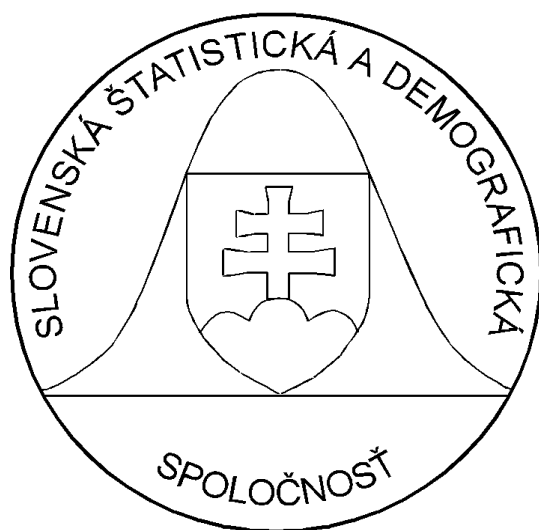


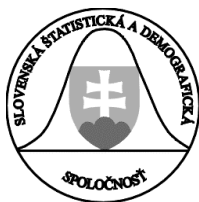
6/2008

FORUM STATISTICUM SLOVACUM



ISSN 1336-7420





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



Naše najbližšie akcie:

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

17. Medzinárodný seminár VÝPOČTOVÁ ŠTATISTIKA

4. – 5. 12. 2008, Bratislava, Infostat

Prehliadka prác mladých štatistikov a demografov

4. 12. 2008, Bratislava, Infostat

Pohl'ady na ekonomiku Slovenska 2009

7. apríl 2009, Bratislava

EKOMSTAT 2009, 23. škola štatistiky

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

31. 5. 2009 – 5. 6. 2009, Trenčianske Teplice

12. SLOVENSKÁ DEMOGRAFICKÁ KONFERENCIA

tematické zameranie: Demografická budúcnosť Slovenska

23. – 25. 9. 2009, Belušké Slatiny, Trenčiansky kraj

Aplikácie metód na podporu rozhodovania vo vedeckej, technickej a spoločenskej praxi

jún 2009, STU Bratislava

Regiónálne akcie

priebežne

ÚVOD

Vážené kolegyně, vážení kolegovia,

šieste číslo piateho ročníka vedeckého časopisu Slovenskej štatistickej a demografickej spoločnosti je zostavené z príspevkov, ktoré sú obsahovo orientované v súlade s tematikou V. ročníka Medzinárodnej konferencie aplikovanej štatistiky FernStat 2008. Táto konferencia sa uskutočnila v dňoch 2. a 3. októbra 2008 v hoteli Lesák v Tajove. Tradičné tematické okruhy konferencie sú: Aplikovaná štatistika, Demografická štatistika, Matematická štatistika, Štatistické riadenie kvality.

Akciu z poverenia Výboru SŠDS, zorganizoval Organizačný a programový výbor: Ing. Vladimír Úradníček PhD. – predseda, Ing. Mária Kanderová, PhD.–tajomník, RNDr. Ján Luha, CSc., doc. Ing. Jozef Chajdiak, CSc., doc. RNDr. Bohdan Linda, CSc., doc. Dr. Jana Kubanová, PhD., RNDr. Peter Mach, Ing. Iveta Stankovičová PhD.

Na príprave a zostavení tohto čísla FORUM STATISTICUM SLOVACUM participovali: Ing. Vladimír Úradníček PhD, Ing. Mária Kanderová PhD., doc. Ing. Jozef Chajdiak, CSc., RNDr. Ján Luha, CSc., Miroslav Kello, Juraj Varga.

Recenziu príspevkov zabezpečili: Ing. Vladimír Úradníček PhD, Ing. Mária Kanderová PhD., doc. Ing. Jozef Chajdiak, CSc., RNDr. Ján Luha, CSc., doc. RNDr. Bohdan Linda, CSc.

Výbor SŠDS

Approaches to volatility estimation based on historical data^{*}

Martin Bod'a

Abstract: Volatility as a quantifiable characteristic is an input to many applications in the area of financial risk management. Even so, it possesses features of abstraction of some degree as for its unequivocal definition. It is because of the unstable mechanics of financial markets and vagueness of its perception in some sense that several ways to estimate volatility have been devised, and the most commonly preferred ones are summarized by the article.

Key words: historical volatility, geometric Brownian motion, close-to-close volatility, Parkinson's high-low estimator of volatility, Garman-Klass estimator of volatility, Rogers-Satchell estimator of volatility, Yang-Zhang estimator of volatility.

1. Introduction

Speaking of risk management, it must be understood that its chief focus is measuring risk and that its attempt to translate risk into a quantifiable characteristic, a number, is utilized throughout virtually entire contemporary financial management. Theory and practice have come to the mutual agreement that risk is a pervading and unavoidable trait of any entrepreneurial activity (to say nothing of any human activity), and that it must be measured. This statement especially holds for the field of financial markets where consequences of risk are possibly felt severer than elsewhere. Financial reality forces traders to take care in building their positions and menaces with the repetition of the well documented events of Black Monday and Black Tuesday of 1929 or Black Monday of 1987. Measuring risk has become a necessity for the engagement in trading at financial markets.

