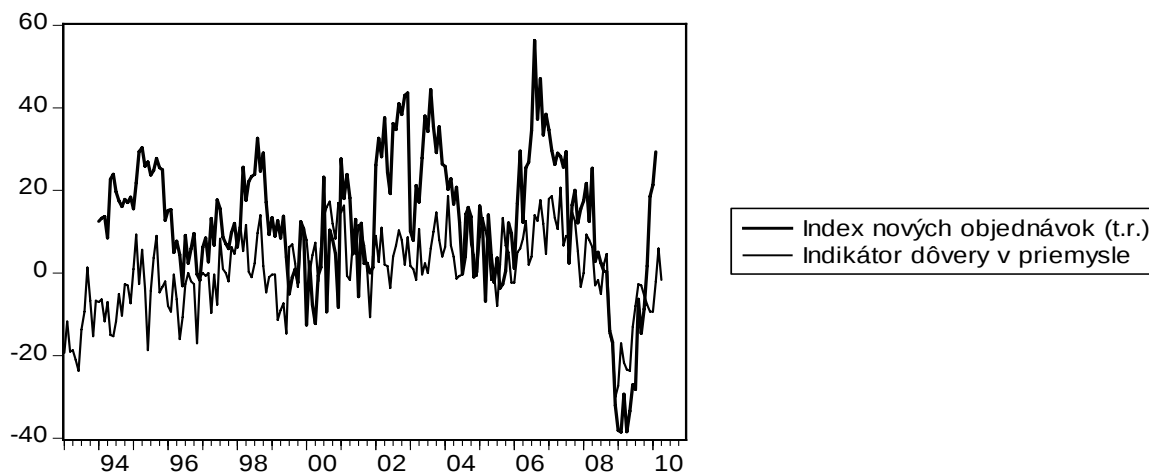
Jednotlivé modely boli odhadnuté pomocou OLS v ekonometrickom programe EViews.

3. Výsledky modelovania a prognostická aplikácia

Predmetom modelovania bol bázičný index nových objednávok v priemyselnej výrobe (v stálych cenách, 2005=100), ktorý zverejňuje Eurostat do 2 mesiacov po ukončení referenčného mesiaca (napr. 20. apríla 2010 boli zverejnené údaje za február). Časový rad indexu nových objednávok je k dispozícii za obdobie január 1993 až február 2010. V tom istom časovom období je časový rad indikátora dôvery v priemysle dostupný za obdobie január 1993 až marec 2010. Korelačný koeficient medziročného tempa rastu indexu nových objednávok a indikátora dôvery v priemysle dosahuje za celé obdobie hodnotu cca 0.47, avšak závislosť oboch ukazovateľov sa v posledných obdobiach zvýšila, pri skrátení analyzovaného obdobia na január 2003 až február 2010 sa koeficient korelácie zvýšil na cca 0.74. Porovnanie vývoja oboch ukazovateľov ilustruje obrázok 2, jedná sa o sezónne neočistené údaje.

Vzhľadom na výsledky korelačnej analýzy sú jednotlivé modely odhadnuté za skrátené obdobie od januára 2003 po február 2010. Pri odhade EC modelov je použitý vyššie popísaný dvojkrokový Engle-Grangerov postup. To znamená, že v 1. kroku je odhadnutý dlhodobý

vzťah v tvare (5), pričom pre vysvetlenie sezónnosti sú použité umelé sezónne premenné SD_i ($i=1,...,12$), ktoré sa v daných mesiacoch ukázali ako štatisticky významné. Indikátor dôvery v priemysle je v modeli použitý s časovým posunom 1 mesiaca, keďže tento postup vedie k získaniu modelu s vyššou výrokovou schopnosťou (IDP zahŕňa aj konjunktúrne saldo očakávanej produkcie na nasledujúce 3 mesiace a z tohto dôvodu môže byť informácia obsiahnutá v ukazovateli IDP tiež mierne časovo posunutá).



Obrázok 2: Porovnanie vývoja medzročného tempa rastu indexu nových objednávok a indikátora dôvery v priemysle

Odhadnutý dlhodobý vzťah pre index nových objednávok predstavuje rovnica (6). Všetky odhadnuté parametre sú štatisticky významné na 5%-nej hladine významnosti. Rezíduá (REZ) odhadnutého dlhodobého rovnovážneho vzťahu sú následne testované pomocou Dickey-Fullerovho testu jednotkového koreňa [1]. Po overení stacionarity rezíduí sa tieto následne použijú v 2. kroku pri odhade vzťahu v tvare modelu s korekčným členom – rovnica (7), pričom sezónnosť v krátkodobom vzťahu reprezentujú diferencie sezónnych premenných.

$$\log(INO)=3.078174749+0.01019303263*time+0.01427929347*IDP(-1)-0.09607095582*SD12 \quad (6)$$

$$d\log(INO)=0.004612487595*d(IDP(-1))-0.3894733377*REZ6(-1)+0.05935775848*d(SD1)+0.06350907023*d(SD3)+0.156267023*d(SD9)+0.149554667*d(SD10)+0.1470933944*d(SD11) \quad (7)$$

Všetky odhadnuté parametre sú štatisticky významné na 5%-nej hladine významnosti aj v prípade modelu odhadnutého v tvare EC. Jednotlivé štatistické charakteristiky modelov sú uvedené v tabuľke 1, kde sú súčasne porovnané aj s charakteristikami odhadnutými pre druhú verziu modelu založenom na latentných premenných.

Prístup založený na latentných premenných bol zvolený z toho dôvodu, aby sa využila aj informácia obsiahnutá v ostatných konjunktúrnych saldách, ktoré nie sú zahrnuté v IDP. Na základe jednoduchých regresných modelov (8) pre medzročné tempo rastu indexu nových objednávok s každými 3 časovými radmi percent odpovedí (p_i) sú identifikované kombinácie s najvyšším vysvetleným rozptylom a najvyššími odhadnutými parametrami: časové rady negatívnych hodnotení pre otázky 2, 3, 4 a 7 (2 – úroveň celkového dopytu, 3 – úroveň zahraničného dopytu, 4 – zásoby hotových výrobkov, 7 – očakávaný počet zamestnancov) a časové rady pozitívnych hodnotení pre otázky 1 a 5 (1 – trend produkcie za posledné 3 mesiace, 5 – očakávaná produkcia).

$$\Delta INO_t = a_1 p_{1t} + a_2 p_{2t} + a_3 p_{3t} + u_t \quad (8)$$

V ďalšom kroku je aplikovaná metóda hlavných komponentov pre vytvorenie samotných latentných premenných. Pre negatívne hodnotenia otázok 2, 3, 4, 7 je identifikovaný 1. hlavný komponent (N4_KOMP1) vysvetľujúci 66% variability premenných a pre pozitívne hodnotenia otázok 1, 5 1. hlavný komponent (P2_KOMP1) vysvetľujúci 53% variability premenných. Oba skonštruované hlavné komponenty sú použité najskôr pri odhade dlhodobého rovnovážneho vzťahu (9) a po otestovaní jeho rezíduí aj pri odhade modelu s korekčným členom (10). Parametre oboch odhadnutých modelov sú štatisticky významné na 5%-nej hladine významnosti.

$$\log(INO) = 3.340608428 + 0.008889008941 \cdot time + 0.08078274578 \cdot N4_KOMP1 + 0.02252081517 \cdot P2_KOMP1 - 0.09789092459 \cdot SD8 - 0.09410534764 \cdot SD12 \quad (9)$$

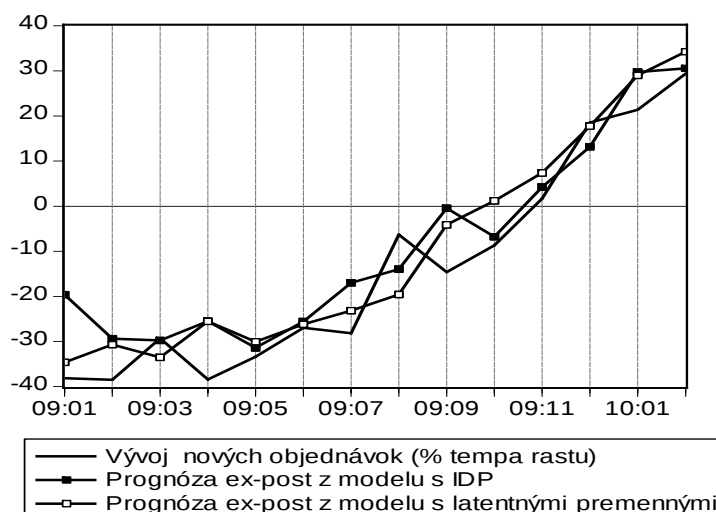
$$d\log(INO) = 0.03467607071 \cdot d(N4_KOMP1) + 0.0299916758 \cdot d(P2_KOMP1) - 0.4957807469 \cdot REZ8(-1) - 0.05071488475 \cdot d(sd2) - 0.06835478773 \cdot d(SD4) + 0.04520904798 \cdot d(SD6) - 0.1176083382 \cdot d(SD8) - 0.1031574392 \cdot d(SD12) \quad (10)$$

Z porovnania jednotlivých štatistických charakteristík všetkých odhadnutých modelov v tabuľke 1 vyplýva, že model s korekčným členom založený na latentných premenných dosahuje vyššie % vysvetleného rozptylu vývoja indexu nových objednávok ako model založený na indikátore dôvery v priemysle.

Tabuľka 1: Porovnanie štatistických charakteristík modelov pre index nových objednávok

vysvetľujúce premenné modelu	IDP	latentné premenné
R^2 – dlhodobý rovnovážny vzťah	0.83	0.87
D-W – dlhodobý rovnovážny vzťah	1.20	1.18
Dickey-Fullerova štatistika pre rezíduá	-3.59	-5.88
Kritická hodnota testu jednotkového koreňa (1%)	-2.59	-2.59
R^2 – model s korekčným členom	0.53	0.62

Prognostická schopnosť oboch odhadnutých modelov je overená ex-post za obdobie od januára 2009 po február 2010 a porovnaná so skutočným medziročným tempom rastu indexu nových objednávok. Prognózy sú výsledkom odhadnutých modelov vždy za aktuálne dostupné obdobie v čase prípravy prognózy a odhadnuté medziročné rasty sú výsledkom statickej simulácie modelov. Aj priemerná štvorcová chyba prognóz MSE (Mean Square Error) [7] potvrdila, že model založený na latentných premenných lepšie vysvetľuje skutočnosť. Pre prognózu z modelu s IDP je MSE za analyzované obdobie 78.04 a pre prognózu z modelu s latentnými premennými 56.62.



Obrázok 3: Porovnanie medziročných rastov (%) indexu nových objednávok v priemyselnej výrobe s prognózovanými hodnotami

4. Záver

Pri prognózovaní vývoja nových objednávok v priemyselnej výrobe sa potvrdzuje vhodnosť použitia modelu s korekčným členom. Odhady modelov aj prognostická aplikácia ex-post ukazujú, že model využívajúci latentné premenné (skonštruované z mikroúdajov z konjunkturálnych prieskumov) má lepšiu vysvetľujúcu schopnosť ako model založený na konjunkturálnych saldách reprezentovaných indikátorom dôvery v priemysle.

5. Literatúra

- [1] ASTERIOU, D. - HALL, S.G. 2007. Applied Econometrics (A Modern Approach using EViews and Microfit). Palgrave Macmillan, Revised edition 2007.
- [2] D'ELIA, E. 2005. Using the results of qualitative surveys in quantitative analysis. Working paper n. 56, September 2005.
- [3] EViews 5 User's Guide, Quantitative Micro Software, LLC, 2004.
- [4] GREEN, W. H. 2003. Econometric Analysis. New Jersey: Prentice – Hall, 2003. ISBN 0-13-066189-9.
- [5] HALUŠKA, J., OLEXA, M., JURIOVÁ, J., KLÚČIK, M. 2008 – Modelový aparát na rýchle odhady vývoja makroekonomických ukazovateľov slovenskej ekonomiky (Využitie konjunkturálnych a spotrebiteľských prieskumov). Bratislava: INFOSTAT, 2008. ISBN 978-80-89398-06-5
- [6] HATRÁK, M. 2007. Ekonometria. Bratislava: Iura Edition, Edícia EKONÓMIA, 2007.
- [7] RUBLÍKOVÁ, E. 2007. Analýza časových radov. Bratislava: Iura Edition, Edícia EKONÓMIA, 2007. ISBN 978-80-8078-139-2.

Adresa autora:

Jana Juriová, Ing.
Dúbravská cesta 3
845 24 Bratislava
juriova@infostat.sk

Produkční a inflační mezera po zavedení eura – vliv společné monetární politiky ECB

Growth and Inflation Gap after joining the Euro – impact of the ECB's Single Monetary Policy

Svatopluk Kapounek

Abstract: The aim of this article is to identify growth and inflation gap in the Eurozone member states after joining the Euro. Author compares growth and inflation gap with the monetary policy character. In the empirical analysis is used Hodrick-Prescott filter, results are discussed in the Taylor rule sense.

Key words: Taylor rule, Hodrick-Prescott filter, single monetary policy

Klíčová slova: Taylorovo pravidlo, Hodrick-Prescott filter, společná monetární politika

1. Úvod

Hlavním cílem příspěvku je identifikovat produkční a inflační mezeru v členských státech eurozóny v krátkém období po zavedení společné měny euro. Autor se tak věnuje analýze očekávání, která byla formulována a prezentována veřejnosti v období před vstupem jednotlivých států do eurozóny v souvislosti s naplněním předpokladů endogenity procesu evropské integrace. Autor hledá odpověď na otázku, zda a jak se změnila výkonnost jednotlivých ekonomik eurozóny ve vztahu k inflačnímu vývoji za předpokladu implementace společné monetární politiky Evropské centrální banky (ECB).

V případě, že Evropská centrální banka nastavuje svou monetární politiku především s ohledem na dosažení primárního cíle (cenové stability), potom mohou mít země po vstupu do eurozóny problémy se zpomalením ekonomického růstu oproti situaci ponechání si autonomie v této oblasti. Snížení dynamiky ekonomického růstu pak může být z pohledu občanů jednotlivých členských zemí jednoznačným nákladem spojeným s existencí společné měny euro. Nespokojenost občanů pak může být zdrojem politické nestability a snah politiků o získání popularity formou „siláckých“ projevů a vystoupení země z HMU, jak se již stalo v případě Itálie.

Možnost ovlivnit nastavení monetární politiky ECB z pohledu členských států je statutem a vysokou nezávislostí ECB velmi limitována. Rozhodovací mechanismus omezuje vliv členských států prosadit politiku „ušitou na míru“ potřeb jedné národní ekonomiky. ECB nedokáže svými nástroji reagovat ani na důsledky asymetrického šoku v jedné ekonomice. ECB se dívá se na eurozónu jako celek a váhu při výpočtu požadované úrokové sazby přiřazuje dle ekonomické výkonnosti daného státu.

2. Metodika

Empirická analýza je založena na odhadu produkční mezery a na jejím vývoji ve členských zemích HMU před a po integraci. Produkční mezeru lze určit jako rozdíl mezi skutečným a potenciálním hospodářským růstem. Vlastní odhad potenciálního produktu je možné provést dvěma způsoby. Pomocí ekonometrického modelu založeném na produkční funkci nebo určením potenciálního produktu pomocí vyhlazení HDP, tedy určením dlouhodobého trendu. Produkční mezera je pak definována jako rozdíl mezi skutečným a trendovým HDP (*Hájek, M., Bezděk, V., 2000*):

$$X_t = Y_t - Y_t^* \quad (1)$$

kde Y_t reprezentuje HDP ve stálých cenách a Y_t^* jeho potenciální, trendovou, úroveň. V analýze je použita druhá metoda, identifikace dlouhodobého trendu v časové řadě (např. *Woodford, 2001*). Vlastní určení dlouhodobého trendu je možné realizovat prostřednictvím filtrů. *Canova (1999)* se zabývá identifikací trendu a cyklické složky v hospodářském cyklu pomocí Kalmanova filtru, Hodrick-Prescott filtru a band-pass filtru. Sami autoři *Hodrick a Prescott (1980)* doporučují použití Hodrick-Prescottova filtru pro identifikaci dlouhodobého trendu v hospodářském cyklu. Stejný názor potvrzuje ve své práci také *Hájek a Bezděk (2000)*. Hodrick-Prescottův filtr je definován jako:

$$\text{Min} \left\{ \sum_{t=1}^T (Y_t - Y_t^*)^2 + \lambda \sum_{t=2}^{T-1} [(Y_{t+1}^* - Y_t^*) - (Y_t^* - Y_{t-1}^*)]^2 \right\} \quad (2)$$

kde λ je parametr určující hladkost trendového vyhlazení. Pro účely této empirické analýzy byla hodnota parametr $\lambda=1600$, což je obecně uznávaná hodnota tohoto parametru pro čtvrtletní pozorování a určení potenciální úrovně HDP.

V souvislosti s definicí primárního cíle ECB v podobě cenové stability a sekundárních cílů prostřednictvím dlouhodobě udržitelného růstu a stabilní zaměstnanosti je v další části analýzy vypočtena inflační mezera jako rozdíl mezi její skutečnou a cílovou hodnotou. Za cílovou hodnotu inflace byla pro všechny členské státy HMU zvolen oficiální inflační cíl ECB, 2% meziroční změna harmonizovaného indexu spotřebitelských cen.¹

Analýza vlivu monetární politiky ECB na růst a inflaci v HMU je na obrázcích 2 až 12 založena na vzájemném vztahu produkční a inflační mezery na straně jedné a krátkodobých úrokových sazeb na mezibankovním trhu na straně druhé. Tento vztah exaktně popisuje Taylorovo pravidlo, které definoval v roce 1993 ekonom ze Stanford University, John Taylor:

$$i = i^* + a(\pi - \pi^*) + b(y - y^*) \quad (3)$$

kde i je výsledná požadovaná úroková sazba, i^* rovnovážná úroková sazba (používají se dlouhodobé vládní cenné papíry), π míra inflace, π^* inflační cíl, y růst HDP, y^* potenciální růst a $(y - y^*)$ mezera růstu. (*Taylor, 1993 a Walsh, 2003*)

Centrální banka zvyšuje úrokové sazby, tedy působí monetární restrikcí, pokud inflace roste nad inflační cíl (anebo v případě cyklického růstu, pokud růst inflace přesahuje trend), případně pokud aktuální hospodářský růst převyšuje jeho potenciální úroveň² a hrozí tak přehřátí ekonomiky a vznik inflačních tlaků.

Velmi důležitým ukazatelem jsou ale váhy, parametry a a b . Tyto představují preference inflační a růstové mezery a specifikují charakter centrální banky. Pokud $b = 0$, lze hovořit o centrální bance silně konzervativní, která sleduje pouze inflační cíl, nikoliv hospodářský růst. Begg, D. a kol. ve své práci analyzoval parametry monetární politiky ECB, váhy inflace a mezery růstu ve výši $a = 2,0$ a $b = 0,8$. Dle těchto údajů lze charakterizovat ECB jako banku, která se soustřeďuje zejména na inflační mezeru. Mezera hospodářského růstu je jejím sekundárním cílem, což plně koresponduje s jejím statutem a definicí jejích cílů v Maastrichtské smlouvě.

Daty vstupujícími do analýzy jsou meziroční změny čtvrtletního HDP a deflátoru HDP v letech 1996 - 2006. Zdrojem vstupních dat je eurostat. Měnovou politiku ECB reprezentují krátkodobé úrokové sazby na mezibankovním trhu HMU. Nedostatečná délka časových řad

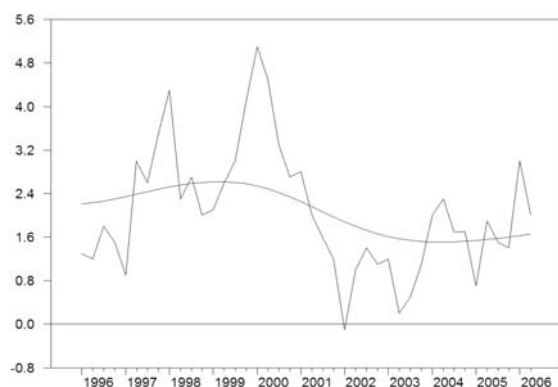
¹ V souvislosti s analýzou produkční mezery není do analýzy zahrnut harmonizovaný index spotřebitelských cen, ale deflátor HDP.

² Vzhledem k tomu, že je potenciální produkt v empirické analýze determinován filtrační technikou, která je často kritizovaná v souvislosti s potenciálem ekonomiky, autor bude dále užívat vhodnější termín dlouhodobý trend ekonomického růstu.

před zavedením společné měny omezuje analýzu pouze na vybrané členské země HMU: Belgii, Německo, Španělsko, Francii, Irsko, Itálii, Lucembursko, Nizozemí, Rakousko a Finsko. Analýza je provedena také za Eurozónu jako celek.

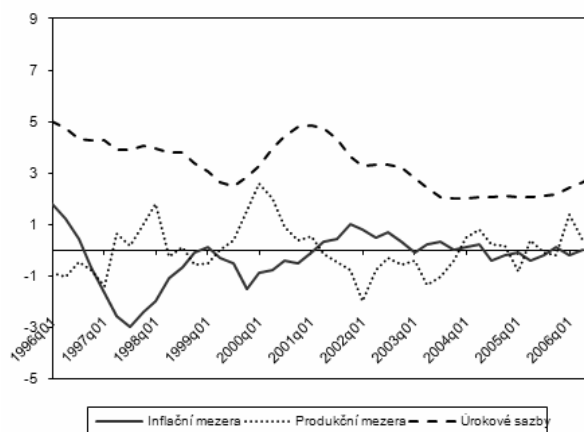
3. Výsledky empirické analýzy

Z obrázku 1 je patrná výrazná tendence k poklesu nejen skutečného HDP, ale zejména dlouhodobého trendu ekonomického růstu po roce 1999 v Eurozóně jako celku.



Obrázek 1: Ekonomický růst v Eurozóně

Zdroj: vlastní výpočet autora

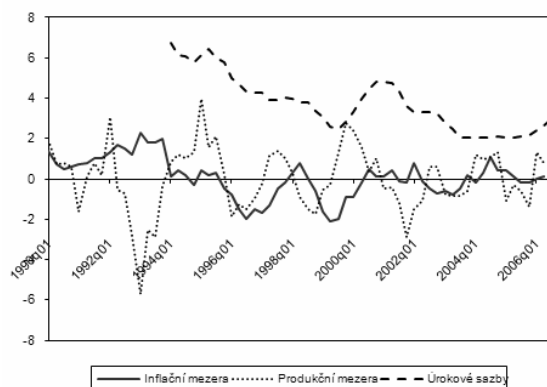


Obrázek 2: Mezera růstu a inflace v Eurozóně

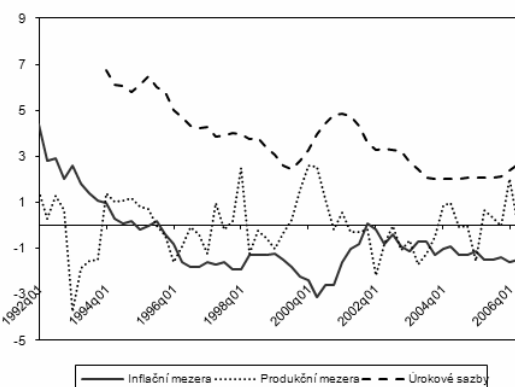
Zdroj: vlastní výpočet autora

Obrázek 2 porovnává produkční a inflační mezera s monetární politikou ECB, reprezentovanou krátkodobými úrokovými sazbami na mezibankovním trhu v HMU. Je patrné, že Eurozóna trpí výraznou produkční mezerou. Inflační mezera lze v období let 1996 až 2006 považovat za negativní. Svůj inflační cíl v podobě 2% inflace tak má ECB tendence spíše podstřelovat, což souvisí s původní definicí jejího inflačního cíle v podobě meziročního růstu harmonizovaného indexu spotřebitelských cen ve výši 0-2%. Např. v roce 2000 je možné pozorovat významné zvýšení úrokových sazeb na mezibankovním trhu i přes nárůst produkční mezery. Mezera inflační byla, i přes mírný nárůst inflace, až do počátku roku 2001 negativní.

Na obrázcích 3 až 12 potvrzují tendence ECB, preferovat inflační mezera nad růstovou, také jednotlivé členské státy HMU. Nejvýraznější kontrast mezi negativní inflační mezerou a pozitivní produkční mezerou vykazuje Německo, kde je také patrný vliv neadekvátní monetární restriktivní politiky ECB na negativní produkční mezera v letech 2001 – 2003.



Obrázek 3: Mezera růstu a inflace v Belgii

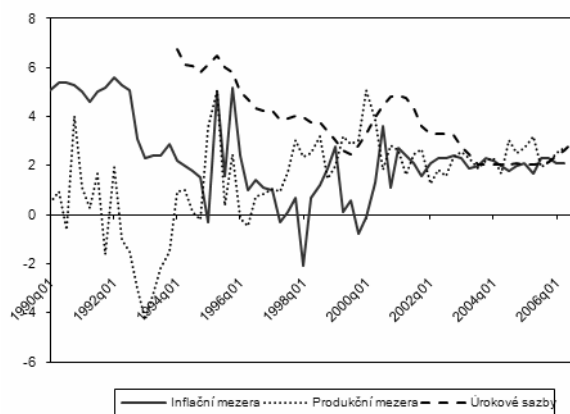


Obrázek 4: Mezera růstu a inflace v Německu

Zdroj: vlastní výpočet autora

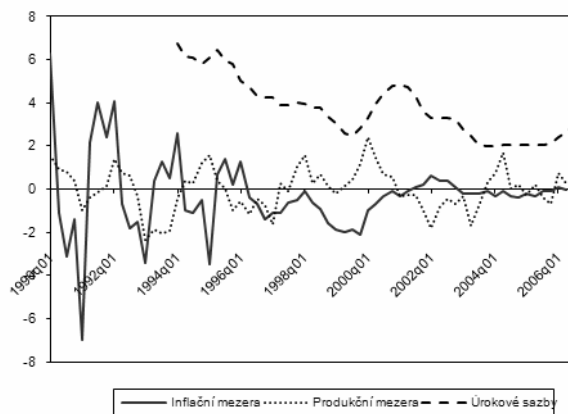
Zdroj: vlastní výpočet autora

Vliv monetární politiky ECB ve Španělsku je možné považovat za adekvátní. Pozitivní produkční mezeru je zde doprovázena pozitivní mezerou inflační. Výrazná monetární restrikce v letech 2001 – 2002 tak odpovídala zvyšujícím se inflačním tlakům v této zemi. Její negativní vliv na hospodářský růst nebyl významný. Pozitivní produkční mezeru si tak Španělsko zachovalo až do roku 2006.



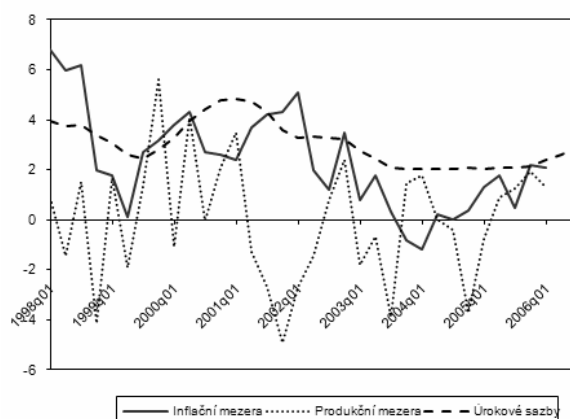
Obrázek 5: Mezera růstu a inflace ve Španělsku

Zdroj: vlastní výpočet autora



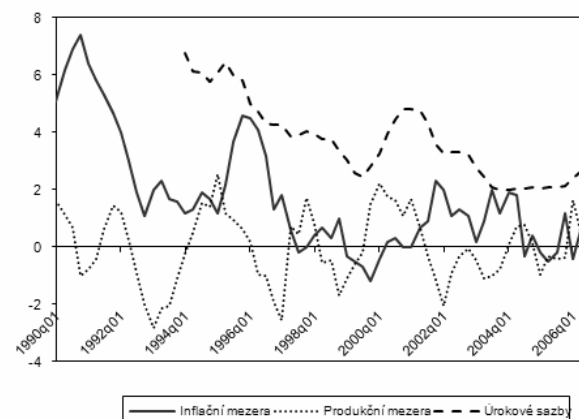
Obrázek 6: Mezera růstu a inflace ve Francii

Zdroj: vlastní výpočet autora



Obrázek 7: Mezera růstu a inflace v Irsku

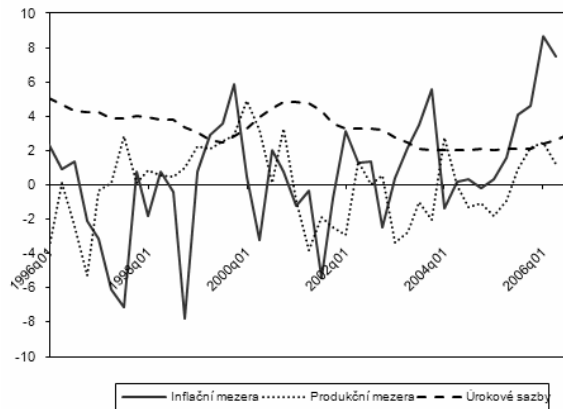
Zdroj: vlastní výpočet autora



Obrázek 8: Mezera růstu a inflace v Itálii

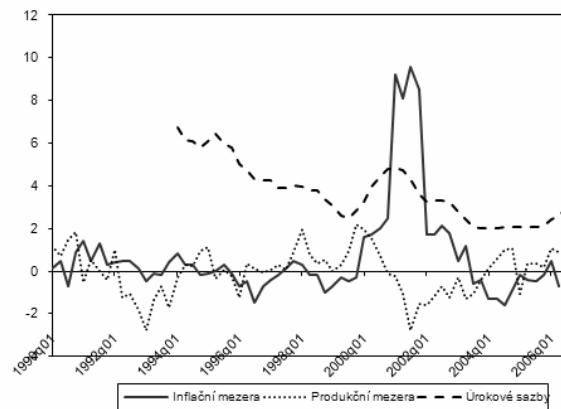
Zdroj: vlastní výpočet autora

Výraznou převahu inflační mezery nad mezerou produkční zaznamenává dlouhodobě Irsko. Výše úrokových sazeb a restriktivní monetární politika se tak v letech do poloviny roku 2003 i přes dlouhodobě klesající ekonomický růst jeví jako adekvátní. Od roku 2004 je ale negativní produkční mezeru, ve srovnání s kladnou mezerou inflační, výraznější.



Obrázek 9: Mezera růstu a inflace v Lucembursku

Zdroj: vlastní výpočet autora

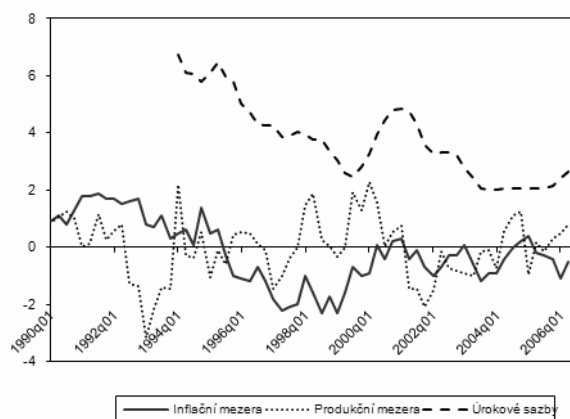


Obrázek 10: Mezera růstu a inflace v Nizozemí

Zdroj: vlastní výpočet autora

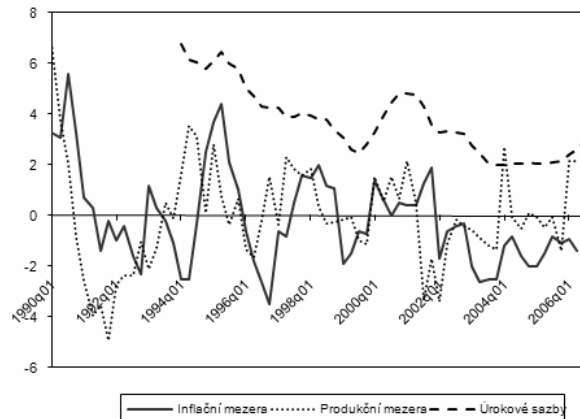
Pozitivní inflační mezera v Lucembursku i Nizozemí oproti mezeře růstové dominuje. Zatímco Lucembursko zaznamenává vysoký nárůst inflační mezery od roku 2005, v Nizozemí byl tento výkyv zaznamenán mezi roky 2000 – 2002. Produkční mezera se zde jeví jako nevýznamná.

V případě Belgie, Francie, Itálie, Rakouska a Finska nelze jednoznačně říci, jaký vliv měla monetární politika ECB reprezentovaná vývojem krátkodobých úrokových sazeb na vývoj produkční a inflační mezery.



Obrázek 11: Mezera růstu a inflace v Rakousku

Zdroj: vlastní výpočet autora



Obrázek 12: Mezera růstu a inflace ve Finsku

Zdroj: vlastní výpočet autora

4. Závěr

Po zavedení společné měny euro lze v členských státech eurozóny i přes veškerá očekávání v souvislosti s endogenitou procesu evropské integrace zaznamenat negativní vývoj dlouhodobého trendu ekonomického růstu. Výjimkami jsou Belgie a Francie, které nezaznamenaly výraznou změnu v dlouhodobém trendu ekonomického růstu po zavedení společné měny euro. U ostatních zemí je možné po roce 1999 pozorovat pokles v ekonomickém růstu. Nejvýrazněji je tomu u Irska, Španělska, Itálie, Nizozemí a Rakouska. Mírná, takřka nepatrná tendence k poklesu se projevila v Německu a Lucembursku. Jedinou zemí, ve které je možné zaznamenat výrazný růst dlouhodobého trendu ekonomického růstu je Finsko.

Jak již vyplývá z definice cílů ECB a výsledků empirické analýzy, charakter společné monetární politiky preferuje inflační mezeru nad mezerou v ekonomickém růstu. Nejvýraznější kontrast a neadekvátnost monetární politiky vzhledem k mezeře v ekonomickém růstu je patrný v Německu v letech 2001 – 2003. V ostatních členských zemích má uvedená tendence pouze mírný charakter.

Jak vyplývá již z učení Keynesa, ekonomický růst je mnohem citlivější na monetární restrikcí ve srovnání s expanzivní monetární politikou. Tím, že společná monetární politika ECB dlouhodobě systematicky střídavě potlačuje svou neadekvátní restrikcí ekonomický růst ve státech, kde se v danou chvíli snížila inflace pod cílovou hodnotu ECB, snižuje tak hospodářský růst v celé HMU. Od zavedení eura v letech 1998 a 1999 se hospodářský růst ze 3% snížil na hodnotu 1,7% v roce 2005. ECB, která by měla připravovat podmínky pro dlouhodobý a udržitelný hospodářský růst tak podporuje dlouhodobou tendenci k jeho poklesu.

Uvedené závěry jsou důležité i z pohledu zemí připravujících se na vstup do eurozóny, mezi které patří i ČR. Jedním z jejích cílů je reálná konvergence s ekonomickou úrovní vyspělých členských zemí EU. V případě, že vstup do eurozóny negativním způsobem ovlivňuje dynamiku ekonomického růstu, by tyto země měly se vstupem do HMU vyčkat až po dosažení vysoké úrovně reálné konvergence, kdy se pravděpodobnost neadekvátnosti monetární politiky ECB sníží na minimum.

Předkládaný příspěvek vznikl za podpory interního grantového projektu 71/2010 “Stabilizační funkce monetární politiky v souvislostech hospodářského cyklu České republiky”.

5. Literatura

1. Hájek, M. - Bezděk, V: Odhad potenciálního produktu a produkční mezery v České republice. Politická ekonomie, Praha : Vysoká škola ekonomická, ISSN 0032-3233, 2001, 2001/IL, č. 4, s. 473-491;
2. Canova, F: Does detrending matter for the determination of the reference cycle and the selection of turning points? The Economic Journal, vol. 109, no. 452, s. 126-150, leden 1999;
3. Hodrick, R.J. – Prescott, E.C: Postwar U.S. business cycles: An Empirical Investigation. Journal of Money, Credit and Banking, vol. 29, s. 1-16, February 1997;
4. Taylor, J.B: Macroeconomic Policy in a World Economy: from econometric design to practical operation. W.W.Norton&Company, New York. ISBN: 0-393-96316-0.
5. Walsh, C.E: Monetary Theory and Policy. Second Edition. London: The MIT Press, 2003. ISBN 0-262-23231-6;
6. Woodford, M: The Taylor Rule and Optimal Monetary Policy, American Economic Review: Papers and Proceedings 91, 232—237.

Adresa autora:

Svatopluk Kapounek, Ph.D., Ing.

Ústav financí, PEF MU v Brně

Zemědělská 1

Česká republika

613 00 Brno

skapounek@mendelu.cz

Vybrané aspekty přijetí eura v souvislosti s ekonomickým růstem v České republice

Selected Aspects of Joining the Euro and Economic Growth in the Czech Republic

Svatopluk Kapounek

Abstract: Joining the euro is often assumed to be structural change in the economy. This change has positive significant potential to affect total factor productivity, international trade and finally income of the inhabitants. However, there is no lab to simulate and estimate benefits and costs of the euro adoption. The author focuses on the selected aspects /factors/ which contribute to the economic activity in long-run. In the second part of the paper are estimated business cycles of selected members of the Eurozone.

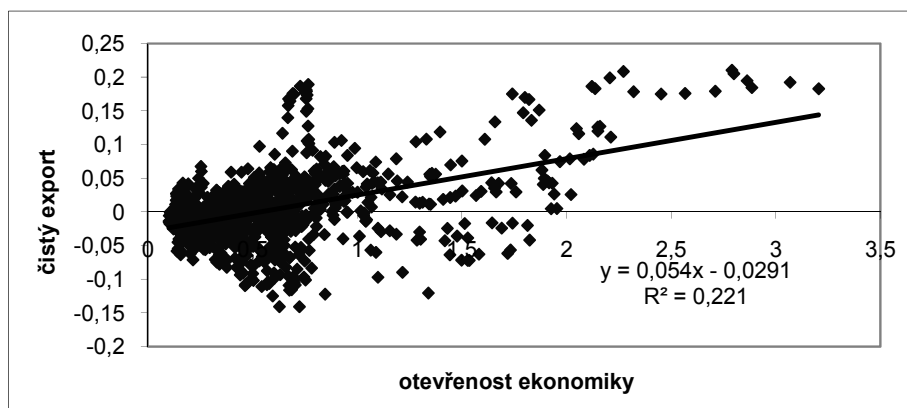
Key words: openness of the economy, inhabitants' income, total factor productivity, investments

Klíčová slova: otevřenost ekonomiky, příjmy obyvatel, produktivita výrobních faktorů, investice

1. Úvod

V souvislosti s plánovaným zavedením společné měny euro nejčastěji diskutovanými tématy dopady na inflaci, zhodnocení české koruny, finanční trhy, obyvatelstvo, mezinárodní obchod či daňový systém. Veškeré vlivy zavedení eura se ovšem na konci slučují do jednoho jediného ukazatele a tím je dlouhodobý ekonomický růst odrážející životní standard obyvatel. Přijetí společné měny euro lze přitom považovat za změnu ve struktuře ekonomiky, která významným způsobem ovlivní produktivitu jejích výrobních faktorů. Za konečný cíl i motiv zavedení eura v České republice tak lze považovat zvýšení dlouhodobého ekonomického růstu a tím i příjmu na obyvatele. V porovnání s ostatními státy Evropy patří Česká republika mezi země s nižším středním příjmem na obyvatele.

Vlastní projekt měnové unie, a tedy i zavedení společné měny, je pokračováním procesu evropské integrace, procesu prohlubování mezinárodní spolupráce a odbourávání bariér zahraničního obchodu. Zvyšování otevřenosti ekonomiky je přitom výhodnější pro malé země oproti zemím větším. Podle teorie reciproční poptávky má vývoz z malých ekonomik, ke kterým patří také Česká republika, mnohem větší šanci najít své spotřebitele v zahraničí, než je tomu u ekonomik velkých. Důvodem je mnohem větší počet potenciálních kupujících, a tedy i vyšší poptávka, kterou nabízí zahraniční trh oproti trhu domácímu. Díky vyšší poptávce v zahraničí je možné, aby malá ekonomika získala z uskutečnění obchodních transakcí ze směny mnohem vyšší hodnotu/cenu, než by dosáhla za nezměněných podmínek na domácím trhu. Čím větší je přitom potenciální poptávka, tím vyšší cenu je možné za tuzemské zboží na zahraničních trzích získat. Výnosy se budou zvyšovat s mírou zapojení do zahraničněobchodních vztahů, a tedy i s mírou odstraňování bariér zahraničního obchodu, ke kterým zavedení společné měny euro patří. Dle tohoto předpokladu by tuzemští spotřebitelé měli využívat mnohem nižších cen při spotřebě dováženého zboží a tuzemští výrobci nižších nákladů při spotřebě dovážených surovin a meziproduktů. Zvýšením výnosů z vývozu a snížením nákladů na dovoz roste čistý vývoz (rozdíl mezi vývozem a dovozem), a tedy i výstup ekonomiky, aniž by docházelo ke zvyšování zahraničního zadlužení. Obecnou platnost uvedené myšlenky dokazuje prostřednictvím vztahu mezi otevřeností ekonomiky a čistým vývozem Obrázek 1 (předpokládá se, že počet malých ekonomik je vyšší než ekonomik velkých).



Obrázek 1: Vztah mezi otevřeností ekonomiky a čistým vývozem

Pozn: Na výběrovém souboru ročních dat 30 států OECD v období let 1970-2005 byla vypočtena otevřenost ekonomiky jako součet absolutních hodnot importu a exportu ku HDP.

2. Vybrané faktory

Bohužel neexistuje „laboratoř“, ve které by bylo možné dopady zavedení eura simulovat a veškeré důsledky následně vyhodnotit a kvantifikovat. Veškeré analýzy dopadů vstupu do měnové unie jsou tak založeny na logických úvahách, domněnkách a předpokladech anebo na zkušenostech ostatních zemí, stávajících členských států eurozóny.

Mezi nevyvratitelné dlouhodobé přínosy patří zejména snížení transakčních nákladů, zlepšení transparentnosti cen a minimalizace kurzového rizika. Zvýšení konkurenčních tlaků v souvislosti s vyšším zapojením naší ekonomiky do mezinárodněobchodních vztahů by mělo vést k efektivnějšímu využívání zdrojů a fungování tržních mechanismů, tedy k růstu produktivity a investic. Nejčastěji zmiňovaným dlouhodobým nákladem je pak ztráta autonomie v oblasti měnové politiky.

V souvislosti se zavedením společné měny se jednoznačně očekává zvýšení produktivity výrobních faktorů. Oproti předpokladům se ale ukazatel růstu souhrnné produktivity výrobních faktorů v zemích tvořících eurozónu letech 2000-2004 nezvýšil. Přitom např. ve Velké Británii, která euro nezavedla, došlo k mnohem vyššímu růstu produktivity než v zemích, které euro zavedly (Tabulka 1).

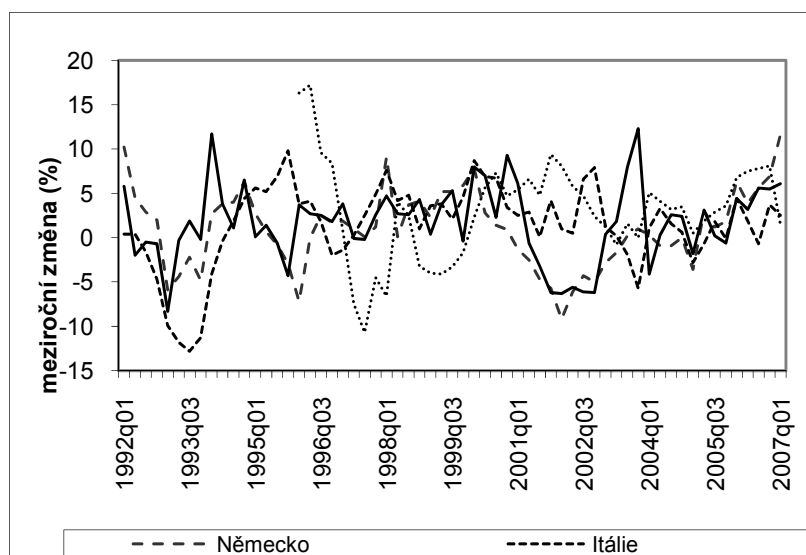
Mezi další důležitý faktor ovlivňující ekonomický růst a zároveň ukazatel, na který má přijetí eura významný vliv jsou investice. Ty tvoří v České republice téměř 30% tvorby HDP. Porovnáme-li meziroční změnu hrubé tvorby fixního kapitálu ve vybraných státech eurozóny, v krátkém období lze hovořit o krátkodobém poklesu růstu tvorby investic v Německu a Rakousku. Po roce 2004 ale dochází k jejich opětovnému nárůstu. Výraznou změnu, kterou bychom mohli spojit s vývojem dlouhodobého ekonomického růstu ale pozorovat nelze (Obrázek 2).

Určité nebezpečí lze ovšem spatřit v asymetrickém vývoji investičních cyklů mezi vybranými státy eurozóny a Českou republikou (zejména období let 1996-2003). Pokud by docházelo v České republice po zavedení eura k odlišnému (asymetrickému) vývoji investičního cyklu a tedy i ekonomického růstu, lze očekávat také odlišný vývoj inflace. Stabilizační funkce měnové politiky by tak (v případě odlišného vývoje inflace a hospodářského cyklu v České republice ve srovnání s eurozónou) byla výrazně omezena ztrátou její autonomie. Evropská centrální banka totiž nemá prostřednictvím společné měnové politiky nástroj, kterým by mohla působit v jednotlivých částech eurozóny různým způsobem.

	1996-2004	1996-1999	2000-2004
EU-15	0,7	0,9	0,5
Belgie	0,8	0,9	0,8
Česká rep.	1,5	0,4	2,3
Dánsko	1	1,2	0,9
Finsko	2,3	2,9	1,9
Francie	0,9	1,3	0,6
Irsko	3,1	4	2,4
Itálie	0,2	0,6	-0,1
Lucembursko	0,7	2,1	-0,3
Německo	0,4	0,4	0,4
Nizozemí	0,8	1,2	0,4
Portugalsko	0,2	1,2	-0,5
Rakousko	0,7	1,1	0,3
Řecko	1,9	1,5	2,3
Španělsko	0,9	1,3	0,1
Švédsko	1,8	2,2	1,5
Vel.Británie	1,2	1,2	1,2

Tabulka 1: Souhrnná produktivita výrobních faktorů

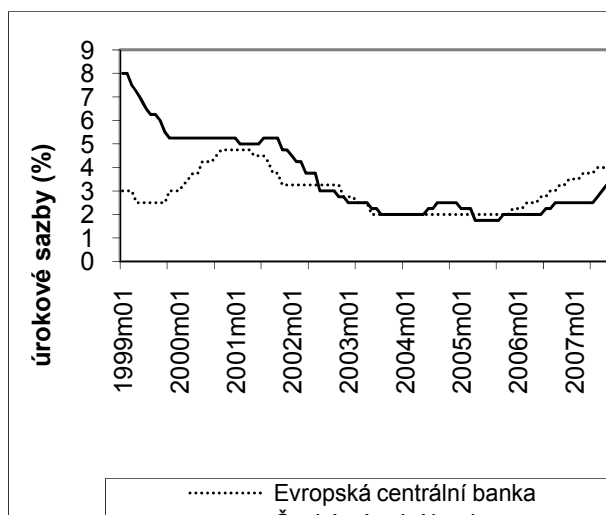
Pramen: Kadeřábková a kol., 2006, s. 35.



Obrázek 1: Hrubá tvorba fixního kapitálu ve vybraných členských státech eurozóny a v ČR

Pramen: Eurostat

Charakter autonomní měnové politiky České národní banky je ale již dnes díky vysoké míře otevřenosti české ekonomiky do značné míry ovlivněn charakterem měnové politiky Evropské centrální banky. To potvrzuje také vývoj úrokových sazeb z hlavních operací obou centrálních bank (Obrázek 3).



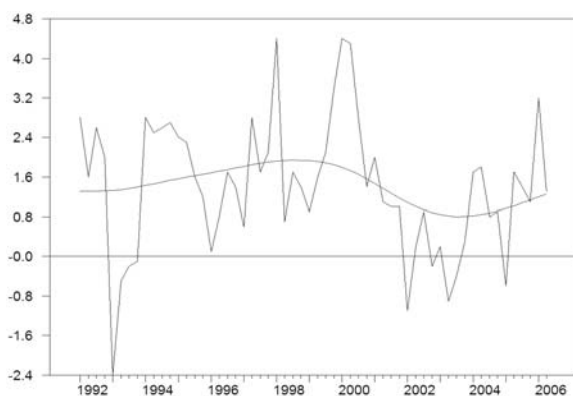
Obrázek 2: Úrokové sazby refinančních operací
Pramen: Eurostat.



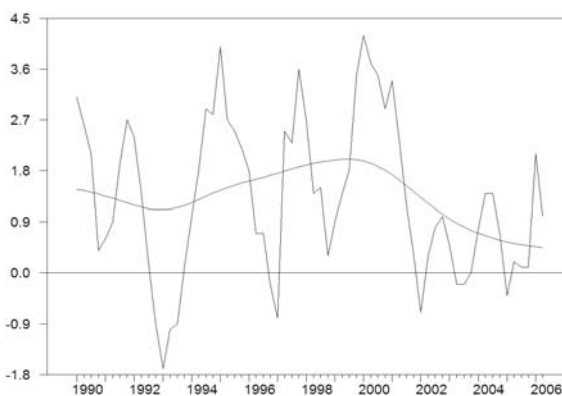
Obrázek 4: Ekonomický růst v Belgii
Pozn. HP filtr, $\lambda=1600$

3. Ekonomický růst po zavedení eura ve vybraných státech Eurozóny

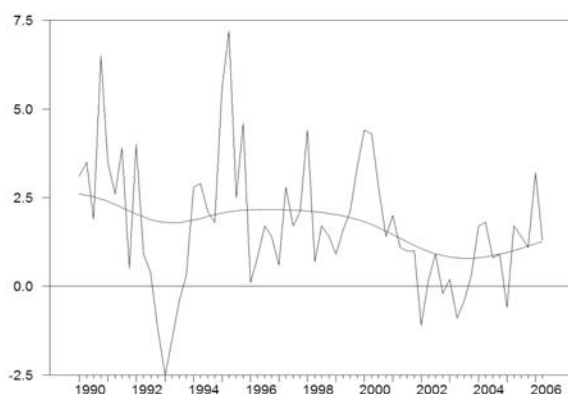
Při analýze dopadu začlenění České republiky do eurozóny na ekonomický růst prostřednictvím zkušeností ostatních zemí je nutno upozornit zásadní problém. Každá ekonomika je odlišná ve struktuře i efektivitě jednotlivých trhů, v produktivitě výrobních faktorů a jejich objemu a tedy i v základních komponentech, zdrojích ekonomického růstu, stejně tak v charakteru hospodářské politiky a efektivitě jejích nástrojů. Zmíněné skutečnosti jsou dále prohloubeny odlišnými podmínkami v době začlenění dané země do eurozóny. Na základě zkušeností ostatních zemí se zavedením eura a dopadu na ekonomický růst se tedy můžeme pouze domnívat, jaký by byl možný dopad na ekonomický růst v ČR. Skutečnost může být zcela jiná.



Obrázek 5: Ekonomický růst v Německu
Pozn. HP filtr, $\lambda=1600$



Obrázek 6: Ekonomický růst v Itálii
Pozn. HP filtr, $\lambda=1600$



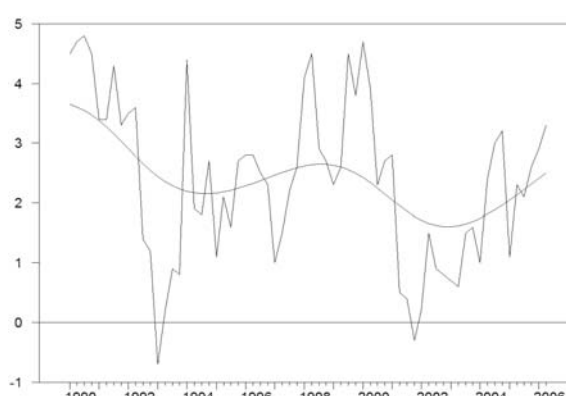
Obrázek 7: Ekonomický růst ve Španělsku

Pozn. HP filtr, $\lambda=1600$ 

Obrázek 8: Ekonomický růst ve Francii

Pozn. HP filtr, $\lambda=1600$ 

Obrázek 9: Ekonomický růst v Nizozemí

Pozn. HP filtr, $\lambda=1600$ 

Obrázek 12: Ekonomický růst v Rakousku

Pozn. HP filtr, $\lambda=1600$

Jak je vidět z obrázků 4 až 13, krátce po zavedení společné měny se objevil výraznější propad v ekonomické aktivitě. Je ovšem diskutabilní, zda uvedený vývoj byl způsoben samotným zavedením společné měny euro či jinými, externími vlivy. Zavedení společné měny by mohlo mít krátkodobý negativní vliv v ekonomické aktivitě díky jednorázovým administrativním nákladům. V dlouhém období ovšem k výrazné změně v ekonomickém růstu u analyzovaných států nedošlo.

4. Závěr

Přijetí společné měny euro lze označit za strukturální zlom, který by měl mít pozitivní vliv na efektivitu využití výrobních zdrojů, jejich produktivitu i přerozdělování příjmů z těchto zdrojů. Zvýšení otevřenosti ekonomiky (reprezentované přijetím společné měny euro, odstraněním kurzového rizika a snížením transakčních nákladů) přispívá k růstu čistého vývozu. K růstu čistého vývozu také může významným způsobem přispět zafixování české koruny na euro, a tím zabránění jejímu dalšímu nominálnímu zhodnocení. Tento předpoklad je podmíněn zachováním cenové stability, což je primárním cílem Evropské centrální banky, která jako silná a důvěryhodná měnová instituce má ke splnění tohoto cíle reálné předpoklady.

Z pohledu zkušeností stávajících členských států eurozóny není ovšem patrné zvýšení růstu produktivity výrobních faktorů, ani zásadní pozitivní zlom ve vývoji dlouhodobého ekonomického růstu. Naopak, výrazná je krátkodobá korekce ekonomického růstu, kterou

ovšem nelze jednoznačně přičítat zavedení společné měny a která byla v řádu několika málo let anulována dalším vývojem.

Závěrem lze shrnout, že přijetí společné měny euro nebude mít pravděpodobně negativní vliv na přizpůsobení reálných příjmů obyvatel České republiky průměrným příjmům členských států eurozóny. Naopak, začlenění České republiky do eurozóny by mělo přispět k dlouhodobému ekonomickému růstu.

Předkládaný příspěvek vznikl za podpory výzkumného záměru „Česká ekonomika v procesech integrace a globalizace a vývoj agrárního sektoru a sektoru služeb v nových podmínkách evropského integrovaného trhu“ MSM 6215648904.

5. Literatura

1. BALDWIN, R., DININO, V. ET AL: Study on the Impact of the Euro on Trade and Foreign Direct Investment. In European Economy [online]. May 2008, Economic Papers 321 [cit 2008-05-30]. Dostupné z: < http://ec.europa.eu/economy_finance/publications/>. ISBN 978-92-79-08246-7;
2. EVROPSKÁ CENTRÁLNÍ BANKA (2002): Recent findings on monetary policy transmission. ECB's Monthly Bulletin, October 2002, ISSN: 1561-0136;
3. FRANKEL, J. ROSE A. K: An Estimate of the Effect of Common Currencies on Trade and Income [online]. 2002. [cit 2008-05-30]. Dostupné z <<http://faculty.haas.berkeley.edu/arose/RecRes.htm>>;
4. HOLMAN, R. a kol: *Dějiny ekonomického myšlení. 2.vydání.* Praha: C.H.Beck, s.105-116. ISBN: 80-7179-631-X;
5. KADERÁBKOVÁ a kol: Ročenka konkurenceschopnosti České republiky. Praha: Linde. 2007, ISBN: 80-86131-77-7;
6. MACH, M: Makroekonomie II. Praha: Melandrium, 2001. s. 317 – 367. ISBN: 80-86175-18-9;
7. NÁRODNÁ BANKA SLOVENSKA: Odhad možných vplyvov zavedenia eura na podnikatelsky sektor v SR [online]. 2006 [2009-06-30]. Dostupné z <http://www.nbs.sk/_img/Documents/PUBLIK/06_kol2.pdf>
8. NÁRODNÁ BANKA SLOVENSKA: Odhad možných vplyvov zavedenia eura na slovenské obyvateľstvo [online]. 2006 [2009-06-30]. Dostupné z <http://www.nbs.sk/_img/Documents/PUBLIK/06_kol3.pdf>

Adresa autora:

Svatopluk Kapounek, Ph.D., Ing.
Ústav financí, PEF MU v Brně
Zemědělská 1
Česká republika
613 00 Brno
skapounek@mendelu.cz

Indikátory klasického cyklu pre priemysel SR Industry Classical Cycle Indicators of Slovakia

Miroslav Kľúčik

Abstract: Recent global economic crisis has stressed the importance of business cycle analysis among the broad tools available for researchers in the area of applied economic statistics. Slovakia has a small and extremely open economy and thus depends very much on development of foreign trade partners. The database is therefore composed of domestic and also foreign data. This papers approaches the analysis from the classical cycles point of view and tries to describe the methodology for construction of lagging, coincident and leading indicators for the reference cycle of industrial production. The final composite indicators are presented in three basic versions: simple average, binary composite variable and as a first component of principal component analysis.

Key words: business cycles, turning point analysis, composite indicators, leading indicator

Kľúčové slová: klasický cyklus, analýza bodov obratu, kompozitné indikátory, predstihový indikátor

1. Úvod

Súčasný globálny hospodársky vývoj je charakterizovaný dynamickými zmenami, na ktoré je z pohľadu tvorcov hospodárskej politiky potrebné dostatočne rýchlo a kvalifikovane reagovať. Úlohou výskumu v oblasti makroekonómie a štatistiky je pomocou matematicko-štatistických metód využívať potenciál údajov získaných prostredníctvom štatistických zisťovaní na odhalenie ekonomických javov a vzťahov medzi nimi. Samostatnou oblasťou je analýza časových radov z hľadiska cyklov (angl. business cycle analysis). Makroekonomické časové rady je možné analyzovať z hľadiska ekonomických cyklov, a tým získať poznatky o fázach vývoja hospodárstva (recesia, expanzia). Identifikácia cyklov v ekonomických časových radoch vyžaduje analýzu časových radov z hľadiska sezónnosti, vývoja trendu, identifikácie náhodnej zložky, metód na vyrovňovanie časových radov, korelačnú analýzu atď. Tento príspevok v skratke vysvetľuje tvorbu kompozitných cyklických indikátorov vývoja priemyslu SR pomocou verejne dostupných údajov. Cieľom je získať indikátory, ktoré dokážu s určitým časovým predstihom signalizovať body obratu ekonomického cyklu priemyslu (predstihové indikátory) a indikátory, ktoré potvrdia s určitým časovým odstupom existenciu fáz ekonomického cyklu (koincidenčné a zaostávajúce indikátory).

2. Identifikácia klasického cyklu

Základom každej analýzy časových radov pri makroekonomickej analýze je vytvorená databáza vstupných dát. V našom prípade pôjde výhradne o mesačné časové rady makroekonomických ukazovateľov získaných z verejne dostupných informačných zdrojov. Mesačné údaje sú preferované pred štvrťročnými hlavne z dôvodu dĺžky časových radov a flexibilnejších informácií o vývoji ekonomiky oproti štvrťročným dátam. Keďže cieľom je maximálne využiť informačný potenciál potrebný pre kvalitnú analýzu, k analýze sa využijú zdroje domáce i zahraničné makroekonomické ukazovatele, ktoré bolo nutné zahrnúť do databázy vzhľadom na vysokú otvorenosť ekonomiky SR [1]. Finálna databáza obsahuje informácie kvantitatívneho (produkcia, tržby, zamestnanosť, mzdy, odpracované hodiny v rôznych odvetviach) i kvalitatívneho charakteru (konjunkturálne prieskumy v odvetviach). Časové rady boli do databázy získané z domácich zdrojov: Štatistický úrad SR (produkcia, tržby, zamestnanosť...), Ministerstvo financií SR (štátny rozpočet, výber daní), Burza cenných

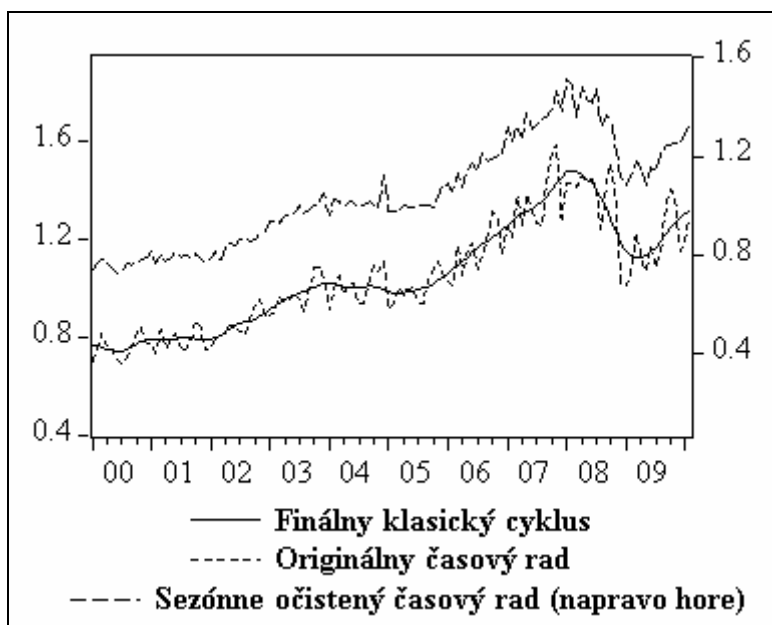
papierov v Bratislave (index SAX, burzové obchody), Národná banka Slovenska (úvery, vklady, úroky, peňažné agregáty). Zo zahraničných zdrojov je to Eurostat a OECD. Spolu tvorí obsah databázy v súčasnosti 1532 časových radov. K dispozícii sú mesačné dáta s rôznym začiatočným obdobím, pričom minimálny začiatok pre dôkladnú analýzu bol stanovený na január 2000, posledná aktualizácia dát bola vykonaná v apríli 2010.

Kľúčové z hľadiska metodológie extrakcie cyklov z časových radov je rozhodnutie o druhu skúmaného cyklu. Rozlišujú sa dva hlavné druhy cyklov, a to klasický a deviačný cyklus. Klasický cyklus predstavuje zmeny vo vývoji sledovaného ukazovateľa v úrovni (level), zatiaľ čo deviačný cyklus predstavuje fluktuáciu ukazovateľa okolo vlastného dlhodobého trendu vývoja. Klasické cykly sú stredobodom pozornosti v obdobiach klasických fáz ekonomického (hospodárskeho) cyklu (recesia, dno, expanzia, vrchol). Pozornosť na deviačné cykly sa sústreďovala hlavne v obdobiach dlhodobého rastu vyspelých západných ekonomík (USA) od osemdesiatych rokov 20. storočia. Témou tohto príspevku bude práve klasický cyklus, vzhľadom na nedávnu recesiю klasického cyklu slovenskej ekonomiky.

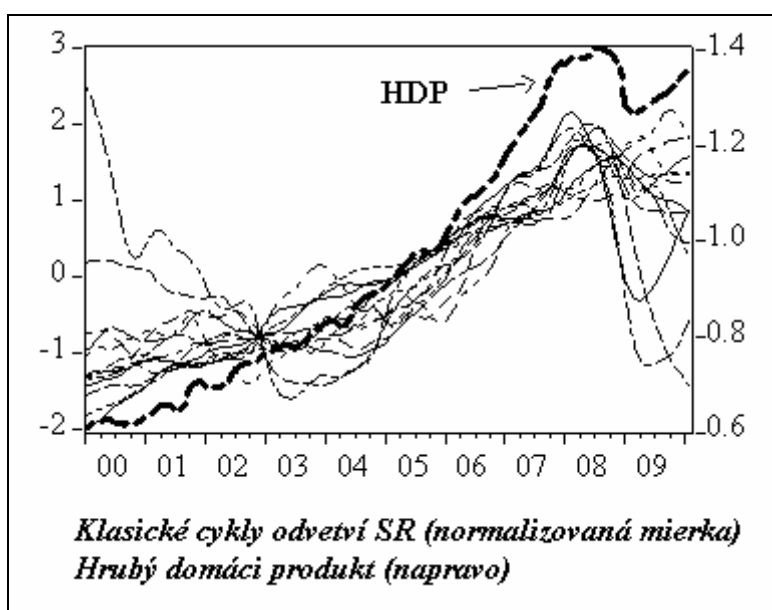
Druhou kľúčovou otázkou pre výskum klasického cyklu je výber metódy na extrakciu cyklu z časových radov. Tieto metódy sa dajú rozdeliť opäť do dvoch hlavných skupín: prvé sa sústreďujú na analýzu časovej domény časových radov, druhé na frekvenčnú doménu ukazovateľov. Z oblasti frekvenčnej domény možno spomenúť spektrálnu analýzu, z oblasti časovej domény ide predovšetkým o modely kľzavého priemeru a pásmové filtre ako Hodrick Prescott filter, Baxter-King filter alebo Christiano-Fitzgerald filter. Pásmové filtre sa využívajú pri deviačných cykloch na identifikáciu trendovej a cyklickej zložky. V tomto príspevku sa pozornosť venuje klasickým cyklom, teda pohybom veličiny v úrovni, takže k identifikácii cyklu je potrebné odstrániť len sezónnu a náhodnú zložku časových radov. Tento výsledný cyklus je zhodný s trendovo-cyklickou zložkou časového radu (angl. trend-cycle component).

Sezónnu zložku možno odstrániť transformáciou originálneho časového radu na časový rad medziročných zmien alebo použiť sezónne filtre. Pri medziročných zmenách by sme stratili ďalší rok, čím by sa skrátili časové rady a poklesla kvalita analýzy, nakoľko i so začiatkom od roku 2000 ide pri klasickej analýze o krátky časový rad. Všetky časové rady v databáze boli očistené od sezónnej zložky pomocou nástroja Tramo\Seats, ktorý je zakomponovaný v programovom balíku softvéru EViews. Prostredníctvom tohto nástroja sme automaticky identifikovali i náhodné výkyvy (angl. outliers) pre všetky časové rady. Výsledný sezónne očistený časový rad možno následne zbaviť i náhodnej zložky.

Ako najvhodnejší nástroj pre tento účel sa ukázal byť Hendersonov kľzavý priemer (Henderson moving average – HMA) [2]. Tento nástroj na jednej strane zabezpečuje odstránenie iregulárnej zložky časových radov a teda vyrovnanie časového radu, a na druhej strane i zachovanie bodov obratu trendovo-cyklickej zložky časového radu. Na Obrázku 1 je zobrazený originálny časový rad indexu priemyselnej produkcie (sezónne neočistený) a finálny klasický cyklus priemyslu po sezónnom očistení pomocou Tramo\Seats a vyrovnaní časového radu pomocou HMA metódy. Ďalší obrázok zobrazuje vývoj hlavných ukazovateľov (vo forme klasického cyklu) najväčších odvetví Slovenska (Obrázok 2). Na tomto obrázku je znázornený ilustračne všeobecný trend viacerých agregátnych ukazovateľov slovenskej ekonomiky, aby bolo umožnené jeho porovnanie s trendom vývoja hrubého domáceho produktu (HDP). Na obrázku sú zobrazené krivky klasického cyklu časových radov tržieb v priemysle, stavebníctve, maloobchode, veľkoobchode, informáciách a komunikácii, doprave, ostatných službách, peňažného agregátu M1, bilancie zahraničného obchodu, zamestnanosti vo vybraných odvetviach (10 odvetví), indexu spotrebiteľských a výrobných cien. Časový rad HDP je transformovaný do mesačnej frekvencie pomocou kvadratickeho priemeru, angl. quadratic match average.



Obrázok 1: Klasický cyklus priemyslu (produkcia)



**Obrázok 2: Vývoj agregátov slovenskej ekonomiky
(klasické cykly, normalizované časové rady)**

3. Indikátory klasického cyklu

V tejto etape výpočtu sú k dispozícii všetky časové rady v databáze vo forme klasického cyklu. Teória podľa NBER (Národný úrad pre ekonomický výskum) z prvej polovice 20. storočia [5] predpokladá existenciu predstihových, koincidenčných a zaostávajúcich indikátorov. Indikátory z hľadiska cyklu teda možno zaradiť do jednej z uvedených tried podľa vývoja ich cyklu v porovnaní s referenčným ukazovateľom. Predstihový indikátor (CLI – angl. composite leading indicator) nesie informáciu o pravdepodobnom vývoji referenčného ukazovateľa s určitým časovým predstihom, spravidla niekoľko mesiacov, koincidenčný ukazovateľ vykazuje body obratu súčasne s referenčným ukazovateľom a zaostávajúcí indikátor je nositeľom tejto informácie až následne s oneskorením niekoľkých mesiacov. Referenčný ukazovateľ je indikátor, ktorý je objektom pozorovania, teda ten, s ktorým sa

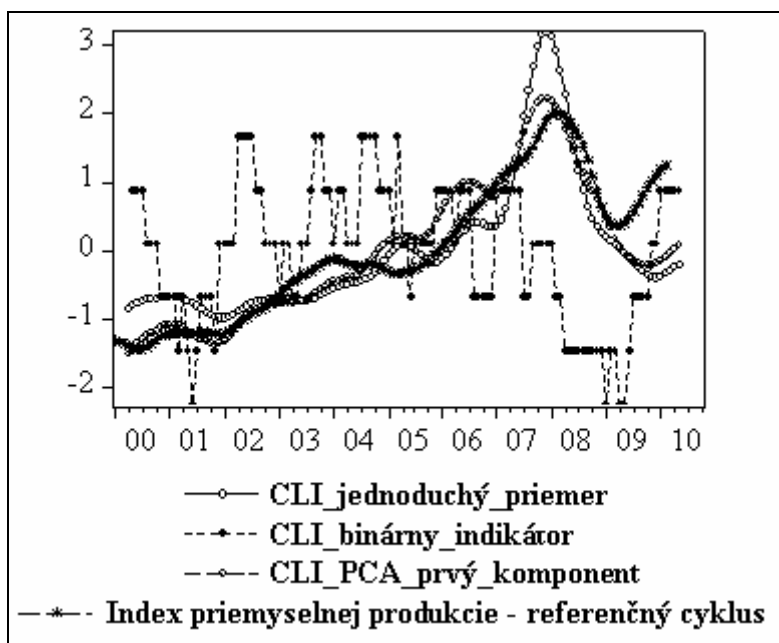
porovnávajú všetky ostatné časové rady. V prípade, že objektom sledovania je vývoj celej ekonomiky, referenčným časovým radom je napr. HDP, v prípade zahraničného obchodu je to jeho bilancia atď. Tento príspevok sa zaoberá iba odvetvím priemyslu, nakoľko je najväčším odvetvím hospodárstva SR (podiel priemyselnej produkcie v stálych cenách sa dlhodobo pohybuje na úrovni vyše 30% z tvorby HDP). Referenčným časovým radom môže byť teda klasický cyklus indexu priemyselnej produkcie alebo tržieb v priemysle. Body obratu (body prechodu medzi cyklami expanzie a recesie) klasických cyklov priemyselných tržieb a priemyselnej produkcie sú takmer identické. Za referenčný časový rad je zvolený nakoniec index priemyselnej produkcie.

Zaradenie časových radov do jednotlivých tried cyklov je možné pomocou analýzy krížovej korelácie, pri ktorej sa otestujú všetky časové posuny časových radov a porovnajú sa na základe výšky korelačného koeficientu daného časového radu s referenčným časovým radom. Za koincidenčný časový rad sa považuje časový rad s posunom ± 2 mesiace oproti referenčnému cyklu. Podľa uvedeného základného kritéria sa každý z časových radov stal členom jednej z troch skupín – predstihové indikátory, koincidenčné indikátory a zaostávajúce indikátory. Z pôvodnej databázy 1532 časových radov bolo identifikovaných 336 predstihových a 327 zaostávajúcich indikátorov, zvyšných 869 časových radov tvorili koincidenčné indikátory.

Z uvedených skupín je možné zložiť po jednom zloženom indikátore z každej skupiny. Keďže každá zo skupín obsahuje aj menej kvalitné indikátory (napr. predstihový indikátor s nízkym korelačným koeficientom), na výber najkvalitnejších indikátorov sú zvolené ďalšie kritéria. Prvým kritériom je výška korelačného koeficienta, ktorá je stanovená na 0.65. Ako ďalšie efektívne kritérium sa javí kritérium zhody bodov obratu jednotlivých indikátorov s referenčným cyklom indexu priemyselnej produkcie. Pre toto kritérium je zadaná potreba zhody minimálne dvoch bodov obratu s referenčným časovým radom pri tolerancii ± 3 mesiace (počet relevantných bodov obratov v referenčnom cykle je 5). Pri predstihových indikátoroch zostalo po zohľadnení spomínaných dvoch kritérií z pôvodného počtu 336 indikátorov už len 18 indikátorov. Z pôvodného počtu 869 ostalo nakoniec 143 koincidenčných indikátorov a rovnako sa znížil i počet zaostávajúcich indikátorov z 327 na 30. Posledným kritériom výberu je subjektívne kritérium, ktoré zohľadňuje ekonomickú interpretovateľnosť odhaleného vzťahu medzi skúmaným časovým radom a referenčným časovým radom. Po zohľadnení posledného kritéria ostalo 5 predstihových, 120 koincidenčných a 7 zaostávajúcich indikátorov.

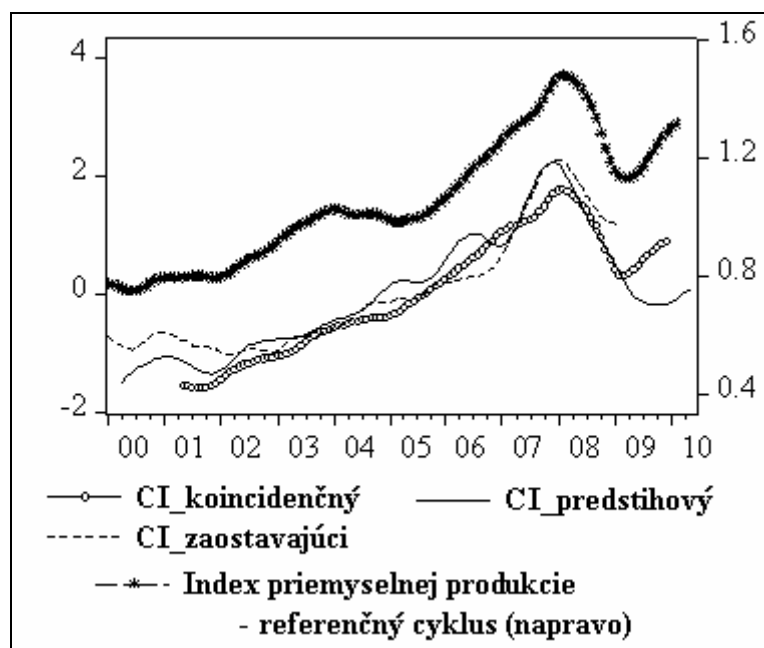
Posledným kľúčovým rozhodnutím pri zostavovaní indikátorov ekonomického cyklu je voľba metódy na konštrukciu kompozitného indikátora z konečného počtu najkvalitnejších indikátorov z každej z troch skupín. Najjednoduchšou metódou je jednoduchý priemer zvolených časových radov (všetky originálne časové rady sú v stálych cenách a vo forme bázičných indexov s rovnakou bázou), druhou metódou je transformácia cyklov na sumárne rastové binárne indexy (0 = medzimesačný pokles, 1 = medzimesačný rast) a nakoniec aplikáciou metódy hlavných komponentov – PCA (principal component analysis) – extrakcia hlavného komponentu z viacerých časových radov. Všetky tri metódy môžeme porovnať na základe Obrázku 3. Okrem týchto metód existujú i omnoho sofistikovanejšie metódy [4]. Z uvedených metód nie je diskvalifikovaná žiadna, naopak všetky tri metódy sa dajú využívať súčasne s porovnávaním výsledkov¹.

¹ Keďže všetky výpočty zostavovania indikátorov ekonomického cyklu uskutočňujeme pomocou automatizovaného prístupu programovacieho jazyka aplikácie EViews, najviac času zaberie dopĺňanie základných databáz a subjektívna voľba ekonomicky interpretovateľných indikátorov a nie samotný výpočet zložených indikátorov rôznymi metódami.



Obrázok 3: Porovnanie metód na vytvorenie kompozitného indikátora

Zložené indikátory sú aktualizované každý mesiac s tým, že sa menia ich hodnoty i späťne. Týmto sa zabezpečí maximálna flexibilita pri analýze vývoja ekonomických ukazovateľov. Na základe Obrázku 4 je možné porovnať predstihový, koincidenčný a zaostávajúci kompozitný index (CI – composite index) s referenčným cyklom indexu priemyselnej produkcie pomocou tretej metódy – PCA. Predstihový indikátor môže slúžiť aj ako nástroj na prognózovanie priemyselnej výroby na minimálne 3 mesiace dopredu, čo sa dá následne využiť i pri rýchlych odhadoch makroekonomických ukazovateľov v referenčnom štvrtroku [3].



Obrázok 4: Porovnanie výsledných indikátorov ekonomického cyklu priemyslu SR

4. Záver

Predstihový indikátor ako významný krátkodobý nástroj na makroekonomické analýzy a prognózy slúži ako zdroj významných informácií pre výskumníkov i predstaviteľov podnikateľského prostredia pri odhade očakávaného vývoja ekonomiky a jej častí. Ako podpora pre cyklickú analýzu sú využiteľné i ďalšie zložené indikátory – koincidenčný indikátor, ktorý potvrdzuje vývoj sledovaného referenčného ukazovateľa a nakoniec zaostávajúci indikátor, pomocou ktorého je možné spätne identifikovať existenciu vzťahov medzi rôznymi makroekonomickými ukazovateľmi.

5. Literatúra

- [1]KLÚČIK, M. 2009. Composite Leading Indicators for the Slovak Economy. Open access document: https://www.ciret.org/news/workshop_budapest/papers/Klucik.pdf. Paper presented at the CIRET/KOF/GKI International Workshop on Sentiment Indicators and the Current Crisis, Budapest, Hungary, 9. – 10. November 2009.
- [2]KLÚČIK, M. 2009. Composite Reference Series and Composite Leading Indicator for Slovakia. Open access document: <http://www.isae.it/MFC/2009/klucik.pdf>. Paper presented at the IFO/INSEE and ISAE International Conference „The First Macroeconomic Forecasting Conference - MFC 2009“, Rome, Italy, 27. March 2009.
- [3]KLÚČIK, M., JURIOVÁ J. 2009. Slowdown or Recession? Forecasts Based on Composite Leading Indicator. Paper presented at the 36th International Conference MACROMODELS 2009, Bochnia, Poland, 2. – 5. December 2009
- [4]MARCELLINO, M. 2006. Leading Indicators. In: HANDBOOK OF ECONOMIC FORECASTING, Volume 1. Amsterdam: Elsevier, 2006. 880-952 s. ISBN 978-0-444-51395-3.
- [5]ZARNOWITZ, V., OZYILDIRIM A. 2002. Time Series Decomposition and Measurement of Business Cycles, Trends and Growth Cycles. Open access document: <http://www.fcsn.gov/03papers/Zarnowitz.pdf>. NBER Working Paper No. W8736.

Adresa autora:

Miroslav Klúčik, Ing.
Dúbravská cesta 3
845 24 Bratislava
klucik@infostat.sk

Výuka statistiky podporovaná excelovskou aplikací Learning statistics endorsed by Excel

Oldřich Kříž, Jiří Neubauer, Marek Sedlačík

Abstract: The contribution is focused on a specially prepared aid made in Excel environment (including English version) which was tried out in the teaching process at the University of Defence. The authors explain the utilization by the way of illustrative solution of real problem and they ponder over other possibilities how to use or modify this aid considering students enquiry results.

Key words: teaching statistics, Excel, basic data processing, normal distribution, point estimates, interval estimates, hypothesis testing.

Klíčová slova: výuka statistiky, Excel, základní zpracování dat, normální rozdělení, bodové odhady, intervalové odhady, testování hypotéz.

1. Úvod

Statistika patří na fakultách s ekonomickým zaměřením mezi stěžejní předměty, protože tvoří nedílnou součást ekonomického vzdělávání. Přestože je v ekonomickém prostředí zcela nenahraditelná, protože se využívá k získávání nejrůznějších informací, studenti ji jako předmět „nemají příliš v lásce“ a považují ji za obtížnou a velmi podobnou matematice. Lze vnímat hned dva důvody této nepříznivé situace. Statistika ve srovnání s matematikou má zcela odlišnou filozofii – užívá zcela odlišný způsob myšlení. Statistika se nezabývá výpočty veličin, ale jejich odhady, což není v časovém prostoru vymezeném základnímu kurzu statistiky vůbec jednoduché pochopit. Druhým důvodem je stavba této disciplíny – statistika je vybudovaná na třech pilířích: pravděpodobnost, náhodná veličina a popisná statistika. Pochopit, jakou roli ve statistice hrají a jaké jsou souvislosti mezi těmito součástmi navzájem a směrem k odhadům, je pro začátečníka často těžko stravitelné. Oproti tomu pozitivně působí skutečnost, že statistika řeší zejména praktické problémy ze života a odpovídá na reálné otázky. Toho však musí učitel statistiky umět využívat a neustále poukazovat na možnosti užití statistiky v reálném světě.

Autoři tohoto příspěvku se již určitou dobu zamýšlí, jak docílit toho, aby studenti lépe a přirozeněji chápali podstatu statistiky, aby pronikali do jejího skutečného fungování, a aby se v konečném důsledku nebáli ji reálně a hlavně smysluplně používat. Jednu z možných cest vidí autoři v počítačové podpoře výuky.

2. Výpočetní technika ve výuce statistiky

Výpočetní technika vstoupila do statistiky již před mnoha lety, pro určité moderní statistické metody je počítačová podpora dokonce nezbytností. Z tohoto důvodu bylo počítačové zpracování dat zařazeno do výuky statistiky také na většině vysokých škol. Na trhu je k dispozici řada softwarových produktů, jako např. Statistica, SPSS, Statgraphics, Unistat, QCExpert, Adstat, Matlab (statistics toolbox), pomocí nichž je možné statistické analýzy provádět. Uvedené produkty se liší rozsahem nabízených metod a analýz, grafickým rozhraním, uživatelskou přístupností, univerzálností apod. Co však mají všechny tyto produkty společné, to je skutečnost, že jsou komerční.

Pokud si učitel statistiky položí otázku, který z uvedeného softwaru by měl ve výuce statistiky používat, bude postavený před dva základní problémy. Předně je to dostupnost zvoleného softwaru pro školu, protože zakoupení kvalitního programu v žádoucí multilicenci určitě není záležitost levná. Stejně tak je však důležitá dostupnost tohoto programu pro

studenty, aby v klidu svého domova či koleje mohli všechno, co učitel ve cvičení ukázal, samostatně vyzkoušet a pronikat tak do tajů statistiky.

Autoři příspěvku už ve výuce vyzkoušeli jak odborné, tak i didaktické přednosti Unistatu, QCExpertu a Matlabu. První dva jmenované softwarové produkty nejsou nikterak náročné na ovládání a orientaci v nabízené struktuře metod resp. jednotlivých technik. Mají pro výuku postačující grafické rozhraní a jsou oba v českém prostředí. Pro podporu výuky statistiky v základním kurzu jsou i vhodné. Matlab je známý a rozšířený produkt s obrovsky širokým záběrem v různých odvětvích exaktních disciplín. Vyžaduje však speciální toolboxy a neobejde se bez znalosti speciálního jazyka. Vhodný je spíše pro podporu výuky v některém z nadstavbových kurzů statistiky. V čem nás však tyto produkty při výuce statistiky neuspokojili, to byla právě jejich omezená využitelnost při samostatné práci studentů v domácím prostředí. Pro studenty jsou v našich podmínkách nedostupné, a tak podpora výuky statistiky končí za dveřmi počítačové učebny.

3. Statistika v Excelu

Tento „handicap“ většiny statistických programů lze alespoň částečně řešit pomocí známého produktu Excel, který je součástí kancelářského balíku MS Office. Značnou výhodou tohoto balíku je právě jeho masová rozšířenost mezi studenty, a také počítačové učebny jsou tímto produktem vybavené. Nejedná se pochopitelně o speciální statistický software, ale tabulkový charakter Excelu umožňuje i ve výuce statistiky využít řady jeho zabudovaných prostředků a zajímavých vlastností.

Především se jedná o široké možnosti užití vlastních numerických výpočtů pomocí rovnic tvořených uživatelem. Excel zvládá i maticové výpočty. Různé numerické výstupy lze velmi jednoduchým způsobem uspořádat do tabulek, kterých statistika hojně využívá. Výhodou je zejména to, že uživatel – v našem případě student – pracuje interaktivně a uspořádá výstupy podle svých potřeb. Zanedbatelné nejsou ani možnosti grafické, Excel nabízí 14 typů grafů, každý ještě v několika modifikacích. Student se podle svých zkušeností může rozhodnout, který z grafů bude pro zobrazení statistické vlastnosti nejvhodnější.

V řadě situací, které v základním kurzu statistiky nelze obejít, je možné využít zabudované procedury z různých oblastí statistiky, označené jako analytické nástroje. Využít je možné balíčky popisné a pořadové statistiky, rozdělení četností, dvouvýběrových testů, analýzy rozptylu, regrese a korelace, a dalších. Jedná se o pevné procedury, které dávají stále stejné typy výsledků, a je možné je proto snadno komentovat. Nejširší nabídku poskytuje Excel v kategorii zabudovaných statistických funkcí. V rámci kurzu statistiky je vhodné je zařazovat postupně a s jistým didaktickým záměrem využívat přednosti práce v excelovském prostředí. To umožňuje studentovi kombinovat různé funkce, vyzkoušet si chování různých dat či modelů a pronikat tak do statistické filozofie.

Nespornou výhodou Excelu je jeho schopnost provést automaticky přepočítání všech definovaných funkcí při změně vstupních hodnot. Tato vlastnost může být při výuce názorně využita, např. se snadno ukáže, co se stane s hodnotou aritmetického průměru a mediánu přidáme-li k naměřeným datům nějakou odlehlou hodnotu apod.

V konkrétních situacích je však třeba studenty upozornit také na jisté neduhy Excelu. Ti, kteří se již se statistickými funkcemi v Excelu setkali, si jistě povšimli poněkud nejasné a místy i zavádějící terminologie v popisu funkcí i v nápovědě k jednotlivým funkcím. Mimo těchto nejasností, nad kterými by mohl mnohý uživatel mávnout rukou – z didaktického hlediska to však není věc vůbec zanedbatelná – se zde objevují i jiné nešvary. Jedná se zejména o nejednotnost při určování hodnot inverzních funkcí hustot některých pravděpodobnostních rozdělení. Tato nejednotnost v zadávání parametrů excelovských funkcí je matoucí a může způsobit studentům jisté problémy, někdy i chyby ve výpočtech. Podobných nepříjemností by bylo možno uvést více.

Autoři příspěvku se v minulých letech netajili ambicí sestavit v excelovském prostředí aplikaci, která by uvedené problémy eliminovala a dostatečně mohla podporovat výuku statistiky, zejména v oblasti popisu naměřených dat a praktického užití odhadů a testů, které se ve statistice váží na konkrétní věcné problémy. Základem této aplikace byl už používaný pracovní sešit [4], který byl rozšířený o datový list a některé další statistické procedury – viz dále.

4. Aplikace STAT1

Vedle využití zabudovaných excelovských nástrojů – grafů, analytických nástrojů a funkcí – lze v Excelu vytvořit prostředí podle vlastních požadavků. Autoři příspěvku připravili v prostředí Excelu aplikaci s názvem STAT1. Je určena především jako podpora základního zpracování dat v podobě popisné statistiky a dále jsou zde implementované nejjednodušší metody v rámci jednorozměrné induktivní statistiky. Aplikace je konstruovaná tak, aby s naprosto minimálními vstupy dávala snadno řadu užitečných výstupů, které již stačí „jen“ správně interpretovat. Z tohoto pohledu se chová jako profesionální programy založené na výběru procedur z menu a stanovení parametrů prováděné analýzy. Protože výuka statistiky probíhá na naší fakultě paralelně také v anglickém jazyce, používá se popsaná aplikace také v anglické verzi.

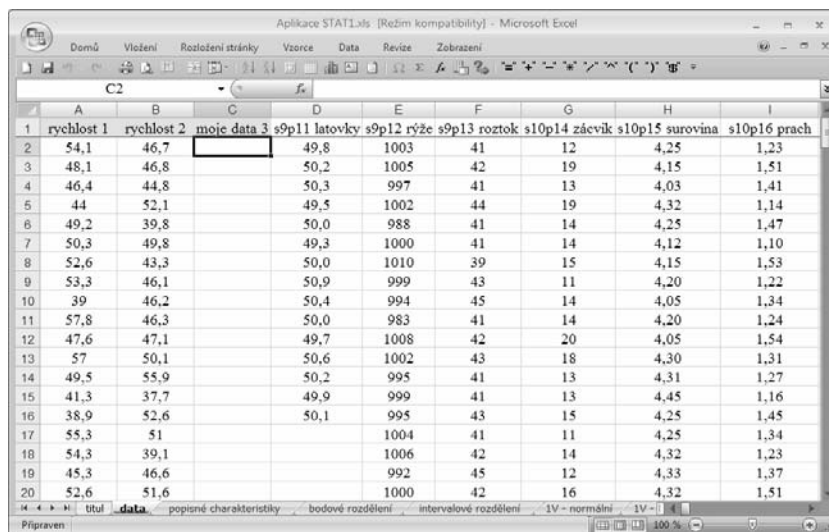
Právě komentáře a interpretace získaných výsledků považujeme z didaktického hlediska za naprosto klíčové, protože umožňují studentovi postupně budovat správnou představu o statistické filosofii. Prostor k tomu je zejména vytvořený tím, jak jednoduché je ovládání této aplikace. Student se tedy soustředí na problémy skutečně statistické a nemusí svojí pozornost věnovat samotnému ovládání programu. V základním kurzu statistiky jsou v podstatě dvě kapitoly, ve kterých je počítačová podpora výuky funkcí: popisná statistika a odhady a testy hypotéz. Obě tyto části tvoří praktickou část statistiky, jejímž východiskem jsou naměřená data.

První list aplikace STAT1 s označením *data* obsahuje datové soubory ze Sbírky úloh ze statistiky [3], která je používána na cvičeních, viz obr. 1. Použitelnost aplikace není však omezena jen na příklady z uvedené sbírky. Je také samozřejmě možné, aby si student do listu *data* vložil svoje vlastní hodnoty, které v záhlaví sloupce označí – pojmenuje. Výběr konkrétních dat pro další zpracování se provádí přímo v listu, ve kterém bude student dále pracovat. V každém výpočetním listu lze provést výběr proměnných, které se zobrazí (uvedený je odkaz na Sbírku a současně klíčové slovo charakterizující úlohu). Dále se vloží požadované parametry úlohy (jsou zvýrazněné červeně), např. hladina významnosti, velikost přípustné chyby apod. Statistické výstupy – výsledky jednotlivých analýz – jsou zobrazené v zelených polích.

Další tři listy – *popisné charakteristiky*, *bodové rozdělení* a *intervalové rozdělení*, viz obr. 2 – umožňují provést exploratorní analýzu dat, pomocí které lze posoudit důležité vlastnosti naměřených dat. Z tabulkového a grafického vyjádření rozdělení četností dokážeme orientačně posoudit, z jakého rozdělení výběr pochází, zda je výběr homogenní a zda neobsahuje odlehlé hodnoty. Dostaneme zde nejpoužívanější výběrové charakteristiky polohy, variability a koncentrace. Ty jsou všechny počítané pomocí excelovských funkcí z původních dat (uvedených v listu *data*).

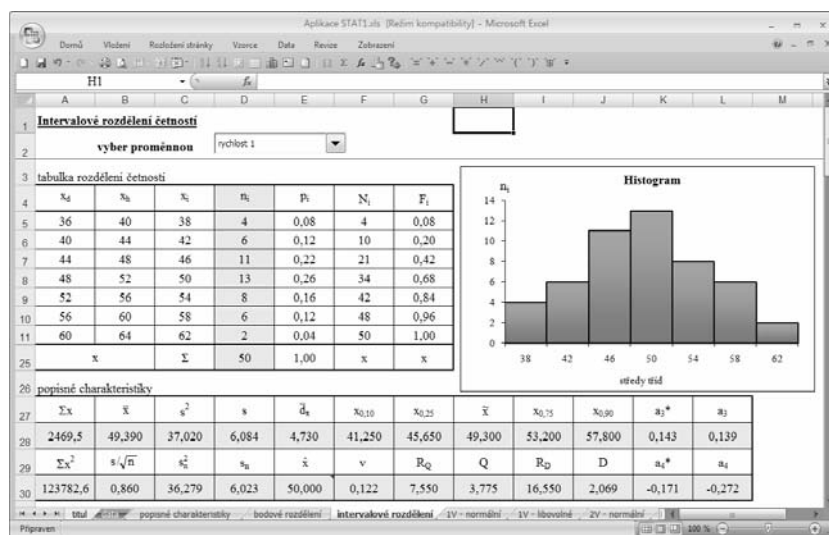
Tyto tři listy navíc shodně obsahují procedury představující testy o nulové šikmosti a nulové špičatosti resp. kombinovaný C-test o šikmosti a špičatosti [1]. Tyto testy se použijí na zvolené hladině významnosti k rozhodnutí, zda zkoumaný výběr pochází z normálního rozdělení. Ve výuce proto představují didakticky názorný prostředek k posouzení normality. Pokud testy normálního rozdělení zamítnou, považujeme data za výběr z libovolného – neznámého rozdělení. To má význam pro výběr dalších listů, na kterých jsou připravené další

statistické metody. S implementací jiných testů tvaru rozdělení, které jsou standardně nabízeny statistickými programy, se počítá v dalším rozšíření aplikace STAT1.



	A	B	C	D	E	F	G	H	I
1	rychlost 1	rychlost 2	moje data 3	s9p11 latovky	s9p12 rýže	s9p13 roztok	s10p14 zácvik	s10p15 surovina	s10p16 prach
2	54,1	46,7		49,8	1003	41	12	4,25	1,23
3	48,1	46,8		50,2	1005	42	19	4,15	1,51
4	46,4	44,8		50,3	997	41	13	4,03	1,41
5	44	52,1		49,5	1002	44	19	4,32	1,14
6	49,2	39,8		50,0	988	41	14	4,25	1,47
7	50,3	49,8		49,3	1000	41	14	4,12	1,10
8	52,6	43,3		50,0	1010	39	15	4,15	1,53
9	53,3	46,1		50,9	999	43	11	4,20	1,22
10	39	46,2		50,4	994	45	14	4,05	1,34
11	57,8	46,3		50,0	983	41	14	4,20	1,24
12	47,6	47,1		49,7	1008	42	20	4,05	1,54
13	57	50,1		50,6	1002	43	18	4,30	1,31
14	49,5	55,9		50,2	995	41	13	4,31	1,27
15	41,3	37,7		49,9	999	41	13	4,45	1,16
16	38,9	52,6		50,1	995	43	15	4,25	1,45
17	55,3	51			1004	41	11	4,25	1,34
18	54,3	39,1			1006	42	14	4,32	1,23
19	45,3	46,6			992	45	12	4,33	1,37
20	52,6	51,6			1000	42	16	4,32	1,51

Obrázek 1: Datové soubory



Obrázek 2: Intervalové rozdělení četností a výběrové charakteristiky

Na dalších listech jsou zpracované jednovýběrové a dvouvýběrové metody (označení listů 1V a 2V) odhadů a testů pro střední hodnoty a rozptyly za předpokladu normality a také pro libovolné rozdělení. Po výběru 1 resp. 2 datových souborů se zobrazí použité výběrové charakteristiky (zpravidla rozsah souborů, výběrové průměry a výběrové odchylky) a na zvolené hladině významnosti se určí intervaly spolehlivosti. Po nastavení potřebné alternativy a u 1V-úloh i hodnoty testovaného parametru, s ohledem na řešení konkrétního problému, se zobrazí kritická hodnota a p -hodnota pro dané nastavení. Veškeré výstupy – výsledky jsou zobrazeny v zelených buňkách. V nejnovější verzi aplikace byl doplněn ještě list s 1V a 2V procedurami o podílech – relativních četnostech. Poslední list aplikace tvoří elektronické statistické tabulky.

Ovládání jednotlivých listů je zcela intuitivní, student má však také oporu ve studijní pomůcce [2], kde po teoretickém výkladu je vždy řešený příklad přímo v prostředí STAT1. Výhodu takto koncipovaného řešení vidíme zejména v tom, že student musí provést vstupní nastavení sám, a to podle vlastního uvážení s ohledem na řešený problém. Neocenitelná je

také možnost učitele ukázat studentům souvislost mezi intervalem spolehlivosti a testem. Tím se učí studenti správně interpretovat výsledky a aktivně vnímat statistickou filozofii. Ve všech listech obsahujících popsání procedury indukční statistiky existuje ještě možnost vedle výpočtů z dat použít vstupy v podobě výběrových charakteristik. Toho lze využít didakticky zejména k ukázkám, jaký dopad na odhad či test bude mít změna rozsahu souboru, velikosti odchylky či hodnoty průměru.

5. Příklad řešený v aplikaci STAT1

Praktické užití aplikace STAT1 si ukažme na konkrétním příkladu. Představme si reálnou situaci, kdy na základě stížností rodičů na chování řidičů v bezprostřední blízkosti základní školy byla provedena dvě měření rychlosti aut. První měření bylo provedené na jaře a po jednání zástupců školy a města bylo schválené opatření ke snížení rychlosti pomocí dvou retardérů. Druhé měření bylo provedené na podzim a mělo za účel posoudit, zda po vybudování retardérů došlo ke snížení rychlosti, jak se očekávalo. První měření představuje 50 pozorování, druhé měření 60 pozorování, obě měření byla provedená za srovnatelných podmínek, v pracovní den mezi 7. a 8. hodinou ranní.

Ke statistickému řešení daného problému uijeme aplikaci STAT1. Oba datové soubory vložíme do datového listu pod jmény *rychlost1* a *rychlost2*, viz obr. 1. Východiskem další analýzy bude exploratorní analýza dat, kterou provedeme na obou souborech. Vzhledem k povaze dat provedeme intervalové rozdělení četností (použitý list *intervalové rozdělení*), ze kterého je zřejmé, že data jsou homogenní, rozdělení téměř symetrické a bez odlehlých hodnot, součástí exploratorní analýzy jsou i výběrové charakteristiky, viz obr. 2. Z obou výstupů je také zřejmé, že oba výběry pochází z normálního rozdělení.