Whereas, not even a century ago, theoretical literature occupied itself with debates on what risk is and sought to define the distinctive terms *risk* and *uncertainty* correctly; nowadays, these efforts *per se* are abandoned and in most literature the terms *risk* and *uncertainty* are interpreted loosely. At financial markets the term *risk* is construed operationally as synonymous with the term *volatility*, and this is where the interest of this paper lies. Volatility, as a measure of market risk, must be estimated and several ways to its estimation are employable – it is the intention of this paper is to summarize some of commonly employed approaches and to demonstrate its practical use.

2. Theoretical and practical perspective on volatility

It is frequently assumed that the process generating returns of an asset follows a geometric Brownian motion. This implies – albeit several writings are possible – that the price of an asset S_t at any time $t \in (0, \infty)$ shall be governed by the formula

$$dS_t = \mu_t S_t dt + \sigma_t S_t dW_t, \quad (1)$$

where W_t is a Wiener process, μ_t denotes the drift of the process and σ_t represents a noise factor of the process. The task is to estimate correctly the noise factor of the process.

Taking μ_t and σ_t to be constant, i. e. μ and σ respectively, and defining W_t as a one-dimensional Gaussian process (in consistency with literature, see Baxter, Rennie, 1999), one obtains instantly that

(A.) for any $t \in (0, \infty)$ the random variable dS_t/S_t is independent of all prices until t , and

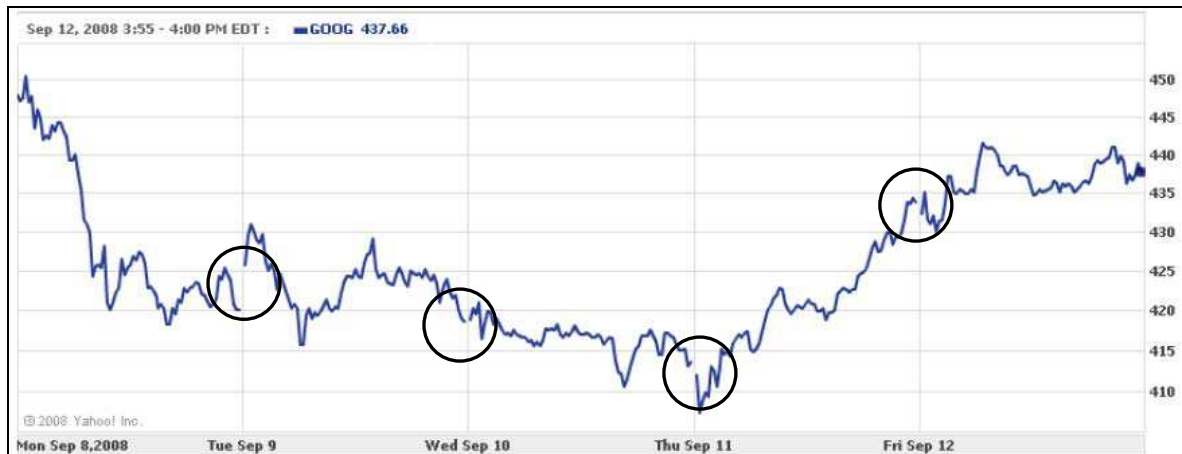
^{*} The paper was prepared under the framework of VEGA project No. 1/4634/07 „Variant methods of prediction of small and medium sized enterprises development after introducing single European currency in the Slovak Republic.“

(B.) for any $t \in (0, \infty)$ the random variable dS_t/S_t is normally distributed with mean μdt and dispersion $\sigma^2 dt$.

It is vital to notice that the relativized increments of the price of an asset having may be treated as the returns of the asset; and that having complied with implication (A.), formula (1) introduces the normal distribution for the asset returns.

The assumption of formula (1) is then utilized in constructing estimators for the noise factor σ , i. e. volatility of the process. In this quest, observations of historical prices are often utilized and this approach leads to estimates of historical volatility.

Theoretically, the process determined by formula (1) is assumed to run continually and continuously at $(0, \infty)$. Never the less, empirically, asset prices are observable only during the periods they are traded. Therefore, there are some time pauses in the series of asset price observations, and the series of observations of one trading day does not link to the series of observations of the day before without intermittence. Furthermore, these observations are discrete, not continuous; all though they are usually so dense that they can be approximated as continuous. These statements are illustrated in scheme 1 by a series of Google Inc. share prices over 5 trading days (from 08/09/2008 until 12/09/2008). High frequency of transactions with the shares of Google Inc. conduces to creating of an impression of a continuous process, however, the differences between closing prices and opening prices between to subsequent trading days are remarkable. It is needless to say that trading days are scarcer than civil days and that between two available observations need not necessarily be the same time gap of, say, 16 hours as it is between two subsequent civil days which are trading days.