Řešení praktické situace má charakter dvouvýběrového problému, v rámci kterého potřebujeme porovnat střední hodnoty odpovídající chování řidičů na jaře a na podzim. Po zavedení retardérů předpokládáme sníženou rychlost aut, čemuž reálně odpovídá nižší druhý výběrový průměr. Proto budeme testovat shodu středních hodnot $\mu_1 = \mu_2$ proti alternativní hypotéze $\mu_1 > \mu_2$. Protože obě proměnné mají normální rozdělení, použijeme list *2V-normální*, v jehož horní části provedeme výběr obou proměnných. Po provedení testu na shodu rozptylů (teorie v [2]) pokračujeme testem o shodě středních hodnot za předpokladu heteroskedasticity. Výsledkem je konstatování, že s 95% spolehlivostí se hypotéza o shodě obou středních hodnot zamítá, viz obr. 3. Prakticky to tedy znamená, že s touto spolehlivostí ke snížení rychlosti aut vlivem instalovaných retardérů skutečně došlo.

The screenshot shows the STAT1 application window with the following data and results:

1. výběr

n_1	$\bar{x}$	$s_1^2(x)$	$s_1^2(x)$
50	49,390	6,084	37,020

2. výběr

n_2	$\bar{y}$	$s_2^2(y)$	$s_2^2(y)$
60	46,348	3,938	15,510

$\alpha = 0,05$

1. Testy hypotéz o shodě rozptylů

Null hypothesis: $\sigma_1^2 = \sigma_2^2$ | Alternative hypothesis: $\sigma_1^2 \neq \sigma_2^2$ (selected), $\sigma_1^2 > \sigma_2^2$, $\sigma_1^2 < \sigma_2^2$

spol.	F	$F \in W_\alpha$	kritická hodnota	p-hodnota	H	A
95%	2,387	ano	0,578	1,707	0,002	zamítá se

heteroskedasticita

3. Testy hypotéz o shodě středních hodnot za předpokladu heteroskedasticity

Null hypothesis: $\mu_1 = \mu_2$ | Alternative hypothesis: $\mu_1 \neq \mu_2$, $\mu_1 > \mu_2$ (selected), $\mu_1 < \mu_2$

spol.	t	$t \in W_\alpha$	krit. hodn.	p-hodnota	H	A
95%	3,043	ano	1,664	0,002	zamítá se	se přijme

kvantily $F_\alpha(v_1, v_2)$

$v_1 = n_1 - 1$	$v_2 = n_2 - 1$
49	59

kvantily $t_\alpha(v^*)$

v^*	$t_{1-\alpha}(v^*)$	$t_{1-\alpha/2}(v^*)$
80	1,664	1,990

Obrázek 3: Dvouvýběrový test o shodě středních hodnot

Učiteli se samozřejmě otevírá prostor pro formulace dalších praktických problémů, které lze na základě našich obou měření řešit. Např.: Jaká byla průměrná rychlost aut ve sledovaném místě na jaře? Jaký byl podíl řidičů, kteří na jaře překračovali povolenou rychlost 50 km/hod.? Snížil se podíl řidičů, kteří na podzim překračovali povolenou rychlost? S využitím aplikace STAT1 se může řešit řada podobných reálných situací a s výhodou tak posilovat důvěru studentů ve statistiku.

6. Závěr

Autory sestavená aplikace v excelovském prostředí se používá na cvičeních vedených učitelem. Odpadají veškeré numerické výpočty a myšlenková kapacita studentů se využívá k postupnému chápání statistické filozofie. Na takto podporovanou výuku může student navázat při své domácí přípravě, protože produkt má volně k dispozici. Také kompletní zápočtový projekt na základě vlastního měření nebo zjišťování řeší student v této aplikaci. Zanedbatelná není ani skutečnost, že studentovi aplikace zůstane k dalšímu použití, po celou dobu studia a vlastně i poté v praxi. Veškeré poznatky získané ve výuce statistiky s touto aplikací se tak mohou v profesním životě ověřovat a využívat.

7. Literatura

- [1] ANDĚL, J. Základy matematické statistiky. 1. vyd. Praha: Matfyzpress, 2005, 358 s. ISBN 80-86732-40-1.
- [2] KŘÍŽ, O. – NEUBAUER, J. – SEDLAČÍK, M. Popisná statistika a výběrová šetření. [Skripta]. 1. vyd. Brno: Univerzita obrany, 2009, 181 s. ISBN 978-80-7231-7073.
- [3] KŘÍŽ, O. Sbírka úloh ze statistiky. [Skripta]. 1. vyd. Vyškov: VVŠ PV, 1999, 115 s. ISBN 80-7231-033-X.
- [4] KŘÍŽ, O. – NEUBAUER, J. Výuka statistiky podporovaná Excelem. In: Sborník XXV. mezinárodního kolokvia řízení osvojovacího procesu: sborník abstraktů a elektronických verzí příspěvků na CD-ROMu [CD-ROM]. Brno: UO, 2007. Adresář: 6clanky/1krizo. pdf. ISBN 978-80-7231-228-3.
- [5] NOVOTNÝ, J. Počítačová podpora cvičení ze statistiky systémem MOODL. In: Sborník příspěvků 6. konference o matematice a fyzice na vysokých školách technických. Brno: UO, 2009. s. 185–188. ISBN 978-80-7321-667-0.

Adresa autorů:

Oldřich Kříž, RNDr.
Univerzita obrany Brno
Fakulta ekonomiky a managementu
Kounicova 65, 662 10 Brno
oldrich.kriz@unob.cz

Jiří Neubauer, Mgr., Ph.D.
Univerzita obrany Brno
Fakulta ekonomiky a managementu
Kounicova 65, 662 10 Brno
jiri.neubauer@unob.cz

Marek Sedlačík, Mgr., Ph.D.
Univerzita obrany Brno
Fakulta ekonomiky a managementu
Kounicova 65, 662 10 Brno
marek.sedlacik@unob.cz

Aby štatistika nebola klamstvo... So statistic not be lie...

Peter Lenčoš

Abstract: In this article we deal with interpretation of results acquired using statistical methods. We focus on hypothesis testing using statistical significance tests

Key words: statistic, significance tests

Kľúčové slová: štatistika, testy významnosti

1. Shaw: Veľké klamstvo, malé klamstvo, štatistika

„Poznám tri druhy klamstva – veľké klamstvo, malé klamstvo a štatistika.“ Výrok nositeľa Nobelovej ceny za literatúru – Georga Bernarda Shawa, známeho svojimi aforizmami, v intenciách satiry evokuje rizikovosť štatistiky ako takej. Faktom je, že štatistika ako matematická vedná disciplína je citlivá na výber vhodných metód, ktoré ponúka za účelom ich aplikácie a na korektnosť interpretácie výsledkov získaných aplikáciou daných metód. Odkaz Shawa k štatistike, ktorý sme ilustrovali uvedeným aforizmom, spočíva vlastne na otázke vhodnosti použitých metód štatistiky v praxi a tiež na interpretácii ich výsledkov.

V rámci článku sa dotkneme otázky korektnej interpretácie výsledkov získaných aplikáciou metód matematickej štatistiky na poli empirie. Explicitne sa zameriame na jednu triedu testov ako nástrojov na testovanie štatistických hypotéz a síce na testy významnosti. Venujeme im zvláštnu pozornosť, pretože nachádzajú široké uplatnenie vo výskume v rozmanitých oblastiach vedy.

2. Do „kuchyne“ testov významnosti

Úloha testovať štatistickú hypotézu je štandardne formulovaná tak, že proti testovanej hypotéze (nazývanej aj nulová hypotéza a označovanej H_0) stojí tzv. alternatívna hypotéza (označovaná H_1). Štatistická hypotéza sa týka buď parametra rozdelenia pravdepodobnosti náhodnej premennej reprezentujúcej nejaký základný súbor alebo typu rozdelenia pravdepodobnosti tejto náhodnej premennej. Pri testovaní nulovej hypotézy sa možno dopustiť dvoch druhov chýb – tzv. chyby 1. druhu a chyby 2. druhu. Chyby 1. druhu sa dopúšťame vtedy, keď v kontexte všeobecných princípov testovania štatistických hypotéz musíme zamietnuť nulovú hypotézu, hoci je táto v skutočnosti pravdivá. Ak nulovú hypotézu nemôžeme zamietnuť, hoci je táto v skutočnosti nepravdivá, dopúšťame sa chyby 2. druhu. Pravdepodobnosť chyby 1. druhu a chyby 2. druhu sa (v uvedenom poradí) označuje α , β .

Zrejme by bolo pri testovaní štatistických hypotéz vhodné eliminovať α , β . Postup založený na súčasnom minimalizovaní α a β by bol však bezvýsledný, pretože sa dá dokázať, že znižovanie α by nikdy nevedlo k znižovaniu β (ale zväčša aj k jej zvyšovaniu) a naopak. V tomto spočíva dôvod, pre ktorý boli navrhnuté tzv. testy významnosti.

Idea testov významnosti je založená na zamietnutí nulovej hypotézy s vopred zvolenou pravdepodobnosťou chyby 1. druhu (α – tzv. hladina významnosti), pričom pravdepodobnosť chyby 2. druhu (β) sa neberie do úvahy. α sa volí ako relatívne veľmi malé číslo, takže pri zamietnutí nulovej hypotézy je pravdepodobnosť omylu (spočívajúceho v tom, že nulovú hypotézu zamietame, hoci v skutočnosti je pravdivá) veľmi malá (rovnajúca sa α). β však môže byť číslo relatívne veľmi veľké, no keďže sa neberie do úvahy, nemáme o ňom vedomosť. Teda, ak nulovú hypotézu nemôžeme zamietnuť, nemôžeme ju ani prijať, resp.

nemôžeme ju považovať dokázanú, iba za prípustnú – totiž, existuje riziko chyby 2. druhu, ktorej pravdepodobnosť (β) môže byť relatívne veľmi veľká.

3. Testy významnosti: nezamietnuť neznamená prijať

Každý matematik si v súvislosti s vymedzenou ideou testov významnosti uvedomuje, že ak v kontexte všeobecných princípov testovania štatistických hypotéz musíme zamietnuť nulovú hypotézu, potom môžeme považovať jej potenciálnu platnosť so spoľahlivosťou $1 - \alpha$ za vylúčenú. Na druhej strane, ak ju nemožno zamietnuť, považujeme ju iba za prípustnú, čo nás neoprávňuje vysloviť sa v prospech jej platnosti. Teda nezamietnutie nulovej hypotézy pri použití testu významnosti neumožňuje jej prijatie, resp. rigidné prijatie jej platnosti, no oprávňuje pripustiť jej potenciálnu platnosť.

Predmetnej skutočnosti sa dotkol vo svojom článku [6] Dalibor Roháč, doktorand na George Mason University vo Virginii. Roháč sa okrajovo zmieňuje o tzv. *t*-teste. Ako aj uvádza, „v praxi sa požíva napr. na porovnanie toho, či sa výsledky meraní z vybratej vzorky štatisticky významne líšia od kontrolnej vzorky.“ Teda hovoril o jednom zo skupiny testov významnosti, ktorý poznáme ako párový *t*-test. Ako ďalej uvádza, „pri *t*-teste sa overuje, či možno hypotézu o tom, že skutočný efekt je nulový, vylúčiť na základe údajov nameraných na vybratej vzorke.“ Čo je dôležité, „ak takzvanú nulovú hypotézu nemôžeme zamietnuť, stále nám to nedáva právo jednoznačne povedať, že skutočný efekt je nulový.“ So stotožnením sa s postrehom Roháča zdôrazňujeme, že v uvedenom kontexte „sa tejto chyby dopúšťa príliš veľa spoločenských vedcov.“ Iste, táto chyba v zmysle považovania nulovej hypotézy za dokázanú v prípade jej nezamietnutia sa netýka len *t*-testu, ale ktoréhokoľvek použitého testu významnosti. A faktom je, že sú ňou poznačené experimentálne overenia, prieskumy, výskumy na rôznych úrovniach.

4. Ukážme si to na párovom *t*-teste

V nasledujúcom texte budeme dosiaľ uvedené skutočnosti týkajúce sa testov významnosti ilustrovať na podklade spomenutého párového *t*-testu. Najskôr uvedieme stručný teoretický výklad k danému testu významnosti (metodika párového *t*-testu je ozrejmenej napr. v [1], [2], [3] a [5]).

Párový *t*-test sa používa, ak na každej z n vybraných štatistických jednotiek sledujeme určitý štatistický znak nezávisle dvakrát. V našom prípade budeme sledovať preferencie aktuálne najsilnejších ôsmich parlamentných a mimoparlamentných politických strán (podľa abecedy HZDS, KDH, Most-Híd, SaS, SDKÚ-DS, SMER-SD, SMK, SNS) za apríl 2010 podľa prieskumu verejnej mienky agentúry POLIS Slovakia (uskutočneného v dňoch 24. – 26. apríla 2010 na reprezentatívnej vzorke 924 respondentov) a agentúry MVK (uskutočneného v dňoch 16. – 23. apríla 2010 na reprezentatívnej vzorke 1022 respondentov). Preferencie jednotlivých strán podľa uvedených agentúr obsahuje tabuľka 1.

Tabuľka 1: Preferencie polit. strán podľa agentúr POLIS Slovakia a MVK za apríl 2010

výsledky prieskumu agentúry POLIS Slovakia (v %)	politická strana	výsledky prieskumu agentúry MVK (v %)
4,0	HZDS	5,2
13,2	KDH	11,4
6,2	Most-Híd	5,1
9,2	SaS	11,6
13,8	SDKÚ-DS	11,7

36,2	SMER-SD	35,1
5,8	SMK	6,0
5,3	SNS	6,2

Na každej štatistickej jednotke teda nameriame dve hodnoty x_i a y_i ($i = 1, 2, \dots, n$), ktoré môžeme formálne považovať za realizácie náhodných premenných (v uvedenom poradí) X_i

a Y_i ($i = 1, 2, \dots, n$). Máme tak náhodný výber $\begin{pmatrix} (X_1, X_2) \\ \dots \\ (X_n, Y_n) \end{pmatrix}$ z dvojrozmerného základného

súboru modelovaného náhodným vektorom (X, Y) , ktorý má v zmysle predpokladov párového t -testu dvojrozmerné normálne rozdelenie pravdepodobnosti $N((\mu_1, \mu_2), (\sigma_1^2, \sigma_2^2))$. V ďalšom budeme preto predpokladať, že údaje v tab. 1 sú realizáciou náhodného výberu z dvojrozmerného normálneho rozdelenia pravdepodobnosti.

Párovým t -testom testujeme nulovú hypotézu $H_0: \mu_1 - \mu_2 = \Delta$ (Δ je číslo) proti dvojstrannej (resp. jednostrannej) alternatíve H_1 . V našom prípade nás bude zaujímať, či sa preferencie politických strán za apríl 2010 (tab. 1) namerané danými agentúrami (POLIS Slovakia a MVK) štatisticky významne odlišujú. Preto v nulovej hypotéze položíme $\Delta = 0$. Ako testovacie kritérium sa v prípade párového t -testu používa štatistika

$$t = \frac{\bar{Z} - \Delta}{S} \cdot \sqrt{n} \sim t(n-1), \quad (1)$$

pričom $\bar{Z} = \frac{1}{n} \cdot \sum_{i=1}^n Z_i$ a $S = \sqrt{\frac{1}{n-1} \cdot \sum_{i=1}^n (Z_i - \bar{Z})^2}$, kde Z_i ($i = 1, 2, \dots, n$), sú náhodné premenné definované nasledovne: $Z_i = X_i - Y_i$ ($i = 1, 2, \dots, n$). Pre údaje z tab. 1 platí: $t = 0,76577$. Kritickým oborom W_α testovacieho kritéria (1) je množina

$$W_\alpha = (-\infty; -t_\alpha(n-1)) \cup (t_\alpha(n-1); \infty), \quad (2)$$

kde $t_\alpha(n-1)$ je kritická hodnota Studentovho rozdelenia. Pre rozsah $n = 8$ nášho náhodného výberu a zvolenú hladinu významnosti $\alpha = 0,05$ máme $W_\alpha = (-\infty; -2,36462) \cup (2,36462; \infty)$. Hodnota $t = 0,76577$ nepatrí do intervalu $(-\infty; -2,36462) \cup (2,36462; \infty)$, preto v kontexte všeobecných princípov testovania štatistických hypotéz hypotézu H_0 na hladine významnosti nemožno zamietnuť. Nevylučujeme teda, že medzi meraniami preferencií agentúr POLIS Slovakia a MVK nie je štatisticky významný rozdiel.

Výsledok aplikácie párového t -testu na údaje tab. 1 by sme mohli interpretovať povedzme aj v kontexte prípadnej zaujatosti agentúr POLIS Slovakia a MVK, pokiaľ ide o prieskum verejnej mienky v otázke podpory politických strán. Keďže nulovú hypotézu sme nemohli zamietnuť, mohli by sme sa vysloviť v prospech ich nezaujatosti. Nemohli by sme však operovať touto hypotézou ako neotrasiteľným faktom. Iste, náš ilustračný príklad (preferencie politických strán) nie je v kontexte chybné interpretácie tak závažný z hľadiska možných negatívnych dopadov, ako by bolo povedzme skúmanie vedľajších účinkov nejakého lieku.

5. Záver

Matematická štatistika ponúka širokú škálu nástrojov umožňujúcich popísať javy, či už prírodné alebo spoločenské. Vždy je však potrebné sledovať otázku korektnosti interpretácie

nimi získaných výsledkov. Platí to aj pre testy významnosti, v hojnej miere využívané vo výskumnej praxi.

6. Literatúra

- [1] KUBANOVÁ, J. 2003. Statistické metódy pro ekonomickú a technickú praxi. Bratislava: STATIS, 2003. ISBN 80-85659-31-X.
- [2] LAMOŠ, F. – POTOCKÝ, R. 1989. Pravdepodobnosť a matematická štatistika – štatistické analýzy. Bratislava: Alfa, 1989. ISBN 80-05-00115-0.
- [3] RIEČAN, P. – LAMOŠ, F. – LENÁRT, C. 1984. Pravdepodobnosť a matematická štatistika. Bratislava: Alfa, 1984.
- [4] TIRPÁKOVÁ, A. – MALÁ, D. 2007. Základy štatistiky pre pedagógov, psychológov a sociológov s popisom postupu práce v programe Excel. Nitra: FPV UKF v Nitre, 2007. ISBN 978-80-8094-220-5.
- [5] VRÁBELOVÁ, M. – MARKECHOVÁ, D. 2001. Pravdepodobnosť a štatistika. Nitra: FPV UKF v Nitre, 2001.
- [6] ROHÁČ, D. 2008. Na veľkosti záleží. In: TREND, 24. 4. 2008, s. 16 – 17.

Adresa autora:

Peter Lenčoš, PaedDr. RNDr.
Tr. A. Hlinku 1
949 74 Nitra
peter.lences@ukf.sk

Analýza odpovedí „neviem“ v batérii otázok The analysis of „do not know“ responses in the battery of questions

Ján Luha, Gabriel Bianchi¹

Abstract: The authors present a specific statistical analysis of responses „do not know“ in the battery of questions in a population surveys. Examples are illustrated by a concrete survey.

Key words: responses “do not know”, analysis, testing, confidence intervals.

Kľúčové slová: odpovede „neviem“, analýza, testovanie, intervaly spoľahlivosti.

1. Úvod

Pri populačných prieskumoch sa často stretávame s problémom odpovede neviem. Podobný je aj problém neodpovedí, teda keď časť respondentov na niektorú z otázok výskumného dotazníka neodpovie.

V príspevku sa zaoberáme špecifickým problémom odpovedí neviem pri batérii otázok. Batéria otázok sa zaoberá rovnakou tematikou, pričom škála odpovedí na jednotlivé otázky v batérii je rovnaká. Odpoveď neviem použije respondent z niekoľkých dôvodov. Ak je variant odpovede neviem súčasťou dotazníka a respondent na otázky odpovedá sám, tak je táto odpoveď frekventovanejšia (úniková voľba) – najmä pri väčšej batérii otázok. Ďalšou príčinou voľby odpovede neviem je určitý nezáujem odpovedať a samozrejme odpoveď neviem môže byť skutočná odpoveď respondenta. Pre batériu otázok výber odpovede neviem môže vyjadrovať aj špecifikum skúmaných oblastí. ***Problematiku budeme ilustrovať pomocou konkrétneho výskumu (Bianchi, 2009).*** Špecifiká odpovedí neviem pri vybraných batériách otázok sú opísané v práci (Bianchi, Luha, 2010).

V tomto príspevku sa zaoberáme štatistickými postupmi overenia hypotézy o rovnakom podiele odpovedí neviem v každej otázke danej batérie otázok. Budeme sa zaoberať tromi metódami overovania uvedenej hypotézy – pomocou exaktných intervalov spoľahlivosti pre podiely, pomocou exaktného binomického testu a pomocou exaktného Fisherovho testu.

2. Podiel odpovedí neviem v batérii otázok

Ako príklad batérie otázok vyberieme z výskumu IVO/COPART-KVSBK, máj 2008 otázku o tom či by spoločnosť mala pomáhať určitým skupinám ľudí. Uvádzame znenie vybranej batérie otázok.

Znenie „spoločnej“ otázky:

18. "Zamyslite sa nad tým, či by mala spoločnosť týmto skupinám ľudí pomáhať riešiť ich problémy, respektíve im vychádzať v ústrety pri uspokojovaní ich špeciálnych potrieb?"

Škála odpovedí pre všetky otázky danej batérie:

1) rozhodne áno 2) skôr áno 3) skôr nie 4) rozhodne nie 9) nevie (NEČÍTAŤ)

¹ Gabriel Bianchi, KVSBK a Centrum excelentnosti SAV pre výskum a rozvoj občianstva a participácie

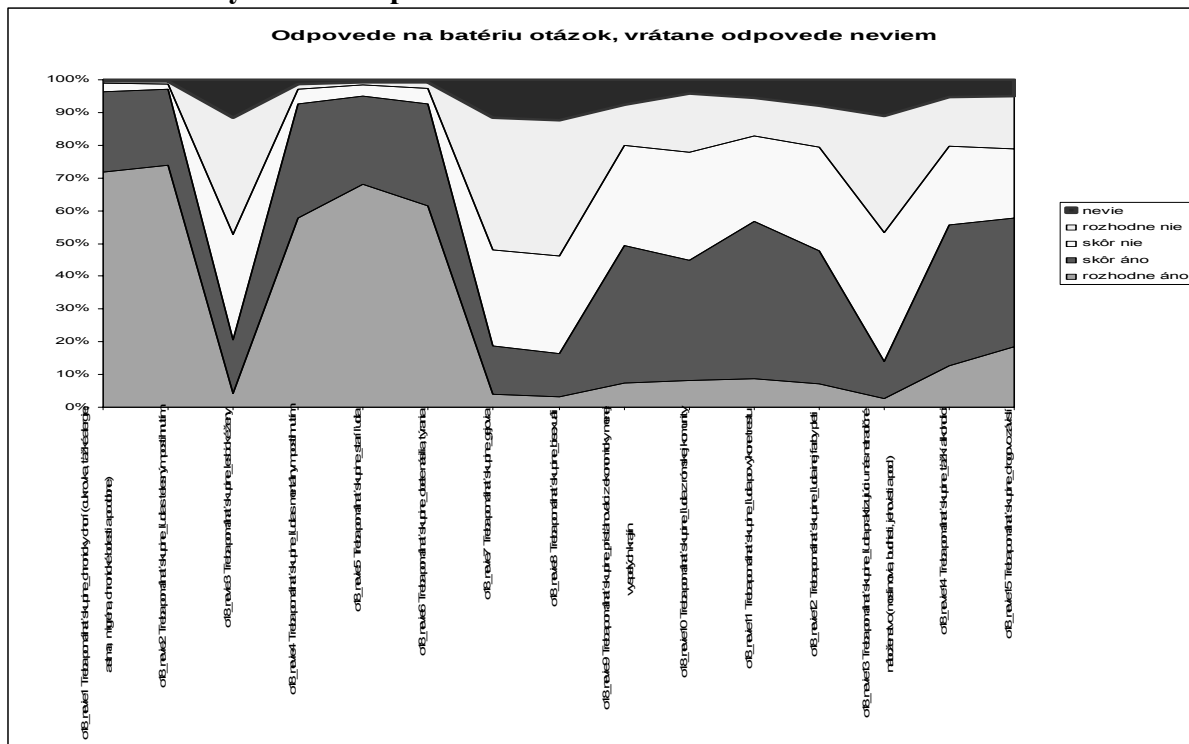
Batéria 15 otázok výskumu s výsledkami v percentách, *kde rozsah výberového súboru bol n=1209*:

Tabuľka č.1: Výsledky odpovedí na batériu otázok

	Otázka	rozhodne áno	skôr áno	skôr nie	rozhodne nie	neviem
18_1	chronicky chorí (cukrovka, ťažké alergie, astma, migréna, chronické bolesti a podobne)	71,83	24,50	2,57	0,66	0,43
18_2	ľudia s telesným postihnutím	73,94	23,04	1,77	0,74	0,51
18_3	lesbické ženy	4,18	16,41	32,26	35,54	11,62
18_4	ľudia s mentálnym postihnutím	57,89	34,68	4,53	1,55	1,35
18_5	starí ľudia	68,17	26,75	3,40	0,93	0,75
18_6	obete násilia, týrania	61,61	30,96	4,77	1,77	0,89
18_7	Gejovia	3,86	14,80	29,36	40,45	11,54
18_8	bisexuáli	3,29	13,11	29,80	41,41	12,39
18_9	prist'ahovalci z ekonomicky menej vyspelých krajín	7,31	41,96	30,79	12,30	7,65
18_10	ľudia z rómskej komunity	8,18	36,75	32,99	17,80	4,27
18_11	ľudia po výkone trestu	8,74	48,01	26,15	11,56	5,53
18_12	ľudia inej farby pleti	7,24	40,51	31,60	12,86	7,80
18_13	ľudia praktizujúci u nás netradičné náboženstvo (moslimovia, budhisti, jehovisti a pod.)	2,75	11,32	39,13	35,78	11,02
18_14	ťažkí alkoholicy	12,73	43,02	24,05	14,93	5,27
18_15	drogovo závislí	18,42	39,29	21,21	16,10	4,98

Graf č.1 názorne zobrazuje štruktúru vecných odpovedí, vrátane odpovede neviem. Výsledky za odpoveď neviem signalizujú vplyv merita otázky na veľkosť percenta tejto odpovede. Názornejšie to vidno z grafickej prezentácie s usporiadaním otázok podľa veľkosti percenta odpovede neviem – Graf. č.2.

Graf č. 1: Graf výsledkov odpovedí na batériu otázok



Znenie otázok sme pre štatistické spracovanie v prostredí SPSS upravili a na skúmanie odpovedí rekódovali do nových premenných o18_nevie1 až o18_nevie15 s kódmi 1=neviem a 0=odpovedal (vzhľadom k tomu, že primárne nás zaujíma jav – neviem). Percento odpovedí neviem v danom výskume je prehľadne v tabuľke:

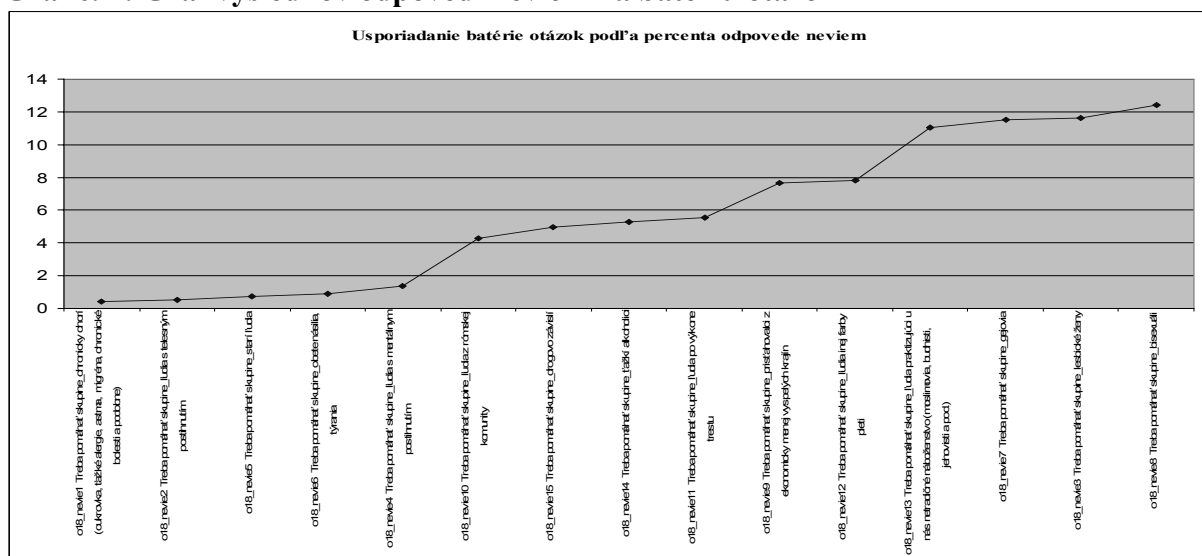
Tabuľka č.2: Výsledky podielu odpovede neviem za batériu otázok

Otázka	percento odpovede neviem
o18_nevie1 Treba pomáhať skupine chronicky chorí (cukrovka, ťažké alergie, astma, migréna, chronické bolesti a podobne)	0,43
o18_nevie2 Treba pomáhať skupine ľudia s telesným postihnutím	0,51
o18_nevie3 Treba pomáhať skupine lesbické ženy	11,62
o18_nevie4 Treba pomáhať skupine ľudia s mentálnym postihnutím	1,35
o18_nevie5 Treba pomáhať skupine starí ľudia	0,75
o18_nevie6 Treba pomáhať skupine obeť násilia, týrania	0,89
o18_nevie7 Treba pomáhať skupine gejovia	11,54
o18_nevie8 Treba pomáhať skupine bisexuáli	12,39
o18_nevie9 Treba pomáhať skupine prisťahovalci z ekonomicky menej vyspelých krajín	7,65
o18_nevie10 Treba pomáhať skupine ľudia z rómskej komunity	4,27
o18_nevie11 Treba pomáhať skupine ľudia po výkone trestu	5,53
o18_nevie12 Treba pomáhať skupine ľudia inej farby pleti	7,80
o18_nevie13 Treba pomáhať skupine ľudia praktizujúci u nás netradičné náboženstvo (moslimovia, budhisti, jehovisti a pod.)	11,02
o18_nevie14 Treba pomáhať skupine ťažkí alkoholicy	5,27
o18_nevie15 Treba pomáhať skupine drogovzo závislí	4,98

Vidno, že najväčší podiel odpovede neviem bol zaznamenaný v otázke na bisexuálov (12,39%), čo je takmer rovnako ako rovnako pri lesbičkách (11,62%) a pri gejoch (11,54%). Nasledujú iné náboženstvá (moslimovia, budhisti, jehovisti a pod. 11,02%). Ďalšiu skupinu otázok s vyšším podielom ako priemer za 15 otázok (5,73%) tvoria ľudia inej farby pleti (7,80%) a prisťahovalci z ekonomicky menej vyspelých krajín (7,65%). Blízko priemeru sa nachádzali ľudia vo výkone trestu (5,53%), ťažkí alkoholicy (5,27%), drogovzo závislí (4,98) a ľudia rómskej komunity (4,27%). Nízke podiely odpovedí neviem boli pri otázkach na pomoc ľuďom s mentálnym postihnutím (1,35%), obetiam násilia, týrania (0,89%), starým ľuďom (0,75%), ľuďom s telesným postihnutím (0,51%) a chronicky chorým ľuďom (0,43%).

V grafickej prezentácii percenta odpovede neviem sme otázky usporiadali podľa veľkosti percenta tejto odpovede.

Graf č. 2: Graf výsledkov odpovedí neviem na batériu otázok



3. Overovanie hypotézy o zhode podielov odpovedí neviem v batérii otázok

Pri analýze podielov odpovedí neviem v batérii otázok nás môže zaujímať či sa tieto podiely signifikantne líšia od určitej hodnoty. Prirodzene sa núka ako hodnota, ku ktorej budeme komparovať jednotlivé podiely, priemerný podiel, ktorý získame ako sumu počtu odpovedí neviem za všetky otázky batérie (1040) k sume počtu všetkých odpovedí (respondentov) za všetky otázky batérie (18142=15.1209). Hodnota priemerného percenta je teda 5,73%.

Budeme teda overovať hypotézy o zhode percenta odpovedí neviem za jednotlivé otázky batérie s hodnotou 5,73%. Označme príslušné percento odpovedí neviem za otázky batérie p_1 , p_2 až p_{15} .

Overujeme teda hypotézy H_i : $p_i=5,73\%$, pre $i=1,2$ až 15.

Ako kritérium na overenie týchto hypotéz môžeme zvoliť exaktné intrervalov spoľahlivosti pre podiely, exaktný binomický test podielov a aj špeciálne zostavené kontingenčné tabuľky a Chí-kvadrát test, prípadne Fisherov exaktný test.

Exaktné intervaly spoľahlivosti pre podiely

Vzťahy pre exaktné hranice (dolnú **pd** resp. hornú **ph**) intervalov spoľahlivosti:

$$pd(p,n)=x/[x+(n-x+1)F1],$$

kde $F1$ je $1-\alpha_1$ kvantil F rozdelenia s počtom stupňov voľnosti: $2(n-x+1)$, $2x$

$$ph(p,n)=[(x+1)F2]/[n-x+(x+1)F2],$$

kde $F2$ je $1-\alpha_2$ kvantil F rozdelenia s počtom stupňov voľnosti: $2(x+1)$, $2(n-x)$.

Obvykle volíme $\alpha_1=\alpha_2=\alpha/2$ a $\alpha=0,05$, resp. v percentách 5%.

V nasledovnej tabuľke uvádzame 95% exaktné hranice spoľahlivosti a vyhodnotenie signifikantnosti oproti hodnote 5,73%. Výsledok je signifikantný, keď sa hodnota 5,73% nenachádza v príslušnom intervale spoľahlivosti, čo je pri otázkach č. 11, 14 a 15. Porovnaním jednotlivých podielov zistíme, ktoré sú signifikantne nižšie ako priemer a ktoré vyššie. Vyššie podiely odpovede neviem signalizujú menšiu znalosť príslušnej problematiky, prípadne signalizujú „vzťah“ respondentov ku skupinám spomínaných v danej otázke.

Tabuľka č.3: Exaktné hranice spoľahlivosti podielov odpovede neviem v batérii otázok

otázka	podiel odpovede neviem	IS_dolna	IS_horna	priemerný podiel	signifikant
o18_nevie1	0,4	0,134	0,962	5,73	sig
o18_nevie2	0,5	0,182	1,077	5,73	sig
o18_nevie3	11,7	9,907	13,607	5,73	sig
o18_nevie4	1,3	0,758	2,140	5,73	sig
o18_nevie5	0,7	0,341	1,408	5,73	sig
o18_nevie6	0,9	0,455	1,622	5,73	sig
o18_nevie7	11,6	9,830	13,519	5,73	sig
o18_nevie8	12,4	10,601	14,398	5,73	sig
o18_nevie9	7,7	6,253	9,341	5,73	sig
o18_nevie10	4,3	3,229	5,602	5,73	sig
o18_nevie11	5,5	4,320	6,985	5,73	nesig
o18_nevie12	7,8	6,328	9,431	5,73	sig
o18_nevie13	11,0	9,292	12,902	5,73	sig
o18_nevie14	5,3	4,100	6,710	5,73	nesig
o18_nevie15	5,0	3,808	6,342	5,73	nesig

Exaktný binomický test pre podiely

Dáva rovnaké výsledky ako exaktné intervaly spoľahlivosti, ako vidno z tabuľky č.4.

Tabuľka č.4: Exaktný binomický test podielov odpovede neviem v batérii otázok

otázka	N	Observed Prop.	Test Prop.	Exact Sig. (1-tailed)
o18_nevie1	5	0,0041	0,0573	0,000
o18_nevie2	6	0,0050	0,0573	0,000
o18_nevie3	141	0,1166	0,0573	0,000
o18_nevie4	16	0,0132	0,0573	0,000
o18_nevie5	9	0,0074	0,0573	0,000
o18_nevie6	11	0,0091	0,0573	0,000
o18_nevie7	140	0,1158	0,0573	0,000
o18_nevie8	150	0,1241	0,0573	0,000
o18_nevie9	93	0,0769	0,0573	0,003
o18_nevie10	52	0,0430	0,0573	0,016
o18_nevie11	67	0,0554	0,0573	0,420
o18_nevie12	94	0,0778	0,0573	0,002
o18_nevie13	133	0,1100	0,0573	0,000
o18_nevie14	64	0,0529	0,0573	0,281
o18_nevie15	60	0,0496	0,0573	0,138

Kontingenčné tabuľky a Fisherov exaktný test

Aby sme mohli aplikovať Fisherov exaktný test kontingenčnej tabuľky je potrebné vhodne kontingenčné tabuľky skonštruovať. Uvažujme 2x2 kontingenčnú tabuľku:

n_{11}	n_{12}	$n_{1.}$
n_{21}	n_{22}	$n_{2.}$
$n_{.1}$	$n_{.2}$	n

Kde n_{11} je počet odpovedí neviem danej otázky a n_{12} doplnok do rozsahu výberového súboru $n=1209$. Druhý súbor bude zakaždým konštrukt odpovedajúci priemernému podielu. Po zaokrúhlení dostávame odpovedajúce počty odpovedí za priemer neviem $n_{21}=69$ a $n_{22}=1140$, za doplnok.

V nasledujúcej tabuľke sú výsledky Fisherovho testu. Takmer rovnaké výsledky signifikantnosti s výsledkami tabuľky č.3 sú za jednostranný Fisherov exaktný test, okrem výsledku za otázku č. 10, kde sme pri exaktných intervaloch spoľahlivosti získali signifikantný výsledok a v tabuľke č. 4 nie je signifikantný, keď P-hodnota je $P=0,068$.

Konštrukcia kontingenčných tabuliek vyžaduje zaokrúhľovanie a tak sú výsledky potenciálne menej presné ako pri exaktných intervaloch spoľahlivosti a pri exaktnom binomickom teste.

Tabuľka č.5: Fisherov exaktný test podielov odpovede neviem v batérii otázok

otázka	nevie	doplnok nevie	priemer nevie	doplnok	Exact Sig. (2- sided)	Exact Sig. (1- sided)
o18_nevie1	5	1204	69	1140	0,000	0,000
o18_nevie2	6	1203	69	1140	0,000	0,000
o18_nevie3	141	1068	69	1140	0,000	0,000
o18_nevie4	16	1193	69	1140	0,000	0,000
o18_nevie5	9	1200	69	1140	0,000	0,000
o18_nevie6	11	1198	69	1140	0,000	0,000
o18_nevie7	140	1069	69	1140	0,000	0,000
o18_nevie8	150	1059	69	1140	0,000	0,000
o18_nevie9	93	1116	69	1140	0,061	0,031
o18_nevie10	52	1157	69	1140	0,135	0,068
o18_nevie11	67	1142	69	1140	0,930	0,465
o18_nevie12	94	1115	69	1140	0,051	0,026
o18_nevie13	133	1076	69	1140	0,000	0,000
o18_nevie14	64	1145	69	1140	0,657	0,361
o18_nevie15	60	1149	69	1140	0,417	0,235

V práci Luha J. (2009) je uvedený príklad kladnej korelácie javov práve pre skúmaný typ javov. Ďalší zaujímavý výsledok sme získali ako výsledok faktorovej analýzy, keď sme hľadali tri faktory. Získali sme rozdelenie otázok batérie presne podľa veľkosti podielu odpovede neviem.

Tabuľka č.6: Výsledok faktorovej analýzy

	1	2	3
o18_nevie7 Treba pomáhať skupine_gejovia	0,89	0,07	0,18
o18_nevie8 Treba pomáhať skupine_bisexuáli	0,88	0,03	0,15
o18_nevie3 Treba pomáhať skupine_lesbické ženy	0,87	0,03	0,17
o18_nevie13 Treba pomáhať skupine_ludia praktizujúci u nás netradičné náboženstvo (moslimovia, budhisti, jehovisti a pod.)	0,69	0,14	0,22
o18_nevie1 Treba pomáhať skupine_chronicky chorí (cukrovka, ťažké alergie, astma, migréna, chronické bolesti a podobne)	0,08	0,80	0,04
o18_nevie5 Treba pomáhať skupine_starí ľudia	0,04	0,78	0,16
o18_nevie2 Treba pomáhať skupine_ludia s telesným postihnutím	0,09	0,78	0,09
o18_nevie4 Treba pomáhať skupine_ludia s mentálnym postihnutím	0,08	0,75	0,27
o18_nevie6 Treba pomáhať skupine_obete násilia, týrania	0,03	0,70	0,13
o18_nevie14 Treba pomáhať skupine_ťažkí alkoholicy	0,10	0,11	0,78
o18_nevie15 Treba pomáhať skupine_drogovo závislí	0,09	0,14	0,78
o18_nevie11 Treba pomáhať skupine_ludia po výkone trestu	0,12	0,13	0,70
o18_nevie12 Treba pomáhať skupine_ludia inej farby pleti	0,33	0,11	0,62
o18_nevie10 Treba pomáhať skupine_ludia z rómskej komunity	0,26	0,19	0,48
o18_nevie9 Treba pomáhať skupine_prisťahovalci z ekonomicky menej vyspelých krajín	0,44	0,10	0,48
Extraction Method: Principal Component Analysis. Rotation Method: Varimax with Kaiser Normalization.			

4. Záver

Podiel respondentov, ktorí neodpovedajú na otázky, spolu s ďalšími matematicko-štatistickými parametrami (ako napr. tvar rozloženia škálových odpovedí), môže byť užitočnou informáciou pri spoločenskovednom výskume – najmä v prípade politicky či morálne citlivých postojov a názorov. Je to najmä v prípade porovnania rôznych mier neodpovedania medzi otázkami v batérii. Hlbšiu analýzu problému podielu odpovede neviem v batérii otázok je možné realizovať pomocou exaktných intervalov spoľahlivosti podielov, pomocou exaktného binomického testu a špeciálnou konštrukciou kontingenčných tabuliek. Na praktickú aplikáciu je vhodnejší postup pomocou exaktných intervalov spoľahlivosti pre podiely a tiež pomocou exaktného binomického testu pre podiely. Pri konštrukcii kontingenčných tabuliek musíme počty odpovedí neviem, a teda aj doplnku do rozsahu výberu, zaokrúhľovať, čím najmä pri menších rozsahoch výberu môže byť aproximácia menej presná.

5. Literatúra

- [1] Bianchi, G. (2009) Inakosť vo verejnom priestore. In: Dušan Selko (Ed.): Psychológia zdravia v praxi. Bratislava, NÚSCH/MAURO, str. 83-91.
- [2] Bianchi G., Luha J. (2010): Náborová neparticipácia: sexuálna a telesná inakosť na okraji záujmu. Sociológia – zadané.
- [3] Chajdiak J. (2003): Štatistika jednoducho. Statis Bratislava 2003, ISBN 808565928-X.
- [4] Chajdiak J. (2005): Štatistické úlohy a ich riešenie v Exceli. STATIS, Bratislava, ISBN 80-85659-39-5.
- [5] Kanderová, M. – Úradníček, V. (2007): Štatistika a pravdepodobnosť pre ekonómov. 1. časť. OZ Financ, Banská Bystrica 2007, ISBN 978-80-969535-5-4.
- [6] Kanderová, M. – Úradníček, V. (2007): Štatistika a pravdepodobnosť pre ekonómov. 2. časť. OZ Financ, Banská Bystrica 2007, ISBN 987-80-696535-6-1.
- [7] Luha J. (1985): Testovanie štatistických hypotéz pri analýze súborov charakterizovaných kvalitatívnymi znakmi. Vydal Odbor Výskumu programov ČST a divákov v SR. Bratislava 1985.
- [8] Luha J. (2003): Matematickoštatistické aspekty spracovania dotazníkových výskumov. Štatistické metódy vo vedecko-výskumnej práci 2003, SŠDS, Bratislava 2003, ISBN 80-88946-32-8.
- [9] Luha J. (2005): Viacrozmerné štatistické metódy analýzy kvalitatívnych znakov. *EKOMSTAT* 2005, Štatistické metódy v praxi. SŠDS Trenčianske Teplice 22. – 27. 5. 2005.
- [10] Luha J. (2008): Prvotná štatistická analýza kvalitatívnych dát. *FORUM STATISTICUM SLOVACUM* 2/2008. SŠDS Bratislava 2008. ISSN 1336-7420.
- [11] Luha J. (2009): Korelácia javov. *FORUM STATISTICUM SLOVACUM* 2/2009. SŠDS Bratislava 2009. ISSN 1336-7420.
- [12] Pecáková I. (2008): Statistika v terénnych průzkumech. Proffessional Publishing, Praha 2008. ISBN 978-80-86946-74-0.
- [13] Řehák J., Řeháková B. (1986): Analýza kategorizovaných dat v sociologii. Academia Praha 1986.
- [14] Řezanková A. (2007): Analýza dat z dotazníkových šetření. Proffessional Publishing, Praha 2007. ISBN 978-80-86946-49-8.
- [15] Stehlíková B., Tirpáková A., Poměnková J., Markechová D. (2009): Metodologie výzkumu a statistická inference. *FOLIA UNIVERSITATIS AGRICULTURAE ET SILVICULRURAE MENDELIANAE BRUNENSIS*. Mendelova zemědělská a lesnická univerzita v Brně, 2009, ISBN 978-80-7375-362-7.
- [16] Stankovičová I., Vojtková M. (2007): Viacrozmerné štatistické metódy s aplikáciami. IURA EDITION, Bratislava 2007, ISBN 978-80-8078-152-1.

Adresa autorov:

Ján Luha, RNDr., CSc.
Ústav lekárskej biológie, genetiky a klinickej
genetiky LF UK a FNsP
Sasinkova 4, Bratislava
jan.luha@fmed.uniba.sk

Gabriel Bianchi, Doc., PhDr., CSc.
Ústav sociálnej a biologickej komunikácie
SAV
Dúbravská cesta 9, Bratislava
bianchi@savba.sk

Analýza vzájemného vztahu mezi výdaji na konečnou spotřebu a spotřebou domácností v ČR

Analysis of the mutual relation of the final consumption expenditure and the household consumption in Czech republic

Radka Martináková

Abstract:

Presented paper is focused on identification of the mutual relation of the business cycle models on the final consumption expenditure and the fluctuation of household consumption around their long-term trend. This analysis focuses on the period 1999–2009 and uses quarterly data provided by the Czech Statistical Office and Eurostat databases.

Key words: household consumption, business cycle, Hodrick-Prescott filter

Klíčové slova: spotřeba domácností, hospodářský cyklus, Hodrick-Prescottův filtr

1. Úvod

Problematikou analýzy hospodářského cyklu se ekonomové zabývají řadu let. V této souvislosti vznikla řada studií, jenž se snaží vysvětlit podstatu a příčiny ekonomických výkyvů. Můžeme například zmínit studii Českého statistického úřadu (2007), jenž se primárně zaměřuje na prozkoumání vztahů v průběhu ekonomických výkyvů na datech České republiky. Jako další studii věnovanou hospodářským cyklům lze uvést studii České národní banky (2005), v níž se autoři zabývají připraveností České republiky na přijetí eura z hlediska dlouhodobých ekonomických trendů, střednědobého vývoje ekonomické aktivity a strukturální podobnosti České republiky s ekonomikou eurozóny a celkovou spotřebu sledují pouze na makroekonomické úrovni. Další oblast týkající se hospodářského cyklu se dotýká problematiky sladění. Uvedme například doktorskou disertační práci Ing. Rozmahela (2006), která se zabývá připraveností kandidátských zemí z pohledu teorie optimálních měnových oblastí. V další studii, jenž je zaměřená na podobné téma, a to synchronizaci hospodářských cyklů uvnitř eurozóny (Weyerstrass et al., 2009), se autoři zmiňují o metodách sloužících pro odstranění trendu z časové řady, například uvádí Baxter–Kingův filtr, Christiano–Fitzgeraldův filtr a také Hodrick–Prescottův filtr, který je v tomto příspěvku využit.

Spotřeba domácností je důležitým ekonomickým ukazatelem, který se podílí na tvorbě hrubého domácího produktu (HDP) více než 50 % při uplatnění výpočtu výdajovou metodou. Od roku 1999 (počátek námi sledovaného období) představoval podíl výdajů domácností na tvorbě HDP celkem 53,9 % a tento podíl se rok od roku zvyšuje, což může být chápáno jako postupné utváření standardního tržního prostředí. Toto tvrzení potvrzuje dominantní postavení spotřebitele ve vyspělé tržní ekonomice, a proto nebude v následující práci pro analýzu hospodářského cyklu používán HDP, ale pouze ukazatel spotřeba domácností (Fuchs, Tuleja, 2003). Pojem spotřeba domácností můžeme definovat jako hodnotu všech výrobků a služeb užitých domácnostmi pro uspokojení individuálních potřeb uhrazených z důchodů domácností a pořízených nákupem, dary i formou naturální spotřeby (ČNB, 2010).

Cílem předkládaného příspěvku je posoudit a vyhodnotit, zda mezi hospodářským cyklem výdajů na konečnou spotřebu v České republice a fluktuacemi výdajů domácností v České republice okolo jejich dlouhodobého trendu existuje významná statistická závislost.

2. Metodika

Abychom mohli z časové řady odstranit trend je nutné použít vhodnou detrendovací techniku. Dle Canovy (1998) lze detrendovací techniky rozdělit do dvou skupin, na ekonomické a na statistické. Statistický přístup předpokládá, že trend a cyklus jsou nepozorovatelné, ale k identifikaci těchto dvou komponent využívají rozdílné statistické předpoklady. Mezi statistické techniky řadíme deterministické modely, první difference nebo například proceduru Beveridge a Nelsona. V ekonomickém přístupu je výběr trendu diktován ekonomickým modelem, preferencemi výzkumníka nebo řešeným problémem. Předpokládá se, že existuje pouze trendová a cyklická složka a použitá data byla nejprve sezónně očištěná nebo že sezónní a cyklická složka jsou považovány jen za složku jednu a nepravidelné fluktuace mají nepatrný význam. Mezi ekonomické metody řadíme například Hodrick-Prescottův (HP) filtr nebo Baxter-Kingův (BK) filtr.