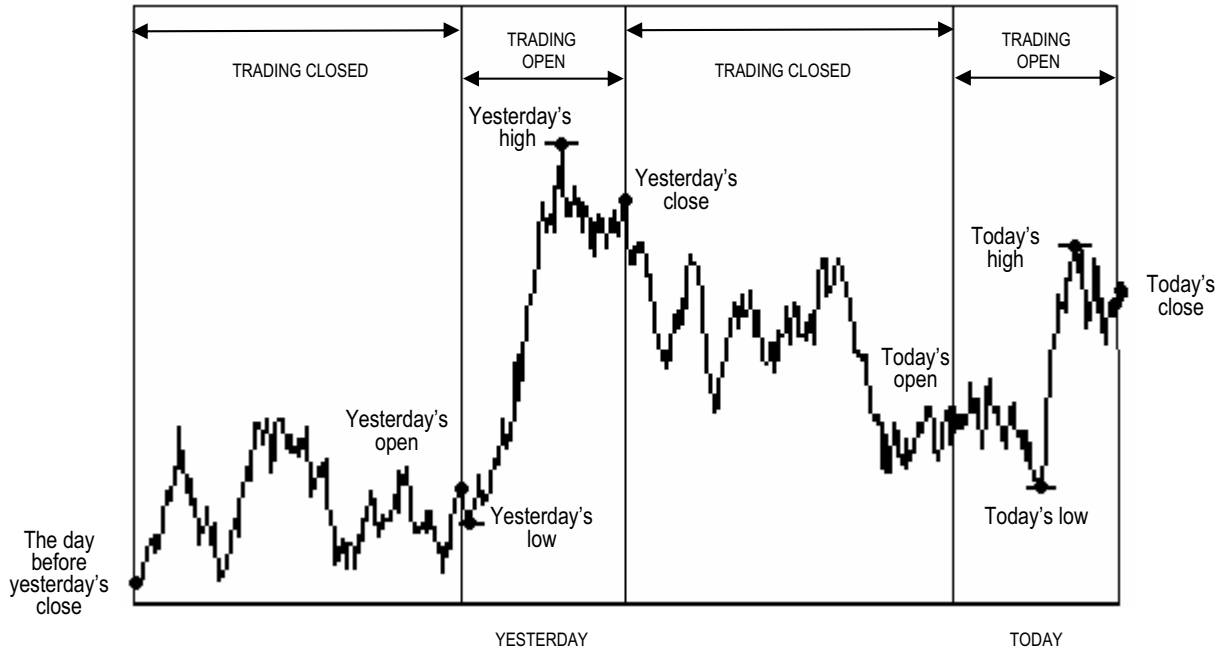
Scheme 1 Observations of the Google Inc. share prices from 08/09/2008 to 12/09/2008 with time interruptions
(Source: <http://finance.yahoo.com>)

Another aspect to make allowance for is the availability of information. All though it is possible, at certain costs, to equip oneself with all price observations during the given period, for day-to-day estimation of volatility of a large portfolio of assets it may amount dear to found the estimation procedure upon all price observations. Rather, as some information is made public free of charge once the trading day has been closed, it may be useful and reasonable to base the estimation procedure upon this information providing that it substitutes all price observations satisfactorily. This information includes opening prices (open), maximum prices (high), minimum prices (low), and closing prices (close) of every trading day as it is reported in table 1 for the prices of the Google Inc. shares from 08/09/2008 up to 12/09/2008 (all denominated in USD). The table moreover contains the volume of transactions (in USD) and the adjusted closing price (in USD) which accounts for dividends and splits.

Table 1 The published free information on the Google Inc. share prices from 08/09/2008 to 12/09/2008
(Source: <http://finance.yahoo.com>)

Date	Open	High	Low	Close	Volume	Adjusted Close
12/09/2008	430.21	441.99	429	437.66	6028000	437.66
11/09/2008	408.35	435.09	406.38	433.75	6471400	433.75
10/09/2008	424.47	424.48	409.68	414.16	6226800	414.16
09/09/2008	423.17	432.38	415	418.66	7229600	418.66
08/09/2008	452.02	452.94	417.55	419.95	9017900	419.95

In accordance with the aforementioned remarks it must be borne in mind that the process generating asset prices is considered a geometric Brownian motion which is continuous regardless of the fact that some of its development is regularly unobservable. During the observable periods, the information on the process comprises the opening price, the highest price, the lowest price and the closing price at which the asset was traded at a given trading period. The alternation of trading periods and non-trading periods with the process running in their background is illustrated in scheme 2. In the following sections are set forth the approaches towards volatility estimation based upon the information on opens, highs, lows and closes.



Scheme 2 The process generating asset prices in the flow of time
(Adjusted according to Garman and Klass, 1992)

3. Notation

Whilst it is assumed that one uses T historical days to estimate volatility, the estimation procedures described later are presented under this adopted notation:

O_t	The opening price on day t ,
H_t	The highest price on day t ,
L_t	The lowest price on day t ,
C_t	The closing price on day t ,
$R_t := \ln(C_t / C_{t-1})$	The logarithmic return on day t ,
$h \in (0, 1)$	The fraction of the civil day when it is traded in the financial market,
M	The number of trading days in a year.

It is obvious that it is convenient to estimate variance of returns (the square of volatility) instead of volatility itself.

4. Procedures for estimating daily volatility

The classical line works on closing prices and employs the unbiased variance estimator in the form

$$\hat{\sigma}_{close-to-close}^2 := \frac{1}{T-1} \sum_{t=1}^{t=T} (R_t - \bar{R})^2, \quad (2)$$

where $\bar{R}$ is the mean return over the T days. This close-to-close is preferred for its simplicity and intuition; the recommendations for its use can be found in most of textbooks (see e. g. Hull, 2000, or Ross, 1999). However, grounded in closing prices only, this estimator ignores much information on intraday volatility which may significantly contribute to its magnitude.