V příspěvku budeme pracovat s časovou řadou y_t , kterou pro účely analýzy hospodářského cyklu budeme transformovat logaritmem $Y_t = \ln(y_t)$, $t=1, \dots, n$, a provedeme rozložení na růstovou komponentu g_t a na cyklickou komponentu c_t , dle vztahu (1) :

$$Y_t = g_t + c_t, \quad t=1, \dots, n. \quad (1)$$

Pro detrendování časové řady bude v příspěvku využito Hodrick-Prescottova filtru (1980). Podstatou jednorozměrného HP filtru je rozklad nestacionární časové řady Y_t na trendovou g_t a cyklickou složku c_t . HP filtr bere při extrakci trendu v úvahu 2 důležitá kritéria: velikost reziduí a míru hladkosti trendu. HP filtr se snaží minimalizovat následující výraz:

$$\min_{\{g_t\}_{t=1}^T} \sum_{t=1}^T (Y_t - g_t)^2 + \lambda \sum_{t=1}^T [(g_{t+1} - g_t) - (g_t - g_{t-1})]^2, \quad (2)$$

kde cyklická komponenta $c_t = Y_t - g_t$ představuje odchylky od dlouhodobého trendu a její hladkost je měřena pomocí kvadrátu druhých diferencí

Literatura uvádí doporučenou hodnotu parametru λ , která však není závazná. Pro roční data je stanovena obvyklá hodnota $\lambda = 100$, pro pololetní data $\lambda = 400$ a pro čtvrtletní data je velmi často užívána hodnota $\lambda = 1600$ (Hodrick, Prescott, 1980). V příspěvku bude využito doporučené hodnoty pro čtvrtletní data.

Mezi výhody HP filtru řadíme fakt, že nepůsobí ztrátu pozorování, to znamená, že očištěná řada má pozorování pro všechny časové okamžiky jako původní řada. Další výhodou je nenáročnost na vstupní data a lze jej poměrně jednoduše aplikovat na jakoukoli časovou řadu. Nevýhodou při použití detrendovací techniky Hodrick-Prescottova filtru je tzv. „problém koncových bodů“. Tedy, že počátek a konec časové řady nejsou vyhlazeny a jsou na obou koncích časové řady vychýleny. Například pokud poslední pozorování vykazuje známky expanze, trend je tažen na konci časové řady nahoru. Tento problém bývá obvykle řešen prodloužením časové řady do nadcházejícího období a tím se „vyhlazovací chyba“ přesune do budoucnosti (Bonenkamp aj., 2001).

Pro identifikaci bodů zvratu, tj. vrcholu a dna, budou v příspěvku využita pravidla dle Canovy (1999). Canova rozlišuje pro určení bodů zvratu 2 části datovacího pravidla. První pravidlo definuje dno, které nastává v případě, kdy po období dvou po sobě následujících čtvrtletí referenčního cyklu následuje čtvrtletí růstu, toto tvrzení lze vyjádřit vztahem: $c_{t-2} > c_{t-1} > c_t < c_{t+1}$. Podobně vrchol nastane v případě, že dvě po sobě jdoucí čtvrtletí růstu jsou následována čtvrtletím poklesu, vztah: $c_{t-2} < c_{t-1} < c_t > c_{t+1}$. Canova dále definuje druhé pravidlo, které by mělo sloužit k minimalizaci rizika tzv. falešných bodů zvratů. Toto pravidlo vybírá dno (vrchol) pokud po sobě následovaly alespoň dva absolutní poklesy (růsty), cyklické komponenty ve dvou po sobě jdoucích čtvrtletích během období tří čtvrtletí, to jest $c_t < (>) 0$ a $c_{t-1} < (>) 0$ nebo pokud $c_{t+1} < (>) 0$ a $c_t < (>) 0$. Zavedení obou pravidel se snaží zabránit předčasné identifikaci bodů zvratů (tzv. falešné signály).

3. Empirická část

Pro účely analýzy hospodářského cyklu výdajů na konečnou spotřebu České republiky byla použita čtvrtletní data výdajů na konečnou spotřebu pro období 1999Q01 – 2009Q03. Data byla již sezónně očištěná, v jednotkách milionů národní měny, vztažná k roku 2000. Zdrojem těchto dat byla databáze Eurostat.

Pro analýzu fluktuací vydání a spotřeby domácností České republiky kolem dlouhodobého trendu byla rovněž použita čtvrtletní data – průměry spotřebního vydání na osobu v Kč za měsíc pro srovnatelné období. Zdrojem dat pro účely analýzy byl Český statistický úřad, který však na rozdíl od databáze Eurostat poskytuje data neočištěná. Bylo tedy třeba nejprve data spotřeby domácností upravit očištěním o inflaci a sezónnost. Pro očištění o sezónnost byla využita triviální metoda empirických sezónních indexů (Hindls aj. 2000). V dalším kroku byla data přepočtena na srovnatelné vyjádření jako celková spotřeba.

Nejprve byl nalezen růstový cyklus a poté identifikovány body zvratu pomocí pravidel definovaných Canovou u obou sledovaných veličin. Jednotlivá čtvrtletí, v nichž byly identifikovány body zvratu jsou uvedeny v tabulce 1 a 2. Tato analýza prokázala, že vývoj obou sledovaných ukazatelů, tj. celkové spotřeby a spotřeby domácností, není v průběhu časové řady vždy shodný. Tuto skutečnost dokládá i obrázek 1.

Tabulka 1: Identifikované body zvratu u celkové spotřeby

dno	vrchol
2001Q4	2003Q3
2005Q1	2007Q1
2008Q4	

Zdroj : vlastní výpočet

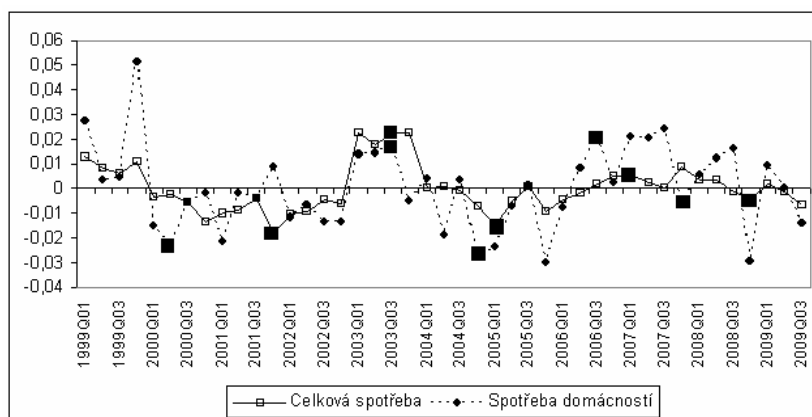
Tabulka 2: Identifikované body zvratu u spotřeby domácností

dno	vrchol
2000Q2	2003Q3
2004Q4	2006Q3
2007Q4	

Zdroj : vlastní výpočet

Za celé sledované období byl zaznamenán pouze jeden shodný bod zvratu, a to vrchol ve 3. čtvrtletí roku 2003. Největší rozdíl můžeme zpozorovat u prvního dna, které u spotřeby domácností nastává ve 2. čtvrtletí roku 2000, ale u celkové spotřeby až o 6 čtvrtletí později, tedy ve 4. čtvrtletí roku 2001. V dalších případech není rozdíl mezi body zvratu tolik výrazný, například v případě dna u spotřeby domácností určeného v posledním čtvrtletí roku 2004 a u celkové spotřeby v 1. čtvrtletí roku 2005 je rozdíl pouze 1 čtvrtletí. U vrcholu stanoveného u spotřeby domácností na 3. čtvrtletí roku 2006 a u celkové spotřeby na 1. čtvrtletí roku 2007 se jedná o rozdíl dvou čtvrtletí a v posledním případě je dna u spotřeby domácností stanoveno na poslední čtvrtletí roku 2007 a u celkové spotřeby na poslední čtvrtletí roku 2008, zde je rozdíl jednoho roku.

Ve všech případech, vyjma shodného vrcholu ve 3. čtvrtletí roku 2003, následoval nejdříve bod zvratu u spotřeby domácností a následně s různými zpožděními následovaly body zvratu u celkové spotřeby.



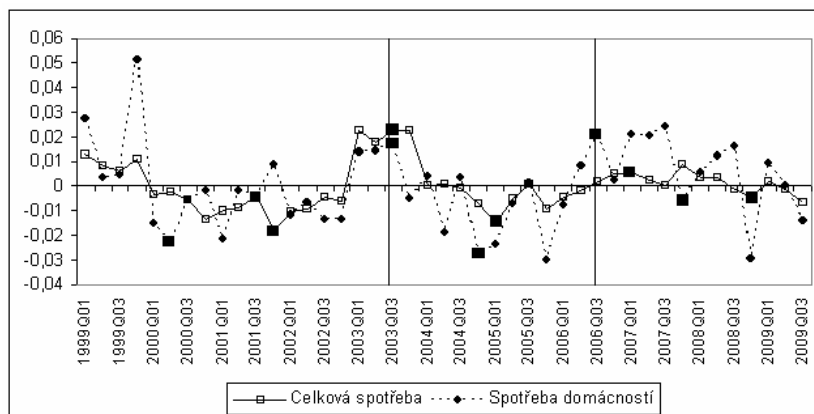
Obrázek 1: Vývoj celkové spotřeby a spotřeby domácností ČR

Pro prověření těsnosti vzájemné vazby celkové spotřeby a spotřeby domácností v letech 1999–2009 byla provedena korelační analýza. Pro testování těsnosti jednotlivých regresních koeficientů byl v příspěvku použit t-test. Existuje více přístupů k vyhodnocení testů, jako například klasický test pomocí kritických hodnot nebo test pomocí p-hodnoty, který nabízí software a bude použit pro účely tohoto příspěvku (Hušek, 2007). P-hodnota pod 0,05 indikuje statisticky významnou nenulovou korelaci na 95% konfidenční úrovni a v našem případě poslouží k rozhodování o statistické významnosti odhadované korelace

Celkový koeficient korelace nabývá hodnoty 0,54 (p -hodnota = 0,0002); lze tedy tvrdit, že na základě provedené korelační analýzy mezi růstovým cyklem modelovaným na výdajích na konečnou spotřebu a fluktuacemi vydání a spotřeby domácností kolem dlouhodobého trendu existuje střední závislost.

Pro bližší rozlišení vzájemné vazby mezi sledovanými ukazateli byla časová řada rozdělena dle „shodných“ bodů zvratu do tří úseků, viz obrázek 2. Jelikož shodný bod zvratu nastal pouze v jednom okamžiku, zbylé „shodné“ body zvratu jsou určeny ve středu mezi jednotlivými body zvratu celkové spotřeby a spotřeby domácností.

První úsek zahrnuje období od počátku sledovaného období tj. 1. čtvrtletí roku 1999 do 3. čtvrtletí roku 2003 (celkem 19 čtvrtletí), kdy nastává shodný bod zvratu u obou sledovaných veličin. Druhý úsek počíná posledním čtvrtletím roku 2003 a trvá do 3. čtvrtletí roku 2006 (celkem 12 čtvrtletí). Poslední úsek (celkem 12 čtvrtletí) začíná 4. čtvrtletím roku 2006, který se nachází ve středu mezi body zvratu spotřeby domácností a celkové spotřeby, tj. vrchol spotřeby domácností nastal ve 3. čtvrtletí roku 2006 a vrchol celkové spotřeby v 1. čtvrtletí roku 2007. Poslední úsek je ukončen koncem sledovaného období 3. čtvrtletím roku 2009.



Obrázek 2: Rozdělení časové řady do tří úseků

Obrázek 2 uvádí vzájemnou závislost celkové spotřeby a spotřeby domácností, svislé čáry pak značí rozdělení do tří úseků dle „shodných“ bodů zvratu.

Tabulka 1: Korelační koeficienty celkové spotřeby a spotřeby domácností

úsek	od	do	n	Kor. koeficient	p-hodnota
1	1999Q01	2003Q03	19	0,62	0,0043
2	2003Q04	2006Q03	12	0,44	0,1536
3	2006Q04	2009Q03	12	0,42	0,1732

Pozn.: p-hodnota $< \alpha$ indikuje statisticky významnou korelaci na hladině významnosti $\alpha\%$
zdroj : vlastní výpočet

Jak je z tabulky patrné nejvyšší hodnoty nabývá korelační koeficient v 1. úseku časové řady a to 0,62. P-hodnota je v tomto úseku rovna 0,0043, je tedy zřejmé, že vzájemnou vazbu posuzovaných veličin v 1. úseku můžeme označit za statisticky významnou. Ve 2. úseku byl vypočten korelační koeficient 0,44 (p-hodnota = 0,1536), ve 3. úseku je těsnost také nižší a to 0,42 (p-hodnota = 0,1732).

3. Závěr

Předkládaný příspěvek zkoumal vztah mezi hospodářským cyklem modelovaným na výdajích na konečnou spotřebu a fluktuacemi vydání a spotřeby domácností v ČR kolem jejich dlouhodobého trendu.

Z provedené analýzy bylo zjištěno, že celková spotřeba má až na určité výjimky podobný průběh se spotřebou domácností. Vypočtený celkový koeficient korelace za celkové období vykazuje střední závislost mezi pozorovanými veličinami. Při bližším zkoumání, kdy byla časová řada rozdělena na 3 úseky, vykazovaly posuzované veličiny nejtěsnější lineární závislost v období 1999Q01 až 2003Q03, kdy bylo sledováno celkem 19 čtvrtletí. V tomto úseku byla prokázána statisticky vysoce významná korelace na 4,3% hladině významnosti. Ve 2. a 3. úseku, kdy bylo pozorováno 12 čtvrtletí, vypovídá koeficient korelace spíše o slabší závislosti mezi pozorovanými veličinami, přičemž p-hodnota poukazuje na statistickou nevýznamnost. V těchto úsecích lze tedy pozorovat určitou asymetrii mezi oběma spotřebami.

Předkládaný příspěvek vznikl za podpory interního grantového projektu 71/2010 “Stabilizační funkce monetární politiky v souvislostech hospodářského cyklu České republiky”.

4. Literatura

- [1] BONENKAMP, J., JACOBS, J., KUPER, G.H. 2001. Measuring Business Cycles in the Netherlands, 1815 - 1913: A comparison of Business Cycle Dating Methods. *SOM Research Report, No. 01C25*, Systems, Organisation and Management, Groningen. University of Groningen [online].
- [2] CANOVA, F. 1998. De-trending and business cycle facts, *Journal of monetary economic*, vol. 41, pp. 533-540.
- [3] CANOVA, F. 1999. Does De-trending Matter for the Determination of the Reference Cycle and Selection of Turniny Points?, *The Economic Journal*, Vol. 109, No. 452 (Jan., 1999), pp. 126-150.
- [4] Česká národní banka [online]. 2005. [cit. 2010-04-28]. Analýzy stupně ekonomické sladěnosti České republiky s Eurozónou. Dostupné

- z WWW:<http://www.cnb.cz/miranda2/export/sites/www.cnb.cz/cs/menova_politika/strategicke_dokumenty/download/analyzy_sladenosti_2005.pdf>.
- [5] Česká národní banka. [online]. 2010. [cit. 2010-03-20]. Indikátory měnového a hospodářského vývoje. Dostupné z WWW: <http://www.cnb.cz/m2export/sites/www.cnb.cz/cs/statistika/dalsi_statistiky/INDIKATORY_CS.PDF>.
- [6] Český statistický úřad [online]. 2007. [cit. 2010-04-20]. Monitorování a analýza hospodářského cyklu. Dostupné z WWW: <<http://www.czso.cz/csu/2006edicniplan.nsf/p/1132-06>>.
- [7] FUCHS, K.; TULEJA, P. 2003. *Makroekonomie*. 1. Brno : Masarykova univerzita, 2003, 283 s. ISBN 80-210-3073-9.
- [8] HINDLS, R.; HRONOVÁ, S.; NOVÁK, 2000. I. *Metody statistické analýzy pro ekonomy*. 2. Praha : Management Press,2000, 259 s. ISBN 80-7261-013-9.
- [9] HODRICK, R.J., PRESCOTT, E.C.1980. Post-war U.S. Business Cycles: An Empirical Investigation, mimeo, Carnegie-Mellon University, Pittsburgh, PA, 24pp.
- [10] HUŠEK, R. 2007. *Ekonometrická analýza*. 1. vyd. Praha : Oeconomica, 367 s. ISBN 978-80-245-1300-3.
- [11] ROZMAHEL, P. 2006. Metodologické aspekty posuzování připravenosti kandidátských zemí pro vstup do eurozóny z pohledu teorie optimálních měnových oblastí. Disertační práce. PEF, MZLU v Brně,2006, 154 s.
- [12] WEYERSTRASS, K., et al. 2009. Business Cycle Synchronisation with(in) the Euro Area:in Search of a Euro Effect.*Open economies review*. 4, 2009, s. 1-20. ISSN 0923-7992.

Adresa autora:

Radka Martináková
Ústav financí, PEF MU v Brně
Zemědělská 1
Česká republika
613 00 Brno
xmartina@node.mendelu.cz

Faktory internacionalizácie malého a stredného podnikania Factors of internationalization of small and medium enterprise

Ladislav Mura

Abstract: The current development of the world economy depends on simultaneous integration and globalization tendencies. In this cross-pressure, businesses about to expand their territory by entering new markets are also facing the competitive environment on foreign markets. Our results indicate that the Slovak firms consider themselves more successful on the markets of Central, Eastern and Southern Europe and other transforming countries. In many cases they use the advantage of similar economic, political and historical development with those countries. Managing directors of the interviewed companies are highly interested in internationalization.

Key words: internationalization, factors of internationalization, SAS

Kľúčové slová: internacionalizácia, faktory internacionalizácie, SAS

1. Úvod

Súčasný rozvoj svetovej ekonomiky závisí v mnohom od postupujúcej globalizácie a na integračných procesoch. [1] V silnom konkurenčnom boji, ktoré je príznačné pre súčasné podnikateľské prostredie, sa podnikateľské subjekty usilujú o budovanie stabilnej pozície s perspektívou ďalšieho rozvoja. Svoj rozvoj čoraz viac smerujú aj mimo domáceho podnikateľského prostredia s cieľom zvýšiť objem predaja a následne svojho zisku, získania konkurenčnej výhody a rastu.

Pri výskume v ekonomických vedách sa čoraz častejšie využíva SAS (Statistical Analysis System). SAS je rozsiahly programový systém pre prácu s dátami, pre ich analýzu a prezentáciu. Jeho hlavná výhoda spočíva v tom, že užívateľské prostredie je až na grafické detaily rovnaké pre rôzne operačné systémy a je možné na rôznych typoch počítačov s rôznymi operačnými systémami zdieľať údaje a aplikácie v SASe. [3] SAS umožňuje prácu s dátami veľkého rozsahu. Užívateľ môže použiť na vyhodnotenie údajov široké spektrum štatistických metód od najjednoduchších po zložité a tiež špeciálne štatistické metódy.

Preto cieľom príspevku je na základe štatistických metód prostredníctvom programu SAS analyzovať faktory internacionalizácie podnikania malých a stredných podnikov.

2. Materiál a metódy

Predmetom vedeckého príspevku je problematika internacionalizácie podnikania malých a stredných podnikov v potravinárstve. V príspevku skúmame faktory internacionalizácie, ktorých poznanie je pre úspešné riadenie podnikateľských subjektov v podnikateľskom prostredí Európskej únie nevyhnutné. Vzhľadom na limitovaný rozsah príspevku uvádzame iba vybrané faktory internacionalizácie s ukážkou štatistického spracovania pomocou SASu.

Splnenie vytýčeného cieľa si vyžiadalo uskutočniť primárny výskum. Hlavnými použitými metódami v primárnom výskume boli dotazníková metóda a technika riadeného rozhovoru s predstaviteľmi manažmentu oslovených podnikov. Dotazník pozostával z 24 otázok. Spracovanie dotazníkov sa uskutočnilo v dvoch etapách. Prvá etapa spracovania získaných údajov po ošetroaní údajovej základne sa uskutočnila prostredníctvom programového balíka MS Office Excel. Druhá etapa spracovania údajov sa uskutočnila za pomoci špeciálneho štatistického programu SAS. Z veľkého množstva metód, ktoré ponúka

SAS, sme pri hodnotení dotazníkového zisťovania použili chí - kvadrát test nezávislosti a Friedmanov test [3].

Kontingenčné koeficienty merajú silu závislosti medzi premennými v kontingenčných tabuľkách v škále od 0 (žiadna závislosť) do 1 (perfektná závislosť). Najčastejšie používanými koeficientmi pre hodnotenie miery závislosti kvantitatívnych znakov sú Pearsonov kontingenčný koeficient C a Cramerov kontingenčný koeficient V. Z testov najpoužívanejší je Chí-kvadrát test. SAS okrem hodnoty testovacej štatistiky uvádza aj p hodnotu, t.j. najnižšiu hladinu významnosti, na ktorej nulovú hypotézu zamietame. Teda ak je p hodnota vyššia ako hladina významnosti 0,05, nulovú hypotézu o závislosti zamietame, t.j. neexistuje štatisticky preukazná závislosť.

Friedmanov test je neparametrickou obdobou dvojfaktorovej analýzy rozptylu s jedným pozorovaním v podtriade. Podstata testu spočíva v porovnaní mediánov jednotlivých skupín podľa úrovne faktora. Test odpovedá na otázku, či vo vzorke zistené rozdiely môžu byť iba náhodné, t. j. medzi premennými nie je vzťah, alebo sú štatisticky významné, t.j. medzi premennými vzťah existuje.

3. Výsledky a diskusia

Dotazníkový výskum sa uskutočnil v Nitrianskom kraji, ktorý možno označiť ako poľnohospodársky región s previazanosťou na infraštruktúrne dobre vybavený potravinársky priemysel [2]. V skúmanom kraji je 190 štatisticky evidovaných podnikateľských subjektov malého a stredného podnikania v potravinárstve. Na otázky v dotazníku odpovedalo 46 podnikov zo všetkých siedmych okresov kraja. Tvoria 24, 21 %, t.j. takmer štvrtinu zo všetkých evidovaných subjektov. Keďže sme neoverovali reprezentatívnosť vzorky vzhľadom na priestorové rozloženie podnikateľských subjektov malého a stredného podnikania v potravinárstve v Nitrianskom kraji ako aj ich veľkostnú štruktúru, považujeme náš výsledok za štúdiu. Sídlo a počet podnikov zúčastnených na výskume prezentuje tabuľka 1.

Tabuľka 1 Rozmiestnenie a počet podnikov podľa sídla

Sídlo okresu	Počet podnikov na výskume
Komárno	9
Levice	7
Nitra	10
Nové Zámky	10
Šaľa	4
Topoľčany	4
Zlaté Moravce	2

Zdroj: vlastný výskum

Z hľadiska kategorizácie podnikov zúčastnených na výskume podľa charakteristiky priemerný evidenčný počet zamestnancov je situácia nasledovná (tabuľka 2):

Tabuľka 2 Podiel podnikov v jednotlivých kategóriách

Kategória podniku	Podiel podnikov na výskume
Mikropodnik (0-9 zamestnancov)	36,96 %
Malý podnik (10 – 49 zamestnancov)	52,17 %
Stredný podnik (50 – 249 zamestnancov)	10,87 %

Zdroj: vlastný výskum

Najväčší počet podnikateľských subjektov na uskutočnenom výskume tvorili malé podniky s počtom zamestnancov v rozmedzí od 10 do 49 zamestnancov (24 podnikov; 52,17 %), druhú najväčšiu skupinu tvorili mikropodniky s počtom zamestnancov v rozmedzí od 0 do 9 zamestnancov (17 podnikov; 36,96 %) a najmenšiu skupinu tvorili stredné podniky s počtom zamestnancov v rozmedzí od 50 do 249 zamestnancov (5 podnikov; 10,87 %). V kategórii stredný podnik počet zamestnancov v analyzovaných podnikoch nepresiahol 75 zamestnancov.

Zo skúmaných podnikateľských subjektov vyše tretina – 17 podnikov (36,96 %) podnikateľsky pôsobí aj na zahraničných trhoch, kým 29 podnikov (63,04 %) neinternacionalizuje svoje podnikanie a pôsobí len v domácom podnikateľskom prostredí.

V ďalšej časti výskumu sme analyzovali príčiny a motívy internacionalizácie podnikania, resp. zotrvania podnikov iba na domácom trhu.

Podnikateľské subjekty, ktoré internacionalizujú svoje podnikanie, najčastejšie uviedli z hľadiska lokalizácie zahraničného trhu podnikateľské prostredie Európskej únie – 38 podnikov (82,61 %), nasledované podnikateľským prostredím Európy, mimo krajín EÚ – 8 podnikov (17,39 %). Z krajín Európskej únie išlo najčastejšie o okolité štáty (Maďarsko – 18 podnikov, Česko – 14 podnikov, Poľsko – 6 podnikov), z ostatných členských štátov Nemecko – 5 podnikov, Taliansko – 3 podniky, Francúzsko – 2 podniky, v ojedinelých prípadoch išlo o ostatné členské štáty EÚ. Podnikateľské prostredie Ameriky, Ázie či iné podnikateľské prostredie neuviedol žiaden podnik.

V rámci internacionalizácie podnikania sme sa zaujímali o výšku podielu zahranično – obchodných aktivít na celkovej podnikateľskej činnosti. Výsledky uvádza tabuľka 3.

Tabuľka 3 Podiel zahranično – obchodných aktivít na podnikateľskej činnosti

Výška podielu	Podiel podnikov
Do 25 %	76,47 %
26 – 50 %	23,53 %

Zdroj: vlastný výskum

Z tabuľky 3 je zrejmé, že najpočetnejšiu skupinu (13 podnikov; 76,47 %) tvoria podniky, ktorých objem zahranično – obchodných aktivít nepresiahne 25 % z celkovej podnikateľskej činnosti. V druhom intervale, teda s podielom od 26 do 50 %, sa nachádzajú 4 podniky (23,53%). Ani u jedného respondenta nepresiahla 50%-tnú hranicu zahranično – obchodná podnikateľská aktivita z celkovej podnikateľskej činnosti.

Tabuľka 4 Spôsob internacionalizácie podnikania

Spôsob internacionalizácie podnikania	Podiel podnikov
priamy export	62,50 %
nepriamy export	8,33 %
dcérsky podnik	12,50 %
subdodávateľ iného podnikateľského subjektu	16,67 %

Zdroj: vlastný výskum

V ďalšej časti výskumu sme zisťovali spôsob internacionalizácie podnikania podnikateľských subjektov malého a stredného podnikania v potravinárskom priemysle. Z tabuľky 4 vyplýva, že najčastejšie využívaným spôsobom internacionalizácie podnikania v skúmaných podnikoch Nitrianskeho kraja, teda prieniku na zahraničné trhy je priamy export. Túto odpoveď uviedlo celkom 15 podnikov. Druhým, ale za prvou možnosťou výrazne zaostávajúcim, spôsobom internacionalizácie podnikania je subdodávateľ pre iný podnikateľský subjekt. Túto možnosť využívajú 4 podniky. Dcérsky podnik v zahraničí založili 3 podniky a nepriamy export využívajú 2 podniky. Niektoré podniky uviedli aj viac spôsobov prieniku na zahraničné trhy. Možnosť licencie, pobočky v zahraničí a franšízing ako spôsoby internacionalizácie podnikania neuviedol žiaden podnik.

Výskum sme zamerali aj na zisťovanie existencie závislosti medzi akceptáciou internacionalizácie podnikania ako súčasného trendu v globalizujúcom sa svete a vývojom objemu predaja podnikov na zahraničné trhy vybraných podnikov v Nitrianskom kraji. Skutočnosť, či podnik internacionalizuje svoje podnikanie, t.j. akceptuje zapojenie sa do zahranično – obchodných aktivít, sme zisťovali otázkou: „*Vykonáva podnikateľský subjekt podnikateľskú činnosť aj na zahraničných trhoch?*“. Respondenti mali na výber z dvoch odpovedí: áno, nie. Vývoj objemu predaja u sledovaných podnikov sme zisťovali otázkou: „*Ako sa zmenil celkový objem Vášho predaja za posledné 3 roky?*“, pričom respondenti mali možnosť odpovedať jednou z týchto možností: rastie, stagnuje alebo klesá.

Pomocou chí-kvadrát testu nezávislosti sme testovali hypotézu, či je závislosť medzi zapojením sa podniku do zahranično-obchodných aktivít a zmenou objemu predaja podniku za posledné tri roky. Na základe výsledkov chí-kvadrát testu ($p = 0,4409$) nulovú hypotézu H_0 o nezávislosti nemôžeme zamietnuť. Čiže v podnikoch, ktoré hodnotíme v rámci našej štúdie neexistuje významná závislosť medzi akceptovaním internacionalizácie ako súčasného trendu v globalizujúcom sa svete a vývojom objemu predaja na zahraničné trhy. Výsledky analýzy v programe SAS prezentuje tabuľka 5.

Tabuľka 5 Testovanie závislosti

Štatistika	Hodnota	P-hodnota
Chí-kvadrát	12.4124	0.4409
Pearsonov koeficient	0.2318	
Kontingenčný koeficient	0.2258	
Cramerov koeficient	0.2318	

Zdroj: dotazníkový výskum – výstup zo softvéru SAS

Prostredníctvom Friedmanovho testu sme testovali, či v skúmanej vzorke respondentov pri odpovediach na jednotlivé otázky ohľadom faktorov, ktoré sú dôvodom internacionalizácie ich podnikania, sú zistené rozdiely iba náhodné alebo sú štatisticky významné. Dôvody internacionalizácie podnikania sme zisťovali otázkou: „*Dôvod podnikania v zahraničnom podnikateľskom prostredí*“. Menovaným faktorom - jednorazový obchod, pravidelné objednávky, skúsenosti manažmentu s medzinárodným podnikaním, silná konkurencia v domácom podnikateľskom prostredí, zvýšenie konkurencieschopnosti podniku - mali respondenti priradiť známku od 1 do 5, pričom známka 1 znamenala najvýznamnejší vplyv a známka 5 najmenší vplyv. Testovali sme hypotézu, či je štatisticky významný rozdiel medzi jednotlivými faktormi, ktoré sú dôvodmi internacionalizácie podnikania jednotlivých podnikov.

Na základe výsledku chí-kvadrát testu ($p = 0,0004$) môžeme hypotézu H_0 o rovnakom vnímaní faktorov, ktoré sú dôvodmi internacionalizácie ich podnikania, zamietnuť s pravdepodobnosťou na zvolenej hladine významnosti 0,05. „Prijmame“ hypotézu H_1 , že dotazované podniky vnímajú jednotlivé faktory, ktoré sú dôvodmi internacionalizácie ich podnikania, rozdielne. Znamená to, že podniky nie sú motivované rovnakými dôvodmi k internacionalizácii ich podnikania.

Z pohľadu manažmentu podnikov bol výskum zameraný na zistenie skutočnosti, či existuje závislosť medzi jednotlivými medzinárodnými skúsenosťami manažmentu predtým, ako začal podnikateľský subjekt uskutočňovať zahranično-obchodné aktivity a vývojom objemu predaja podnikov na zahraničné trhy. Medzinárodné skúsenosti manažmentu boli modelované otázkou či „*Mal manažment podniku skúsenosti s medzinárodným obchodom skôr, ako podnik začal uskutočňovať zahranično-obchodné aktivity*“. Možné odpovede boli áno a nie. Testovali sme hypotézu, že neexistuje závislosť medzi jednotlivými medzinárodnými skúsenosťami manažmentu a vývojom objemu predaja na zahraničné trhy.

Aj v tomto prípade bola p hodnota chí-kvadrát testu väčšia ($p = 0,5357$) ako hladina významnosti 0,05. Na základe výsledkov testu nezamietame hypotézu H_0 a teda neexistuje v skúmaných podnikoch štatisticky významná závislosť medzi jednotlivými medzinárodnými skúsenosťami manažmentu a vývojom objemu predaja skúmaných podnikov na zahraničné trhy.

4. Záver

Proces internacionalizácie sa dotýka bezprostredne alebo okrajovo každého potravinárskeho podniku vzhľadom na skutočnosť, že možnosti expanzie na domácom trhu sú obmedzené. Možnou alternatívou na zhodnotenie zdrojov je podnikateľsky pôsobiť na zahraničných trhoch. Internacionalizácia podnikania však okrem zaujímavých možností získať nové trhy prináša so sebou aj riziká. Manažment podniku musí zvážiť svoje rozhodnutie o internacionalizácii podnikania, adekvátne sa pripraviť a až následne prenikáť na nové trhy. V skúmanom súbore podnikov sme zistili, že manažmenty podnikov majú záujem o rozvoj internacionalizácie svojho podnikania. Na niekoľkých príkladoch testovania hypotéz sme dokázali, že výpočtová technika s adekvátnym softvérovým vybavením – programom SAS – výrazne uľahčuje prácu s údajmi získanými z realizovaného výskumu.

5. Literatúra

- [1] HORSKÁ, E. a kol. 2008. Internacionalizácia agropotravinárskych podnikov SR. Nitra: SPU, 2008, 234 s., ISBN 978-80-552-0136-8
- [2] MURA, L. 2010. Internacionalizácia podnikania malých a stredných podnikov na Slovensku. Dizertačná práca. Nitra, 2010, 157 s.

[3] STEHLÍKOVÁ, B. 2003. Štatistická analýza systémom SAS. Nitra: SPU, 2003, 127 s., ISBN 80-8069-221-1

Adresa autora:

Ladislav Mura, Ing. et Bc.
Dubnický technologický inštitút v Dubnici nad Váhom
Ul. Sládkovičova 533/20
018 41 Dubnica nad Váhom
E mail: ladislav.mura@gmail.com

Komparácia regiónov SR a ČR z pohľadu vybraných demografických ukazovateľov

Comparison of regions of the Slovak Republic and Czech Republic from the viewpoint of chosen demographic indicators

Zuzana Poláková, Peter Obtulovič

Abstrakt

The issue of a demographic development and different views of courses, causes and also consequences of the demographic processes consider it very important to analyze and compare the population progress with other countries or regions. The main aim is the comparison of chosen regions in Slovak Republic and Czech Republic in the area of demography. The development in particular regions was monitored and analyzed in the period from 2000 to 2007. In these regions we were studying the progress of a population and the primary demographic processes including natality, mortality and natural growth increment. Also, the indexes of economic and social status – the youth economic dependence index, the old economic dependence index, the economic load index.

Key words: demographic processes, natality, mortality, development of population

Kľúčové slová: demografické procesy, pôrodnosť, úmrtnosť, vývoj populácie

1. Úvod

Demografický vývoj každej krajiny je v istom zmysle ukazovateľom zmien, ktoré sa uskutočňujú v rámci jej ekonomického a sociálneho rozvoja. Dôležitú a neoddeliteľnú časť demografie tvorí regionálna demografia, ktorá sa zaoberá štúdiom demografickej reprodukcie v rôzne vymedzených regiónoch. Pre spoločnosť je dôležité poznať demografické javy a procesy nielen na úrovni štátu, ale musí sa zaoberať aj menšími územnými jednotkami na úrovni krajov, okresov a obcí. Pri skúmaní regiónov do popredia vystupujú ich všeobecné a charakteristické vlastnosti, ale aj spoločné znaky, ktoré sa opakujú v mnohých vyčlenených územiach.

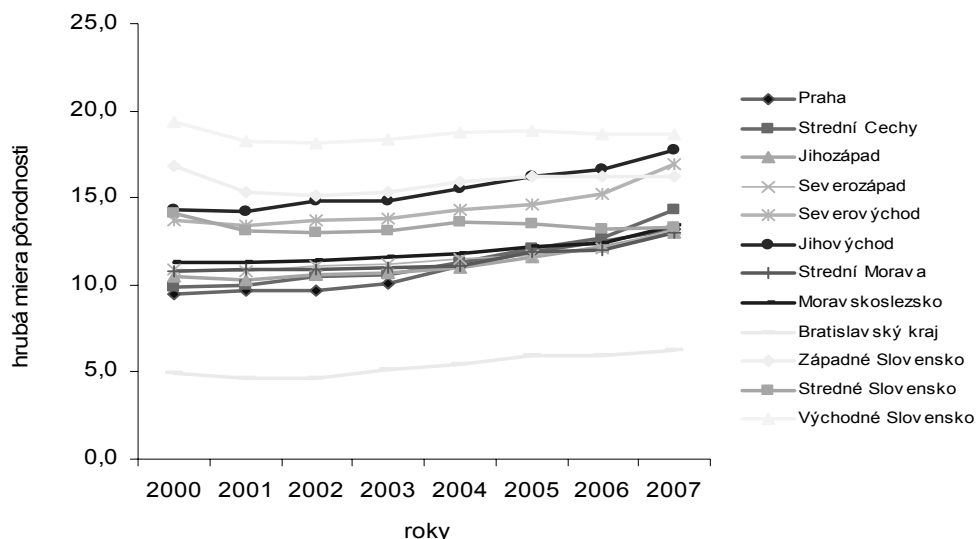
Okolo roku 2000 oproti predchádzajúcemu obdobiu nastalo výrazné zníženie pôrodnosti v celoeurópskom meradle. Populácia začala rapídne starnúť, menila sa veková štruktúra. Cieľom príspevku je analyzovať a porovnať súčasnú demografickú situáciu, t.j. ako sa situácia v demografickom správaní zmenila po tomto nepriaznivom období. Analyzované budú vybrané regióny Slovenskej a Českej republiky a bude poukázané na ich diverzitu z pohľadu vybraných demografických ukazovateľov. Porovnávaných bude 12 regiónov Slovenskej a Českej republiky. Českú republiku tvorí 8 regiónov: Praha, Stredné Čechy, Juhozápad, Severozápad, Severovýchod, Juhovýchod, Stredná Morava a Moravskosliezsko. Na Slovensku porovnáваме 4 regióny: Bratislavský kraj, Západné Slovensko, Stredné Slovensko a Východné Slovensko. Cieľom bude zhodnotiť a porovnať pôrodnosť, úmrtnosť a prirodzený prírastok obyvateľstva v jednotlivých regiónoch a poukázať na vývojové tendencie indexov závislosti I, II a indexu ekonomického zaťaženia. V ČR sa podľa Dufeka, J. a Minaříka, B.(2007) prechod na nový štýl života spolu s niektorými ďalšími faktormi prejavil negatívne hlavne u pôrodnosti. V roku 1994 bol prvýkrát v histórii Českej republiky zaznamenaný úbytok obyvateľstva prirodzenou mierou (Rychtaříková, J., 1996). Je potešiteľné, že v súčasnosti dochádza k miernemu rastu pôrodnosti, hoci najväčší podiel rodičiek zaujímajú vydaté ženy, narastá podiel detí narodených mimo manželstva, tvrdia Dufek, J, Minařík B. (2007).

2. Metodika

Údajovú základňu pre analýzu tvorili databázy získané zo štatistických úradov SR a ČR. Pre vyššie uvedenú komparáciu boli použité ukazovatele hrubá miera pôrodnosti, hrubá miera úmrtnosti, prirodzený prírastok, index závislosti mladých (I), index závislosti starých (II), index ekonomického zaťaženia. Podľa Klufovej, R.(2007) sa pre štúdium regionálneho potenciálu ako veľmi vhodné javí využitie moderných metód priestorovej štatistiky v spojení s nástrojmi GIS, najmä nástroje pre analýzu priestorovej autokorelácie a identifikácia zhlukov podobných hodnôt študovanej premennej v priestore (tzv. hot spot analýza).

3. Dosiahnuté výsledky

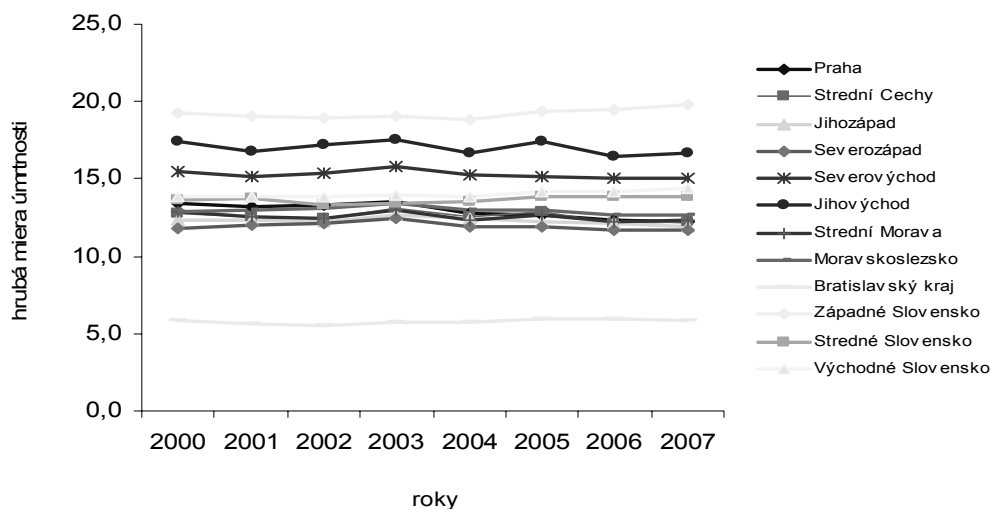
Demografické procesy boli hodnotené za roky 2000 až 2007. Pri hodnotení pôrodnosti sme použili ukazovateľ hrubá miera pôrodnosti. Na základe jeho vývoja môžeme tvrdiť, že ani v jednom z regiónov nenastal výrazný výkyv hodnôt v počte živonarodených detí na 1000 obyvateľov v celom sledovanom období. Východné Slovensko sa radí na prvé miesto, kde hrubá miera pôrodnosti dosiahla v priemere hodnotu 18,6 ‰. O toto vedúce postavenie sa stará najmä rómske etnikum, ktoré je charakteristické mnohočlennými rodinami. Naopak najnižšie hodnoty vykazuje Bratislavský kraj v priemere 5,3 ‰. Tento fakt možno pripísať tomu, že ľudia žijúci v tomto regióne majú vytvorené lepšie pracovné podmienky, ktorých sa nechcú vzdať, viac možností vzdelávania, ktoré chcú využiť a tiež aj využívanie zdravotnej starostlivosti, ktorú si môžu dovoliť je na oveľa vyššej úrovni ako v prípade Rómov, ktorým zdravie a hygiena takmer nič nehovorí. Vo všetkých českých regiónoch možno konštatovať, že počet živonarodených detí na 1000 obyvateľov má rastúcu tendenciu, t.j. každým rokom sa rodí viac detí (obrázok1).



Obr.1: Vývoj hrubej miery pôrodnosti (v ‰) vo vybraných regiónoch v rokoch 2000 až 2007, Zdroj: Eurostat, vlastné prepočty

Úmrtnosť v regiónoch bola hodnotená ukazovateľom hrubá miera úmrtnosti. Regiónom s najvyšším počtom zomretých ľudí na 1000 obyvateľov je Západné Slovensko, ktorého priemerná hodnota za celé hodnotené obdobie predstavuje hodnotu 19,2‰. Presný opak možno vidieť v Bratislavskom kraji, ktorý dosiahol najnižšiu priemernú hodnotu hrubej miery úmrtnosti 5,7‰. Hodnoty v ostatných regiónoch sa pohybujú v intervale od 11,8 do 17,5‰. Z českých regiónov je na tom najlepšie Severozápad, kde počet zomretých ľudí na 1000

obyvateľov predstavuje v približne 12 ľudí a najvyššiu hodnotu hrubej miery úmrtnosti dosiahol Juhovýchod a to 17‰ (obrázok 2).



Obr.2: Vývoj hrubej miery úmrtnosti (v ‰) vo vybraných regiónoch v rokoch 2000 až 2007, Zdroj: Eurostat, vlastné prepočty

Z hodnotenia jednotlivých regiónov z hľadiska prirodzeného prírastku možno tvrdiť, že všetky české regióny do roku 2005 vykazovali záporné hodnoty v intervale od -0,2 do -3,9‰ čo predstavuje úbytok obyvateľstva, ale v roku 2007 naopak prišlo k nárastu obyvateľstva a prirodzený prírastok nadobúdala kladné hodnoty od 0,7 do 2,1‰. Úplne iný vývoj mal prirodzený prírastok v slovenských regiónoch. Extrémy boli zaznamenané na Západnom Slovensku, ktoré počas sledovaného obdobia nadobúdalo záporné hodnoty v rozmedzí od -2,5 do -3,8‰ a Východné Slovensko, ktoré sa od ostatných regiónov líši tým, že ako jediné vykazuje počas celého hodnoteného obdobia prírastok obyvateľstva s maximálnymi hodnotami v intervale od 4,2 do 5,5‰ (tabuľka 1).

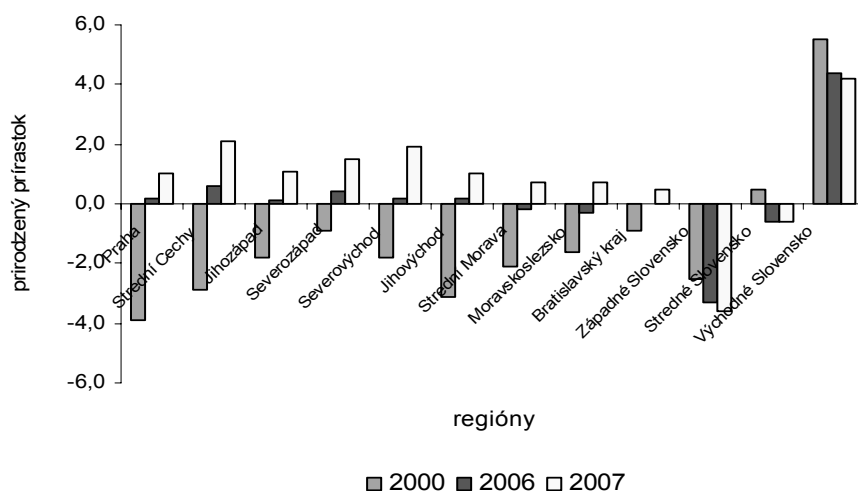
Tabuľka 1 Vývoj prirodzeného prírastku (v ‰) vo vybraných regiónoch SR a ČR

región	Prirodzený prírastok v ‰							
	2000	2001	2002	2003	2004	2005	2006	2007
Praha	-3,9	-3,5	-3,6	-3,4	-1,7	-0,8	0,2	1,0
Středné Čechy	-2,9	-2,6	-1,9	-2,5	-1,3	-0,8	0,6	2,1
Juhozápad	-1,8	-2,0	-1,6	-2,0	-1,4	-0,6	0,1	1,1
Severozápad	-0,9	-1,2	-1,0	-1,2	-0,4	-0,2	0,4	1,5
Severovýchod	-1,8	-1,8	-1,7	-2,0	-1,0	-0,5	0,2	1,9
Juhovýchod	-3,1	-2,6	-2,4	-2,7	-1,2	-1,2	0,2	1,0
Středná Morava	-2,1	-1,7	-1,6	-2,0	-1,2	-0,8	-0,2	0,7
Moravskoslezsko	-1,6	-1,7	-1,7	-1,8	-1,2	-0,8	-0,3	0,7
Bratislavský kraj	-0,9	-1,0	-0,9	-0,6	-0,3	0,0	0,0	0,5
Západné Slovensko	-2,5	-3,7	-3,8	-3,8	-2,9	-3,2	-3,3	-3,6
Stredné Slovensko	0,5	-0,6	-0,3	-0,3	0,1	-0,4	-0,6	-0,6
Východné Slovensko	5,5	4,4	4,3	4,3	4,9	4,7	4,4	4,2

Zdroj: Eurostat, vlastné prepočty

Z grafického zobrazenia prirodzeného prírastku je evidentné, že v roku 2000 nastal takmer vo všetkých regiónoch úbytok obyvateľstva, ktorý pretrvával až do roku 2006, kedy sme zaznamenali vo viacerých regiónoch mierny nárast počtu ľudí, ale oveľa priaznivejšie sa vyvíjal tento ukazovateľ v roku 2007, v ktorom okrem Stredného a Západného Slovenska

prírodný prírastok vykazoval kladné hodnoty. Zaujímavú skupinu tvoria Západné Slovensko, v ktorom aj v roku 2007 zomrelo viac ľudí ako sa narodilo a presný opak je na Východnom Slovensku, kde bol oveľa vyšší príliv narodených ako zomretých, čo je podmienené tým, že tu žije početné rómske etnikum (obrázok3).



Obr.3: Vývoj prirodzeného prírastku (v ‰) vo vybraných regiónoch v SR a ČR, Zdroj: Eurostat, vlastné prepočty

Pod pojmom ekonomická štruktúra obyvateľstva sa rozumie štruktúra ľudí podľa ekonomickej aktivity, zárobkovej činnosti ako aj zdrojov príjmu. K ekonomicky aktívnym osobám zaraďujeme tých, ktorí svojou činnosťou prispievajú k hospodárskemu výsledku spoločnosti a tých, ktorí museli svoju činnosť z rôznych dôvodov pozastaviť.