First advances in the improvement of volatility estimation are due to Parkinson (1980) who derived the formula using only the high and low prices (which he regards as extreme values, whilst he calls his method the extreme value method)

$$\hat{\sigma}_{Parkinson's\ high-low}^2 := \frac{1}{T 4 \ln 2} \sum_{t=1}^{t=T} \left(\ln \frac{H_t}{L_t} \right)^2. \quad (3)$$

As Garman and Klass (1992) remark, “intuition would then tell us that high/low prices contain more information regarding volatility than do open/close prices”. The drawback of Parkinson’s high-low estimator rests with the assumptions. He (implicitly) assumes that the process in (1) has zero drift (i. e. $\mu = 0$) and that trading is conducted ceaselessly over the entire day. Garman and Klass (1992) assist in improving Parkinson’s formula by remedying the second said imperfection and they suggest the estimator

$$\hat{\sigma}_{improved\ Parkinson's\ high-low}^2 := \frac{1}{T} \left[\frac{0.83}{h 4 \ln 2} \sum_{t=1}^{t=T} \left(\ln \frac{H_t}{L_t} \right)^2 + \frac{0.17}{1-h} \sum_{t=1}^{t=T} \left(\ln \frac{O_t}{C_{t-1}} \right)^2 \right], \quad (4)$$

in which the coefficients 0.83 and 0.17 were determined in the fashion providing (4) with the highest efficiency. Criticism of Parkinson’s formula is established on the fact that formulæ (3) and (4) do not allow for the joint effects between the open, high, low and close of a given day.

Garman and Klass (1992), too, assume zero drift and no interruptions in trading, however, they seek the best analytic scale-invariant estimator, whereas “best” stands for unbiased with minimum variance. They derive their best analytic scale-invariant estimator of historical volatility in the simplified form

$$\hat{\sigma}_{Garman-Klass}^2 := \frac{1}{T} \left[\frac{1}{2} \sum_{t=1}^{t=T} \left(\ln \frac{H_t}{L_t} \right)^2 - (2 \ln 2 - 1) \sum_{t=1}^{t=T} \left(\ln \frac{C_t}{O_t} \right)^2 \right], \quad (5)$$

and suggest its improvement after considering interruptions in trading in the form

$$\hat{\sigma}_{improved\ Garman-Klass}^2 := \frac{0.88}{h} \hat{\sigma}_{Garman-Klass}^2 + \frac{0.12}{1-h} \frac{1}{T} \sum_{t=1}^{t=T} \left(\ln \frac{O_t}{C_{t-1}} \right)^2, \quad (6)$$

where, again, the values 0.88 and 0.12 were set to bring forth the highest efficiency.

Rogers and Satchell in 1991 take into account non-zero drift of the process in (1), albeit they operate under the assumption of no interruptions in trading. Their estimator is given by

$$\hat{\sigma}_{Rogers-Satchell}^2 := \frac{1}{T} \sum_{t=1}^{t=T} \left[\ln \frac{H_t}{O_t} \ln \frac{H_t}{C_t} + \ln \frac{L_t}{O_t} \ln \frac{L_t}{C_t} \right]. \quad (7)$$

A very interesting result is by Yang and Zhang (2000) who were the first to construct an historical volatility estimator that – besides the property of minimum variance – is independent of the drift and is independent of interruptions in trading. It thus does not matter whether there be zero or nonzero drift or how long trading be closed. Its efficiency is 14 times higher than that of the close-to-close estimator in (2). The advantages of the estimator of Yang and Zhan are at the expense of simplicity. The formula is ruled by

$$\hat{\sigma}_{Yang-Zhang}^2 := \hat{\sigma}_{open}^2 + k \hat{\sigma}_{close}^2 + (1-k) \hat{\sigma}_{Rogers-Satchell}^2, \quad (8a)$$

where

$$\hat{\sigma}_{open}^2 := \frac{1}{T-1} \sum_{t=1}^{t=T} \left(\ln \frac{O_t}{C_{t-1}} - \overline{\ln \frac{O}{C_{-1}}} \right)^2, \quad (8b1)$$

$$\hat{\sigma}_{close}^2 := \frac{1}{T-1} \sum_{t=1}^{t=T} \left(\ln \frac{C_t}{O_t} - \overline{\ln \frac{C}{O}} \right)^2, \quad (8b2)$$

$$k = 0.17 \left(1 - \frac{1}{T} \right), \quad (8b3)$$

with $\overline{\ln \frac{O}{C_{-1}}} := \frac{1}{T} \sum_{t=1}^{t=T} \ln \frac{O_t}{C_{t-1}}$ and $\overline{\ln \frac{C}{O}} := \frac{1}{T} \sum_{t=1}^{t=T} \ln \frac{C_t}{O_t}$.

The formula in (8a) is suggestive of some weighted average of the open-to-close volatility, the close-to-open volatility and the cross-product open-high-low-close estimator of Rogers and Satchell.