Tabuľka 2 Vývoj indexu ekonomickej závislosti I, II v regiónoch SR a ČR (v ‰)

	Index ekonomickej závislosti I		Index ekonomickej závislosti II	
	2000	2007	2000	2007
Praha	19,90	16,86	23,42	21,72
Stredné Čechy	23,52	20,85	20,77	19,91
Juhozápad	23,83	20,22	19,90	20,48
Severozápad	24,48	21,34	17,02	18,01
Severovýchod	24,69	20,94	20,04	20,41
Juhovýchod	24,43	20,41	20,42	21,08
Stredná Morava	24,39	20,27	19,55	20,67
Moravskosliezsko	25,15	20,50	17,40	19,18
Bratislavský kraj	22,85	17,38	15,48	16,56
Západné Slovensko	26,67	20,05	13,31	14,99
Stredné Slovensko	29,47	22,97	18,74	19,78
Východné Slovensko	33,37	26,90	14,23	15,35

Zdroj: Eurostat, vlastné prepočty

Je preukazné, že v roku 2000 bol vo všetkých hodnotených regiónoch index ekonomickej závislosti I vyšší ako v roku 2007. V roku 2000 dosiahla najnižšiu hodnotu indexu Praha (20 ľudí v predproduktívnom veku pripadajúcich na 100 ľudí v produktívnom veku). Presným opakom bolo Východné Slovensko (33,37). Spomínané regióny si udržali minimálne a maximálne hodnoty indexu ekonomickej závislosti aj v roku 2007, ale hodnoty boli

poznačené poklesom v počte predproduktívnych ľudí pripadajúcich na 100 produktívnych ľudí. V Prahe dosiahol index hodnotu 16,86 a na Východnom Slovensku 26,9 (tabuľka2).

Z hodnotenia indexu ekonomickej závislosti II vyplýva, že najvyšší podiel zaťaženia produktívnej zložky obyvateľstva poproduktívnou v porovnaní s ostatnými regiónmi je v Prahe. V roku 2000 predstavoval hodnotu 23,42, do roku 2007 klesol na 21,72. Naopak najnižší podiel je na Západnom Slovensku, kde v roku 2000 vykazoval hodnotu 13,31, v roku v roku 2007 zaznamenal nárast na 14,99. Rok 2000 je charakteristický tým, že iba regióny Praha a Stredné Čechy sa vyznačujú pozitívnou zmenou hodnoty indexu ekonomického zaťaženia starých v porovnaní s rokom 2007 (tabuľka2).

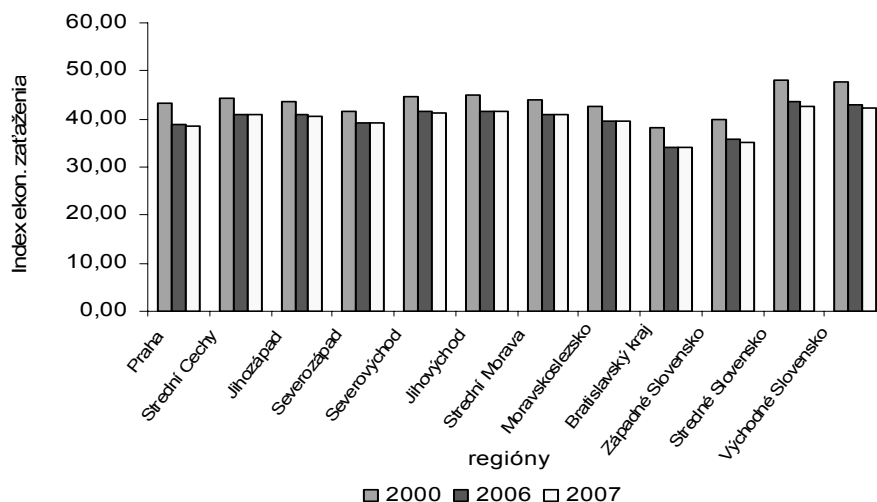
Zaťaženosť produktívnej populácie mladou menej ako dôchodcovskou je v Prahe, v Juhozápadnom, Juhovýchodnom regióne a v Strednej Morave. V ostatných regiónoch je stav opačný, čo nasvedčuje priaznivému vývoju – zvyšujúcej sa natalite.

Tabuľka 3 Vývoj indexu ekonomického zaťaženia v regiónoch SR a ČR (v %)

	2000	2007
Praha	43,32	38,58
Stredné Čechy	44,29	40,76
Juhozápad	43,73	40,71
Severozápad	41,50	39,35
Severovýchod	44,73	41,35
Juhovýchod	44,85	41,50
Stredná Morava	43,94	40,94
Moravskosliezsko	42,54	39,69
Bratislavský kraj	38,32	33,94
Západné Slovensko	39,98	35,04
Stredné Slovensko	48,22	42,75
Východné Slovensko	47,60	42,26

Zdroj: Eurostat, vlastné prepočty

Celkové zaťaženie produktívnej zložky obyvateľstva zložkami v predproduktívnom a poproduktívnom veku dosahuje vyššie hodnoty v roku 2000 v porovnaní s rokom 2007. Regióny Stredné a Východné Slovensko musia čeliť situácii, kedy vo všetkých sledovaných rokoch zaznamenali najvyššie hodnoty indexu ekonomického zaťaženia. Najmenej ekonomicky zaťaženým regiónom je Bratislavský kraj, ktorý v analyzovaných rokoch vykazoval najnižšie hodnoty. Najlepším regiónom v ČR je Praha, kde v roku 2007 v priemere muselo 100 ľudí v produktívnom veku pracovať na 39 ľudí v neproduktívnom veku. Najhorší scenár je na Strednom Slovensku (tab.3, obrázok 4).



Obr.4: Vývoj indexu ekonomického zaťaženia v regiónoch SR a ČR, Zdroj: Eurostat, vlastné prepočty

4. Záver

Tendencie vývoja na Slovensku ako aj v Čechách boli v podstate rovnaké, ale vďaka menším odlišnostiam v spoločenskej, ekonomickej a politickej situácii sme zaznamenali rôzne disparity.

Z hľadiska jednotlivých demografických ukazovateľov v sledovanom období je zrejmé, že vo veľkosti prirodzeného prírastku existujú výrazné regionálne disparity. Napríklad v Stredných Čechách (2,1 ‰) dosahuje prirodzený prírastok najvyššie hodnoty v ČR, v SR je to Východoslovenský kraj (4,2 ‰). Opakom je Západné Slovensko, kde ukazovateľ v roku 2007 nadobudol hodnotu -3,6 ‰. V ČR je najhoršia situácia v Strednej Morave a Moravskosliezsku, kde bola hodnota ukazovateľa 0,7 ‰ v roku 2007.

Ekonomicky je najviac zaťažená populácia v SR na Strednom Slovensku (42,75 %, pokles oproti roku 2000 o 11,34 %), v ČR Juhovýchod (41,50%, pokles oproti roku 2000 o 7,47%). Najmenej v SR je zaťažená populácia v Bratislavskom kraji (33,94%), v ČR v Prahe (35,58%).

Vzhľadom na danú demografickú situáciu v budúcnosti bude mnoho závisieť aj od ďalšieho politického a spoločenského vývoja a od napredovania procesu transformácie. Je potrebné uvažovať aj o koncepčných zmenách v oblasti týkajúcej sa rodinnej politiky – ako prostriedku starostlivosti štátu o rodinu a podporu pôrodnosti.

5. Literatúra:

- [1] DUFEK, J., MINAŘÍK, B. 2007. Analýza demografického vývoje České republiky a krajů regionu Jihovýchod. Brno, MZLU, s. 166, ISBN 978-80-7375-063-3.
- [2] DUFEK, J., MINAŘÍK, B. 2007. Přirozený pohyb obyvatelstva v Jihovýchodním regionu České republiky podle kraje. In: Sborník příspěvku z mezinárodní vědecké konference INPROFORUM 2007, 27.-28.11.2007, EF JU v Českých Budějovicích, ISBN 978-80-7394-016-4.
- [3] Klufová R. 2007. Sociálnegeografická exponovanost obcí Jihočeského kraje. In: Sborník příspěvku z mezinárodní vědecké konference INPROFORUM 2007, 27.-28.11.2007, EF JU v Českých Budějovicích, ISBN 978-80-7394-016-4.
- [4] RYCHTAŘÍKOVÁ, J. 1996. Současné změny charakteru reprodukce v České republice a mezinárodní situace. Demografie, 1996, roč. 38, č. 2, s. 77-89, ISSN 0011-8265 (Print)

Adresa autorov:

Zuzana Poláková, Ing. PhD.
Tr. A. Hlinku 2
949 76 Nitra
SR
Zuzana.Polakova@uniag.sk

Peter Obtulovič, doc. Ing. CSc.
Tr. A. Hlinku 2
949 76 Nitra
SR
Peter.Obtulovic@uniag.sk

Spektrální analýza cyklické struktury průmyslové výroby ČR

Spectral analysis of the cyclical behaving of the Czech Republic industrial production

Jitka Poměnková, Roman Maršálek

Abstract:

Presented paper is focused on evaluation of cycle behaving of growth business cycle in the Czech Republic. Growth business cycle is identified on the data of industrial production in 1996/Q1 – 2008/Q using Baxter-King band-pass filter and Hodrick-Prescott High-pass filter. For this aim, spectral analysis using autoregression process is applied. Obtained results are compared with the results from the time domain analysis, for which business cycle dating via naive rules and Bry-Boschan algorithm were applied.

Key words: spectral analysis, AR process, business cycle, dating

Klíčová slova: spektrální analýzy, AR proces, hospodářský cyklus, datování

1. Úvod

Standardní přístup empirické analýz hospodářského cyklu transitivity ekonomik vychází z růstového pojetí cyklu zkoumaného v časové doméně. Pro získání hodnot růstového hospodářského cyklu je možné využít řady metod. Mezi běžně používané metody patří filtry typu pásmová propust nebo filtry typu horní propust (Baxter a King, 1999). Analýza hospodářského cyklu v časové doméně však není schopna poskytnout podrobnější informace o cyklickém chování ukazatele zvoleného. Vyjdeme-li z aplikace obvyklých datovacích technik, kterými jsou v odborné ekonomické veřejnosti naivní pravidla upravená Canovou (1999) nebo Bry-Boschanův algoritmus (1971), zpravidla identifikujeme dva až čtyři cykly o různých délkách závislosti na zvoleném typu ukazatele. Z hlediska popsání cyklické struktury je to však nedostačující. Proto je na místě aplikace spektrální analýzy ve frekvenční doméně, která umožňuje lépe popsat toto cyklické chování.

Metody odhadu spektra lze dělit na parametrické a neparametrické. Neparametrické metody odhadují spektrum přímo ze zadané časové řady, zde řadíme periodogram (Hamilton, 1994). Parametrické metody oproti tomu fungují tak, že jsou nejprve odhadnuty parametry modelu popisující chování vstupní časové řady a teprve poté, s využitím odhadnutých koeficientů modelu, je odhadováno samotné spektrum (Proakis, 2002). Modelem se zpravidla rozumí procesy AR, MA nebo ARMA (Wooldridge, 2003).

Cílem předkládaného příspěvku je popsat a vyhodnotit cyklickou strukturu průmyslové výroby České republiky prostřednictvím odhadu spektra autoregresním procesem. Zjištěné výsledky analýzy cyklického chování ve frekvenční doméně dále porovnat a vyhodnotit s výsledky datování plynoucí z časové domény. Hodnoty průmyslové výroby v letech 1996/Q1 - 2008Q4 jsou chápány jako ukazatele hospodářského vývoje země. Pro potřeby analýzy bude pro nalezení růstového cyklu použito Baxter-Kingova filtru typu pásmová propust a Hodrick-Prescottova filtru (Hodrick, Prescott; 1980) typu horní propust.

2. Metodika

Spektrum časové řady Y_t je možné vyjádřit Fourierovou sumou (Dejong, Chetan; 2007)

$$S_Y(\omega) = \frac{1}{2\pi} \sum_{j=-\infty}^{+\infty} \gamma_j e^{-i\omega j}, \quad (1)$$

kde $\gamma_j = \text{cov}(Y_t, Y_{t+j}) = E(Y_t - \mu_t)(Y_{t+j} - \mu_{t+j})$ je autokovariance mezi Y_t a Y_{t+j} , $i = \sqrt{-1}$. Pro úhlovou frekvenci ω v radiánech platí $\omega = 2\pi/n = 2\pi f$, kde n je rozsah souboru, f je frekvence. Výběrové spektrum lze odhadovat několika způsoby. Při použití neparametrických metod je základní metodou odhadu spektra metoda periodogramu, v případě parametrických metod lze odhadovat spektrum prostřednictvím autoregresního procesu (Hamilton, 1994).

Odhad spektra pomocí AR modelu můžeme vypočítat jako (Proakis, 2002):

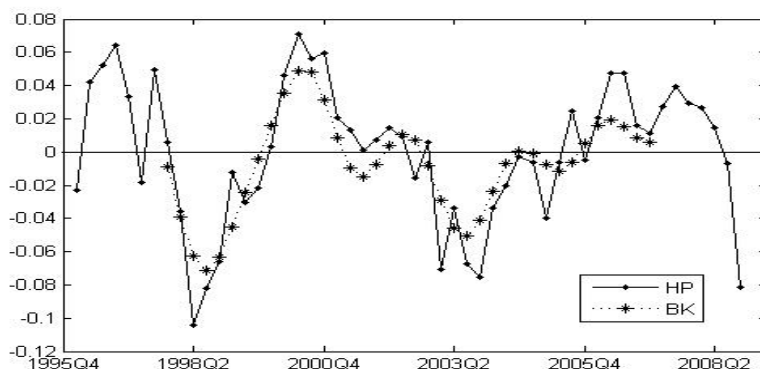
$$\hat{S}_Y = \frac{\sigma_w^2}{|H_w(e^{j\omega})|^2} = \frac{\sigma_w^2}{\left|1 - \sum_{i=1}^p a_i e^{-j\omega i}\right|^2}, \quad (2)$$

kde a_i jsou koeficienty AR modelu řádu p , n je rozsah datového souboru, resp. délka vstupní časové řady. Pro odhad parametrů modelu $AR(p)$ s optimalizovaným řádem zpoždění lze využít Yule-Walkerovu metodu (Proakis, 2002). Pro optimalizaci řádu zpoždění autoregresního procesu lze využít Akaiikovo nebo Schwartz-Cauchyho informačního kritéria (Seddighi, 2000).

Pro datování růstového hospodářského cyklu lze využít několika přístupů. Mezi běžně používaný orientační způsob patří využití naivních pravidel, jak je definoval Canova (1999). Canova označuje za „dno“ situaci, kdy ve dvou po sobě jdoucích čtvrtletích poklesu referenčního cyklu následuje růst, tj. $c_{t-2} > c_{t-1} > c_t < c_{t+1}$. Podobně, „vrchol“ je situace, kdy ve dvě po sobě následující čtvrtletí rostou a poté dochází k poklesu, tj. $c_{t-2} < c_{t-1} < c_t > c_{t+1}$. Druhá část pravidla vybírá čtvrtletí t jako dno hospodářské aktivity, resp. vrchol, pokud po sobě následovaly alespoň dva poklesy, resp. růsty, cyklické komponenty ve dvou po sobě jdoucích čtvrtletích během období tří čtvrtletí. Tedy, pokud platí $c_t < (>) 0$ a $c_{t-1} < (>) 0$ nebo pokud $c_{t+1} < (>) 0$ a $c_t < (>) 0$. Důvod zavedení obou pravidel pramení ze snahy zmenšit riziko předčasné identifikace bodů zvratu. Pokročilejší způsob nabízí Bry-Boschanův algoritmus (Bry, Boschan, 1971), který je obdobou tradiční posloupnosti, kdy je nejprve provedena identifikace hlavních cyklických pohybů, poté vymezeno okolí maxima a minima a nakonec zúžení hledání bodů zlomu vedoucí ke specifikaci okamžiku bodu zlomu. Tento algoritmus přihlíží k individuálnímu charakteru časové řady. Komplexní analýza s různými statistickými nástroji může totiž vést ke ztrátě stálosti výsledků v čase.

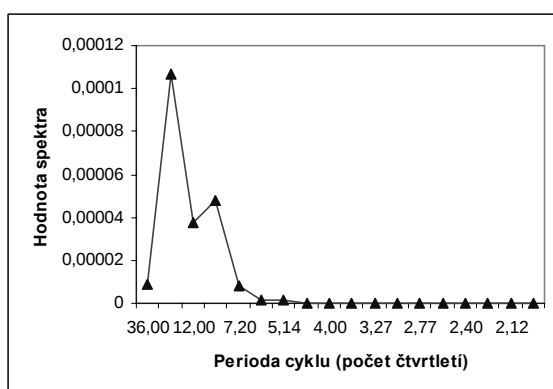
3. Empirická část

Pro empirickou analýzu byly zvoleny sezónně očištěné absolutní hodnoty celkové průmyslové výroby (bez stavebnictví) v České republice v konstantních cenách v období 1996Q1–2008Q4, které budou před analýzou transformovány přirozeným logaritmem, neboť růstové charakteristiky lépe vystihují povahy cyklů. Pro získání hodnot růstového cyklu byl využit Hodrick-Prescottův (HP) filtr s parametrem $\lambda = 1600$ (Kapounek, 2009) pro čtvrtletní hodnoty (Hodrick, Prescott; 1980) a Baxter-Kingův dvoupásmový filtr BK s délkou klouzavé části $K = 7$, který prochází složkami časové řady mezi šesti čtvrtletími až osmi lety (Baxter, King; 1999). Získaná rezidua byla označena jako hodnoty růstového cyklu ČR (obr. č. 1). Tvar získaných růstových cyklů je zobrazen na obr. č. 1. Aplikací rozšířeného Dickey-Fulerova testu (ADF) bylo zjištěno, že hodnoty růstového cyklu jsou stacionární v úrovni při 1, 5, 10 % a zpoždění $p = 1$ při detrendování Hodrick-Prescottovým filtrem, $p = 4$ při detrendování Baxterovým-Kingovým filtrem (Seddighi, Lawler, Katos; 2000). Dále byla pomocí Jarque-Berova testu zjištěna normalita těchto reziduí pro 1, 5% riziko.

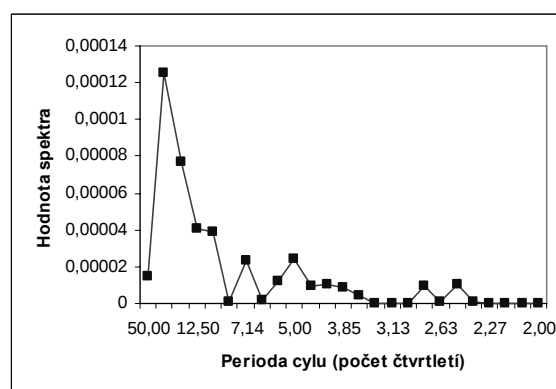


Obrázek 1: Růstový cyklus hodnot průmyslové výroby v ČR

Pro optimalizaci řádu zpoždění autoregresního procesu bylo využito Akaiikova (AIC) a Schwartz-Cauchyho (SC) informačního kritéria doplněného testem bílého šumu získaných reziduí Dickey-Fullerovým (DF) a Jarque-Bera testem normality (Seddighi, 2000). Nalezené optimální hodnoty řádu zpoždění byly pro HP filtr $p=18$ a pro BK filtr $p=11$. Odhad spektra pomocí optimalizovaného $AR(p)$ procesu růstového cyklu získaného detrendováním BK filtru je zobrazen na obr. č. 2, pro HP filtr na obr. č. 3.



Obrázek 2: Spektrum růstového cyklu hodnot průmyslové výroby v ČR (detrendování BK filtrem)



Obrázek 3: Spektrum růstového cyklu hodnot průmyslové výroby v ČR (detrendování HP filtrem)

Tabulka 1: Statistická významnost hodnot spektra

	Periody cyklu	51	25,50	17	12,75	10,20	8,50	7,29
HP	Významnost		**	**	**			
	Periody cyklu	37	18,50	12,33	9,25	7,40	6,17	5,29
BK	Významnost	**	***	**	**	**	**	**

Pozn. Statisticky významné na 1%(***), 5%(**), 10%(*)

Nalezené hodnoty odhadu spektra obou růstových cyklů byly testovány testem R. A. Fishera (Anděl, 1976; Poměnková, Maršálek; 2010) a byla posuzovaná jejich statistická významnost (tab. č. 1). Poznamenejme, že při práci s Baxter-Kingovým filtrem dochází ke ztrátě dat vlivem nastavení délky klouzavé části K . Proto i možné typy period cyklu, které se odvíjejí od délky časové řady, jsou jiné, než při práci s růstovým cyklem získaným aplikací Hodrick-Prescottova filtru. Tohoto lze využít pro identifikaci významných vnořených cyklů,

neboť spektrální rozlišitelnost je pro pomalu se pohybující složky časové řady (cykly dlouhé délky) spíše malá. Různost typů period cyklu pak umožní dodatečný pohled na možné typy cyklů.

Z nastavení frekvenčního pásma Baxter-Kingova filtru od 6 do 32 čtvrtletí a z výsledků spektrální analýzy růstového cyklu získaného detrendováním tímto filtrem (tab. 1) je patrné, že identifikace významné délky periody v trvání 5,29 čtvrtletí je pod dolní hranicí nastavení filtru. Tento fakt je způsoben propustností Baxter-Kingova filtru vlivem nepřesné aproximace ideálního filtru v hraničních oblastech, který je v časové doméně dán nekonečným klouzavým průměrem (Baxter, King; 1999). Protože růstový cyklus získaný detrendováním vstupních hodnot Hodrick-Prescottovým filtrem nevykázal tuto složku jako významnou, nebudeme ji uvažovat. V případě složky délky 37 čtvrtletí dochází rovněž k propuštění nad hranici frekvenčního pásma. Spektrální rozlišitelnost však při aplikaci na růstový cyklus získaný pomocí Hodrick-Prescottova filtru neumožnil potvrzení nebo vyvrácení významnosti této složky, proto i tato složka nebude uvažována. Zbýlé složky označené za statisticky významné říkají, že cyklická struktura průmyslové výroby České republiky obsahuje vnořené cykly střední délky s délkou trvání od pěti do sedmi let, krátké délky s délkou trvání od tří do pěti let a cykly velmi krátké cykly s délkou trvání do tří let.

Datování obou růstových cyklů bylo v časové doméně provedeno pomocí Bry-Boschanova algoritmu a byly stanoveny následující body zlomu (tab. 2). Poznamenejme, že v případě datování naivními pravidly nastala pro růstový cyklus získaný pomocí Hodrick-Prescottova filtru shoda s výsledky datování Bry-Boschanovým algoritmem. V případě růstového cyklu získaného aplikací Baxter-Kingova filtru bylo datováním Bry-Boschanovým algoritmem odhaleno o jedno dno (2005Q2) a vrchol (2004Q3) více, než při orientačním datování pomocí naivních pravidel. Bližší grafickou i numerickou analýzou zjišťujeme, že hodnota růstového cyklu v období 2004Q3 dosahuje úrovně 0,0008, což je při obvyklé aplikaci naivních pravidel chápáno jako hodnota velmi malá, rovna nule. Pak i přes splnění podmínky dvou rostoucích hodnot před bodem zlomu a jedné klesající po bodu zlomu nebyla splněna podmínka druhého pravidla, tj. hodnota bodu zlomu nad nulovou úrovní. Zohledníme-li, že získání hodnot růstových cyklů dosáhneme odstraněním odhadu dlouhodobého trendu aplikací Baxter-Kingova filtru typu pásmová propust, je na místě připustit i existenci intervalu spolehlivosti pro získané růstové hodnoty. Nelze tedy vyloučit, že hodnota odhadu růstového cyklu v okamžiku 2004Q3 by mohla být i větší, a tedy mohla být označena i naivními pravidly za bod zlomu.

Porovnáme-li počet identifikovaných bodů zlomu růstového cyklu získaného detrendováním Hodrick-Prescottovým filtrem s výsledky při aplikaci Baxter-Kingova filtru, vidíme, že počet extrému je při aplikaci Baxter-Kingova filtru větší. Podrobnějším grafickým a numerickým rozbořem růstového cyklu získaného při detrendování Hodrick-Prescottovým filtrem nalézáme situaci podobnou té, která se vyskytla v období 2004Q3, a to období 2001Q3 (potenciální dno), kdy hodnota odhadu růstového cyklu vykazovala téměř nulovou úroveň (0,0011). Pokud bychom připustili existenci intervalu spolehlivosti pro hodnoty růstového cyklu, pak bychom mohli připustit okamžik 2001Q3 jako dno a následně 2002Q1 jako vrchol. Analogicky pro okamžik potenciálního vrcholu v 2004Q3 (-0,0027) a následně dna v 2005Q1. Volatilitu a nejednoznačnost výsledků při aplikaci Hodrick-Prescottova filtru můžeme spatřovat také v samotném filtru, který je na rozdíl od Baxter-Kingova filtru (pásmová propust) koncipován jako filtr typu horní propust, který z časové řady odstraní pouze pomalu se pohybující složky (trend). Vzhledem k vývoji hodnot růstového cyklu získaného detrendováním Hodrick-Prescottovým filtrem v úvodním časovém období, tj. v době 1996Q1-1998Q2 můžeme bod zlomu (vrchol) v okamžiku 1996Q4 považovat za skutečný. Při aplikaci Baxter-Kingova filtru nebylo vlivem nastavené délky klouzavé části možné tento bod zlomu identifikovat.

Tabulka 2: Body zlomu růstových cyklů

HP		BK	
Vrchol	Dno	Vrchol	Dno
1996Q4			
	1998Q2		1998Q3
2000Q2		2000Q2	
			2001Q3
		2002Q2	
	2003Q4		2003Q3
		2004Q3	
			2005Q2
2006Q2		2006Q2	

Body zlomu, na jejichž základě bude vyčíslena délka cyklů identifikovaná v časové doméně, budou body zlomu identifikované Bry-Boschanovým algoritmem na růstovém cyklu získaném aplikací Baxter-Kingova filtru doplněné o bod zlomu v okamžiku 1996Q4 identifikovaný na růstovém cyklu při aplikaci Hodrick-Prescottova filtru. Doba trvání cyklu od nalezeného vrcholu až do období jednoho čtvrtletí před následující vrchol je vyjádřena v tabulce 3.

Tabulka 3: Doba a trvání cyklu identifikovaných v časové doméně

Doba trvání cyklu	Délka cyklu ve čtvrtletích	Délka cyklu v letech
1996Q4-2000Q1	14	3,50
2000Q2-2002Q1	8	2
2002Q1-2004Q2	9	2,25
2004Q3-2006Q1	7	1,75

Porovnáme-li výsledky zjištěné analýzou v časové doméně s výsledky zjištěné ve frekvenční doméně, vidíme, že analýza v časové doméně není dostačující. Ve frekvenční doméně byla odhalena řada dalších cyklů, a to kratších (pod hranici sedmi čtvrtletích) i delších (nad hranici čtrnácti čtvrtletí).

4. Závěr

Předkládaný příspěvek se zabýval cyklickou strukturou indexu průmyslové výroby České republiky ve frekvenční doméně, a to prostřednictvím odhadu spektra autoregresním procesem. Zjištěné výsledky byly dále porovnávány a vyhodnocovány s výsledky plynoucí z datování hospodářského cyklu v časové domény. Růstový cyklus byl identifikován pomocí filtračních technik, a to Baxter-Kingova a Hodrick-Prescottova filtru.

Pomocí odhadu spektra optimalizovaným $AR(p)$ procesem bylo zjištěno, že růstový cyklus průmyslové výroby České republiky obsahuje vnořené cykly jak velmi krátké s dobou trvání do tří let, tak krátké cykly s dobou trvání od tří do pěti let, ale i střední cykly s délkou trvání od pěti do sedmi let. Délky cyklů zjištěné pomocí datování růstového hospodářského cyklu v časové doméně ukázaly existenci pouze krátkých cyklů s dobou trvání do tří a půl let.

Porovnáním obou přístupů vidíme, že pomocí spektrální analýzy je možné lépe nahlédnout do struktury cyklického chování hospodářského cyklu lépe, než je tomu při analýze v časové doméně. Ta sice umožní identifikovat počátek a konec cyklu, avšak neodhalí existenci vnořených cyklů.

Předkládaný příspěvek vznikl za podpory interního grantového projektu 71/2010 "Stabilizační funkce monetární politiky v souvislostech hospodářského cyklu České republiky".

5. Literatura

- [1] ANDĚL, J., 1976. Statistická analýza časových řad, Praha: SNTL, 271 pp.
- [2] BAXTER, R., KING, R. G. 1999. Measuring Business Cycles: Approximate Band – Pass Filters for Economic Time Series. *Review of Economic and Statistics*, vol. 81, no. 4, p. 575-593
- [3] BRY, G., BOSCHAN, C. 1971. Cyclical Analysis of Time Series: Selected Procedures and Computer Programs, Technical Paper 20, National Bureau of Economic Research, New York.
- [4] CANOVA, F. 1999. Does De-trending Matter for the Determination of the Reference Cycle nad Selection of Turniny Points?, *The Economic Journal*, Vol. 109, No. 452 (Jan., 1999), p. 126-150.
- [5] HAMILTON, J. D. 1994. Time Series Analysis. Princeton University Press, 799 pp., ISBN -10:0-691-04289-6
- [6] HODRICK, R. J., PRESCOTT, E. C. 1980. Post-war U.S. Business Cycles: An Empirical Investigation, mimeo, Carnegie-Mellon University, Pittsburgh, PA. 1980, 24 pp.
- [7] KAPOUNEK, S. 2009. Estimation of the Business Cycles - Selected Methodological Problems of the Hodrick-Prescott Filter Application. *Polish Journal of Environmental Studies*. 2009. Vol. V, No. 6, p. 2. ISSN 1230-1485.
- [8] KING, R. G., REBELO, S. T. 1993. Low frequency Filtering and Real Business Cycles. *Journal of Economic Dynamics and Control*. vol. 17, p. 207-231
- [9] POMĚNKOVÁ, J., MARŠÁLEK, R. 2010. Testování významnosti periodicity průmyslové výroby pomocí R. A. Fisherova testu, *Acta univ. agric. et silvic. Mendel. Brun.*, (v tisku)
- [10] PROAKIS, J. G., RADER, Ch. M., Ling, F.L., NIKIAS, Ch. L., MOONEN, M., PROUDLER, J.K. 2002. Algorithms for Statistical Signal Processing, Prentice Hall, ISBN 0-13-062219-2
- [11] SEDIGGI, H. R., LAWLER, K. A., KATOS., A. V., 2000. Econometrics. A practical approach. New York 2000, p. 262 – 287.
- [12] WOOLDRIDGE, J. M. (2003). Introductory Econometrics: A modern approach. Ohio, 863 pp., ISBN 0-324-11364-1.

Adresy autorů:

Jitka Poměnková, Ph.D., RNDr.
Ústav financí
PEF, MENDELU v Brně
Zemědělská 1
613 00 Brno
pomenka@mendelu.cz

doc. Ing. Roman Maršálek, Ph.D.
Ústav radioelektroniky
Fakulta elektrotechniky a komunikačních
technologií, VUT v Brně
Purkyňova 118
612 00 Brno
marsaler@feec.vutbr.cz

Alternativní indexy predikční síly credit scoringových modelů Alternative indexes of predictive power for credit scoring models

Martin Řezáč

Abstract: Standard procedure for assessing the predictive power of credit scoring model is to calculate indexes such as the Gini coefficient, KS, Lift, or Informational value. Despite the great popularity of Gini coefficient in assessing the predictive power of scoring models it is clear that this index is not able to capture some essential features of these models. Alternative indexes of predictive power are proposed in this paper. All of them are based on Lift function. First, it is QLift, which indicates how many times the model is better than the random model at given level of rejection. Furthermore, original concepts of indexes Lift ratio (LR), the relative lift (RLift) and integrated relative lift (IRL) are proposed and discussed.

Key words: Predictive power, Gini index, Lift, Lift ratio, Relative Lift, IRL.

Klíčová slova: Predikční síla, Giniho index, Lift, Liftový poměr, relativní Lift, IRL.

1. Úvod

Obvyklým postupem pro posouzení predikční síly credit scoringového modelu je výpočet indexů jako jsou Giniho koeficient [1], [2], [6], K-S [1],[5], Lift [3],[4] nebo Informační statistika [2], [3], [4]. Máme-li posoudit jak bude klient finanční instituce splácet své závazky, bude predikční síla credit scoringového modelu daná zmíněnými indexy reprezentovat jeho schopnost rozlišit mezi klienty, kteří své finanční závazky bez problému plní, a klienty, kteří jsou defaultní.

Předpokládejme, že máme k dispozici skóre S představující výstup nějaké scoringové funkce. Toto skóre je typicky odhadem pravděpodobnosti defaultu klienta. Dále máme ke každému klientovi dáno, zda u něj nastal nebo nenastal default. Zavedeme následující označení:

- $F_{BAD}(s)$ - distribuční funkce skóre špatných (defaultních) klientů
- $F_{GOOD}(s)$ - distribuční funkce skóre dobrých (nedefaultních) klientů
- $F_{ALL}(s)$ - distribuční funkce skóre všech klientů

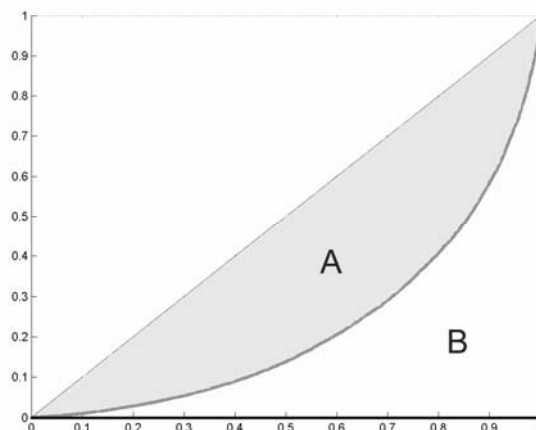
Relativní četnost špatných, resp. dobrých, klientů označíme p_B , resp. p_G .

Zřejmě nejznámějším a v praxi nejpoužívanějším indexem predikční síly scoringového modelu je Giniho koeficient. Definuje se vztahem:

$$Gini = 1 - 2 \int_0^1 F_{GOOD}(F_{BAD}^{-1}(s)) ds \quad (1)$$

S využitím Lorenzovy křivky zobrazené na následujícím obrázku 1 jej lze definovat jako

$$Gini = \frac{A}{A+B} = 2A \quad (2)$$



Obrázek 1: Giniho index. Zdroj: vlastní konstrukce.

Jde o index vyjadřující globální predikční sílu modelu. I přes jeho velikou popularitu je třeba říci, že má jednu zásadní nevýhodu. Není totiž schopen postihnout fakt, že největší predikční síly by scoringový model měl dosahovat v oblasti očekávané cutoffové hodnoty skóre. Typicky jde o hodnoty skóre odpovídající 10 – 20 % nejhorších klientů. Tento problém řeší použití např. Liftu [3], [4], definovaného vztahem:

$$Lift(s) = \frac{F_{BAD}(s)}{F_{ALL}(s)} \quad s \in [L, H] \quad (3)$$

S použitím tohoto indexu predikční síly je totiž možné posoudit lokální chování daného scoringového modelu.

2. QLift

Pro praktické účely je mnohem praktičtější výpočet hodnoty Liftu pro 10 %, 20 %, ..., 100 % klientů s nejhorším skóre [1]. Důvodem je interpretace této hodnoty. Zatímco hodnota Liftu pro danou hodnotu skóre je v praxi těžko interpretovatelná, hodnota Liftu pro 10 % nejhorších klientů má jasnou ekonomickou interpretaci. Tato hodnota totiž udává kolikrát více špatných klientů zamítne scoringový model na úrovni zamítání 10 % ve srovnání se zamítáním náhodným. V této souvislosti definujeme

$$QLift(q) = \frac{F_{BAD}(F_{ALL}^{-1}(q))}{F_{ALL}(F_{ALL}^{-1}(q))} = \frac{1}{q} F_{BAD}(F_{ALL}^{-1}(q)) \quad q \in (0,1], \quad (4)$$

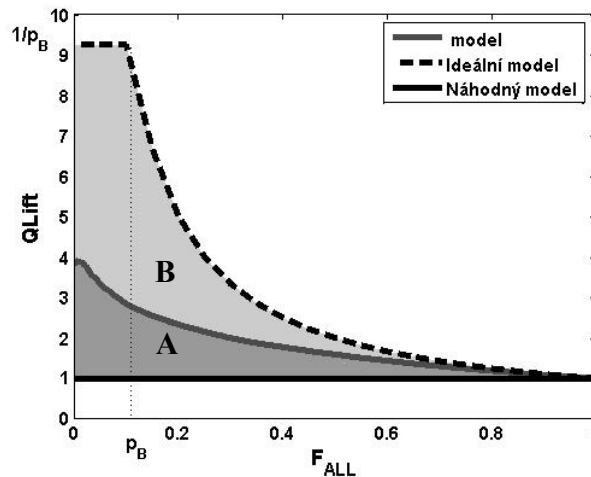
kde q reprezentuje úroveň 100 q % nehorších klientů podle skóre a $F_{ALL}^{-1}(q)$ je kvantilová funkce skóre všech klientů. Tu lze empiricky odhadnout pomocí vzorce

$$F_{ALL}^{-1}(q) = \min\{s \in [L, H] \mid F_{ALL}(s) \geq q\} \quad (5)$$

Dá se snadno ukázat, že pro ideální model, tj. scoringový model bezchybně odhadující dobré a špatné klienty, platí

$$QLift_{ideal}(q) = \begin{cases} \frac{1}{p_B} & q \in (0, p_B] \\ \frac{1}{q} & q \in (p_B, 1] \end{cases} \quad (6)$$

Na následujícím obrázku 2 je zobrazen příklad funkce QLift pro ideální, náhodný a posuzovaný reálný scoringový model.



Obrázek 2: QLift pro reálný, ideální a náhodný model. Zdroj: vlastní konstrukce.

3. Liftový poměr, relativní Lift, IRL

S využitím předchozího obrázku definujeme index LR (Lift Ratio) jakožto analogii k Giniho koeficientu

$$LR = \frac{A}{A+B} = \frac{\int_0^1 QLift(q) dq - 1}{\int_0^1 QLift_{ideal}(q) dq - 1} \quad (7)$$

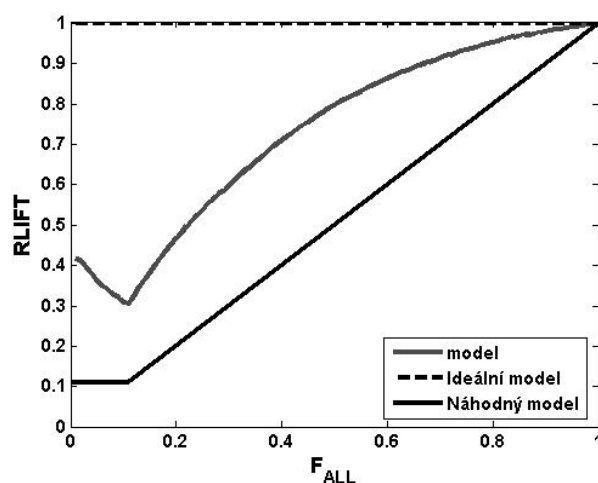
Je zřejmé, že jde o globální míru kvality modelu a že nabývá hodnot od 0 do 1. Hodnota 0 odpovídá náhodnému modelu, hodnota 1 pak modelu ideálnímu. Význam indexu je jednoduchý. Čím vyšších hodnot nabývá, tím je daný model lepší. Podstatnou výhodou tohoto indexu je fakt, že umožňuje korektní porovnání modelů vyvinutých na různých datech, což není možné pomocí hodnot funkce QLift.

Zatímco LR porovnává plochy pod funkcí Liftu pro daný model a model ideální, následující myšlenka je založena na porovnání přímo těchto funkcí samotných. Definujme relativní Lift funkci pomocí

$$RLift(q) = \frac{QLift(q)}{QLift_{ideal}(q)}, \quad q \in (0,1] \quad (8)$$

Příklad této funkce je zobrazen na obrázku 3. Definiční obor této funkce je interval $(0,1]$, obor hodnot je podinterval $[0,1]$. Počátečním bodem grafu je bod $[q_{min}, p_B \cdot QLift(q_{min})]$, kde q_{min} je kladné číslo blízké nule. Následně graf pokračuje do bodu $[p_B, p_B \cdot QLift(p_B)]$ a poté

roste do bodu $[1, 1]$. Je zřejmé, že čím je model lepší, tím více se graf této funkce blíží k úsečce $[[0, 1], [1, 1]]$, která reprezentuje graf pro ideální model.



Obrázek 3: RLIFT pro reálný, náhodný a ideální model. Zdroj: vlastní konstrukce.

Nyní je přirozené dokončit celý koncept pomocí hodnoty integrálu funkce RLIFT. Definujme tedy integrální relativní Lift jako

$$IRL = \int_0^1 RLIFT(q) dq \quad (9)$$

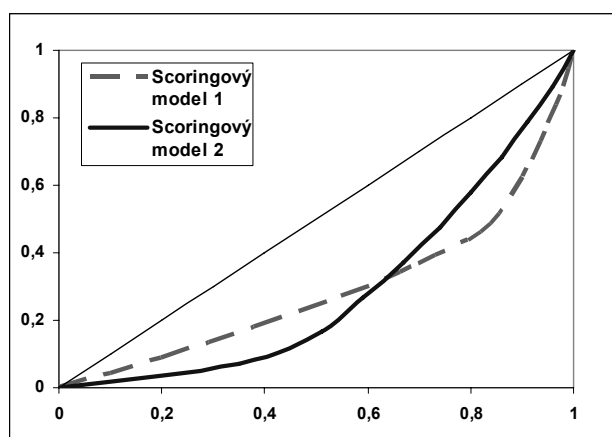
Nabývá hodnot od $0,5 + \frac{p_B^2}{2}$ pro náhodný model po 1 pro model ideální. Opět platí, že čím vyšší je hodnota tohoto indexu, tím lépe. Tento globální index má zajímavou vlastnost. Dá se ukázat, že IRL je přibližně rovno c-statistice v případě rovnosti rozptylů skóre dobrých a špatných klientů. Na druhou stranu se od této statistiky může velmi výrazně lišit pokud jsou tyto rozptyly odlišné.

Uvažujme nyní následující příklad dvou scoringových modelů:

Tabulka 1: Počty špatných klientů a QLIFT podle decilů skóre pro scoringové modely 1 a 2. Zdroj: vlastní konstrukce.

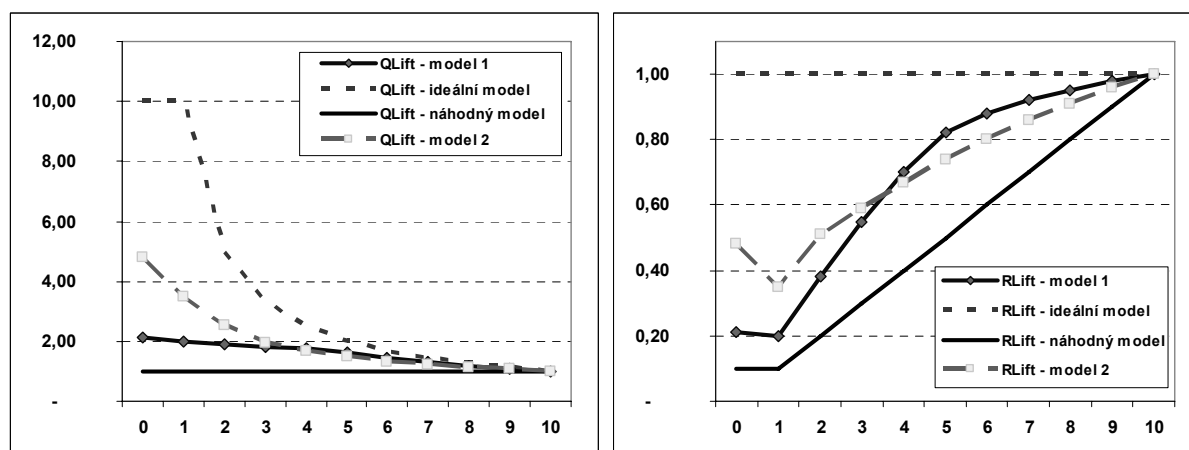
decil skóre	# klientů	scoringový model 1				scoringový model 2			
		# špatných klientů	# kumul. špatných klientů	# kumul. Def. Rate	QLIFT	# špatných klientů	# kumul. špatných klientů	# kumul. Def. Rate	QLIFT
1	100	20	20	20,0%	2,00	35	35	35,0%	3,50
2	100	18	38	19,0%	1,90	16	51	25,5%	2,55
3	100	17	55	18,3%	1,83	8	59	19,7%	1,97
4	100	15	70	17,5%	1,75	8	67	16,8%	1,68
5	100	12	82	16,4%	1,64	7	74	14,8%	1,48
6	100	6	88	14,7%	1,47	6	80	13,3%	1,33
7	100	4	92	13,1%	1,31	6	86	12,3%	1,23
8	100	3	95	11,9%	1,19	5	91	11,4%	1,14
9	100	3	98	10,9%	1,09	5	96	10,7%	1,07
10	100	2	100	10,0%	1,00	4	100	10,0%	1,00
All	1000	100				100			

Scoringový model 1 (ScM 1) reprezentuje situaci, kdy jsou lépe rozpoznáni dobří klienti od těch ještě lepších, ovšem není schopen rozlišit mezi špatnými a nejhoršími klienty. Největší predikční síly tedy dosahuje v oblasti vysokých skóre, která odpovídá dobrým klientům. Naproti tomu scoringový model 2 (ScM 2) reprezentuje přesně opačnou situaci. To znamená, že největší predikční síly je dosaženo v oblasti nízkých skóre odpovídající špatným klientům. Pro oba tyto modely je tato situace zobrazena pomocí Lorenzovy křivky na obrázku 4.



Obrázek 4: Lorenzovy křivky pro ScM 1 a ScM 2. Zdroj: vlastní konstrukce.

Na následujícím obrázku 5 jsou zobrazeny hodnoty QLiftu a RLiftu pro uvažované dva scoringové modely, model ideální a náhodný. Je zřejmé, že lepších hodnot dosahuje model ScM 2.



Obrázek 5: QLift a RLift pro ScM 1, ScM 2, ideální a náhodný model. Zdroj: vlastní konstrukce.

Souhrn uvedených indexů pro posuzované scoringové modely udává tabulka 2. Pokud bychom modely porovnávaly pouze pomocí Giniho koeficientu, došli bychom k závěru, že jsou oba stejně kvalitní, tj. že mají stejnou predikční sílu. Jak ale ukazují všechny ostatní uvedené indexy, lepší je evidentně model č. 2. Byla-li by předpokládána míra zamítnutí scoringového modelu 10 %, je tento fakt podpořen hodnotou QLift(0,1). Uvažujeme-li globální indexy LR a IRL, dostáváme stejný výsledek.

Tabulka 2: Počty špatných klientů a QLift podle decilů skóre pro scoringové modely 1 a 2.
Zdroj: vlastní konstrukce.

	scoringový model 1	scoringový model 2
GINI	0,420	0,420
QLift(0,1)	2,000	3,500
LR	0,242	0,372
IRL	0,699	0,713

4. Závěr

I přes velkou popularitu Giniho koeficientu při posuzování predikční síly scoringových modelů se ukazuje, že tento index není schopen postihnout některé podstatné vlastnosti těchto modelů. Jakožto alternativní indexy predikční síly byly v tomto článku navrženy indexy založené na Liftové funkci. Jednak jde o QLift, který udává kolikrát je daný model na dané úrovni předpokládaného zamítání klientů lepší než model náhodný. Dále byly představeny zcela původní koncepty indexů Liftového poměru (LR), relativního Liftu (RLift) a integrálního relativního Liftu (IRL).

Uvedený příklad demonstruje opodstatněnost navržených indexů. Vyplyvá z něj totiž, že posuzování predikční síly scoringového modelu pouze na základě Giniho koeficientu může vést k situaci, kdy dva evidentně rozdílné modely jsou posouzeny jako stejně kvalitní. Dokonce se může stát, že model s nejvyšším Giniho koeficientem bude nasazen do praxe ačkoli existoval model s mnohem lepší predikční silou v oblasti předpokládané cutoffové hodnoty skóre. Přesně tento problém řeší použití navržených indexů.

5. Literatura

- [1] ANDERSON, R. 2007. *The Credit Scoring Toolkit: Theory and Practice for Retail Credit Risk Management and Decision Automation*, Oxford University Press, Oxford. ISBN 9780199226405.
- [2] HAND, D.J. and HENLEY, W.E. 1997. Statistical Classification Methods in Consumer Credit Scoring: a review. In: *Journal. of the Royal Statistical Society, Series A*, 160, No.3:523-541.
- [3] ŘEZÁČ, M. 2009. Indexy kvality normálně rozložených skóre. In: *Forum Statisticum Slovaceum*, 2009, č. 2., 144-148.
- [4] ŘEZÁČ, M., ŘEZÁČ, F. 2009. Quality indexes of predictive models in risk and portfolio management. In: *IMAEF 2009 Proceedings*. Ioannina, Greece.
- [5] SIDDIQI, N. 2006. *Credit Risk Scorecards: developing and implementing intelligent credit scoring*, Wiley, New Jersey. ISBN 9780471754510.
- [6] THOMAS, L.C., EDELMAN, D.B., CROOK, J.N. 2002. *Credit Scoring and Its Applications*, SIAM Monographs on Mathematical Modeling and Computation, Philadelphia. ISBN 9780898714838.

Adresa autora:

Martin Řezáč, Mgr. Ph.D.
Ústav statistiky a operačního výzkumu PEF MENDELU
Zemědělská 1
613 00 Brno
rezac@mendelu.cz

Handicapy a model výpočtu pravdepodobnosti z kurzu Handicaps and model of computing probability from rates

Lubomír Rybanský

Abstract: The article deals with statistical evaluation of bet game and model of probability computation related to data from NBA basketball game.

Key words: handicap, probability model, hypothesis testing

Kľúčové slová: handicap, pravdepodobnostný model, testovanie hypotéz

1. Úvod

Existuje veľa stávkových spoločností, ktoré nám denne ponúkajú možnosť zúčastniť sa rôznych hier. My sme túto ponuku využili na získanie reálnych údajov, na základe ktorých sme sa pokúsili vytvoriť jednoduchý model prevodu známych kurzov na pravdepodobnosti. Zaoberali sme sa tiež výsledkami jednej ponúkanej hry s názvom handicap, v ktorej sme analýzou výsledkov dospeli k zaujímavému zisteniu.

2. Handicapy

Počas jednej sezóny 2009/2010 basketbalovej súťaže NBA sme okrem výsledkov jednotlivých stretnutí zaznamenávali aj jednu zaujímavú štatistiku a to rozdiel skóre. Porovnávali sme bodové rozpätie stanovené stávkovou kanceláriou so skutočným výsledkom zápasu.

Keď je jeden z tímov favorizovaný, tak stávková kancelária vypisuje také bodové rozpätie pre víťazstvo, aby prilákala zhruba rovnaký počet stávok na oba tímy. To im umožňuje, aby sa hry zúčastnilo čo najviac hráčov a súčasne sa tým eliminujú možné finančné straty.

Napríklad basketbalový tím Cleveland je lepší ako New Jersey (v sledovanej sezóne to tak bolo). Je možná stávka, pri ktorej sú na oba tímy vypísané rovnaké kurzy, ale stávky na Cleveland kancelária vyplatí iba vtedy, keď porazí New Jersey aspoň o 10 bodov. Takýto druh stávky sa nazýva handicap. V záujme každej stávkovej kancelárie je vypísať handicap tak, aby čo najlepšie vystihoval reálne možnosti oboch súperov.

Dôležitým ukazovateľom je premenná *rozpätie*, ktorá vyjadruje rozdiel medzi skutočným bodovým rozdielom v stretnutí a handicapom z pohľadu favorita. Napríklad stretnutie Cleveland - New Jersey malo vypísaný handicap +9,5 bodu v prospech New Jersey a stretnutie sa skončilo výsledkom 100:90. To znamená, že hodnota premennej *rozpätie* je 0,5. V tabuľke 1 sú uvedené popisné štatistiky premenných *odhadovaný rozdiel*, *skutočný rozdiel* a *rozpätie*.

Najskôr vyhodnotíme odhadovaný rozdiel (handicap v prospech slabšieho tímu) a skutočného bodového rozdielu vzhľadom na favorizovaný tím.

Tabuľka 1: Popisné štatistiky

	N	Priemer	Medián	Modus	Min.	Max.	Sm.odch.	Šikmost'	Špicatost'
odhad. rozdiel	766	7,34	6,5	6,5	1,5	16,5	3,09	0,58	-0,41
skutočný rozdiel	766	4,87	6,0	0,0	-43,0	50,0	12,93	-0,24	0,44
rozpätie	766	-0,56	-0,5	-3,5	-45,5	46,5	11,66	0,00	0,61

Z popisných štatistík náhodnej premennej *rozpätie* sme zistili, že priemerná hodnota je -0,56, čo naznačuje, že vypísané handicapy sú v porovnaní so skutočnými rozdielmi mierne podhodnotené, približne o pol bodu na zápas. Smerodajná odchýlka je 11,66, čo znamená, že približne v 2/3 skúmaných prípadov sa bodové rozpätie líšilo od skutočného rozdielu bodov dosiahnutých v hre o menej než 11,66 bodu.

V ďalšom testujeme hypotézu, že stredná hodnota premennej *rozpätie* je 0, proti alternatívnej obojstrannej hypotéze. Použijeme t - test. Dosiahnutá p - hodnota je 0,185. Rovnaký výsledok sme získali aj použitím jednovýberového Wilcoxonovho testu. Dosiahnutá p - hodnota bola 0,171. Na hladine významnosti 0,05 prijímame hypotézu, že stredná hodnota náhodnej premennej *rozpätie* je rovná 0, čo znamená, že vypísané handicapy žiaden tím ani nepreceňujú, ani nepodceňujú.

Napriek tomu, že vypísaná hodnota handicapu dobre vyrovnáva rozdiely medzi tímami, budeme skúmať, či je to tak aj v prípade, ak stretnutia rozdelíme do určitých skupín.

Množinu všetkých zápasov, v ktorých bol vypísaný handicap, rozdelíme na podmnožiny podľa kurzu na víťazstvo favorizovaného tímu.

My sme sa rozhodli vytvoriť podmnožiny zápasov tak, aby rozpätie kurzu na favorita bola jedna desatina. V každej podmnožine sme zisťovali počet zápasov, v ktorých favorizovaný tím zvíťazil vyšším rozdielom ako bol handicap - kód 1, počet zápasov v ktorých zvíťazil, ale nižším rozdielom ako bol handicap - kód 0, počet zápasov, v ktorých favorit prehral - kód -1 a zápasy, ktoré sa skončili nerozhodne - kód -2. Vzhľadom na predchádzajúci výsledok by sme sa mohli domnievať, že počet zápasov s kódom 1 je približne rovnaký ako počet zápasov s kódmi 0, -1, -2.

Testujeme hypotézu, že početnosť kódu 1 a súčet početností kódov -2, -1, 0 sú v rámci danej podmnožiny kurzov rovnaké. Platnosť hypotézy sme overovali χ^2 testom dobrej zhody a testom hypotézy o strednej hodnote binomického rozdelenia. V prípade druhého testu uvažujeme dva možné výsledky: buď nastane kód 1 alebo niektorý z kódov -2, -1, 0. Testujeme tak hypotézu $H_0: p=1/2$ proti dvojstrannej alternatíve.

Zistené početnosti jednotlivých kódov sú spolu s výsledkami χ^2 testu dobrej zhody a testom hypotézy o strednej hodnote binomického rozdelenia uvedené v tabuľke 2.

Tabuľka 2: Početnosti kódov a výsledky testov

kurz		kódy				χ^2 test			hyp. o binom. rozdelení	
od	do	-2	-1	0	1	víťazstvá nad hcp (%)	χ^2	p -hodnota	U	p -hodnota
1,05	1,15	1	8	42	31	37,80	4,891	0,027	-2,209	0,028
1,16	1,25	3	15	35	57	51,82	0,863	0,353	0,381	0,354
1,26	1,35	7	18	23	52	52,00	0,024	0,878	0,400	0,880
1,36	1,45	9	34	15	60	50,85	0,061	0,805	0,184	0,806
1,46	1,55	6	24	10	42	51,22	0,026	0,872	0,221	0,873
1,56	1,65	6	30	10	34	42,50	3,106	0,078	-1,342	0,076
1,66	1,75	5	10	2	15	46,88	0,044	0,834	-0,354	0,835
1,76	1,85	2	2	1	4	44,44	0,143	0,705	-0,333	0,706

Z výsledkov vyplýva, že v kategórii kurzov 1,05 - 1,15 je v porovnaní s teoretickou hodnotou 50% prípadov víťazstiev nad handicap iba 37,80% , čo je štatisticky významný rozdiel (na hladine významnosti 0,05). Dosiahnuté p - hodnoty sú 0,027 resp. 0,028. Do tejto kategórie patria zápasy, v ktorých hrali výrazní favoriti proti slabším tímom. V žiadnej inej podmnožine sme nezistili podobnú neznhodu medzi teoretickými a empirickými hodnotami.

Zaujímavou skutočnosťou je, že v spomínanej kategórii kurzov je väčší počet víťazstiev pod handicap (kód 0) ako nad handicap (kód 1). Tento výsledok sa dá vysvetliť buď nevhodne

nastavenou hodnotou handicapu, t.j. pre favorita bol príliš veľký, alebo malou vôľou zvíťaziť vyšším rozdielom zo strany lepšieho tímu.

3. Odhad pravdepodobnosti z kurzov

Zaoberajme sa otázkou, ako zo známych kurzov určiť pravdepodobnosti nastania týchto udalostí. V našom prípade budeme uvažovať tri možné, vzájomne sa vylučujúce realizácie náhodnej udalosti, ktorou je výsledok stretnutia.

Nech X je náhodná premenná, ktorá vyjadruje výsledok stretnutia. Môže nadobúdať hodnoty 1 - zvíťazí domáci tím, 0 - zápas sa skončí nerozhodne a 2 - zvíťazí hosťujúci tím. Označme $P(X=1)$ pravdepodobnosť víťazstva domáceho tímu, $P(X=0)$ pravdepodobnosť nerozhodného výsledku a $P(X=2)$ pravdepodobnosť víťazstva hosťujúceho tímu.

Nech a je kurz vypísaný na víťazstvo domáceho tímu, b kurz na remízu a c kurz vypísaný na víťazstvo hosťujúceho tímu.

Pri hľadaní vhodného modelu, ako zo známych kurzov určiť príslušné pravdepodobnosti, je dôležité uvedomiť si, akým spôsobom sa určuje kurz v prípade, že by sme chceli tipovať viacero udalostí. Ak predpokladáme, že uvedené udalosti sú nezávislé, tak pravdepodobnosť toho, že všetky udalosti nastanú tak ako sme tipovali, je daná súčinom pravdepodobností nastania jednotlivých udalostí.

Uvažujme napríklad dve stretnutia s rovnakými kurzami: Cleveland – New Jersey 1,5 – 10,0 – 2,25 a Boston – Chicago 1,5 – 10 – 2,25. Ak by sme chceli tipovať na kurz 2,25, tak to môžeme urobiť dvoma spôsobmi:

- vsadíme na New Jersey alebo Chicago,
- vsadíme na Cleveland a Boston.

Pretože v oboch prípadoch je výsledný kurz rovnaký, predpokladáme, že nastanú s rovnakou pravdepodobnosťou.

To nás privádza k tomu, že medzi kurzom a pravdepodobnosťou musí platiť

$$\frac{P^{k_1}(X=1)}{P(X=0)} = \frac{a^{k_1}}{b} = 1, \text{ kde } k_1 \in R. \quad (1)$$

Reálne číslo k_1 vyjadruje, koľko udalostí s kurzom a dá rovnaký kurz, ako udalosť s kurzom b . Rovnako to platí pre pravdepodobnosti. Z toho je tiež vidieť, že čím má udalosť nižší kurz, tým je väčšia pravdepodobnosť jej nastania.

Jednoduchou úpravou zo vzťahu (1) vyjadríme k_1 a $P(X=1)$

$$k_1 = \frac{\ln b}{\ln a} \text{ a } P(X=1) = \sqrt[k_1]{P(X=0)}. \quad (2)$$

Prirodzené logaritmy vo vzťahu (2) existujú vždy, pretože pre hodnoty kurzov sú vždy väčšie ako jedna ($a > 1$, $b > 1$, $c > 1$).

Rovnakou úvahou získame aj vzťah (3) z ktorého vyplýva vzťah (4)

$$\frac{P^{k_2}(X=2)}{P(X=0)} = \frac{c^{k_2}}{b} = 1, \text{ kde } k_2 \in R, \quad (3)$$

$$k_2 = \frac{\ln b}{\ln c} \text{ a } P(X=2) = \sqrt[k_2]{P(X=0)}. \quad (4)$$

Pri určovaní pravdepodobností $P(X=i)$, $i=0,1,2$ využijeme uvedené vzťahy a to, že uvedené udalosti tvoria úplný systém vzájomne sa vylučujúcich udalostí:

$$P(X=1) + P(X=0) + P(X=2) = 1. \quad (5)$$

Dosadením z (2) a (4) do (5) dostaneme:

$$\sqrt[k_1]{P(X=0)} + P(X=0) + \sqrt[k_2]{P(X=0)} = 1,$$

$$\sqrt[\frac{\ln b}{\ln a}]{P(X=0)} + P(X=0) + \sqrt[\frac{\ln b}{\ln c}]{P(X=0)} = 1. \quad (6)$$

Riešením rovnice (6) je pravdepodobnosť toho, že sa zápas skončí nerozhodne. Pravdepodobnosti $P(X=1)$, resp. $P(X=2)$ získame dosadením $P(X=0)$ do vzťahov (2) resp. (4).

Rovnica (6) sa vo všeobecnosti nedá riešiť analytickými metódami a je preto nutné použiť vhodné numerické metódy, prípadne matematický softvér, ktorý dokáže túto rovnicu vyriešiť. My sme použili program Mathematica 7.

Pred testovaním modelu sme museli vyriešiť ten problém, že ak aj boli napríklad rovnaké kurzy na víťazstvo jedného tímu, tak na ostatné dve udalosti sa už kurzy rovnať nemuseli a ani sa vo väčšine prípadov nerovnali. Tento problém sme vyriešili tak, že množinu všetkých stretnutí sme rozdelili na podmnožiny podľa kurzu na favorita stretnutia, tak ako je to uvedené v tabuľke 3.

Tabuľka 3: Stredné kurzy a ich pravdepodobnosti

kurz		stredné kurzy			pravdepodobnosti		
od	do	F	R	S	F	R	S
1,05	1,20	1,144	14,814	5,806	0,847	0,036	0,117
1,21	1,30	1,248	13,860	4,080	0,769	0,044	0,187
1,31	1,40	1,352	13,040	3,349	0,704	0,050	0,246
1,41	1,50	1,444	12,260	2,907	0,654	0,055	0,291
1,51	1,60	1,549	11,860	2,580	0,604	0,059	0,337
1,61	1,70	1,649	11,170	2,370	0,565	0,063	0,372
1,71	1,80	1,761	10,441	2,157	0,525	0,068	0,407

V každej podmnožine sme vypočítali „stredný“ kurz na favorita stretnutia F, stredný kurz na remízu R a stredný kurz na víťazstvo slabšieho tímu S. Stredný kurz na favorita sme vypočítali ako geometrický priemer všetkých kurzov na favorita v príslušnej podmnožine. Analogicky sme postupovali aj pri výpočte stredného kurzu na remízu a víťazstvo slabšieho tímu.

V tabuľke 3 sú tiež uvedené stredné kurzy v jednotlivých podmnožinách a im prislúchajúce pravdepodobnosti.

Vhodnosť tohto modelu sme testovali na údajoch získaných zo stretnutí basketbalovej súťaže NBA, kde sme počas celej sezóny 2009/2010 spolu s vypísanými kurzami sledovali aj výsledky týchto stretnutí.

V tabuľke 4 sú uvedené empirické početnosti jednotlivých udalostí F, R, S a ich teoretické početnosti odhadnuté použitím modelu. Teoretické početnosti sme vypočítali ako súčin celkového počtu prípadov N v danej kategórii a príslušnej pravdepodobnosti.

V poslednom stĺpci sú p -hodnoty χ^2 testu dobrej zhody, ktorým sme na hladine významnosti $\alpha = 0,05$ testovali zhodu medzi teoretickými a empirickými hodnotami.

V každom zo skúmaných prípadov je táto hodnota väčšia ako zvolená hladina významnosti, a preto môžeme povedať, že uvedený model odhadu pravdepodobností zo známych kurzov dobre popisuje získané údaje.

Tabuľka 4: Výsledky testovania modelu

kurz		početnosti prípadov				teoretické početnosti			χ^2 test	
od	do	N	F	R	S	F	R	S	χ^2	p -hodnota
1,05	1,20	170	149	4	17	138,22	5,96	18,82	1,661	0,434
1,21	1,30	121	99	2	20	93,18	5,20	22,62	2,671	0,263
1,31	1,40	148	97	10	41	104,20	7,40	36,40	1,993	0,369
1,41	1,50	113	71	12	30	73,80	6,20	33,00	5,837	0,054
1,51	1,60	119	68	7	44	71,87	6,90	40,23	0,610	0,737
1,61	1,70	104	62	8	34	58,77	6,55	38,68	1,079	0,583
1,71	1,80	91	51	5	35	47,78	6,09	37,13	0,540	0,763

4. Záver

Na základe zistených údajov môžeme konštatovať, že aj jednoduchá analýza reálnych dát nám môže pomôcť poodhaliť skutočnosti, ktoré by sme neočakávali a zrejme ani nepostrehli. Sledovaním kurzov a stávkovaním pravdepodobne nezbohatneme (ak nie sme prevádzkovatelia), ale to nemusí byť našim cieľom. Hľadanie vzťahov medzi tým, čo pozorujeme a mali by sme pozorovať, je azda lákavejšie.

5. Literatúra

- [1]ANDĚL, J. 2003. Statistické metody. Praha: MATFYZPRESS, 2001, ISBN 80-86732-08-8.
- [2]HEBÁK, P. – HUSTOPECKÝ, J.- JAROŠOVÁ, E.- PECÁKOVÁ, I. , 2007. Vícerozměrné statistické metody [1]. Praha: Informatorium, 2007. ISBN 978-80-7333-056-9.
- [3]ZVÁRA, K. – ŠTĚPÁN, J. , 2001. Pravděpodobnost a matematická statistika. Praha: MATFYZPRESS, 2001. ISBN 80-2240736-4.

Adresa autora:

Ľubomír Rybanský, RNDr.
Tr. A. Hlinku 1
949 01 Nitra
lubomir.rybansky@ukf.sk

Cílování inflace v České republice Inflation targeting in the Czech Republic

Kristina Soukupová

Abstract:

Czech Nation Bank implemented an Inflation targeting regime in 1998 and obliged to maintain the price stability. The main features of inflation targeting are using an inflation forecast and the explicit public announcement of an inflation target or sequence of targets. Decisions, how to change the settings of monetary policy instruments – interest rates to achieve its target, are based on inflation forecast. This article deals with Inflation targeting regime in the Czech Republic and ability of Czech National Bank generate credible inflation forecast as an elementary task to make correct decisions. Main aim of this article is an empirical analysis which judges the effectiveness of Czech National Bank to stimulate price stability through the changes of interest rates settings.

Key words: inflation targeting, inflation forecast, interest rates, Granger causality test.

Klíčová slova: cílování inflace, predikce inflace, úrokové sazby, Grangerův test kauzality.

1. Úvod

Režim cílování inflace byl poprvé zaveden novozélandskou centrální bankou v roce 1990. Postupně režim cílování inflace začaly zavádět centrální dalších zemí, které nejčastěji opouštěly režim cílování měnového kurzu. První centrální bankou v transformující se ekonomice, která zavedla strategii cílování inflace, byla Česká národní banka.

Podstatou cílování inflace je určení žádoucí trajektorie vývoje inflace v budoucnosti v krátkém až středně dlouhém období. Žádoucí hodnoty jsou definovány jako bodové hodnoty nebo jako interval, kdy střední hodnota intervalu může být považována za vlastní cíl. Po stanovení inflačního koridoru a zveřejnění střední hodnoty představující inflační cíl, sestaví centrální banka skutečnou predikci inflace (Kvasnička, 2000). Stoupenci cílování inflace zdůrazňují, že k získání této predikce musí být využity všechny dostupné informace a nástroje. Centrální banka se tak neomezuje na využití jediného zprostředkujícího kritéria, ani na jedinou ekonomickou teorii. Centrální banka naopak využívá mnoha kvantitativních modelů, expertních odhadů a dalších nástrojů v rámci svého predikčního aparátu. Inflační predikce je dosaženo vzájemným porovnáním výsledků nebo jejich agregací. Dalším krokem je srovnání inflační predikce s inflačním cílem. Pokud je predikovaná inflace vyšší než cílová (inflace přestřeluje cíl), centrální banka zpřísní svou monetární politiku, a to například tím způsobem, že sníží množství peněz v oběhu zvýšením úrokových sazeb. Reakce ekonomických subjektů na změny v úrokových sazbách závisí na rychlosti jejich reakce, efektivnost nastavení úrokových sazeb a především na transmisním mechanismu. Ten představuje kauzální vztahy mezi použitými měnovými nástroji centrální banky a cílem, jehož má být těmito nástroji dosaženo. Jinými slovy je transmisní mechanismus definován jako řetězec ekonomických vazeb, který umožňuje, aby změny v charakteru monetární politiky vedly k naplnění konečných cílů centrální banky.

Za hlavní výhodu cílování inflace bývá považována její transparentnost, srozumitelnost a jasný závazek udržet inflaci v určitých mezích. Předpokládá se, že inflační závazek centrální banky podaný ve srozumitelné podobě veřejnosti výrazně ovlivní očekávání ekonomických subjektů. Závaznost a transparentnost usnadní kontrolu jednání centrální banky a umožní bance bránit se proti politickým tlakům. Za další výhodu bývá považována nezávislost

inflačního cílování na jediném zprostředkujícím cíli. Inflační cílování umožňuje využití maximálního množství informací, využití většího množství modelů a teoretických přístupů. Význam inflačního cílování však především tkví v jeho krátkodobém až střednědobém charakteru.

Česká národní banka zavedla cílování inflace v lednu 1998 a jejím základním cílem se stala péče o cenovou stabilitu. Tento cíl je primární a nadřazený ostatním funkcím a činnostem, které banka vykonává. Česká národní banka rovněž podporuje obecnou hospodářskou politiku vlády, pokud není tento sekundární cíl v rozporu s cílem primárním. První inflační cíle České národní banky byly stanoveny pomocí tzv. čisté inflace, jež představuje podmnožinu celkové spotřebitelské inflace očištěnou o administrativní zásahy státu. Od roku 2002 Česká národní banka cíluje celkovou inflaci s tím, že ji odhaduje na 12-14 měsíců dopředu, tzv. období nejúčinnější transmise. Navíc od července 2002 Česká národní banka přechází od vyhlášení podmíněné inflační predikce, fungující za předpokladu neměnnosti měnové politických nástrojů, ke stanovování predikce nepodmíněné. Ta spočívá v aktivní měnové politice centrální banky.

Cílem předkládaného příspěvku je posouzení účinnosti strategie cílování inflace v České republice a schopností České národní banky vytvářet důvěryhodnou predikci jako stěžejní úkol pro její rozhodování. Hodnotí zda Česká národní banka dokáže efektivně kontrolovat cenovou hladinu a cílenými změnami v úrokových sazbách ji stimulovat či nikoliv.

2. Metodika

Východiskem k vymezení proměnných pro empirickou analýzu je myšlenka Taylorova pravidla, které popisuje relaci mezi inflací, hospodářským růstem a monetární politikou centrální banky. Taylorovo pravidlo je elegantním nástrojem pro popis chování centrálních bank, vycházejícího z reakcí na odchylky inflace a ekonomického růstu z požadovaných cílovaných úrovní. V tomto příspěvku je využito Taylorova pravidla, které je pro případ otevřené ekonomiky ve tvaru:

$$i = i^* + a(p - p^*) + b(y - y^*) + dER_t \quad (1)$$

kde i je výsledná požadovaná úroková sazba, i^* rovnovážná úroková sazba (používají se dlouhodobé vládní cenné papíry), p míra inflace, p^* inflační cíl, y růst HDP, y^* potenciální růst, $(y - y^*)$ mezera růstu, dER_t je devizový kurz, a a b parametry monetární politiky (Taylor, 1993).

V ekonometrických modelech se můžeme setkat s explicitně definovanou kauzalitou mezi proměnnými u Grangerova testu kauzality. Ve smyslu Grangerovy kauzality říkáme, že „ x_t ovlivňuje y_t , pokud minulé hodnoty x_t pomáhají lépe vysvětlit chování y_t (v regresi y_t na minulé hodnoty y_t a x_t)“ (Kočenda, 1998). Pokud zpožděné hodnoty proměnné x_t nemají vysvětlovací schopnost žádné proměnné v systému, pak lze na x pohlížet jako na slabě exogenní. S přihlédnutím ke směru kauzality můžeme rozlišit následující případy:

- jednosměrná kauzalita: x_t determinuje y_t , ne však naopak.
- obousměrná kauzalita: x_t a y_t jsou vzájemně determinovány.

Uvažujme model VAR(k) o dvou proměnných x_t a y_t :

$$y_t = \alpha_{10} + \sum_{j=1}^k \alpha_{1j} x_{t-j} + \sum_{j=1}^k \beta_{1j} y_{t-j} + \varepsilon_{1t} \quad (2)$$

$$x_t = \alpha_{20} + \sum_{j=1}^k \alpha_{2j} x_{t-j} + \sum_{j=1}^k \beta_{2j} y_{t-j} + \varepsilon_{2t} \quad (3)$$

Lze rozlišit následující případy:

Tabulka 1: Typy kauzalit pro model VAR(k) o dvou proměnných x_t a y_t :

$\{\alpha_{11}, \alpha_{12}, \dots, \alpha_{1k}\}$	$\{\beta_{21}, \beta_{22}, \dots, \beta_{2k}\}$	typ kauzality
$\neq 0$	$= 0$	jednosměrná kauzalita $x \rightarrow y$
$= 0$	$\neq 0$	jednosměrná kauzalita $y \rightarrow x$
$\neq 0$	$\neq 0$	obousměrná kauzalita $x \leftrightarrow y$

Zdroj: Seddighi, 2000

Pro testování kauzality odkazující k významnosti parametrů výše popsaného modelu VAR je použito Wald FW statistiky ve tvaru

$$FW = \frac{(RSS_r - RSS_u) / k}{RSS_u / (n - 2k - 1)} \sim F(k, n - 2k - 1), \quad (4)$$

kde RSS_u je reziduální součet čtverců bez omezení, RSS_r je reziduální součet čtverců s předpokladem, že množina proměnných je omezená, k je počet omezení a n je počet pozorování (Seddighi, 2000). Testované hypotézy jsou:

H_0 : x neovlivňuje v Grangerově smyslu y , $\{\alpha_{11}, \alpha_{12}, \dots, \alpha_{1k}\} = 0$, neprůkazné, $FW < F$,

H_1 : x ovlivňuje v Grangerově smyslu y , $\{\alpha_{11}, \alpha_{12}, \dots, \alpha_{1k}\} \neq 0$, průkazné, $FW \geq F$,

a dále

H_0 : y neovlivňuje v Grangerově smyslu x , $\{\beta_{21}, \beta_{22}, \dots, \beta_{2k}\} = 0$, neprůkazné, $FW < F$,

H_1 : x ovlivňuje v Grangerově smyslu y , $\{\beta_{21}, \beta_{22}, \dots, \beta_{2k}\} \neq 0$, průkazné, $FW \geq F$.

3. Výsledky empirické analýzy

Pro analýzu cílování inflace v České republice prostřednictvím identifikace vazeb mezi proměnnými, plynoucími z reakční funkce centrální banky, byly zvoleny čtvrtletní údaje o dvoutýdenních úrokových sazbách na mezibankovním trhu (i), hrubém domácím produktu ve stálých cenách (y), cenové hladině vyjádřené přírůstkem spotřebitelských cen (p) a reálném devizovém kurzu (dER) za období 1998Q1 – 2009Q3. Vzhledem k velké otevřenosti ekonomiky je reakční funkce centrální banky v tomto příspěvku reprezentována Taylorovým pravidlem rozšířeným o devizový kurz.

Součástí empirické analýzy bylo učení délky zpoždění pro VAR, proto bylo využito Akaikeho informačního kritéria (AIC), Schwartzova kritéria (SC) a Hannan – Quinnova kritéria (HQ) (Seddighi, 2000). Výsledky délky zpoždění udává tabulka č. 2. Z této tabulky vyplývá optimální délka zpětného zpoždění jako osm čtvrtletí, což odpovídá nejnižším hodnotám všech tří testových kritérií.

Tabulka 2: Výsledky identifikace délky zpětného zpoždění, 1998Q1-2009/Q3 ČR

Délka zpoždění	AIC	HQ	SC
1Q	0.0805	0.0829	0.0875
2Q	0.0755	0.0804	0.0896
3Q	0.0718	0.0791	0.0929
4Q	0.0740	0.0838	0.1021
5Q	0.0631	0.0754	0.0983
6Q	0.0588	0.0736	0.1010
7Q	0.0438	0.0610	0.0931
8Q	0.0147	0.0343	0.0710

Výsledky Grangerova testu kauzality s ohledem na VAR(8), uvádí tabulka č. 3.

Tabulka 3: Výsledky testu Grangerovy kauzality v období 1998Q1-2009/Q3 ČR

Závisle proměnná	Nezávisle proměnná	FW-hodnota	p-hodnota	Významnost
i	i	3.6694	0.0653	(*)
	y	3.4533	0.0742	(*)
	dER	1.2267	0.4131	
	p	1.6196	0.2868	
y	i	1.1796	0.4322	
	y	12.651	0.0031	(***)
	dER	3.8828	0.0578	(*)
	p	4.6830	0.0380	(**)
dER	i	5.3004	0.0285	(**)
	y	4.1382	0.0502	(**)
	dER	3.0617	0.0949	(*)
	p	3.4176	0.0751	(*)
p	i	2.1905	0.1775	
	y	4.1263	0.0506	(**)
	dER	2.4766	0.1427	
	p	8.6769	0.0084	(***)

Pozn.: Hladina významnosti na 1 % (***), na 5 % (**) a na 10 % (*)

Veškeré uvedené dosažené výsledky jsou v tabulce č. 4 transformovány tak, aby bylo zřejmé, jaký mají kauzality charakter. Následně bude možné celkově zhodnotit empirickou analýzu a formulovat závěry o efektivitě měnové politiky České národní banky.

Tabulka 4: Identifikace kauzalit proměnných

Závisle proměnná i	Závisle proměnná y	Závisle proměnná dER	Závisle proměnná p
y: $y \rightarrow i$	i: x	i: $i \rightarrow dER$	i: x
dER: x	dER: $dER \rightarrow y$	y: $y \rightarrow dER$	y: $y \rightarrow p$
p: x	p: $p \rightarrow y$	p: $p \rightarrow dER$	dER: x

Na základě Grangerova testu kauzality byla zjištěna jednostranná vazba směřující od ekonomického růstu k úrokovým sazbám, tedy změny v ekonomickém růstu předcházejí změnám v úrokových sazbách na hladině významnosti 10 %. S ohledem na prokázanou vazbu vyplývá, že centrální banka spíše reaguje až na skutečný vývoj hospodářství a úrokové sazby mění v závislosti na změnách v ekonomickém růstu. Grangerův test kauzality dále identifikoval oboustrannou kauzalitu mezi ekonomickým růstem a inflací na hladině významnosti 5 %. Oboustranná kauzalita mezi devizovým kurzem a ekonomickým růstem, prokázána na hladině významnosti 5 %, resp. 10 %, je odrazem otevřenosti ekonomiky a významnosti mezinárodního obchodu v České republice, kdy výše devizových kursů a jejich

změny působí na ceny dovozu a vývozu zboží a služeb. Neprokázání kauzality směřující od úrokových sazeb k inflaci jako hlavního předpokladu fungování cílování inflace bylo překvapujícím výsledkem, což značně zredukovalo představy o efektivitě prováděné monetární politiky České národní banky.

4. Závěr

Předkládaný příspěvek posuzuje zda Česká národní banka dokáže stimulovat cenovou hladinu za pomoci svých měnově politických nástrojů – krátkodobých úrokových sazeb – a efektivně tak provádět svoji monetární politiku. Cílování inflace v České republice má dnes již dvanáctiletou zkušenost a je prováděno s pomocí sofistikovaného predikčního aparátu schopného odhadnout budoucí vývoj inflace, který naznačí jaké změny v úrokových sazbách jsou adekvátní k dosažení konečného cíle, nízké hodnoty inflace.

Aplikace Grangerova testu kauzality však neprokázal, že změny v úrokových sazbách předchází změnám v cenové hladině, naopak prokázal kauzalitu směřující od ekonomického růstu ke změnám v úrokových sazbách. Výsledek značí, že Česká národní banka aktivně reaguje změnami v úrokových sazbách na aktuální vývoj v hospodářství, oproti cílenému ovlivňování budoucích hodnot inflace. Tento závěr prokazuje, že Česká národní banka neplní svůj primární cíl, ke kterému se zavázala. Vezmeme-li však v úvahu, že se jedná o jednoduchý model, i to, že reakční funkci centrální banky reprezentovalo Taylorovo pravidlo, které bylo původně vytvořeno pro odlišný typ ekonomiky, nelze s jistým přesvědčením zamítnout hypotézu o efektivním fungování cílování inflace. Vzhledem k tomu, že ve sledovaném období došlo i ke změně metodiky v predikčním aparátu, mohlo dojít k určité distorzi v datech.

Z výsledků empirické analýzy tohoto příspěvku je možné říci, že Česká národní banka v letech 1998–2009 byla velmi aktivní v provádění monetární politiky a důrazně sledovala vývoj v ekonomice na úkor cíleného ovlivňování hodnot inflace. Do budoucna by se tak měla zaměřit pouze na svůj primární cíl, snížit objem reakcí na šoky v ekonomice nebo využívat více institutu výjimek.

Předkládaný příspěvek vznikl za podpory interního grantového projektu 71/2010 “Stabilizační funkce monetární politiky v souvislostech hospodářského cyklu České republiky”.

5. Literatura

- [1]GREEN, W.H. 1997. *Econometric Analyses*. London: Prentice – Hall, 1997. 1076 s. ISBN 0-13-7246659-5.
- [2]HANOUSEK, J. – KOČENDA, E. 1998. Monetární vazby na českém mezibankovním trhu. *Finance a úvěr*, č.2. roč. 48. s. 99 – 109.
- [3]KVASNIČKA, M. 2000. Cílování inflace: teorie a praxe. *Politická ekonomie*, s.647 – 658.
- [4]POMĚNKOVÁ, J. – KAPOUNEK, S. 2009. Interest Rates and Causality in the Czech Republic – Granger Approach. *Agricultural Economics*, č.7. roč.55. s. 347 – 356.
- [5]SEDDIGHI, H.R. *Econometrics: A Practical Approach*. London: Routledge, 2000. 396 s. ISBN 0-415-15645-9.
- [6]TAYLOR, J.B. 1993. *Macroeconomics Policy in a World Economy: from econometric design to a practical operation*. New York: WW.Norton&Company, ISBN 0-393-96316-0.

Adresa autora:

Kristina Soukupová, Bc.

Ústav financí

PEF, MENDELU v Brně

Zemědělská 1

613 00 Brno

xsoukup4@node.mendelu.cz

**Štatistika pre 21. storočie
vo vedeckom výskume a vzdelávaní
Statistics for the Twenty-First Century
in the Scientific Research and in the Education**

Beáta Stehlíková, Dagmar Markechová

Abstract: On May 6-8, 2002 approximately fifty statisticians from around the world gathered at the workshop of the National Science Foundation to identify the future challenges and opportunities for the statistics profession. In this contribution we cite the basic ideas presented by the participants of this workshop.

Key words: statistics in scientific research, statistics in education, theoretical and applied research

Kľúčové slová: štatistika vo vedeckom výskume, štatistika vo vzdelávaní, teoretický a aplikovaný výskum

1. Úvod

V snahe zhodnotiť možnosti a stanoviť budúce výzvy štatistiky výbor National Science Foundation¹ zvolal workshop s názvom Statistics: Challenges and Opportunities for the Twenty-First Century, ktorý sa konal v dňoch 6. až 8. mája 2002. Workshopu o súčasnosti a budúcnosti štatistiky sa zúčastnilo približne 50 významných štatistikov z celého sveta. Seminár bol z veľkej časti zameraný na vedecko-výskumnú oblasť, ale účastníci boli požiadaní, aby diskutovali aj o úlohe vzdelávania a o odbornej príprave na dosiahnutie dlhodobých cieľov, ako aj o postavení štatistikov. Seminár sa uskutočnil aj na základe potreby objasniť úlohu štatistiky v oblasti vedy a výskumu. Zo seminára bola vypracovaná správa [1], v ktorej sú rozpracované základné problémy štatistiky a sú v nej stanovené hlavné oblasti štatistiky, na ktoré by sa štatistici mali v budúcnosti zamerať. Bolo vymedzených týchto päť hlavných oblastí:

- *Podpora jednoznačnej identifikácie štatistiky.* Ide o vymedzenie postavenia štatistiky voči ostatným vedám, a to najmä pri objasňovaní rozdielov medzi štatistikou a inými matematickými oblasťami. Toto vymedzenie umožní štatistike používať nové netradičné nástroje, ktoré pomôžu jej rastu ako vednej disciplíny.
- *Posilnenie výskumnej oblasti.* Je potrebné posilniť a zachovať jednak vedomostný rast a potom aj rozširovanie financovania rôznych programov v tejto oblasti.
- *Posilnenie multidisciplinárneho výskumu.* To znamená rešpektovanie rovnoprávneho postavenia štatistikov a odborníkov iných vedných disciplín pri spolupráci vo vedeckých výskumoch v danej vednej disciplíne.

¹ National Science Foundation (NSF) je nezávislá federálna agentúra, ktorú vytvoril Kongres USA v roku 1950 „na podporu pokroku vedy, národného zdravia, prosperity a blahobytu, na zabezpečenie národnej obrany ...“ S ročným rozpočtom približne 6,9 dolárov miliardy eur je zdrojom financovania pre približne 20 percent federálne podporovaného základného výskumu na amerických školách. V mnohých oblastiach, ako je matematika, počítačové vedy a spoločenské vedy, NSF je hlavným zdrojom federálnej pomoci.

- *Vývoj nových modelov pre štatistické vzdelávanie.* To by malo byť realizované na všetkých úrovniach, a to zvlášť pre veľké spoločnosti a zvlášť pre malé spoločnosti. Na dosiahnutie tohto cieľa bude potrebné čo najviac zjednotiť spoločný prístup vzdelávacích inštitúcií.
- *Získavanie novej generácie štatistikov.* Presvedčiť mladých o potrebe a možnostiach štatistiky.

Uvedené oblasti záujmu a tiež aj spôsoby na dosiahnutie vytýčených cieľov sú podrobne popísané v spomínanej správe.

Cieľom tohto príspevku je priblížiť hlavné myšlienky, ktoré na workshope o súčasnosti a budúcnosti štatistiky a štatistikov odzneli.

2. Čo je to štatistika

Táto otázka rezonovala počas celého workshopu a účastníci vo svojich vystúpeniach hľadali na ňu odpoveď. Hlavným rečníkom workshopu bol významný profesor doktor D. R. Cox z Oxford University². Vo svojej prednáške s názvom "Čo je to štatistika?" hľadal odpovede na kľúčové otázky štatistiky. Okrem iného skonštatoval, že štatistika je veda, ktorá získava zmysluplné informácie z údajov všetkých typov, pričom štatistici v spolupráci s odborníkmi z iných vedných disciplín sa podieľajú na celej vedeckovýskumnej práci od začiatku experimentu cez analýzu dát až k záverom.

Workshop bol zameraný na šesť hlavných oblastí: jadro štatistiky a päť hlavných oblastí použitia štatistiky v:

- (1) biologických vedách,
- (2) inžinierskych a priemyselných vedných oblastiach,
- (3) geologických a environmentálnych vedách,
- (4) informačných technológiách a fyzikálnych vedách,
- (5) sociálnych a ekonomických vedách.

Tieto kategórie boli vybrané tak, aby približne zodpovedali oblastiam, ktorých výskum je finančne podporovaný agentúrou National Science Foundation.

Hlavnou činnosťou štatistikov - okrem spolupráce s inými vednými oblasťami - je tvorba matematických a koncepčných nástrojov, ktoré môžu byť použité na extrakciu informácií. Aj keď má štatistika svoj teoretický výskum, jej konečným cieľom je poskytovať aplikovateľné výsledky. Štatistika bola vždy spojená s aplikáciou výsledkov. Význam jej výsledkov, a to aj v teoretickej štatistike, je silne závislý na tom, aby výsledky boli pre danú aplikáciu relevantné. V tomto aspekte sa štatistika líši od ostatných odborov matematických vied, s výnimkou výpočtovej matematiky. Jeden zo zásadných rozdielov medzi matematikou a štatistikou je v rovnováhe medzi vytváraním koncepcie v základnej matematike a interakciou s oblasťami, ktoré používajú matematiku, ako sú veda, technika, technológie, financie, národná bezpečnosť a pod.

Jednou z kľúčových zásad štatistiky by mala byť opatrnosť pri hľadaní vedeckej pravdy v údajoch. Pri hľadaní súvislostí v dátach je potrebná opatrnosť pri formulovaní a overovaní hypotéz. Pripomeňme si výrok Marka Twaina, ktorý povedal: "Existujú tri druhy klamstiev: „lži, veľké lži a lži štatistiky“. V skutočnosti sú štatistici tí odborníci, ktorí pomáhajú oddeľovať vedeckú pravdu od vedeckej fikcie. Zásadnou úlohou štatistického výskumu je preto vyvinúť také nové metódy, ktoré by boli použiteľné pre iné vedné disciplíny.

² Sir David R. Cox, Oxford University, UK

Jednou z kľúčových oblastí, ktorou sa zaoberali odborníci na seminári bola aj výchova a pracovné prostredie štatistikov. Keďže semináru sa zúčastnili prevažne odborníci z popredných amerických univerzít, spôsob vzdelávania, ktorý je v správe popísaný sa týka prevažne amerického školského systému. Na prevažnej väčšine univerzít je Pravdepodobnosť a štatistika jedným z oddelení resp. katedier matematiky podobne ako Algebra alebo Topológia. Účastníci seminára však poukázali na to, že štatistika sa čím ďalej tým viac odlišuje od ostatných matematických disciplín. V štatistike sú počítač a nástroje príslušnej vednej disciplíny, kde štatistiku aplikujeme, rovnako dôležité ako teória pravdepodobnosti. Oddelenia resp. katedry štatistiky už existujú aj na univerzitách humanitného zamerania, na ekonomických univerzitách, medicínskych univerzitách a pod. Okrem akademickej pôdy štatistici pracujú tiež vo výskume, vládnych agentúrach, vo farmaceutických odvetviach, v priemysle a v ďalších odvetviach.

3. Súčasný stav

V dôsledku nástupu výkonných počítačov a senzorov dochádza k veľkému nárastu objemu dát a tiež aj k zmenám v zbere dát. Tieto zmeny vytvárajú nové možnosti pre štatistiku ako účinného nástroja pre organizovanie týchto dát a extrahovanie podstatných informácií z dát. Zo všetkých matematických vied práve štatistická veda je jednoznačne zameraná na zber a analýzu experimentálnych údajov. Ďalším motívom pre rozvoj súčasnej štatistiky je tlak iných vedných disciplín, keďže ich rozvoj je úzko spätý aj s rozvojom štatistiky.

Výskum v štatistike sa rozdeľuje na výskum v „jadre“ štatistiky a aplikovaný výskum viazaný na iné vedné disciplíny. Výskum v jadre štatistiky je zameraný na rozvoj štatistických modelov a metód súvisiacich s matematickou teóriou. Cieľom tohto výskumu je vytvoriť jednotiacu filozofiu, koncepciu, štatistické metódy a výpočtové nástroje. Tieto základné metodiky môžu byť použité súčasne v širokom spektre vedných disciplín a v aplikáciách a v dôsledku toho sú užitočné pre všetky vedy. Výskum v „jadre“ štatistiky môže byť v kontraste s aplikovaným výskumom, špecifickým pre daný riešený problém. Môže čerpať zo základných vedomostí a používať tradičné nástroje. Je základom pre ďalší rozvoj štatistiky.

Pre štatistiku viac ako pre iné vedné disciplíny platí, že jej rozvoj je nepredvídateľný a preto je potrebné udržať základnú filozofiu dostatočne pružnú, aby sa prispôbila zmenám a neprerástla v nesúrodú zbierku techniky.

Riešenie reálnych problémov kladie na štatistikov stále väčšie požiadavky: zvyšujú sa požiadavky na ich matematické vedomosti. Zároveň štatistik musí mať potrebné vedomosti z výpočtovej techniky, aby mohol prispieť k rozvoju algoritmov a počítačových programov potrebných pre analýzu dát. Potreba stále sa zvyšujúcich technických zručností je ďalším argumentom pre zachovanie „jadra“ štatistiky ako dôležitého miesta pre integráciu štatistických nápadov. Motiváciou pre rozvoj „jadra“ štatistiky je jednak skutočnosť, že väčšina dát má spoločné črty a tiež fakt, že štatistické metódy platné vo všeobecnosti, je možné aplikovať v rôznych odboroch.

V oblasti biologických vied štatistický výskum bude závisieť od nových problémov, ktoré v biologických vedách vzniknú. Ide o rozvoj štatistických metód pre klinické štúdie, laboratórne problémy a experimenty. Výpočtové metódy a metódy štatistiky hrajú a budú hrať dôležitú úlohu napríklad v genetike, v genetickej epidemiológii a v mapovaní génov, v populačnej genetike, ekológii a v neurovede. Jednou z oblastí, ktorá podnietila rýchly rozvoj nových štatistických metód, je aj genetika, kde sa štatistické metódy používajú napríklad na meranie génového záznamu v normálnych a chorých bunkách s cieľom, aby sa diagnostikovala choroba. Toto je typický problém, kde treba posúdiť rozdiel medzi kontrolnou a experimentálnou (liečenou) skupinou pre veľký počet (v tisícoch) génov z

relatívne malej vzorky jednotlivcov. Pre posúdenie signifikantnosti rozdielu bola vďaka spolupráci štatistikov a výskumných pracovníkov z onkológie a biochémie vypracovaná metóda SAM (Significance analysis of microarrays). V oblastiach ako sú molekulárna medicína a bunková a vývojová biológia sa v súčasnosti výskum orientuje na využívanie genetických údajov na identifikáciu rizikových faktorov pre drogovú závislosť, na klasifikáciu rôznych chorôb, na proteínové profily a na rozvíjanie individuálnej terapeutickú intervencie založenej na prediktívnych modeloch, ktoré používajú rôzne testovacie štatistiky. Výskum v tomto smere závisí a bude závisieť od výsledkov klinicky orientovanej bioštatistiky. Takýchto príkladov by sme mohli uviesť oveľa viac. Stručne povedané, veľké súbory dát, ktoré sú získavané v moderných biologických experimentoch, vyvolávajú dopyt po štatistikoch, ktorí vedú komunikovať s biológmi a dokážu navrhnúť nový experimentálny dizajn a biologické analýzy dát. Z veľkého počtu štatistických metód, ktoré boli vypracované špeciálne pre potreby biologického výskumu, spomeňme experimentálny dizajn pre genetické mapovanie a zhukové analýzy pre analýzu evolučných vzťahov.

Pokrok v oblasti prieskumu a zberu dát sa prejavil aj v technológiách, ktoré umožňujú zhromažďovať rozsiahle súbory údajov. Tieto údaje majú často zložitú štruktúru v podobe časových radov, priestorových procesov, textov a obrázkov veľmi veľkých rozmerov s hierarchickou štruktúrou. Zber, modelovanie a analýza týchto údajov predstavuje široké spektrum výskumu aj v inžinierskych a priemyselných vedných oblastiach. Potreby obrany, vývoj elektroniky a nových typov lietadiel tiež stimulovali rozvoj nových oblastí štatistiky, akými sú sekvenčná analýza a spoľahlivosť odhadov. Tiež monitoring, diagnostika a vývoj nových výrobných procesov vyžadujú nové metódy pre kompresiu dát a vývoj inteligentných diagnostik. Napríklad, výskumný projekt na zvýšenie výkonu polovodičov viedol k vyvinutiu nových metód na analýzu a vizualizáciu priestorových údajov, vrátane metód pre sledovanie priestorových procesov. Riešenie týchto problémov je veľmi dôležité aj z hospodárskeho hľadiska. Štatistika hrá významnú úlohu aj v softvérovom inžinierstve, kde sú štatistické metódy významným pomocníkom pri vytváraní softvérových metrik a efektívnom využití údajov získaných v experimentoch (napríklad pri redukovaní počtu prípadov potrebných na testovanie softvérovej efektívnosti).

Aplikácie štatistických metód v geofyzikálnych a environmentálnych vedách sú späté s takými rôznorodými oblasťami, ako je napríklad poľnohospodárstvo, základná biológia, stavebníctvo, atmosférická chémia a ekológia. V súčasnosti sa kladie dôraz na využitie štatistických modelov v deterministických procesných modeloch, kde sú štatistické modely motiváciou pre ich ďalší rozvoj. Na základe takéhoto výskumu - kombináciou štatistických modelov a danej vednej disciplíny - možno dospieť k záverom, ktoré budú podnetom k ďalšiemu výskumu. Napríklad, pochopenie mnohých geologických a ekologických procesov môže vo vzájomnej kombinácii so štatistickými modelmi viesť k ďalším novým poznatkom, ktoré budú základom pre nový výskum. Meteorológovia pre štúdium klimatických zmien použili štatistické modely, ktoré boli pôvodne vyvinuté pre problémy ekonómie a aplikovali ich na problematiku spojenú s rýchlosťou vetra. To dokazuje, že mnohé zo štatistických metód sú univerzálne.

Vývoj internetu a exponenciálne rastúce schopnosti počítačových systémov predstavujú netušené možnosti výmeny informácií, zhromažďovania a analyzovania extrémne veľkých súborov rôznych dát z prírody a z iných zdrojov. Kontakty štatistikov s odborníkmi z iných vedeckých oblastí vedú k rozvíjaniu nových štatistických metód, ktoré je možné aplikovať aj v iných oblastiach a pri zbere dát.

V telekomunikáciách je každú minútu denne generované obrovské množstvo komunikačných záznamov. Každý z bezdrôtových a drôtových hovorov je zaznamenaný, zaznamenáva sa miesto, čas, dĺžka hovoru a tiež jeho cena. Každý užívateľ žiadosti na

stiahnutie súboru z internetovej stránky je zaznamenaný v tzv. log súbore, je zaznamenaný aj každý príspevok k on-line chatu vo verejnom fóre. Tieto komunikačné záznamy sú predmetom záujmu sieťových inžinierov, ktorí navrhujú design sietí a rozvoj nových služieb. Takisto sú predmetom záujmu sociológov, ktorí sa zaoberajú tým, ako ľudia komunikujú a aké tvoria sociálne skupiny. Ďalej sú predmetom záujmu odborníkov, ktorí vyšetrujú podvody, a takto sledujú a hľadajú rôzne zločinecké a teroristické skupiny a ich aktivity. Uvedené skutočnosti možno na základe vhodných štatistických metód popísať pravdepodobnostnými modelmi. Pravdepodobnostný model potom umožňuje odhadnúť správanie sa každého jednotlivca, ale aj celej skupiny. K dispozícii je teda celý rad náročných štatistických problémov, ktoré vyžadujú riešenie.

Štatistika sa vždy usilovala o nájdenie abstraktnej podstaty problému a rozvíjala štatistické metódy, ktoré riešili tento abstraktný problém. Aj z tohto dôvodu metódy vyvinuté pre aplikáciu v jednej oblasti vedeckého výskumu sú užitočné aj v inej oblasti vedeckého výskumu. Napríklad, modely vyvinuté pre vedy o živej prírode teraz využívajú ekonómovia.