It is worth mentioning that Magdon-Ismail and Atiya establish their derivation of the volatility estimator on the maximum likelihood approach. They set up the conditional probability density $f(H, L | \mu, \sigma, S)$ for the high price (H) and the low price (L) over a period of T days if the process undergoes a geometric Brownian motion with the parameters μ and σ respectively and begins at a certain value S . The estimate for volatility is determined by

$$\hat{\sigma}_{MLE} := \arg \max_{\sigma} L(\mu, \sigma), \quad (9)$$

wherein $f(H, L | \mu, \sigma, S) \equiv L(\mu, \sigma)$. In the article, the presentation of the likelihood function (or the conditional density) is omitted for its pure technical character.

5. Procedures for obtaining yearly volatility

Having the estimate of daily volatility at hand, one can use the square-root-of-time rule, which is under the assumption of geometric Brownian motion true, to obtain the estimate of yearly volatility. This requires the employment of the conversion relationship

$$\hat{\sigma}_{yearly}^2 = M \hat{\sigma}_{daily}^2. \quad (10)$$

6. Summary

The article gives a brief insight on some of the ways to estimate historical volatility in financial markets. Over 30 years, several methods (more or less adequate to the true spirit of financial markets) have been invented since literature commenced to occupy itself with the *correct* estimation of volatility. The most spread method is based upon closing prices solely – so called close-to-close volatility. Notwithstanding, its simplicity warrants not its precision. It

is to be taken under advisement that closing prices solely do not exhaust the entire information on price dynamics, realizing that there are also opening prices, maximum prices and minimum prices available. In general, assuming a geometric Brownian motion for the process generating asset prices it was arrived at volatility estimators also considering opens, highs and lows. At present, the theoretically and empirically soundest estimator of volatility is perhaps ascribable to Yang and Zhang and is free of the shortcomings of the other volatility estimators overviewed in the article.

7. References

- [1.] BAXTER, Martin, RENNIE, Andrew. 1999. *Financial calculus. An introduction to derivative pricing*. Cambridge : Cambridge University Press, 1999. ISBN 0-521-55289-3.
- [2.] BRANDT, Michael W., KINLAY, Jonathan. 2004. *Estimating historical volatility*. Workpaper.
- [3.] GARMAN, Mark B., KLASS, Michael J. 1992. *The estimation of security price volatility from historical data*. An updated version of work originally titled "On the estimation of security price volatilities from historical data," published in *Journal of Business*. 1980. Vol. 53. No. 1. Pp. 67-78.
- [4.] HULL, John C. 2000. *Options, Futures & Other Derivatives*. Fourth Edition. Upper Saddle River [USA] : Prentice Hall 2000. ISBN 0-13-015822-4.
- [5.] MAGDON-ISMAIL, Malik, ATIYA, Amir F. *Volatility estimation using high, low, and close data – a maximum likelihood approach*. Workpaper
- [6.] PARKINSON, Michael, 1980. *The extreme value method for estimating the variance of the rate of return*. In *Journal of Business*. 1980. Vol. 53. No 1. Pp. 61-65.
- [7.] ROSS, Sheldon M 1999. *An introduction to mathematical finance: options and other topics*. Cambridge : Cambridge University Press 1999. ISBN 0-521-77043-2.
- [8.] YANG, Dennis, ZHANG, Qiang 2000. *Drift-Independent volatility estimation based on high, low, open, and close prices*. In *Journal of Business*. 2000. Vol. 73. No 3. Pp. 477-491.

The author's address

Martin Bod'a, Ing. et Bc.
Univerzita Mateja Bela v Banskej Bystrici
Ekonomická fakulta
Katedra kvalitatívnych metód a informačných systémov
Tajovského 10
975 90 Banská Bystrica
martin.boda@umb.sk

Overovanie efektívnosti trhového portfólia v priestore mean-variance: model CAPM¹

Martin Bod'a² Mária Kanderová³

Abstract: In this contribution we present some results from Markowitz model of the selection of efficient portfolio in space mean-variance. The usual CAPM equation is direct implication of the mean-variance efficiency of the market portfolio. We focus on the econometric analysis of Sharpe and Lintner version of the CAPM. The contribution presents the framework for the estimation of the parameters of the CAPM and some statistical tests for verification of its validity.

Key words: mean-variance efficient portfolio, market portfolio, expected excess return, riskfree asset, zero-beta portfolio, Sharpe ratio, Wald test statistic, MacKinlay statistic

1. Úvod

Markowitzová teória výberu portfólia položila základy pre odvodenie modelov oceňovania kapitálových aktív, ktoré vysvetľujú vzťah medzi očakávaným výnosom každého rizikového aktíva a jeho rizikom za podmienok rovnováhy na trhu. Podľa tejto teórie investori budú držať efektívne portfólio v priestore mean-variance, to znamená také, ktoré má najvyšší očakávaný výnos pre danú úroveň rizika. Sharpe (1964) a Lintner (1965) rozšírili výsledky Markowitzovho modelu. Ukázali, že ak investori majú homogénne očakávania a držia portfóliá, ktoré sú efektívne v priestore mean-variance, potom za predpokladu, že neexistujú trhové obmedzenia, trhové portfólio bude tiež portfólio efektívne v priestore mean-variance.