4. Záver

V predložennom príspevku prezentujeme hlavné myšlienky, ktoré odznali na workshope o súčasnosti a budúcnosti štatistiky a štatistikov, ktorý sa konal v dňoch 6. až 8. mája 2002 a ktorý bol organizovaný nezávislou federálnou agentúrou National Science Foundation. Túto agentúru vytvoril Kongres USA v roku 1950 „na podporu pokroku vedy, národného zdravia, prosperity a blahobytu, na zabezpečenie národnej obrany ...". Účastníci seminára diskutovali o vedecko-výskumných úlohách súčasnej štatistiky, o hlavných smeroch jej rozvoja a tiež o úlohe vzdelávania a o postavení štatistikov.

5. Literatúra

- [1] Statistics: Challenges and Opportunities for the Twenty-First Century, edited by: Jon Kettering, Bruce Lindsay and David Siegmund, 6 April 2003

Adresy autorov:

Beáta Stehlíková, prof. RNDr., CSc.
Katedra matematiky FPV UKF
Tr. A. Hlinku 1
949 01 Nitra1
bstehlikova@ukf.sk

Dagmar Markechová, Doc. RNDr., CSc.
Katedra matematiky FPV UKF
Tr. A. Hlinku 1
949 01 Nitra1
dmarkechova@ukf.sk

Zhluková analýza štruktúry slovenských a anglických vyučovacích hodín **Cluster analysis of the slovak and english lessons**

Ján Šunderlík, Štefan Havrlent, Ľubomír Rybanský

Abstract: This article presents one approach to analysis of the lesson structure within and between countries. We try to adapt the cluster analysis to analyse the data from pre-service student teaching in the last semester of their study. We decide to use the Ward's method that minimize the inner distances between the units. We also discuss the strengths and limitations of this approach.

Key words: Cluster analysis, lesson structure, comparative studies, Ward's method.

Kľúčové slová: Zhluková analýza, štruktúra vyučovacej hodiny, komparatívne štúdie, Wardova metóda.

1. Úvod

V našom príspevku sa budeme venovať štúdiu a analýze štruktúry vyučovacích hodín vedených študentmi učiteľstva matematiky zo Slovenska a z Anglicka. Analyzované vyučovacie hodiny boli odučené počas pedagogickej praxe, ktorú študenti absolvovali na konci svojho prípravného programu štúdia na svojich univerzitách v Anglicku a na Slovensku.

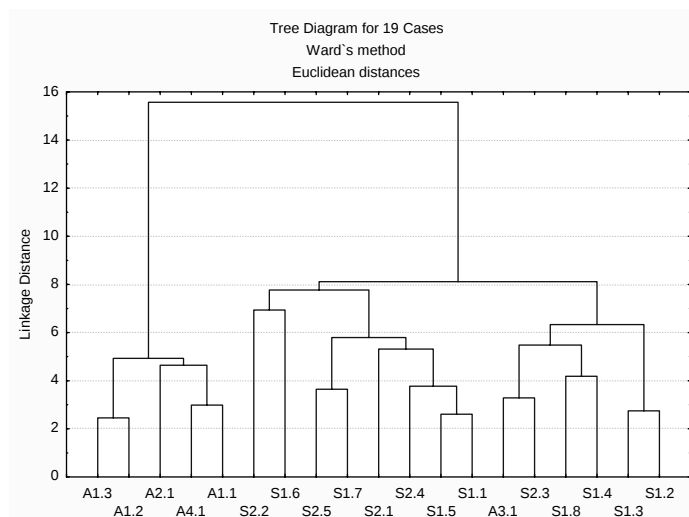
2. Prístup k analýze

Na základe pozorovania devätnástich vyučovacích hodín (kódovanie hodín je nasledovné: prvé písmeno určuje, či išlo o anglickú alebo slovenskú hodinu, prvé číslo určuje učiteľa a druhé číslo je poradové číslo vyučovacej hodiny), sme celkový čas každej vyučovacej hodiny rozdelili do 15 kategórií (premenných) podľa povahy vykonávanej činnosti, tak ako je to uvedené v tabuľke 2.

Získané hodnoty v rámci každej premennej sme štandardizovali t.j. od každej hodnoty príslušnej premennej sme odrátali jej aritmetický priemer a vydělili smerodajnou odchýlkou. Túto transformáciu údajov sme vykonali z dôvodu rôznej škály sledovaných premenných.

Zaujímá nás, či existujú skupiny (zhluky) podobných vyučovacích hodín, ktoré sme mali k dispozícii. Ako vhodný nástroj sa ukazuje použiť zhlukovú analýzu. Pri interpretácii získaných výsledkov budeme vychádzať z dendrogramu získaného Wardovou metódou. Táto metóda minimalizuje celkový súčet štvorcov odchýliek všetkých hodnôt od príslušných zhlukových priemerov. Inak povedané minimalizuje vnútrozhlukový rozptyl.

Charakter sledovaných premenných umožňuje používať viacero metód merania vzdialenosti medzi objektmi. My sme použili euklidovskú vzdialenosť.



Obrázok 1: Dendrogram pre 19 vyučovacích hodín

Z dendrogramu na obrázku 1 vyčítame, že najbližšie sú si objekty A1.2 a A1.3, ktoré sa spájajú v prvom cykle na hladine 2,45 v ďalšom kroku sa na hladine 2,60 spájajú objekty S1.1 a S1.5. Ďalší postup spájania je zrejmy z obrázka 1.

Porovnaním dendrogramov vytvorených rôznymi metódami (ktorých výstupy tu neuvádzame) sme zistili, že štruktúru dát najlepšie vystihuje 5 zhlukov, ktoré získame rezom v dendrograme na hladine vzdialenosti 6,50. Jednotlivé zhluky sú tvorené objektmi tak ako je to uvedené v tabuľke 2.

Tabuľka 1: Zaradenie vyučovacích hodín do zhlukov

zhluk 1	A1.1	A1.2	A1.3	A2.1	A4.1	
zhluk 2	S1.1	S1.5	S1.7	S2.1	S2.4	S2.5
zhluk 3	A3.1	S1.2	S1.3	S1.4	S1.8	S2.3
zhluk 4	S1.6					
zhluk 5	S2.2					

Pri charakterizovaní zhlukov budeme vychádzať z tabuľky 1 v ktorej sú priemerné hodnoty jednotlivých premenných pre každý zhluk. Vzhľadom na štandardizáciu údajov je priemerná hodnota každej premennej 0, čo nám uľahčí interpretáciu výsledkov.

Tabuľka 2: Podmienené priemery pre premenné

	Anglické a slovenské hodiny	zhluk 1	zhluk 2	zhluk 3	zhluk 4	zhluk 5
		priemer	priemer	priemer	priemer	priemer
	nematematická činnosť	0,82	-0,04	-0,43	-0,63	0,31
	kontrola a zadanie DÚ	-0,50	-0,34	0,03	2,96	-0,49
	opakovanie učiva	-0,60	-0,49	0,38	2,28	0,68
	skúšanie	-0,25	-0,25	-0,23	-0,25	4,12
	predstavenie témy, cieľa hodiny	0,96	-0,37	-0,14	-0,69	-0,69
	kroky na riešenie problému prezentované učiteľom	-0,52	0,98	-0,33	-0,26	-0,52
	vysvetľovanie učiva	-0,81	0,02	0,60	-0,94	-0,24
	precvičovanie príkladov, spoločne na tabuľu	-0,31	1,34	-0,51	-0,85	-0,57

	Anglické a slovenské hodiny	zhluk 1	zhluk 2	zhluk 3	zhluk 4	zhluk 5
		priemer	priemer	priemer	priemer	priemer
	žiaci samostatne riešia úlohu	1,38	-0,78	-0,26	-0,78	-0,78
	žiaci prezentujú svoje riešenie	0,39	-0,71	0,55	-0,71	-0,71
	diskusia riešenia	-0,42	-0,12	0,82	-0,77	-0,77
	žiacka otázka	-0,18	0,02	-0,21	3,09	-0,50
	zopakovanie hlavných bodov učiva	-0,53	0,20	-0,52	0,99	1,30
	predstavenie problému dňa	1,23	-0,62	-0,22	-0,62	-0,62
	rozcvička	1,57	-0,56	-0,56	-0,56	-0,56

3. Interpretácia výsledkov

Na základe týchto hodnôt sa pokúsime charakterizovať jednotlivé zhluky a porovnať jednotlivé typy hodín. Naším cieľom nie je vytvorenie akéhosi modelu národného typu hodiny, pretože si uvedomujeme, že samotná štruktúra vyučovacej hodiny je podmienená tým, kde sa vyučovacia hodina v rámci celku nachádza. Na druhej strane na základe hodnôt získaných zo zhlukovej analýzy môžeme konštatovať, v čom sú si jednotlivé vyučovacie hodiny podobné a v čom sú rozdielne.

3.1 Zhluk 1

Pre prvý zhluk, ktorý je tvorený piatimi anglickými vyučovacími hodinami je typické, že prevažná časť vyučovacej hodiny je tvorená *predstavením cieľa* ako aj *problému dňa*, *samostatnou prácou žiakov* a *prezentáciou žiackych výsledkov*. Tieto štyri premenné nadobúdajú spomedzi všetkých zhlukov najväčšie hodnoty. Preto by sme tento zhluk mohli pomenovať ako organizované bádanie žiakov alebo „objavné vyučovanie“.

Študent učiteľstva organizuje vyučovanie tak, aby žiaci mohli samostatne pracovať na matematickom probléme. Tejto práci predchádza rozcvička, ktorá neslúži len na zopakovanie predošlej látky, ale skôr ako zahrievacia činnosť pre ďalšiu prácu žiakov. Ďalšia premenná, ktorá dosahuje nadpriemerné hodnoty je *nematematická činnosť* ktorá zahŕňa napríklad upozorňovanie žiakov. Jedno z možných vysvetlení je, že zvládnutie výučby objavnou metódou si vyžaduje dlhodobejšie úsilie a prax v spôsobe organizácie triedy. Druhým možným vysvetlením je kultúrna podmienenosť správania sa žiakov. Na druhej strane je veľmi cenné, že študenti učiteľstva pracujú spolu s cvičnými učiteľmi, ako aj univerzitnými tútormi, ktorí im v ťažkých začiatkoch pomáhajú nielen overenými riešeniami, ale hlavne spätnou väzbou a pomáhajú im identifikovať kritické momenty a oblasti v ktorých sa študenti potrebujú zlepšovať.

3.2 Zhluk 2

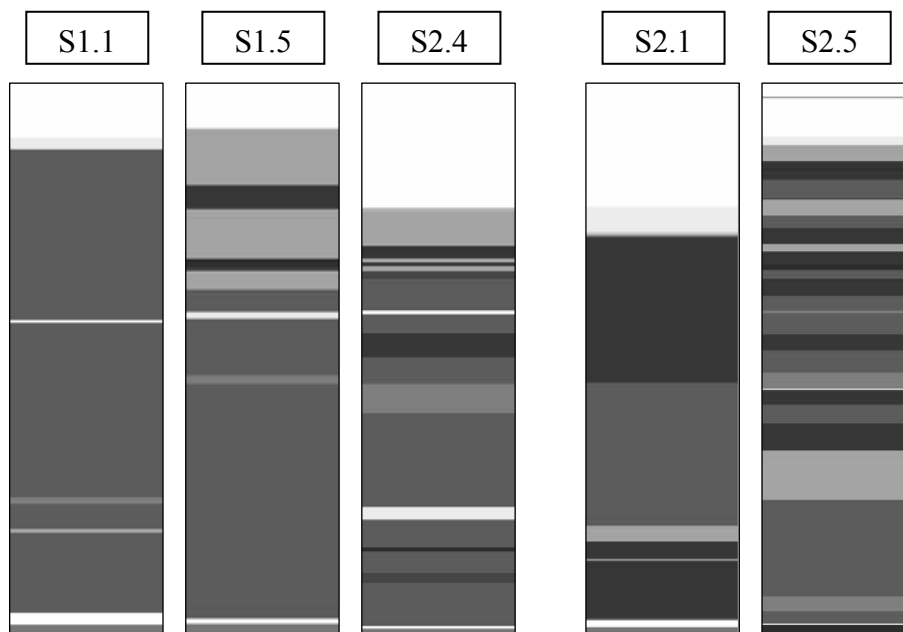
Druhý zhluk nadobúda najväčšie hodnoty pre premenné *kroky na riešenie problému prezentované učiteľom* a *precvičovanie príkladov*. Na základe týchto porovnaní by sme mohli tento zhluk pomenovať ako „utvrdzovanie nového učiva“.

Jedná sa o vyučovacie hodiny na ktorých študenti učiteľstva prezentovali kroky na riešenie nového príkladu, a zadávali úlohy na precvičenie novej látky. To odzrkadľuje aj nadpriemerný čas venovaný týmto dvom činnostiam. Tieto dve činnosti, sú doplnené priemerným časom *vysvetľovania*.

Za nedostatok analýzy považujeme to, že nezohľadňuje v ktorej časti hodiny sa jednotlivé premenné nachádzajú, ale iba ich celkové zastúpenie. Preto napríklad aj v rámci

druhého zhluku by sme mohli oddeliť hodinu S1.1 a hodinu S1.5 kde v prvej z týchto hodín je potrebné si uvedomiť aj následnosť jednotlivých činností v priebehu vyučovacej hodiny

Obrázok 2 znázorňuje časovú následnosť jednotlivých činností na vyučovacích hodinách spolu s časom venovaním tejto činnosti. Legenda k zafarbeniu grafov je v uvedená v tabuľke 1.



Obrázok 2: Grafické znázornenie štruktúry vyučovacej hodiny v zhluke 2

Z uvedených grafov na obrázku 2 môžeme vidieť, ako sa hodnoty prezentované v rámci jedného zhluku odlišujú. V hodine S1.1 je väčšina času venovaná precvičovaniu na príkladoch. Počas hodiny S1.5 je na úvod vyučovacej hodiny zaradené vysvetľovanie a príklad vyriešený učiteľom, za ktorým nasleduje precvičovanie nového učiva na príkladoch.

Iný prípad je charakteristický pre hodiny S2.4 a S2.5, kde študent učiteľstva strieda jednotlivé časti hodiny častejšie, čo má za následok väčšiu členitosť ako aj postupné predstavovanie nového učiva. Nakoľko hodiny S2.4 a S2.5 boli odučené v treťom ročníku, tak tento spôsob by sme mohli charakterizovať ako snahu prebrať väčšie množstvo učiva za kratší čas.

3.3 Zhluk 3

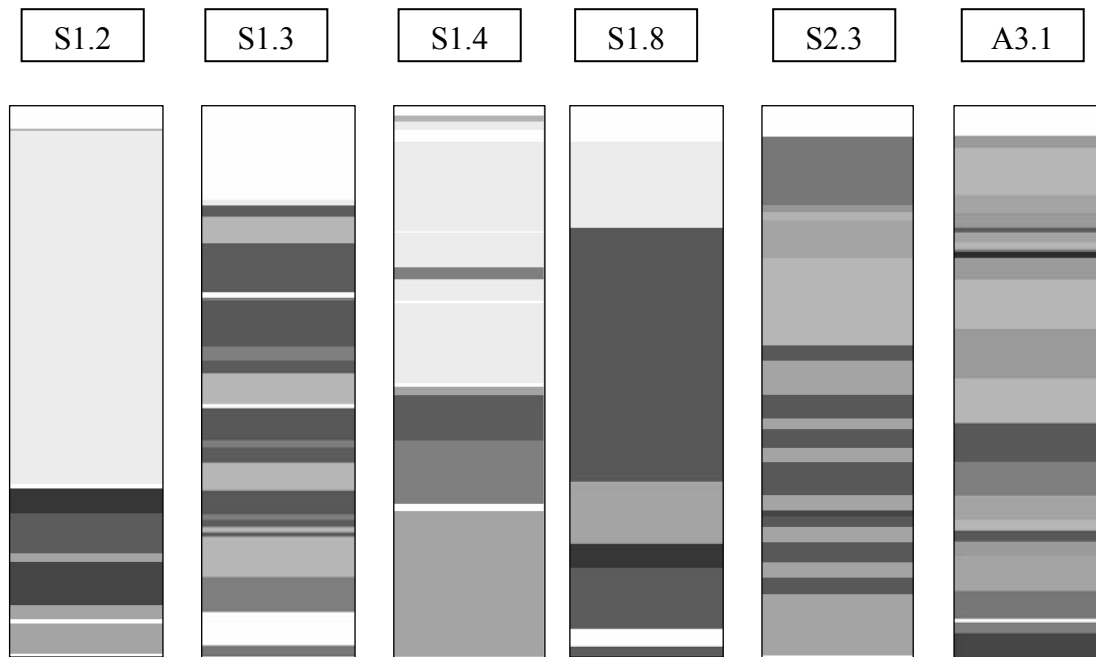
Nadpriemerné hodnoty v zhluke 3 dosahujú premenné *opakovanie učiva*, *vysvetľovanie riešenia*, *žiacka prezentácia riešenia* a najväčšiu hodnotu dosahuje *diskusia riešenia*.

Z týchto hodnôt je ťažké interpretovať spoločné charakteristiky. Preto sme sa bližšie pozreli na to v akom poradí boli tieto vyučovacie hodiny zaradované do zhlukov. Ako prvé boli zoskupené vyučovacie hodiny S1.2 a S1.3. Tieto hodiny boli paralelnými hodinami.

Vo významnej miere bolo zastúpené *opakovanie*, *príklad riešený študentom učiteľstva* ako aj *vysvetľovanie a diktovanie poznámok*. Druhá dvojica hodín je S2.3 a A3.1. Typické pre túto dvojicu je striedanie sa inštrukcií vyučujúceho so samostatnou prácou žiakov. Štruktúra hodín je veľmi podobná, avšak obsah vyučovania odlišný.

Na hodine A3.1 sa študentka učiteľstva venovala otvoreným problémom a bádaniu Möbiovho pásiku, a počas hodiny S2.3 boli zadávané žiakom úlohy uzatvoreného charakteru, s jediným riešením.

Posledná dvojica hodín v tomto zhltuku sú hodiny S1.4 a S1.8. Obe hodiny majú veľmi silné zastúpenie *prezentácie žiackych riešení*. Rozdiel medzi hodinami je opäť v tom, v akej časti hodiny žiaci svoje riešenie prezentujú. Počas hodiny S1.4 ho prezentujú po skupinovej práci v rámci opakovania a počas hodiny S1.8 ho prezentujú po samostatnej práci ako súčasť preberania nového učiva.



Obrázok 2: Grafické znázornenie štruktúry vyučovacích hodín v zhltuku 3

Zhltuk 3 sa nedá pomenovať spoločným názvom na základe prevládajúcich činností. Charakteristické preň je nízke zastúpenie *precvičovania nového učiva*. To je však spôsobené rôznymi zameraniami vyučovacích hodín.

3.4 Zhltuk 4

Zhltuk 4 je tvorený iba jednou vyučovacou hodinou. Táto hodina je špecifická v porovnaní s inými hodinami v tom, že veľa času bolo venované *kontrole domácej úlohy, opakovaniu učiva a žiackym otázkam*.

Táto hodina by sa dala charakterizovať ako opakovacia hodina prípadne „doučovanie“, kde žiaci riešia príklady na zopakovanie jedného typu príkladov.

3.5 Zhltuk 5

Zhltuk 5 je opäť tvorený iba jednou vyučovacou hodinou. Táto hodina sa odlišuje od ostatných tým, že žiaci počas nej písali písomnú prácu počas dvadsiatich minút. Po odovzdaní riešení študent učiteľstva venoval niekoľko minút na preriešenie si príkladov z písomky. Na záver hodiny študent učiteľstva preberal novú látku. Takúto hodinu by sme mohli pomenovať ako „overovanie vedomostí“ a pokračovanie ďalej vo vyučovaní.

4. Záver

Prezentovaná metóda poskytuje zoskupenie vyučovacích hodín do jednotlivých zhlukov, ktoré sme charakterizovali na základe spoločných vlastností. Zaujímavým výsledkom je rozdelenie slovenských a anglických vyučovacích hodín do rozdielnych zhlukov. Vzhľadom na pozorované premenné tak vieme, relatívne ľahko, určiť hlavné rozdiely a podobnosti týchto hodín. Na druhej strane obmedzením prezentovanej zhlukovej analýzy je nezohľadňovanie následností jednotlivých premenných počas vyučovacej hodiny. To má za následok združovanie vyučovacích hodín s iným priebehom a cieľom vyučovania. Preto pri interpretácii výsledkov analýzy je potrebné zohľadniť ten fakt, že cieľ, štruktúra ako aj forma vyučovacej hodiny nie sú navzájom nezávislé a jedna druhú vnútorne ovplyvňuje.

5. Literatúra

- [1]CLARKE, D. – MESITI, C. – JABLONKA, E. – SHIMIZU, Y. 2006. Addressing the Challenge of Legitimate International Comparisons: Lesson structure in the USA, Germany and Japan. In: Making Connections: Comparing Mathematics Classrooms Around The World , 2006, Rotterdam: Sense Publishers, 2006 s. 26 – 45.
- [2]HEBÁK, P. – HUSTOPECKÝ, J. – PECÁKOVÁ, I. – PRÚŠA, M.– ŘEZANKOVÁ, H. – SVOBODOVÁ, A.– VLACH, P. 2007. Vícerozměrné statistické metody (3). Praha: Informatorium, 2007. ISBN 978-80-7333-001-9.

Adresa autorov:

Ján Šunderlík, PaedDr.
Trieda A. Hlinku 1
949 74 Nitra
jan.sunderlik@ukf.sk

Ľubomír Rybanský, RNDr.
Trieda A. Hlinku 1
949 74 Nitra
lubomir.rybansky@ukf.sk

Štefan Havrlent, Mgr.
Trieda A. Hlinku 1
949 74 Nitra
stefan.havrlent@ukf.sk

Grafické vyhodnotenie pôdnej heterogenity Graphic plotting soil heterogeneity

Ján Tirpák

Abstract: Soil, as a basic component of environment, has irreplaceable importance in this process. The main aim of this article is to evaluate soil heterogeneousness by software Oasis Montaj.

Key words: agriculture, heterogeneity, software Oasis Montaj

Kľúčové slová: poľnohospodárstvo, pôdna heterogenita, grafický softvér Oasis Montaj

1. Úvod

Pôda má pre ľudí mimoriadny význam. V procese ekonomickej produkcie je pôda jednou z najdôležitejších faktorov. Pôda, ako základný výrobný faktor zároveň predstavuje aj základnú zložku životného prostredia. Pôda sa tiež považuje za rozhodujúcu zložku prírodného prostredia a to z toho dôvodu, že determinuje bohatstvo krajiny [2]. Pôda nielenže tvorí územie krajiny, ale zároveň je základom pre ekonomicko – sociálne využitie priestoru. Pôdu nezaraďujeme medzi obnoviteľné prírodné zdroje, pretože jej regenerácia či obnova trvá stovky až tisícky rokov. Napriek tomu, na európskej úrovni jej nie je venovaná dostatočná pozornosť [3].

Z tohto dôvodu je potrebné kontrolovať a zabezpečovať kvalitu pôdy. Bez udržateľného vývoja pôdy, nie je možné zabezpečiť udržateľný rozvoj prírodného prostredia, ale ani udržateľný rozvoj ekonomických a sociálnych parametrov súčasnej spoločnosti [1].

Z ekonomického hľadiska je pre pôdu rozhodujúca jej základná vlastnosť, ktorou je jej úrodnosť. Od úrodnosti je odvodená základná funkcia pôdy v poľnohospodárstve – jej produkčná funkcia. Schopnosť pôdy zásobovať rastliny živinami, potrebným množstvom vody a vzduchu v priebehu vegetácie – úrodnosť, je jedinečnou pôdnou funkciou, ktorá nie je ničím nahraditeľná. Úrodnosť rozoznávame *prirodzenú, hospodársku a ekonomickú* [2].

Vyrovnanosť pôdnej úrodnosti sa odhaduje na základe kontrolných odrôd, ktoré plnia funkciu miery kontroly pôdnej heterogenity.

2. Materiál a metódy

Možnosti prezentovať vyhodnotenie pôdnej heterogenity sú v súčasnosti možné viacerými softvérovými programami akými sú napr. Surfer od firmy Golden Software, USA alebo Oasis Montaj od firmy Geosoft, Kanada. V príspevku sme použili softvérový balík Oasis Montaj, ktorý obsahuje import dát, ich spracovanie, vizualizáciu, vytváranie máp a možnosti ich integrácie. Softvér je vybavený základnými aj pokročilými metódami gridizácie a vykresľovacími funkciami. Softvér sa používa hlavne pre spracovanie geovedných dát v rámci geologických prieskumov, ťažby ropy a plynu, vrátane ekologických projektov, archeologických nedeštruktívnych výskumov a detekcie nevybuchnutej munície.

Metódu grafického znázornenia pôdnej skutočnosti pomocou výpočtovej techniky publikovali Stehlíková, B. a Žofajová, A. [4]. Metóda bola založená na vyrovnávaní príslušných hodnôt kontrolných pásov bikubickým splinom. Aplikovateľnosť uvedenej metódy je podmienená splnením nasledujúcich predpokladov:

- a) na pokusnom poli je pravidelne rozmiestnená kontrolná odroda;
- b) pokusná plocha má štvorcový tvar (v prípade obdĺžnika je členená na štvorce, ktorých obrazy sa pospájajú);
- c) hodnoty meraného znaku (zvyčajne úroda) musia byť zisťované s rovnakou presnosťou a v rovnakých jednotkách.

3. Výsledky a diskusia

Pre grafické znázornenie pôdnej heterogenity použijeme softvér Oasis Montaj, pričom budeme vychádzať z práce [4]. V práci [4] Stehlíková, B. a Žofajová, A. popísali metódu grafického znázornenia pôdnej skutočnosti použitím výpočtovej techniky. Uvedenú metódu Stehlíková, B. a Žofajová, A. ilustrovali na nasledujúcom príklade.

Bol realizovaný pokus bez opakovania s ozimnou pšenicom, v ktorom je v piatich pásoch za každou 5. parcelkou skúšaných novo vyšľachtených odrôd umiestnená kontrolná odroda. Dosiahnuté úrody sú prehľadne zapísané v nasledujúcej tabuľke (Tabuľka 1).

Tabuľka 1: Namerané úrody kontrolnej odrody ozimnej pšenice

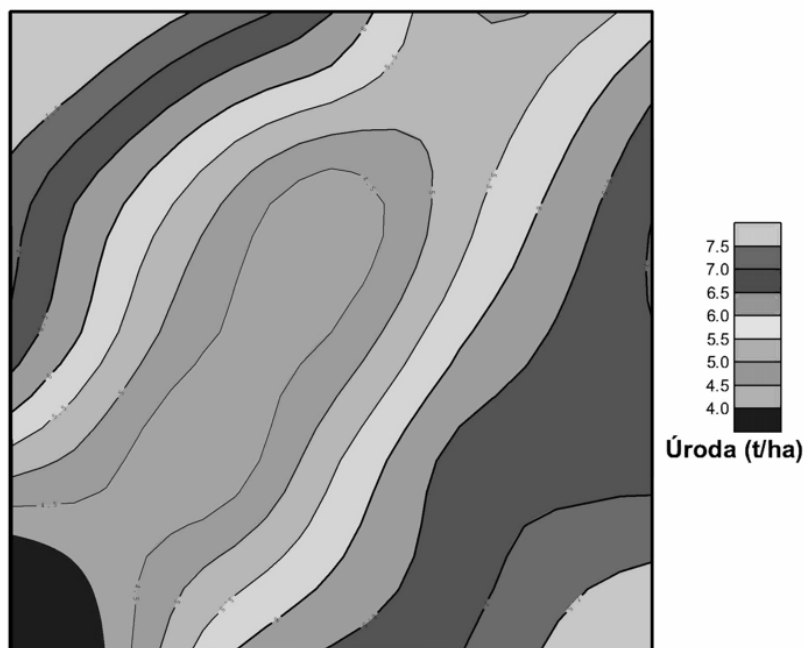
Pás	Úroda				
1	7,5	6,8	6,7	6,4	5,8
2	6,8	6,7	5,9	5,3	5,1
3	6,1	5,4	4,5	4,1	5,9
4	5,2	4,0	4,6	5,2	6,6
5	3,8	4,9	6,1	6,5	7,7

Zdroj: [4]

Uvedený príklad v našom príspevku sme spracovali použitím softvérového balíka Oasis Montaj. Postupovali sme nasledovne: údaje z tabuľky 1 boli importované do databázy tak, aby boli formátované do súboru s koncovkou XYZ. Takto získané údaje z databázy boli následne gridizované pomocou bi-directional numerickej techniky, ktorá je vhodná pre paralelné usporiadané údaje.

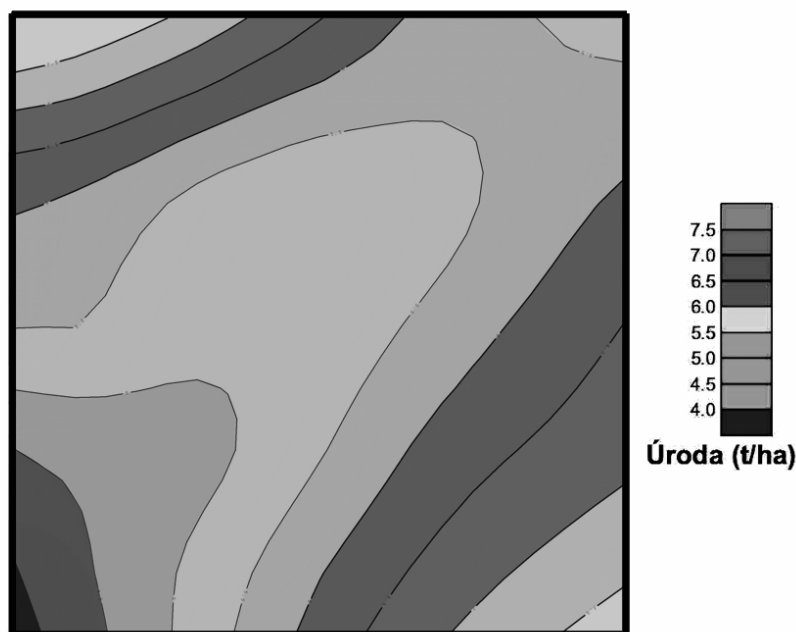
Z takto vytvorených gridových súborov bola zostrojená mapa pôdnej heterogenity nameraných úrod kontrolnej odrody ozimnej pšenice v pásoch, úroda je v tonách na hektár (obrázok 1).

Získané údaje boli spracované aj za použitia filtra B-spline a následne gridizované pomocou bi-directional numerickej techniky. Zo získaného gridizovaného súboru bola zostrojená mapa pôdnej heterogenity nameraných úrod kontrolnej odrody ozimnej pšenice v pásoch, úroda je v tonách na hektár (Obrázok 2).



Zdroj: [4] a vlastne zobrazenie

Obrázok 1: Mapa izoanomálií nameraných úrod kontrolnej odrody ozimnej pšenice v pásoch, úroda je v tonách na hektár



Zdroj: [4] a vlastne zobrazenie

Obrázok 2: Mapa izoanomálií s použitím filtra B-spline pri interpolácii dát nameraných úrod kontrolnej odrody ozimnej pšenice v pásoch, úroda je v tonách na hektár

4. Záver

V príspevku je prezentované grafické znázornenie pôdnej heterogenity použitím počítačového softvéru Oasis Montaj. Uvedená metóda môže byť vhodným doplnkom metód hodnotenia pokusov bez opakovania a tiež ukážkou použitia výpočtovej techniky pri analýze pôdnej heterogenity.

5. Literatúra

- [1] HRONEC, O. et al.: Ekologické základy poľnohospodárskej výroby. Nitra : Slovenská poľnohospodárska univerzita, 2001. 165 s. ISBN 80-7137-956-5.
- [2] HUTTMANOVÁ, E.: Prírodné podmienky ako základný výrobný a lokalizačný faktor a ich význam pre produkčný potenciál a udržateľnosť regiónov. Súbor vedeckých štúdií projektu VEGA č. 1/4638/07 a Centra excelentnosti výskumu kognícií CEVKOG / Róbert Štefko (Ed.). - Prešov : Fakulta manažmentu Prešovskej univerzity, 2008. ISBN 978-80-8068-890-5. S.65-67. - VEGA č. 1/4838/07. FM 181/08
www.pulib.sk/elpub2/FM/Kotulic11/pdf_doc/07.pdf
- [3] KANIANSKA, R.: Pôda a jej stav v Európskej únii. 4/2007 *ENVIROMAGAZÍN*
- [4] STEHLÍKOVÁ, B. – ŽOFAJOVÁ, A. 1985: Grafické znázornenie pôdnej heterogenity počítačom. In: Poľnohospodárstvo, roč.31, čís. 11, s.1049-1052
- [5] www.geosoft.com: Quick start tutorials Oasis Montaj

Adresa autora:

RNDr. Ján Tirpák, CSc..
Katedra zoológie a antropológie FPV UKF
Nábřežie mládeže 91,
949 74 Nitra
jan.tirpak@savba.sk

Štatistické metódy a umelá inteligencia na kapitálových trhoch

Statistical methods and artificial intelligence on kapital markets

Marta Urbaníková

Abstract: The article deals with the statistical methods of analysis of capital markets. It describes the principles of fundamental and technical analysis of stock markets too. Share index is considered to be the indicator of the development of the stock markets. Besides the Box – Jenkins methodology and Exponential smoothing, methods of artificial intelligence are used for the analysis of the share index.

Key words: Fundamental analysis, Technical analysis, Exponential smoothing, Box-Jenkins methodology, Neural networks.

Kľúčové slová: Fundamentálna analýza, technická analýza, exponenciálne vyrovňovanie, Box-Jenkinsonova metodológia, neurónové siete.

1. Úvod

V súvislosti so svetovou hospodárskou a finančnou krízou majú investori čoraz viac dôvodov zaoberať sa vývojom kapitálových trhov. Cieľom príspevku je prezentácia metód analýzy akciových trhov. Okrem fundamentálnej a technickej analýzy je pozornosť venovaná štatistickým analýzám časových radov a umelej inteligencii. Ako indikátor vývoja akciových trhov je braný akciový index. Na analýzu vývoja akciového indexu DJIA sú okrem Box-Jenkinsonovej metodológie a exponenciálneho vyrovňovania použité aj neurónové siete.

2. Analýza akciových trhov

Metódy analýzy akciových trhov prešli dlhodobým vývojom. Hnacím motorom hľadania nových metód a prognóz boli najmä burzové krachy a finančné krízy. V súčasnosti je najviac rozvinutá fundamentálna a technická analýza, menej uznávanou metódou je psychologická analýza. Do popredia sa dostávajú aj analýzy časových radov a metódy umelej inteligencie.

Fundamentálna analýza sa snaží o kvantifikáciu vnútornej hodnoty akcií. Zaoberá sa základnými faktormi globálneho, politického, ekonomického, odvetvového a firemného charakteru, ktoré významne ovplyvňujú kurz resp. vnútornú hodnotu akcie. Fundamentálna analýza sa uskutočňuje na troch úrovniach:

- Makroekonomická úroveň, ktorá analyzuje ekonomiku ako celok a jej vplyv na akciové kurzy
- Odvetvová úroveň, ktorá skúma špecifiká odvetvia a ich dopad na akciové kurzy
- Mikroekonomická úroveň, ktorá uskutočňuje finančnú analýzu jednotlivých spoločností.

Na makroekonomickej úrovni sa posudzujú faktory ako inflácia, nezamestnanosť, hrubý domáci produkt, výška úrokovej sadzby a obchodná bilancia. Nesmie sa zabúdať ani na vplyv vládnej politiky, vplyv ekonomických a politických šokov a na medzinárodný pohyb kapitálu v rámci svetových trhov. Pri dennom obchodovaní môžu hýbať trhom správy ohľadom makroekonomy. Predovšetkým sa treba sústrediť na zisky spoločností, ktoré sú základom tvorby HDP. V prípade vysokých úrokových mier investori predávajú akcie, investujú do bezpečnejších dlhopisov a ceny akcií klesajú. Pri odvetvovej analýze sa posudzujú jednotlivé odvetvia a sektory. Netreba však zabudnúť na súvis s fázou ekonomického cyklu. Pri oživení

ekonomiky sa darí hlavne cyklickým odvetviam (stúpa dopyt v automobilovom priemysle, strojárstve, stavebníctve, dopyt po veciach dlhodobej spotreby a pod.). V období recesie sa spotreba znižuje a darí sa tzv. anticyklickým odvetviam (káblková TV, časopisy, noviny). Obyčajne stúpajú ceny drahých kovov a dopyt po akciách baní na zlato, striebro, diamanty a pod. To možno vidieť aj dnes, keď ceny zlata lámu rekordy. Na hospodársky cyklus nereagujú tzv. neutrálne odvetvia. Ide o odvetvia po ktorých je stály dopyt (potravinarstvo, liečivá, tabak, veci krátkodobej spotreby a pod.). Na mikroekonomickej úrovni sa analytici zameriavajú na analýzu firmy. Používajú rôzne metódy, ktoré skúmajú zisky, rast dividend, cash flow, bilanciu, aktíva, pasíva, rôzne historické pomery hodnôt a podobne. Snažia sa stanoviť vnútornú hodnotu akcie a porovnávajú ju s trhovou cenou. Niektorí investori vyhľadávajú podhodnotené akcie (typickým predstaviteľom je Warren Buffett), iní sa sústreďujú na firmy, ktoré vykazujú nadpriemerný rast alebo vysoké dividendy. Pri výbere jednotlivých akcií je nutné zvažovať aj riziko investície.

Technická analýza na rozdiel od fundamentálnej analýzy sa sústreďuje iba na dianie na kapitálovom trhu. Skúmaním minulých a súčasných informácií o objemoch a pohyboch kurzov akcií sa snaží predpovedať vývoj jednotlivých akcií respektíve celkový vývoj trhu. Technická analýza vedie tak k podstatne rýchlejšim rozhodnutiam. Podobne ako fundamentálna analýza uplatňuje štruktúrny prístup. To znamená, že sa aplikuje najskôr na celý trh, ktorý väčšinou reprezentuje trhoví index, ale tiež na príslušné odvetvie, kde využíva najmä odvetvové indexy a celkový podiel odvetvia v rámci hospodárstva a nakoniec taktiež analyzuje jednotlivé akcie. Technická analýza pozostáva z analýzy grafov kurzov akcií a z analýzy technických indikátorov (Momentum, MACD, RSI, SSto, ADX, CCI, W%R, kombinácie kľzavých priemerov atď.), ktoré na základe vývoja cien akcií a objemov obchodov v minulosti predpovedajú vývoj cien akcií v budúcnosti. Technická analýza vychádza z troch základných postulátov: Trhové ceny odrážajú a zahŕňajú všetky informácie, ceny sa pohybujú v určitých trendoch, v ktorých aj zotrvávajú a nemenia smer svojho pohybu a vývoj cien na trhu má tendenciu opakovať sa.

Okrem uvedených najčastejšie používaných metód sa ponúkajú aj ďalšie možnosti využitia metód štatistickej analýzy. Štatistické algoritmy možno využívať v rôznych oblastiach analýz kapitálových trhov. Najčastejšie sa používajú na:

- Predpovede budúcich pohybov časových radov
- Určeniu vplyvu premenných na časový rad resp. akciový index

3. Box-Jenkinsonova metodológia, exponenciálne vyrovnávanie a neurónové siete

Metódy analýzy časových radov možno rozdeliť na dve skupiny. Jednu skupinu tvoria metódy, keď vývoj modelovaného ukazovateľa závisí len od času. Ide o dekompozičné metódy, Box-Jenkinsonové modely a spektrálnu analýzu. Druhú skupinu tvoria ekonometrické metódy modelovania.

Dekompozičné metódy zahŕňajú regresné metódy a metódy exponenciálneho vyrovnávania. Podstatou regresných metód je vyjadrenie závislosti hodnôt skúmaného ukazovateľa Y od času. Regresné metódy sú vhodné, ak charakter vývoja časového radu Y je jasný a zreteľný.

Podstatou metód *exponenciálneho vyrovnávania* je používanie rôznych váh jednotlivých pozorovaní v časovom rade. Čím je pozorovanie vzdialenejšie do minulosti, tým má menšiu váhu. Metódy exponenciálneho vyrovnávania dokážu pomerne veľmi presne vyrovnať časový rad. Exponenciálne vyrovnávanie nie je výpočtovo zložité a možno ho použiť na analýzu širokého spektra časových radov. Na rozdiel od metódy kľzavých priemerov, kde sme na výpočet vyrovnanej hodnoty časového radu použili krátke úseky určitej dĺžky minulých pozorovaní časového radu (ktorých dĺžku sme museli dopredu určiť), metóda exponenciálneho vyrovnávania využíva na výpočet každej vyrovnanej hodnoty

časového radu všetky dostupné minulé pozorovania radu. Váhy priradené jednotlivým pozorovaniám exponenciálne klesajú smerom do minulosti. Základná metóda najmenších štvorcov sa modifikuje tak, že váhy jednotlivých štvorcov v minimalizovanom súčte smerom do minulosti exponenciálne klesajú. Nech $\hat{y}_t$ je vyrovnaná hodnota radu v čase t . Potom minimalizujeme výraz

$$(y_t - \hat{y}_t)^2 + \alpha(y_{t-1} - \hat{y}_{t-1})^2 + \alpha^2(y_{t-2} - \hat{y}_{t-2})^2 + \dots = \sum_{i=0}^{\infty} \alpha^i (y_{t-i} - \hat{y}_{t-i})^2,$$

kde sčítujeme cez všetky známe hodnoty radu. Parameter α je vyrovnávacia konštanta a spĺňa podmienku $0 < \alpha < 1$. Používa sa jednoduché, dvojité, trojité exponenciálne vyrovnávanie a Holtova-Wintersova metóda exponenciálneho vyrovnávania. Vo väčšine softvérových produktov sa tieto metódy nachádzajú.

Box-Jenkinsonové modely predstavujú nový prístup k modelovaniu časových radov. Kým pri predchádzajúcich metódach sa vychádzalo z trendovej a sezónnej zložky časového radu, pri Box-Jenkinsonových modeloch sa vychádza z náhodnej zložky. Pri ich konštrukcii sa vychádza z autoregresných modelov AR a modelov kĺzavých súčtov MA. Keďže pri jasných časových radoch je vhodnejšie použiť regresné modely, Box-Jenkinsonové modely sa používajú na modelovanie nejasných časových radov, kde okrem ich nejasnosti je aj zložitejšie intuitívne predpovedať ich budúci vývoj. Box-Jenkinsonovou metódou možno modelovať iba stacionárne časové rady. Tento predpoklad však nie je veľmi obmedzujúci, nakoľko pomocou rôznych transformácií je možné veľa nestacionárnych radov z praxe previesť na stacionárne rady. Analýza časového radu v rámci Box-Jenkinsonovej metodológie sa robí systematicky podľa dopredu daného kľúča. Postup možno rozdeliť do troch základných krokov:

- identifikácia modelu
- odhad parametrov modelu
- testovanie vhodnosti modelu, prípadne modifikácia modelu.

Po overení vhodnosti modelu môžeme tento model použiť na prognózy ďalšieho vývoja sledovaného časového radu.

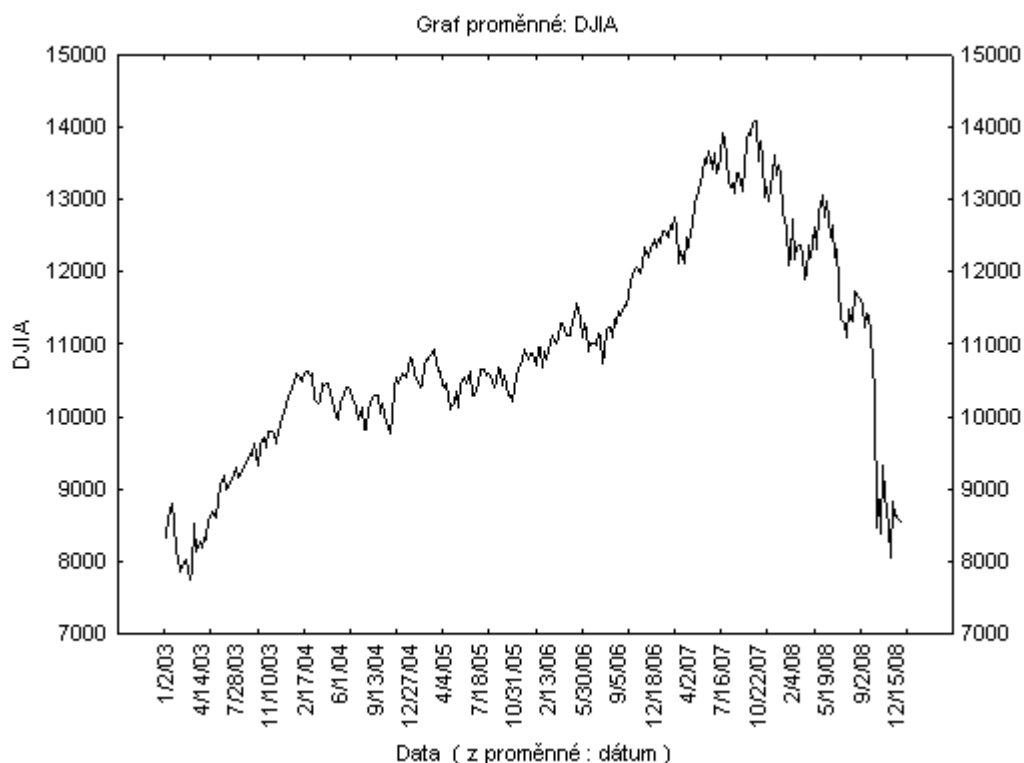
Spektrálna analýza vychádza z idey, že každý časový rad možno rozdeliť na nekonečné množstvo sínusoid. Využíva sa Fourierova transformácia časového radu a na jej základe sa snažíme odhadnúť budúci vývoj. Od časového radu sa vyžaduje stacionárnosť. Spektrálna analýza je pomerne náročná na dĺžku časového radu.

Neurónové siete ponúkajú investorom riešenie klasifikačných a predikčných úloh. Na kapitálových trhoch ide predovšetkým o prácu s časovými radmi. Predpovede neurónových sietí nám umožňujú doplniť resp. nahradiť štatistické odhady robené klasickými štatistickými metódami. V súčasnosti existuje množstvo softvérových produktov, ktoré umožňujú riešiť ekonomické úlohy pomocou neurónových sietí (napr. *Statistica Neural Networks*, *Neural Connection*, *Mathematica Neural Network...*). Hlavnou výhodou nasadenia technológie neurónových sietí je možnosť zachytenia nelineárnych vzťahov a riešení komplexných a zložitých úloh. Kvalita dát predurčuje kvalitu riešení i v prípade technológie neurónových sietí. Neurónové siete sú všeobecne odolnejšie voči vstupným dátam (neúplné časové rady a podobne), ako bežné štatistické nástroje. Ak je kvalita dát dobrá, t.j. ak je zaistená ich vysoká validita a reliabilita (platnosť a spoľahlivosť), výsledky obvykle nevykazujú pri užití adekvátnych štatistických metód výraznejšie rozdiely v porovnaní s výstupom z neurónovej siete. Stále však platí, že vhodnosť použitých prostriedkov závisí predovšetkým na riešenej úlohe. Neurónové siete majú však v porovnaní s inými metódami spracovania dát niektoré nepríjemné vlastnosti. Predovšetkým ide o naplnenie požiadaviek identickej reprodukovateľnosti výsledkov. Rôzne softvéry neurónových sietí, rôzne typy neurónových sietí a tiež použitie rôznych parametrov pre tréning siete, doba tréningu neurónovej siete

atď. nezabezpečujú vždy, aj po použití rovnakej množiny dát na vstupe, rovnaký výstup. Ďalším obmedzujúcim faktorom je, že stav voľných parametrov natrénovanej neurónovej siete neumožňuje vo väčšine prípadov matematicky zachytiť vzťahy medzi premennými.

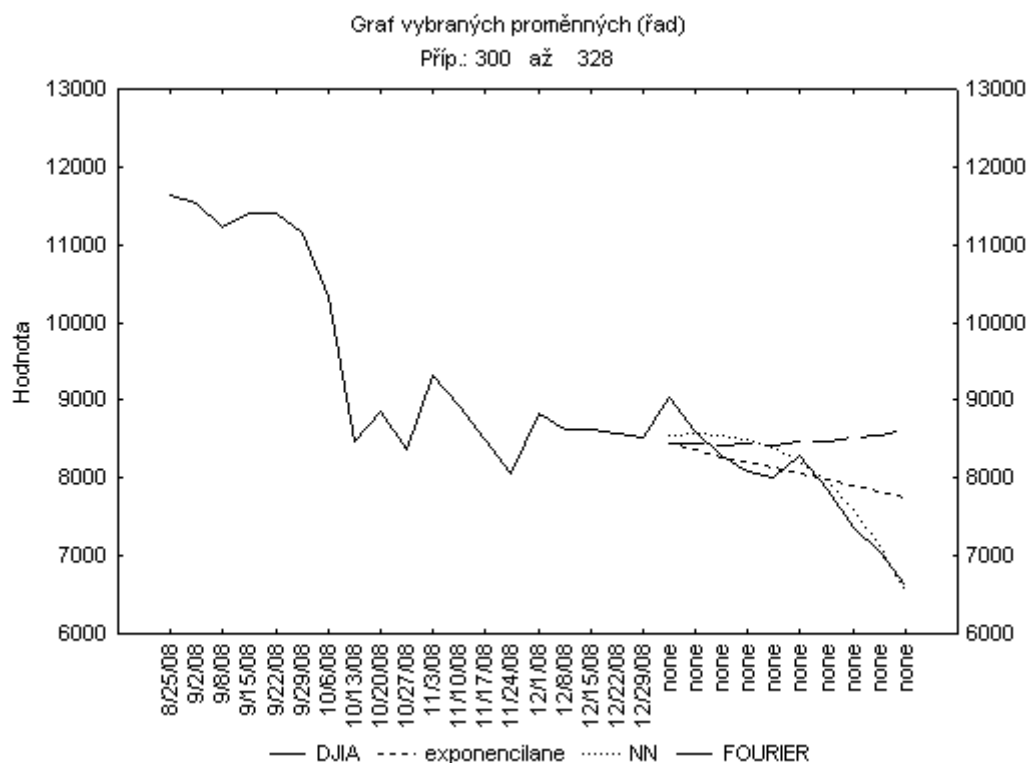
V nasledujúcom príklade použitím softvéru *Statistica CZ trial version* s plugin-om *Statistica Neural Networks* analyzujeme časový rad týždenných hodnôt indexu DJIA od 2.1.2003 do 23.2.2009.