2. Výber portfólia v priestore mean-variance

Modely oceňovania kapitálových aktív sú založené na Markowitzovej teórii portfólia. Markowitz vysvetlil fenomén diverzifikácie portfólia v priestore mean-variance.

Nech portfólio obsahuje N rizikových aktív s $(N \times 1)$ vektorom očakávaných výnosov μ a $(N \times N)$ kovariančnou maticou Σ . Nech w je $(N \times 1)$ vektor váh aktív v portfóliu.

Portfólio P je portfólio s minimálnym rozptylom zo všetkých portfólií so stredným výnosom μ_P ak vektor váh w je riešením nasledujúcej úlohy kvadratického programovania

$$\min \sigma_P^2 = w' \Sigma w \quad (1)$$

za podmienok

$$w' \mu = \mu_P \quad (2)$$

$$w' e = 1 \quad (3)$$

Úloha vyjadruje vzťah medzi rozptylom portfólia σ_P^2 s minimálnym rozptylom a jeho očakávaným výnosom μ_P . Grafické vyjadrenie vzťahu medzi smerodajnou odchýlkou portfólia s minimálnym rozptylom a jeho očakávaným výnosom predstavuje hranicu portfólií s minimálnym rozptylom (obrázok 1). Portfóliá s minimálnym rozptylom s výnosom väčším alebo rovným, ako je očakávaný výnos portfólia s globálne minimálnym rozptylom G , sú efektívne portfólia.

Nech P a Z sú portfólia s minimálnym rozptylom. Ak platí, že ich kovariancia je rovná nule

$$w_P' \Sigma w_Z = 0 \quad (4)$$

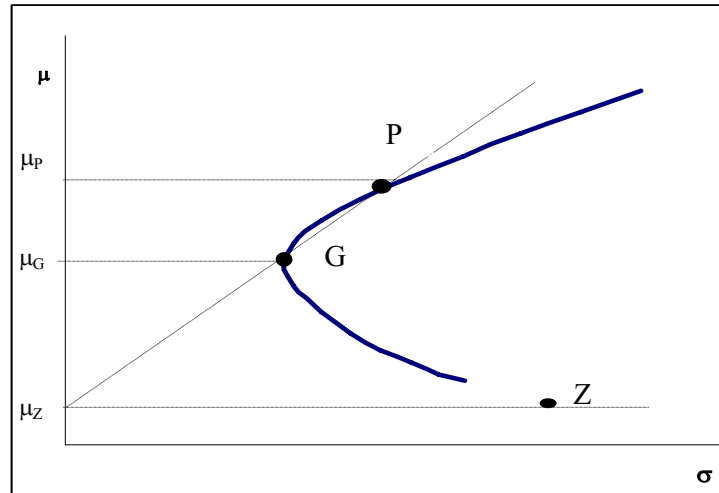
¹ Tento príspevok bol spracovaný v rámci riešenia grantovej úlohy KEGA3/5214/07

² Martin Bod'a, Katedra kvantitatívnych metód a informatiky, Ekonomická fakulta UMB v Banskej Bystrici

³ Mária Kanderová, Katedra kvantitatívnych metód a informatiky, Ekonomická fakulta UMB Banskej Bystrici

takéto portfólia sú ortogonálne. Portfólio Z sa nazýva zero-beta portfólio vzhľadom na portfólio P .

Obrázok 1 Hranica portfólií N -rizikových aktív s minimálnym rozptylom



Dôležitou vlastnosťou úlohy výberu portfólia je veta o vyčlenení dvoch fondov.

Nech P a R sú dve portfólia s minimálnym rozptylom, s očakávaným výnosom μ_P a μ_R , kde $\mu_P \neq \mu_R$. Potom

1. Každé portfólio S s minimálnym rozptylom je lineárnou kombináciou portfólií P a R .
2. Hranica portfólií s minimálnym rozptylom môže byť generovaná z ľubovoľných dvoch rôznych portfólií s minimálnym rozptylom.
3. Každé portfólio, ktoré je lineárnou kombináciou portfólií s minimálnym rozptylom, je tiež portfólio s minimálnym rozptylom.
3. Nech portfólia P a R sú efektívne portfólia s minimálnym rozptylom. Potom portfólio S , ktoré je ich lineárnou kombináciou, je tiež efektívne portfólio v priestore mean-variance.
4. Ku každému portfóliu P s minimálnym rozptylom okrem portfólia s globálne minimálnym rozptylom existuje jedno portfólio Z s minimálnym rozptylom s očakávaným výnosom μ_Z , ktoré má nulovú kovarianciu s portfóliom P . V špeciálnom prípade, ak existuje bezriziková úroková sadzba R_f , potom $\mu_Z = R_f$.

Z uvedených tvrdení vyplýva, že akékoľvek efektívne portfólio v priestore mean-variance možno generovať ľubovoľnými dvoma portfóliami, ktoré sú efektívne v tomto priestore.

3. Portfólio v priestore mean-variance a bezrizikové aktívum.

Predpokladajme, že portfólio je zložené z bezrizikového aktíva a z N rizikových aktív. Nech $\gamma = \mu - R_f e$ je $(N \times 1)$ vektor dodatočných výnosov portfólia, R_f je bezriziková úroková sadzba a e je $(N \times 1)$ jednotkový vektor. Pre dodatočný očakávaný výnos portfólia γ_P platí

$$\gamma_P = w^T \mu + (1 - w^T e) R_f - R_f = w^T \gamma. \quad (5)$$

Úloha výberu portfólia N rizikových aktív s bezrizikovým aktívom je riešením nasledovnej úlohy kvadratického programovania

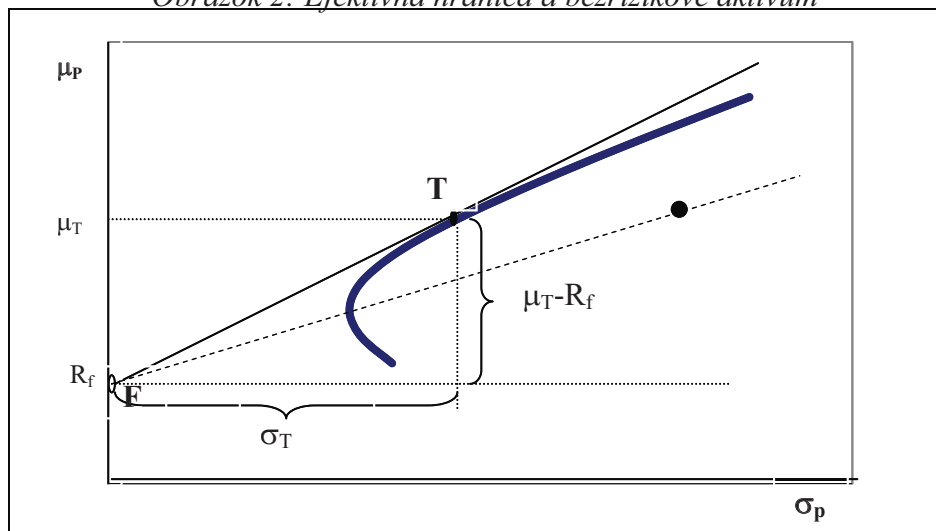
$$\min \sigma_P^2 = w^T C w \quad (6)$$

za podmienok

$$w^T \gamma = \gamma_P. \quad (7)$$

Efektívne portfólia, ktoré vzniknú kombináciou bezrizikového aktíva a mean - variance portfólia rizikových aktív ležia na priamke, ktorá prechádza bodmi $[R_f, 0]$ a $[\mu_T, \sigma_T]$ a má maximálny sklon zo všetkých priamok, ktoré sú kombináciou bezrizikového aktíva a mean-variance portfólia. Priamka určená danými bodmi je priamka kapitálového trhu a portfólio T je dotýčnicové portfólio.

Obrázok 2: Efektívna hranica a bezrizikové aktívum



Sklon tejto priamky charakterizuje Sharpovú mieru portfólia v tvare

$$sr_p = \frac{\mu_p - R_f}{\sigma_p}. \quad (8)$$

Sharpova miera portfólia definuje dodatočný očakávaný výnos na jednotku rizika. Očakávaný dodatočný výnos na jednotku portfólia je vhodný ako východisko pre ekonomickú interpretáciu testov CAPM modelu. Testovanie efektívnosti portfólia v priestore mean-variance je ekvivalentné testovaniu, či Sharpova miera portfólia (8) je maximálna zo všetkých prípustných portfólií.

4. Model oceňovania kapitálových aktív

Základy CAPM modelu položil Markowitz (1959), ktorý v svojej pôvodnej práci argumentoval, že investori držia efektívne portfólio v priestore mean-variance, to znamená portfólio s najvyšším očakávaným výnosom pri danej úrovni rizika.

Sharpe (1964) a Lintner (1965) ukázali, že ak investori majú homogénne očakávania a každý z nich drží optimálne mean-variance efektívne portfólio, pri neexistencii trhových obmedzení, portfólio celého investovaného bohatstva alebo trhové portfólio bude tiež efektívne portfólio v priestore mean-variance. Sharpe a Lintner odvodili CAPM za predpokladu existencie bezrizikovej zapožičiavacej a požičiavacej sadzby. V tejto verzii CAPM pre očakávaný výnos i -tého aktíva platí:

$$E[R_i] = R_f + \beta_{im}(E[R_m] - R_f), \quad (9)$$

$$\beta_{im} = \frac{\text{Cov}[R_i, R_m]}{D[R_m]}, \quad (10)$$

kde R_m je výnos trhového portfólia,
 R_f je výnos bezrizikového aktíva.

Sharpova-Lintnerova verzia CAPM modelu môže byť vysvetlená cez pojem *dobatočný výnos*. Nech γ_i je *dobatočný výnos* i -tého aktíva nad výnosom bezrizikového aktíva, $\gamma_i = R_i - R_f$. Potom pre CAPM model platí

$$E[\gamma_i] = \beta_{im}E[\gamma_m], \quad (11)$$

$$\beta_{im} = \frac{\text{Cov}[\gamma_i, \gamma_m]}{D[\gamma_m]}, \quad (12)$$

kde γ_m je *dobatočný výnos* trhového portfólia.