Na obrázku 1 je vykreslený časový rad indexu DJIA.



Obrázok 1: Časový rad indexu DJIA

Predikcie uskutočnené využitím neurónových sietí porovnáme s predikciami uskutočnenými Box-Jenkinsonovou metodológiou a exponenciálnym vyrovnávaním. Empirickou skúsenosťou sa ukáže, ktorá metóda je vhodnejšia na analýzu a predikcie daného časového radu. Na obrázku 2 vidíme skutočný vývoj indexu DJIA a predikcie hodnôt indexu. Z dôvodu lepšej prehľadnosti sú zobrazené iba prípady 300-328 t.j. august 2008 až marec 2009. Odhad Box-Jenkinsonovou metodológiou bol robený pomocou modelu ARIMA (1,1,1), ešte predtým bol však rad vyrovnávaný pomocou fourierovej analýzy kvôli cyklickej zložke. Pri exponenciálnom vyrovnávaní bolo použité jednoduché exponenciálne vyhladzovanie bez sezónnej zložky s automatickým odhadom parametrov. Neurónová sieť v softvéri SNN (*Statistica Neural Networks*) bola trénovaná pomocou voľby *Intelligent problem Solver* a bola nájdená najlepšia sieť RBF 1:1-67-1:1, ktorá podľa programu mala „veľmi dobrý výkon“. Kvôli prehľadnosti sú výsledky predpovedí všetkých testovaných modelov a reálnych údajov uvedené v tabuľke 1.



Obrázok 2: Index DJIA a predikcie

Tabuľka 1: Predikované a reálne hodnoty

B-J	Exponencial	NN	DJIA
8437,63	8434,52	8545,54	9027,13
8458,60	8356,29	8563,04	8599,26
8415,67	8278,80	8555,37	8279,63
8443,17	8202,02	8503,94	8078,04
8430,18	8125,95	8397,17	8000,62
8467,62	8050,59	8219,10	8281,38
8479,95	7975,92	7956,50	7845,63
8528,60	7901,95	7597,41	7365,99
8561,97	7828,67	7131,85	7056,48
8621,46	7756,06	6551,26	6625,74

Pre zistenie „najoptimálnejšieho“ modelu, je nutné vykonať porovnanie predikcií modelov s reálnymi údajmi a na základe relatívnej percentuálnej chyby (MAPE) určiť najlepší model.

$$MAPE = \frac{1}{N} \sum_{t=1}^N \left| \frac{y_t - \hat{y}_t}{y_t} \right| \cdot 100\%$$

kde N je počet predpovedí, y_t sú pozorované hodnoty a $\hat{y}_t$ sú predpovede.

Výsledky sú uvedené v tabuľke č. 2. Vidíme, že najmenšiu odchýlku 2,68% vykazuje model neurónovej siete RBF 1:1-67-1:1, ktorý označíme ako najlepší.

Tabuľka 2: Hodnoty MAPE

B-J	Exponencial	NN
0,065303	0,065647664	0,053349
0,016357	0,028254757	0,004212
0,016431	0,000100246	0,033303
0,0452	0,015347782	0,052723
0,053691	0,015665036	0,049565
0,22489	0,027868544	0,00752
0,08085	0,016606697	0,014131
0,157835	0,072761435	0,031417
0,213349	0,109429914	0,010681
0,301207	0,170595284	0,011241
0,972712	0,522277359	0,268144
9,727119	5,222773593	2,681439

Najhoršie predikcie nám poskytol model založený na Fourierovej analýze a metóde Box-Jenkins, kde odchýlka od skutočných hodnôt predstavovala skoro 10%.

4. Záver

Ďalší vývoj akciových trhov nebude jednoduchý, pretože situácia v eurozóne je neistá a fiškálne problémy má viac štátov. Pri ďalšom zvyšovaní rizika na trhoch môže prísť k presunu peňazí do bezpečnejších prístavov, akými sú zlato a drahé kovy. Treba zdôrazniť, že nemožno preferovať použitie žiadnej z metód analýzy akciových trhov oproti iným technológiám spracovania dát, ako je napríklad štatistická analýza. Komparácia výsledkov získaných z rôznych metód a analytických nástrojov je preto základnou požiadavkou serióznej analytickej práce.

5. Literatúra

- [1]ARTL, J. 2003. Finanční časové rady. Grada Publishing, Praha, 219s. ISBN 80-247-0330-0.
- [2]FANTA, J. 1999. Technologie umělé inteligence na kapitálových trzích. Karolinum, Praha. ISBN 80-247-0024-7.
- [3]URBANÍKOVÁ, M., VRÁBELOVÁ, M. 2006. Modelovanie ekonomických a finančných procesov. Nitra, 96s. ISBN 80-80-50-929-8.
- [4]SKŘIVÁNKOVÁ, V., SKŘIVÁNEK, J. 2006. Kvantitativne metody finančních operací. Iura Edition, Bratislava, ISBN 80-8078-074-9

Adresa autora :

Marta Urbaníková, Doc., RNDr. CSc.
Tr. Andreja Hlinku 1
949 01 Nitra
murbanikova@ukf.sk

Multinominálny logitový model pravdepodobností prístupov na webové časti portálu

The Multinomial Logit Model of the Probabilities of Accesses to the Web Parts of Portal

Marta Vrábelová, Michal Munk, Jozef Kapusta

Abstract: We model how the probabilities of choice of the web parts by visitors depend on the day hour and on the week day. The multinomial logit model is used.

Key words: web log mining, web parts, multinomial logit model, probabilities of choice.

Kľúčové slová: web log mining, webové časti, multinominálny logitový model, pravdepodobnosti výberu.

1. Úvod

Cieľom príspevku je modelovanie pravdepodobností prístupov na kategórie webových častí portálu. Získané znalosti môžu byť využité ako podklady pre analýzu návštevnosti, optimalizáciu, resp. personalizáciu portálu. Za týmto účelom boli zozbierané údaje o prístupoch na kategórie webových častí portálu FPV UKF v Nitre. V tomto článku sa budeme zaoberať pravdepodobnosťami prístupu zamestnancov, študentov a iných návštevníkov na jednotlivé kategórie obsahu skúmaného portálu a to v závislosti na dni v týždni a v závislosti na hodine dňa. Tieto pravdepodobnosti odhadneme pomocou multinominálneho logitového modelu [1], [2] a to zvlášť pre každý typ návštevníka.

2. Príprava dát

Naším zdrojom dát sú automaticky ukladané dáta v logovacích súboroch. Z tohto dôvodu sa táto oblasť často označuje aj ako web log mining. V takýchto dátach sledujeme postupnosti – sekvencie v návšteve jednotlivých stránok používateľom, ktorého identifikuje za určitých podmienok IP adresa. Logovací súbor vo svojej štandardnej štruktúre nazývanej ako Common Log File [3] zaznamenáva každú transakciu, ktorú prehliadač pri prístupe k webu vykonal. Každý riadok predstavuje záznam o IP adrese, čase a dátume návštevy, prístupovanom objekte a odkazovanom objekte. V prípade, že použijeme jeho rozšírenú podobu môžeme zaznamenávať aj údaje o verzii prehliadača používateľa, tzv. User-Agent.

Log súbor webového servera je zdrojom anonymných dát o používateľovi. Tieto anonymné dáta zároveň predstavujú problém jednoznačnej identifikácie návštevníka webu. Pre použiteľnosť log súboru pre ďalšie analýzy je potrebné vykonanie viacerých úprav s cieľom identifikovať konkrétného používateľa. Za týmto účelom sa vykonávajú nasledujúce úpravy:

- **Čistenie dát od crawlerov** vyhľadávacích služieb, ktoré prístupujú na portál [4].
- **Identifikácia návštevníkov** na základe rôznej verzie internetového prehliadača [5].
- **Identifikácie sedení**, kde sedenie môže byť definované ako postupnosť krokov, ktoré vedú k naplneniu určitej úlohy [6] alebo ako postupnosť krokov, ktoré vedú k dosiahnutiu určitého cieľa [7]. Najjednoduchšou metódou je, ak za sedenie považujeme **sériu kliknutí za určitý čas**, napr. 30 minút [8].
- **Rekonštrukcia aktivít** návštevníka webu. Taucher a Greenberg [9] dokázali, že viac ako 50% prístupov na webe je pohyb späť. Tu nastáva problém s cache prehliadača. Pri pohybe späť sa nevykonáva dotaz na webový server, teda neexistuje o tom ani záznam v

logovacou súboru. Riešením tohto problému je dopĺňanie cesty. **Dopĺňaním cesty** pridávame do log súboru tieto chýbajúce riadky [5].
Všetky tieto techniky sú závislé na prvotnom kroku zberu a čistenia dát.

3. Použité premenné

Skúmanou kategorickou premennou je premenná KATEGORIA, ktorá má úrovne: *úvod, štúdium, oznamy, o fakulte, informácie pre, informácie o, predpisy vyhlášky, dokumenty, veda výskum, ostatné, vyhľadávanie, udalosti, konferencie*, kategórie boli zvolené na základe ich príslušnosti k obsahu. Nezávisle premennou je hodina dňa HOD_DEN s hodnotami 0-23. Dni v týždni použijeme ako umelé premenné, ktoré označujeme PÖ, UT, STR, STVR, PIA, SO (stačí zaviesť len 6 umelých premenných). Použité údaje popisujú individuálne prístupy na internetovú stránku FPV UKF v Nitre v priebehu dvoch týždňov zo stredu letného semestra.

4. Popis modelu

Pravdepodobnosť, že si návštevník zvolí kategóriu j , ak vyberá v čase i označme π_{ij} , pričom $j = 1, 2, \dots, J$, kde J je počet kategórií www stránky a $i = 0, 1, 2, \dots, 23$. Pretože platí $\sum_{j=1}^J \pi_{ij} = 1$, máme len $J - 1$ parametrov.

Nech Y_{ij} označuje počet prístupov v čase i na kategóriu j s pozorovaniami y_{ij} . Označme $n_i = \sum_j y_{ij}$ počet prístupov v čase i . Pravdepodobnostné rozdelenie počtov Y_{ij} , ak je daný súčet n_i je multinomické

$$P[Y_{i1} = y_{i1}, Y_{i2} = y_{i2}, \dots, Y_{iJ} = y_{iJ}] = \frac{n_i!}{y_{i1}! y_{i2}! \dots y_{iJ}!} \pi_{i1}^{y_{i1}} \pi_{i2}^{y_{i2}} \dots \pi_{iJ}^{y_{iJ}}. \quad (1)$$

Pravdepodobnosti π_{ij} výberu kategórií v závislosti na čase i vypočítame z multinominálnych logitov, ktoré budeme modelovať. Multinominálne logity sú to logaritmy

$$\ln \frac{\pi_{ij}}{\pi_{iJ}}, j = 1, 2, \dots, J - 1,$$

pričom π_{iJ} je pravdepodobnosť poslednej kategórie, ktorú si zvolíme za referenčnú. Predpokladáme, že platí nasledovný model

$$\eta_{ij} = \ln \frac{\pi_{ij}}{\pi_{iJ}} = \alpha_j + \mathbf{x}_i' \beta_j, j = 1, 2, \dots, J - 1, \quad (2)$$

kde $\mathbf{x}_i'$ riadok matice $\mathbf{X}$, ktorá neobsahuje stĺpec jednotiek, α_j sú konštanty a β_j je vektor regresných koeficientov. Máme $J - 1$ rovníc, ktoré popisujú kontrasty medzi kategóriou j a poslednou kategóriou J (za referenčnú kategóriu sme mohli zvoliť ľubovoľnú inú kategóriu). Pravdepodobnosti π_{ij} vypočítame z logitov podľa vzorca

$$\pi_{iJ} = \frac{1}{1 + \sum_{k=1}^{J-1} e^{\eta_{ik}}}, \quad \pi_{ij} = e^{\eta_{ij}} \pi_{iJ}, \quad j = 1, 2, \dots, J - 1.$$

Maximálne virohodný odhad parametrov modelu (2) sa robí maximalizáciou multinominálnej funkcie virohodnosti (1), pričom pravdepodobnosti π_{ij} sú vyjadrené ako funkcie α_j a β_j . Odhad parametrov možno urobiť pre individuálne údaje v štatistickom systéme.

5. Určenie modelu a odhad parametrov

Na základe kontingenčných tabuliek pre premenné KATEGORIA, HOD_DEN a DEN_TYZD sme zistili, že významné počty (viac ako 10, až na niektoré odľahlé hodnoty) prístupov sú

- pre zamestnancov len v pracovných dňoch PO-PIA, v hodinách 8-15 a to na kategórie *úvod, štúdium, oznamy, o fakulte*,
- pre študentov v dňoch PO-SO, v hodinách 7-22, na kategórie *úvod, štúdium, oznamy, o fakulte*,
- pre iných návštevníkov v dňoch PO-NE, v hodinách 0-1 a 6-23, na kategórie *úvod, štúdium, oznamy, o fakulte, informácie pre, informácie o, predpisy vyhlášky, veda výskum, vyhľadávanie*.

Nevýznamné prístupy sme vylúčili. Celkovo sme skúmali počty prístupov uvedené v tabuľke 1.

Tabuľka 1: Počty prístupov na jednotlivé kategórie

KATEGÓRIA	Počet prístupov		
	zamestnancov	študentov	iných návštevníkov
úvod	758	2245	5991
štúdium	818	972	15252
oznamy	111	164	871
o fakulte	287	119	5271
informácie pre			591
informácie o			234
predpisy vyhlášky			333
veda výskum			374
vyhľadávanie			309
celkom	1974	3500	29225

Použili sme nasledovné modely (premenná HOD_DEN je označená ako t).
Pre zamestnancov:

$$\eta_{ij} = \alpha_j + \beta_{1j}t_i + \beta_{2j}t_i^2 + \gamma_{1j}PO_i + \gamma_{2j}UT_i + \gamma_{3j}STR_i + \gamma_{4j}STVR_i$$

Pre študentov:

$$\eta_{ij} = \alpha_j + \beta_{1j}t_i + \beta_{2j}t_i^2 + \gamma_{1j}PO_i + \gamma_{2j}UT_i + \gamma_{3j}STR_i + \gamma_{4j}STVR_i + \gamma_{5j}PIA_i$$

Pre iných návštevníkov:

$$\eta_{ij} = \alpha_j + \beta_{1j}t_i + \beta_{2j}t_i^2 + \gamma_{1j}PO_i + \gamma_{2j}UT_i + \gamma_{3j}STR_i + \gamma_{4j}STVR_i + \gamma_{5j}PIA_i + \gamma_{6j}SO_i$$

Parametre modelov sme odhadli pre individuálne údaje v programe *STATISTICA*, uvedené sú v tabuľke 2. Významné parametre sú podfarbené. Významnosť parametrov bola testovaná pomocou Waldovho testu.

Tabuľka 2: Odhady parametrov pre jednotlivé modely

Odhady parametrov						
	KATEGÓRIA	pre zamestnan.	pre študentov	pre iných	KATEGÓRIA	pre iných
Intercept 1	uvod	6,839662	1,34835	3,116701	informacie_pre	0,856933
HOD_DEN	uvod	-0,977737	0,34182	-0,033368	informacie_pre	-0,057496
HOD_KV	uvod	0,041187	-0,01563	0,000389	informacie_pre	0,002289
PO	uvod	0,258197	-0,63353	0,343872	informacie_pre	0,414386
UT	uvod	0,208238	1,10307	-0,155347	informacie_pre	-0,336326
STR	uvod	-0,641734	0,11376	0,182570	informacie_pre	-0,149373
STVR	uvod	-0,916287	-0,49911	0,395581	informacie_pre	-0,248156
PIA	uvod		0,68682	0,402306	informacie_pre	0,418208
SO	uvod			0,954307	informacie_pre	0,463285
Intercept 2	studium	6,136931	-2,84924	4,109729	informacie_o	-0,095363
HOD_DEN	studium	-0,926228	0,38448	-0,053337	informacie_o	-0,046607
HOD_KV	studium	0,038839	-0,00867	0,001855	informacie_o	0,001711
PO	studium	0,424357	0,95518	0,309280	informacie_o	0,574992
UT	studium	0,949406	3,02822	-0,283642	informacie_o	-0,614693
STR	studium	-0,154207	1,54827	0,138681	informacie_o	-0,210748
STVR	studium	-0,011270	1,22148	0,185091	informacie_o	0,316949
PIA	studium		1,20767	0,256058	informacie_o	-0,358549
SO	studium			0,586272	informacie_o	0,522521
Intercept 3	oznamy	2,691152	1,51459	1,331879	predpisy_vyhlaskey	-0,866023
HOD_DEN	oznamy	-0,820874	-0,38342	-0,146409	predpisy_vyhlaskey	0,221840
HOD_KV	oznamy	0,030078	0,01248	0,004713	predpisy_vyhlaskey	-0,009925
PO	oznamy	2,266394	1,65098	1,359966	predpisy_vyhlaskey	-0,238335
UT	oznamy	2,984632	3,04913	0,480137	predpisy_vyhlaskey	-0,589468
STR	oznamy	0,550496	1,59056	0,489341	predpisy_vyhlaskey	-0,099511
STVR	oznamy	-0,610526	0,14253	0,418918	predpisy_vyhlaskey	-0,479032
PIA	oznamy		-1,24192	0,388898	predpisy_vyhlaskey	1,192950
SO	oznamy			0,733100	predpisy_vyhlaskey	-0,486781
Intercept 4	o_fakulte			3,425147	veda_vyskum	0,839007
HOD_DEN	o_fakulte			-0,091528	veda_vyskum	-0,187123
HOD_KV	o_fakulte			0,002458	veda_vyskum	0,007003
PO	o_fakulte			0,122460	veda_vyskum	0,045298
UT	o_fakulte			-0,312283	veda_vyskum	-0,076282
STR	o_fakulte			0,324361	veda_vyskum	0,750511
STVR	o_fakulte			0,563993	veda_vyskum	0,996957
PIA	o_fakulte			0,221580	veda_vyskum	0,582450
SO	o_fakulte			0,717273	veda_vyskum	1,036882

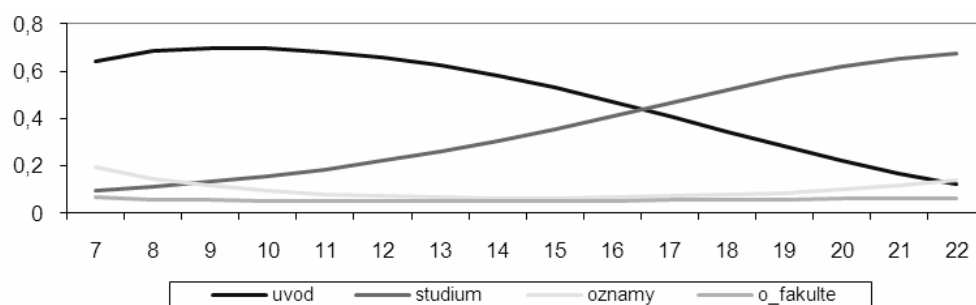
V tabuľke 2 vidíme, že logity pre kategóriu *úvod* sú u zamestnancov i u študentov významne závislé od hodiny prístupu aj od jej štvorca, ale u iných užívateľov od hodiny významne nezávisia. Na hodnoty týchto logitov u zamestnancov významne vplýva pondelok a štvrtok, u študentov pondelok, utorok a u iných návštevníkov len sobota. Takto by sme mohli interpretovať odhady parametrov pre ďalšie kategórie. Všimnime si ešte, že logity pre kategórie *informácie o* a *informácie pre* významne nezávisia ani od hodiny, ani od dňa.

Pomocou odhadnutých parametrov možno vypočítať odhady logitov a z nich pravdepodobnosti výberu jednotlivých kategórií v danej hodine dňa. Vypočítať sa tiež dajú teoretické počty prístupov na jednotlivé kategórie. Ako ukážku uvádzame tieto

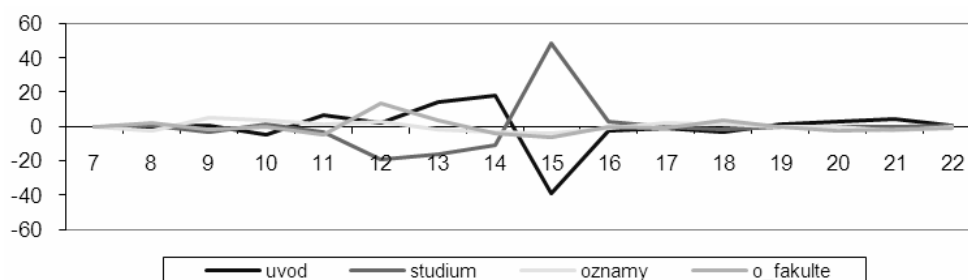
pravdepodobnosti výberu kategórií študentmi v pondelok (tabuľka 3), ktoré sú zobrazené na obrázku 1. Na obrázku 2 sú zobrazené rozdiely medzi empirickými a teoretickými početnosťami prístupov študentov na jednotlivé kategórie v pondelok.

Tabuľka 3: Pravdepodobnosti prístupu študentov na jednotlivé kategórie

Pravdepodobnosti prístupu študentov v pondelok				
Hodina	úvod	štúdium	oznamy	o fakulte
7	0,656748	0,095908	0,19717	0,066083
8	0,685003	0,110749	0,145081	0,059167
9	0,696699	0,132313	0,115221	0,055767
10	0,694756	0,157163	0,094908	0,053173
11	0,681157	0,186109	0,081305	0,051428
12	0,657113	0,219892	0,072498	0,050497
13	0,623378	0,259068	0,067247	0,050307
14	0,580593	0,303861	0,06478	0,050766
15	0,529619	0,353958	0,064654	0,051769
16	0,471824	0,408318	0,066665	0,053194
17	0,409257	0,46504	0,070796	0,054908
18	0,344627	0,521393	0,07721	0,05677
19	0,28104	0,574051	0,086263	0,058646
20	0,221536	0,619495	0,098547	0,060423
21	0,168585	0,65444	0,114969	0,062006
22	0,123728	0,676117	0,136838	0,063317



Obrázok 1: Pravdepodobnosti prístupu študentov na jednotlivé kategórie v pondelok



Obrázok 2: Rozdiel empirických a teoretických početností prístupov študentov v pondelok

Na obrázku 1 vidíme, že pravdepodobnosť prístupu na kategóriu úvod je ráno pomerne vysoká, najvyššia je v 9-tej hodine dňa, keď dosiahne hodnotu 0,697 a k večeru postupne klesá. Naopak, pravdepodobnosť výberu štúdium je ráno malá a v priebehu dňa sa zvyšuje až na 0,676. Pravdepodobnosti prístupu na kategórie oznamy a o fakulte sú v priebehu pondelka pomerne nízke. Pomocou obrázka 2 možno posúdiť kvalitu modelu. Vidíme, že väčšie rozdiely sa objavili v 15. hodine dňa. Nedostali sme veľmi dobré odhady početností medzi 12. a 15. hodinou.

6. Záver

V článku sme chceli poukázať na možnosti modelovania rozdelenia kategorickej premennej pomocou multinominálneho logitového modelu. Ako príklad kategorickej premennej sme vzali premennú KATEGORIA, ktorej úrovňami sú kategórie webových častí portálu FPV. Vytvorili sme model pre zamestnancov, študentov a iných návštevníkov, pričom sme ako prediktívne premenné použili deň v týždni a hodinu dňa. Treba však poznamenať, že tento model vyžaduje nezávislosť individuálnych prístupov na jednotlivé kategórie, čo v našom prípade zrejme nie je zabezpečené. Z tohto dôvodu sme ani neskúmali vhodnosť získaných modelov pomocou testov.

7. Literatúra

- [1] ANDĚL, J. 2007. *Základy matematické statistiky*. MATFYZPRESS, Praha 2007. ISBN 80-7378-001-1.
- [2] RODRÍGUEZ, G. 2007. *Generalized Linear Models*. 2007. [cit. 2010-05-11]. Available from: <http://data.princeton.edu/wws509/stata>
- [3] W3C. 1995. *Configuration File of W3C httpd* [online]. 1995. [cit. 2009-03-10]. Available from: <http://www.w3.org/Daemon/User/Config/Logging.html>
- [4] LOURENÇO, A. G., BELO, O. O. 2006. Catching web crawlers in the act. *Proceedings of the 6th international Conference on Web Engineering*, 2006, ICWE '06, Vol. 263, ACM, New York, NY, pp. 265-272. ISBN 1-59593-352-2.
- [5] COOLEY, R., MOBASHER, B., & SRIVASTAVA, J. 1999. Data Preparation for Mining World Wide Web Browsing Patterns. *Knowledge and Information System*, 1999, Springer-Verlag, Vol. 1, ISSN 0219-1377.
- [6] SPILIOPOULOU, M., FAULSTICH, L.C. 1999. WUM: A Tool for Web Utilization Analysis. *Extended version of Proc. EDBT Workshop WebDB'98*, 1999, Springer Verlag, pp 184–203.
- [7] CHEN, M., PARK, J.S., & YU, P.S. 1996. Data mining for path traversal patterns in a web environment. *ICDCS*, 1996, pp. 385–392. ISBN 0-8186-7398-2.
- [8] BERENDT, B., SPILIOPOULOU, M. 2000. Analysis of navigation behaviour in web sites integrating multiple information systems. *The VLDB Journal*, 2000, Vol. 9, No. 1, pp. 56-75. ISSN 1066-8888.
- [9] TAUCHER, L., GREENBERG, S. 1997. Revisitation patterns in world wide web navigation. *Proc. of Int. Conf. CHI'97*, 1997, Atlanta.

Adresy autorov:

Marta Vrábelová, doc. RNDr.,
CSc.
KM FPV UKF, Tr. A. Hlinku 1
949 74 Nitra
mvrabelova@ukf.sk

Michal Munk, RNDr., PhD., Jozef Kapusta, PaedDr.,
PhD.
KI FPV UKF, Tr. A. Hlinku 1
949 74 Nitra
mmunk@ukf.sk, jkapusta@ukf.sk

Vzt'ah medzi demokraciou a vzdelaním The relationship between democracy and education

Lubomír Zelenický, Ondrej Šedivý

Abstract: The contribution describes the relationship between Democracy index and Education index by using the mathematical-statistical method: Pearson's coefficient of correlation.

Key words: Democracy index, Education index, Pearson's coefficient of correlation

Kľúčové slová: Index demokracie, Index vzdelania, Pearsonov koeficient korelácie

1. Úvod

Vzdelanie je nevyhnutné pre demokraciu a rovnako je demokracia potrebná pre akékoľvek vzdelávanie. Vzdelávanie podľa Deweyho (1916) je dôležité pre demokraciu aj preto, že umožňuje "kultúru demokracie" a v konečnom dôsledku vedie k väčšej prosperite.

Životná úroveň a kvalita života je najvyššia v najslobodnejších a najdemokratickejších krajinách. Popredné miesta dlhodobo patria krajinám Škandinávie, kde je osobná sloboda, sloboda tlače, ale i ekonomická sloboda najvyššia na svete.

Sloboda je základom rastu a iba 14 % ľudí sveta žije v demokracii. Občania členských krajín Európskej únie patria medzi 14 % ľudí na svete, ktorí žijú v plnej demokracii. Každý občan tak má slobodnú voľbu rozhodovať sám za seba a je obmedzený iba dodržiavaním základných zákonov.

2. Demokracia a vzdelanie

The Economist pravidelne zostavuje *Index demokracie*, ktorým preveruje stav demokracie v 167 štátoch. Výsledný Index demokracie je zostavený na základe nasledujúcich 5 kritérií: volebný proces a pluralita, občianska sloboda, fungovanie vlády, politická spoluúčasť a politická kultúra. V každej kategórii je samozrejme hodnotená celá škála aspektov. Napríklad občianska sloboda je najvyššia v tých krajinách, kde je súdna moc nezávislá na vládnej moci, kde majú občania právo sa zhromažďovať či štrajkovať, kde je zaručená sloboda slova, kde sú nezávislé a slobodné médiá, kde štát nezakazuje prístup k vybraným internetovým stránkam, kde majú cudzinci zhodné práva s domácimi občanmi, kde sa toleruje odlišná viera, kde nie je diskriminácia na základe farby pleti, či vyznania, kde nie je problém spísať a odovzdať verejnej správe petíciu a ďalšie. Obdobným spôsobom sú hodnotené i ďalšie ukazovatele. Pre ilustráciu uvádzame, že za rok 2008 je index demokracie najvyšší vo Švédsku, ktorému tak patrí prvé miesto. Česku patrí 19. miesto z celkového počtu 167 krajín. Krajiny s úplnou demokraciou dosahujú hodnotu indexu demokracie v rozpätí 8-10 skóre, krajiny s nefunkčnou demokraciou 6-7,9, s hybridnou demokraciou 4-5,9 a krajiny s autoritatívnym režimom majú skóre menšie ako 4. Napríklad Index demokracie má hodnotu vo Švédsku 9,88 a v Severnej Kórei má hodnotu 0,86.

Organizácia spojených národov každoročne uverejňuje Index ľudského rozvoja (Human development index - HDI), ktorý sa skladá z Indexu vzdelania (Education index), Indexu hrubého domáceho produktu – HDP (GDP Index) a Indexu priemernej dĺžky života (Life Expectancy Index) (Human Development Reports).

Index vzdelania ako súčasť indexu ľudského rozvoja, je mierou vedomostí, meraných gramotnosťou dospelého obyvateľstva (dve tretiny váhy hodnoty údaju) a počtom prihlásených na školy prvého, druhého a tretieho stupňa (tretina váhy hodnoty údaju). Gramotnosťou dospelých sa rozumie schopnosť čítať a písať. Index vzdelávania môže

nadobúdať hodnoty od 0 po 1. 1 je najvyššie možné teoretické skóre, čo znamená perfektné vzdelanie. Všetky krajiny považované za vyspelé krajiny majú minimálne skóre 0,8 alebo vyššie, aj keď väčšina vyspelých krajín má skóre 0,9 alebo vyššie. V tabuľke 1 sú pre porovnanie uvedené krajiny s najvyššou a najnižšou hodnotou Indexu vzdelania vrátane Slovenska vo vybraných rokoch: v roku 2000 a 2008.

Tabuľka 1: Hodnoty Indexu vzdelania vybraných krajín (rok 2000 a 2008)

Poradie v roku 2000	Krajina	Index vzdelania rok 2000	Index vzdelania rok 2008
1	 Austrália	0,993	0,993
1	 Belgicko	0,993	0,974
1	 Fínsko	0,993	0,993
1	 Švédsko	0,993	0,974
42	 Slovensko	0,902	0,928
173	 Čad	0,274	0,293
174	 Mali	0,236	0,300
175	 Burkina Faso	0,158	0,274
176	 Niger	0,118	0,286

Zdroj: http://en.wikipedia.org/wiki/Education_Index

3. Metódy a interpretácia výsledkov

Na výpočet miery závislosti medzi Indexom demokracie a Indexom vzdelania v štátoch EU sme použili *Pearsonov – Bravaisov koeficient korelácie*, ktorý popíšeme nasledovne. Nech hodnoty znaku X sú $x_1, x_2, \dots, x_n$ a hodnoty znaku Y sú $y_1, y_2, \dots, y_n$. Najpoužívanejšou číselnou charakteristikou štatistickej závislosti dvoch kvantitatívnych znakov je korelačný koeficient (tzv. *Pearsonov – Bravaisov koeficient*), ktorý je definovaný takto (Anděl, 2003):

$$r = \frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum_{i=1}^n (x_i - \bar{x})^2} \sqrt{\sum_{i=1}^n (y_i - \bar{y})^2}},$$

Korelačný koeficient nadobúda hodnoty z intervalu $\langle -1, 1 \rangle$.

V našom prípade sme vypočítali hodnoty koeficientov korelácie nielen medzi Indexom vzdelania a Indexom demokracie, ale aj medzi Indexom vzdelania a každým z piatich kritérií Indexu demokracie. Všetky uvedené korelačné koeficienty sú pre ilustráciu vypočítané za rok 2000 a to pre novoasociované štáty EÚ, pre 15 štátov EÚ a pre EÚ celkom. Vypočítané hodnoty koeficientov korelácie sú prehľadne zapísané v tabuľke 2.

Tabuľka 2: Koeficienty korelácie medzi Indexom vzdelania a Indexom demokracie, medzi Indexom vzdelania a piatimi kritériami Indexu demokracie za rok 2000

	Index vzdelania		
	Novosociované štáty EÚ	15 štátov EÚ	EÚ celkom
Index demokracie (celkom)	-0,31068	0,666634	0,581879
Volebný proces a pluralita	0,62222	0,598601	0,647853
Fungovanie vlády	-0,27159	0,727629	0,578808
Politická spoluúčasť	-0,14822	0,588955	0,504737
Politická kultúra	-0,26759	0,432288	0,437741
Občianska sloboda	-0,25048	0,538379	0,440739

Poznámka. Kritická hodnota Pearsonovho koeficienta korelácie (s 10 stupňami voľnosti) pre novosociované krajiny je 0,576; pre 15 štátov EÚ (13 stupňov voľnosti) je 0,514 a pre EÚ celkom (25 stupňov voľnosti) je kritická hodnota Pearsonovho koeficienta korelácie 0,381.

Zdroj: vlastný výpočet

Na základe vypočítaných koeficientov korelácie môžeme konštatovať, že medzi Indexom vzdelania a Indexom demokracie (celkom) je koeficient korelácie štatisticky významný ako v 15 štátoch EÚ ($r = 0,666635$) tak aj v štátoch EÚ celkom ($r = 0,581879$). V novosociovaných štátoch sa koeficient korelácie štatisticky významne nelíši od nuly, t.j. neexistuje väzba medzi Indexom demokracie a Indexom vzdelania. Rovnako aj koeficienty korelácie medzi Indexom vzdelania a jednotlivými zložkami Indexu demokracie sa štatisticky významne nelíšia od nuly až na zložku Volebný proces a pluralita, kde koeficient korelácie je štatisticky významný ($r = 0,62222$). Záporné korelačné koeficienty, ale štatisticky preukazné nevýznamné medzi Indexom vzdelania a Indexom demokracie v novosociovaných štátoch sú dôsledkom povinnej školskej dochádzky. V 15 štátoch EÚ je štatisticky významný koeficient korelácie medzi Indexom vzdelania a všetkými zložkami Indexu demokracie okrem zložky politická kultúra. V tomto prípade sa koeficient korelácie štatisticky významne nelíši od nuly ($r = 0,432288$), t.j. medzi politickou kultúrou a Indexom vzdelania nie je štatisticky významná lineárna súvislosť. V štátoch EÚ celkom je závislosť Indexu vzdelania a Indexu demokracie štatisticky významná, rovnako ako sú štatisticky významné všetky koeficienty korelácie medzi Indexom vzdelania a zložkami Indexu demokracie.

4. Záver

Na vytvorenie demokratickej spoločnosti musia byť splnené podmienky nevyhnutné na existenciu demokracie. Jednou z nich je aj to, že občania musia mať na spoločenské otázky vytvorený vlastný názor. Občania musia byť vzdelaní. Musia mať k dispozícii všetky známe informácie potrebné na to, aby si mohli vytvoriť na spoločenské otázky vlastný a objektívny názor. Navyše sa musia o veci verejné dostatočne zaujímať. Cieľom však nie je samoúčelné vzdelávanie, ale výsledné znalosti občanov a všeobecný prehľad. Preto by vzdelávanie malo byť časovo a finančne efektívne a pre študentov zaujímavé. Vzťah ku vzdelaniu a ku demokracii by mal byť vstrepovaný žiakom už v základných školách, aby si uvedomovali, že vzdelávanie sa má zásadný význam pre demokraciu a demokratické hodnoty slobody. Vzdelanie je hlavnou zložkou prosperity a je dôležitým komponentom miery hospodárskeho rozvoja a kvality života. Kvalita života je kľúčovým faktorom, ktorý rozhoduje, či je krajina rozvinutá, rozvojová, alebo nedostatočne rozvinutá.

5. Literatúra

- [1]ANDEĽ, J. 2003. Statistické metody. MATFYZPRESS, Praha, 299 s. ISBN 80-86732-08-8
- [2]DEWEY, J. 1916. Demokracia a výchova, New York, Macmillan Company
- [3]Human Development Reports, <http://hdr.undp.org/en/>
- [4]KEKIC, L. 2007. The Economist Intelligence Unit's index of democracy.
http://www.economist.com/media/pdf/democracy_index_2007_v3.pdf
- [5]LESAGE, J.P. – KELLY PACE, R. 2005. Spatial and Spatiotemporal Econometrics, Volume 18 (Advances in Econometrics) [Hardcover], JAI Press; 1 edition, 340 p. ISBN-10: 0762311487, ISBN-13: 978-0762311484

Adresy autorov:

Ľubomír Zelenický, prof. RNDr., CSc.
Katedra fyziky FPV UKF
Trieda A. Hlinku 1
949 74 Nitra
lzelenicky@ukf.sk

Ondrej Šedivý, prof. RNDr., CSc.
Katedra matematiky FPV UKF
Trieda A. Hlinku 1
949 74 Nitra
osedivy@ukf.sk

OBSAH

	Úvod	1
Antoch J.	Change Point Detection	2
Riečan B.	Pravdepodobnosť na algebraických štruktúrach	16
Bauerová M., Brindza J., Stehlíková B., Tirpáková A.	Kvantifikácia biodiverzity	21
Brindza J., Bauerová M., Stehlíková B., Tirpáková A.	Porovnávanie mier biodiverzity s využitím štatistických metód	27
Broďáni J.	Prognózovania v športe	32
Fabová Ľ.	Vývoj nominálnych a reálnych úrokových sadzieb v SR	38
Hájková R.	Aktivita versus pasivita vybraných centrálnych bank	41
Chajdiak J.	Stupeň dôležitosti zdrojov informácií pri inovačných aktivitách	47
Janiga I., Žák K., Valent A.	Spôsobilosť procesu úpravy úžitkovej vody vzhľadom k vybranému faktoru	52
Janová J.	Validace optimalizačního modelu pro plánování produkce v zemědělství	57
Juriová J.	Prognózovanie vývoja nových objednávok v priemyselnej výrobe SR	63
Kapounek S.	Produkční a inflační mezera po zavedení eura – vliv společné monetární politiky ECB	69
Kapounek S.	Vybrané aspekty přijetí eura v souvislosti s ekonomickým růstem v České republice	75
Klůčik M.	Indikátory klasického cyklu pre priemysel SR	81
Kříž O., Neubauer J., Sedlačík M.	Výuka statistiky podporovaná excelovskou aplikací	87
Lenčes P.	Aby štatistika nebola klamstvo...	93
Luha J., Bianchi G.	Analýza odpovedí „neviem“ v batérii otázok	97
Martináková R.	Analýza vzájemného vztahu mezi výdaji na konečnou spotřebu a spotřebou domácností v ČR	105
Mura L.	Faktory internacionalizácie malého a stredného podnikania	111
Poláková Z., Obtulovič P.	Komparácia regiónov SR a ČR z pohľadu vybraných demografických ukazovateľov	117
Poměnková J., Maršálek R.	Spektrální analýza cyklické struktury průmyslové výroby ČR	123
Řezáč, M.	Alternativní indexy predikční síly credit scoringových modelů	129
Rybanský Ľ.	Handicapy a model výpočtu pravdepodobnosti z kurzu	135
Soukupová K.	Cílování inflace v České republice	140
Stehlíková B., Markechová D.	Štatistika pre 21. storočie vo vedeckom výskume a vzdelávaní	146

Šunderlík J., Havrlent Š., Rybanský Ľ.	Zhluková analýza štruktúry slovenských a anglických vyučovacích hodín	151
Tirpák J.	Grafické vyhodnotenie pôdnej heterogenity	157
Urbaníková M.	Štatistické metódy a umelá inteligencia na kapitálových trhoch	161
Vrábelová M., Munk M., Kapusta J.	Multinominálny logitový model pravdepodobností prístupov na webové časti portálu	167
Zelenický Ľ., Šedivý O.	Vzťah medzi demokraciou a vzdelaním	173

CONTENTS

	Introduction	1
Antoch J.	Change Point Detection	2
Riečan B.	Probability on Algebraic Structures	16
Bauerová M., Brindza J., Stehlíková B., Tirpáková A.	Quantification of biodiversity	21
Brindza J., Bauerová M., Stehlíková B., Tirpáková A.	Comparing measures of biodiversity using statistical methods	27
Broďáni J.	Forecasting in sport	32
Fabová Ľ.	Nominal and real interest rates development in Slovak Republic	38
Hájková R.	The activity vs. passivity of selected central banks	41
Chajdiak J.	Degree of Importance Sources of Information for Innovation Activities	47
Janiga I., Žák K., Valent A.	Process capability of supply water modification with respect to chosen factor	52
Janová J.	Validation of production planning optimization model in agriculture	57
Juriová J.	Forecasting new orders in manufacturing in the SR	63
Kapounek S.	Growth and Inflation Gap after joining the Euro – impact of the ECB's Single Monetary Policy	69
Kapounek S.	Selected Aspects of Joining the Euro and Economic Growth in the Czech Republic	75
Kľúčik M.	Industry Classical Cycle Indicators of Slovakia	81
Kříž O., Neubauer J., Sedlačík M.	Learning statistics endorsed by Excel	87
Lenčes P.	So statistic not be lie...	93
Luha J., Bianchi G.	The analysis of „do not know“ responses in the battery of questions	97
Martináková R.	Analysis of the mutual relation of the final consumption expenditure and the household consumption in Czech republic	105
Mura L.	Factors of internationalization of small and medium enterprise	111
Poláková Z., Obtulovič P.	Comparison of regions of the Slovak Republic and Czech Republic from the viewpoint of chosen demographic indicators	117
Poměnková J., Maršálek R.	Spectral analysis of the cyclical behaving of the Czech Republic industrial production	123
Řezáč, M.	Alternative indexes of predictive power for credit scoring models	129
Rybanský Ľ.	Handicaps and model of computing probability from rates	135
Stehlíková B., Markechová D.	Statistics for the Twenty-First Century in the Scientific Research and in the Education	140

Soukupová K.	Inflation targeting in the Czech Republic	146
Šunderlík J., Havrlent Š., Rybanský Ľ.	Cluster analysis of the slovak and english lessons	151
Tirpák J.	Graphic plotting soil heterogeneity	157
Urbaníková M.	Statistical methods and artificial intelligence on kapital markets	161
Vrábelová M., Munk M., Kapusta J.	The Multinomial Logit Model of the Probabilities of Accesses to the Web Parts of Portal	167
Zelenický Ľ., Šedivý O.	The relationship between democracy and education	173

Pokyny pre autorov

Jednotlivé čísla vedeckého časopisu FORUM STATISTICUM SLOVACUM sú prevažne tematicky zamerané zhodne s tematickým zameraním akcií SŠDS. Príspevky v elektronickej podobe prijíma zástupca redakčnej rady na elektronickej adrese uvedenej v pozvánke na konkrétne odborné podujatie Slovenskej štatistickej a demografickej spoločnosti. Názov word-súboru uvádzajte a posielajte v tvare: **priezvisko_nazovakcie.doc**

Forma: Príspevky písané výlučne len v textovom editore MS WORD, verzia 6 a vyššia do verzie 2003, písmo Times New Roman CE 12, riadkovanie jednoduché (1), formát strany A4, všetky okraje 2,5 cm, strany nečíslovať. Tabuľky a grafy v čierno-bielom prevedení zaradiť priamo do textu článku a označiť podľa šablony. Bibliografické odkazy uvádzať v súlade s normou STN ISO 690 a v súlade s medzinárodnými štandardami. Citácie s poradovým číslom z bibliografického zoznamu uvádzať priamo v texte.

Rozsah: Maximálny rozsah príspevku je 6 strán.

Príspevky sú recenzované. Redakčná rada zabezpečí posúdenie príspevku členom redakčnej rady alebo externým oponentom.

Štruktúra príspevku: (Pri písaní príspevku využite elektronickú šablónu: <http://www.ssds.sk/> v časti Vedecký časopis, Pokyny pre autorov.)

Názov príspevku v slovenskom jazyku (štýl Názov: Time New Roman 14, Bold, centrovať)

Názov príspevku v anglickom jazyku (štýl Názov: Time New Roman 14, Bold, centrovať)

Vynechať riadok

Meno1 Priezvisko1, Meno2 Priezvisko2 (štýl normálny: Time New Roman 12, centrovať)

Vynechať riadok

Abstract: Text abstraktu v anglickom jazyku, max. 10 riadkov (štýl normálny: Time New Roman 12).

Vynechať riadok

Key words: Kľúčové slová v anglickom jazyku, max. 2 riadky (štýl normálny: Time New Roman 12).

Vynechať riadok

Kľúčové slová: Kľúčové slová v jazyku v akom je napísaný príspevok, max. 2 riadky (štýl normálny: Time New Roman 12).

Vynechať riadok

Vlastný text príspevku v členení:

1. **Úvod** (štýl Nadpis 1: Time New Roman 12, bold, zarovnať vľavo, číslovať)
2. **Názov časti 1** (štýl Nadpis 1: Time New Roman 12, bold, zarovnať vľavo, číslovať)
3. **Názov časti 1. . .**
4. **Záver** (štýl Nadpis 1: Time New Roman 12, bold, zarovnať vľavo, číslovať)

Vlastný text jednotlivých častí je písaný štýlom Normal: písmo Time New Roman 12, prvý riadok odseku je odsadený vždy na 1 cm, odsek je zarovnaný s pevným okrajom. Riadky medzi časťami nevynechávajúte.

5. **Literatúra** (štýl Nadpis 1: Time New Roman 12, bold, zarovnať vľavo, číslovať)

[1] Písať podľa normy STN ISO 690

[2] GRANGER, C.W. – NEWBOLD, P. 1974. Spurious Regression in Econometrics. In: Journal of Econometrics, č. 2, 1974, s. 111 – 120.

Adresa autora (-ov) (štýl Nadpis 1: Time New Roman 12, bold, zarovnať vľavo, adresy vpísať do tabuľky bez orámovania s potrebným počtom stĺpcov a s 1 riadkom):

Meno1 Priezvisko1, tituly1
Ulica1
970 00 Mesto1
meno1.priezvisko1@mail.sk

Meno2 Priezvisko2, tituly2
Ulica2
970 00 Mesto2
meno2.priezvisko2@mail.sk

FORUM STATISTICUM SLOVACUM

vedecký recenzovaný časopis Slovenskej štatistickej a demografickej spoločnosti

Vydavateľ

Slovenská štatistická a demografická
spoločnosť
Miletičova 3
824 67 Bratislava 24
Slovenská republika

Redakcia

Miletičova 3
824 67 Bratislava 24
Slovenská republika

Fax

02/39004009

e-mail

chajdiak@statis.biz
jan.luha@fmed.uniba.sk

Registráciu vykonalo

Ministerstvo kultúry Slovenskej republiky

Registračné číslo

3416/2005

Evidenčné číslo

EV 3287/09

Tematická skupina

B1

Dátum registrácie

22. 7. 2005

Objednávky

Slovenská štatistická a demografická
spoločnosť
Miletičova 3, 824 67 Bratislava 24
Slovenská republika
IČO: 178764
DIČ: 2021504276
Číslo účtu: 0011469672/0900

ISSN 1336-7420

Redakčná rada

RNDr. Peter Mach – *predseda*

Doc. Ing. Jozef Chajdiak, CSc. – *šéfredaktor*

RNDr. Ján Luha, CSc. – *tajomník*

členovia:

Ing. František Bernadič
Mgr. Branislav Bleha, PhD.
Ing. Mikuláš Cár, CSc.
Ing. Ján Cuper
Ing. Pavel Flák, DrSc.
Ing. Edita Holíčková
Doc. RNDr. Ivan Janiga, CSc.
Ing. Anna Janusová
Doc. RNDr. PaedDr. Stanislav Katina, PhD.
Prof. RNDr. Jozef Komorník, DrSc.
RNDr. Samuel Koróny, PhD.
Doc. Ing. Milan Kovačka, CSc.
Doc. RNDr. Bohdan Linda, CSc.
Prof. RNDr. Jozef Mládek, DrSc.
Doc. RNDr. Oľga Nánásiová, CSc.
Doc. RNDr. Karol Pastor, CSc.
Prof. RNDr. Rastislav Potocký, CSc.
Doc. RNDr. Viliam Páleník, PhD.
Ing. Iveta Stankovičová, PhD.
Prof. RNDr. Beata Stehlíková, CSc.
Prof. RNDr. Anna Tirpáková, CSc.
Prof. RNDr. Michal Tkáč, CSc.
Ing. Vladimír Úradníček, PhD.
Ing. Boris Vaňo
Doc. MUDr. Anna Volná, CSc., MBA.
Ing. Mária Vojtková, PhD.
Prof. RNDr. Gejza Wimmer, DrSc.
Mgr. Milan Žirko

Ročník

VI.

Číslo

2/2010

Cena výtlačku 20 EUR

Ročné predplatné 80 EUR