Empirické testy Sharpe-Lintnerovho CAPM sú zamerané na tri implikácie modelu (11):

1. Intercept (úrovňová konštanta) je rovná nule.
2. Beta úplne vyčerpáva variabilitu očakávaných *dobatočných výnosov*.
3. Očakávaná *trhová riziková prémie* (očakávaný *dobatočný výnos* trhu $E[\gamma_m]$) je kladná.

V našom príspevku budeme testovať prvú implikáciu, že úrovňová konštanta je rovná nule.

Black (1972) odvodil zovšeobecnenú verziu CAPM modelu za predpokladu neexistencie bezrizikovej úrokovej sadzby. V tejto verzii *dobatočný očakávaný výnos* aktíva i nad výnosom *zero-beta* portfólia je v lineárnom vzťahu k jeho beta.

$$E[R_i] = E[R_{0m}] + \beta_{im}(E[R_m] - E[R_{0m}]), \quad (13)$$

$$\beta_{im} = \frac{\text{Cov}[R_i, R_m]}{D[R_m]}, \quad (14)$$

kde R_m je výnos trhového portfólia,

R_{0m} je výnos *zero-beta* portfólia, ktoré je ortogonálne k trhovému portfóliu m . *Zero-beta* portfólio je definované ako portfólio, ktoré má minimálny rozptyl zo všetkých portfólií nekorelovaných s trhovým portfóliom m .

Ekonometrická analýza Blackovej verzie CAPM považuje výnos *zero-beta* portfólia ako nemerateľnú veličinu. Blackova verzia môže byť testovaná ako obmedzenie na reálnom modeli výnosov trhu. Pre reálny výnos modelu trhu platí

$$E[R_i] = \alpha_i + \beta_{im}E[R_m], \quad (15)$$

kde podľa (13)

$$\alpha_i = E[R_{0m}](1 - \beta_{im}). \quad (16)$$

CAPM v tvare (11) a (15) je model pre jedno obdobie, nemá časovú dimenziu. Pre ekonometrickú analýzu je nevyhnutné pridať predpoklad, ktorý zohľadňuje vývoj výnosov v čase a odhadnúť model v čase.

5. Ekonometrický model Sharpe- Lintnerovej verzie CAPM.

Predpokladáme že výnosy sú nezávislé a rovnako rozdelené v čase a majú združené N -rozmerné normálne rozdelenie. Tieto predpoklady aplikujeme na *dobatočné výnosy* pre testovanie Sharpovej-Lintnerovej verzie CAPM. Ďalej predpokladajme, že investori majú možnosť zapožičiavať a vypožičiavať za bezrizikovú výnosovú mieru. Nech γ je $(N \times 1)$ vektor *dobatočných výnosov* N aktív v portfóliu. Pre týchto N aktív môže byť *dobatočný výnos* vyjadrený trhovým modelom v tvare

$$\gamma_t = \alpha + \beta \gamma_{mt} + \varepsilon_t, \quad (16)$$

$$E[\varepsilon_t] = 0, E[\varepsilon_t \varepsilon_t^T] = \Sigma, E[\gamma_{mt}] = \mu_m, E[(\gamma_{mt} - \mu_m)^2] = \sigma_m^2, \text{Cov}[\gamma_{mt}, \varepsilon_t] = 0, \quad (17)$$

kde β je $(N \times 1)$ vektor koeficientov beta, γ_{mt} je dodatočný výnos trhového portfólia za obdobie t , α je $(N \times 1)$ vektor koeficientov alfa a ε_t je $(N \times 1)$ vektor náhodných chýb.

Dôsledkom verzie CAPM (8) je, že všetky prvky vektora α sú nulové. Z toho vyplývajú základné hypotézy pre testovanie tohto modelu. Ak všetky prvky vektora α sú nulové, potom portfólio m je dotyčnicové portfólio.

Na odhad parametrov modelu (16) použijeme metódu maximálnej virohodnosti. Rozdelenie dodatočných výnosov aktív je podmienené dodatočnými výnosmi trhu. Za predpokladu, že dodatočné výnosy γ_t majú združené normálne rozdelenie, potom pre ich funkciu hustoty rozdelenia pravdepodobnosti platí

$$f(\gamma_t | \gamma_{mt}) = (2\pi)^{-N/2} |\Sigma|^{-1/2} e^{\left[-\frac{1}{2} (\gamma_t - \alpha - \beta \gamma_{mt})^T \Sigma^{-1} (\gamma_t - \alpha - \beta \gamma_{mt}) \right]}. \quad (18a)$$

Keďže dodatočné výnosy sú v čase IID, potom pre T pozorovaní je združená pravdepodobnostná funkcia hustoty

$$f(\gamma_1, \gamma_2, \dots, \gamma_T | \gamma_{m1}, \gamma_{m2}, \dots, \gamma_{mT}) = \prod_{t=1}^T (2\pi)^{-N/2} |\Sigma|^{-1/2} e^{\left[-\frac{1}{2} (\gamma_t - \alpha - \beta \gamma_{mt})^T \Sigma^{-1} (\gamma_t - \alpha - \beta \gamma_{mt}) \right]}. \quad (18b)$$

Odhad parametrov získame, ak parciálne derivácie logaritmu združenej funkcie hustoty (18b) podľa α , β , Σ položíme rovné 0. Maximálny virohodný odhad pre parametre a ich podmienené rozdelenia

$$\hat{\alpha} = \bar{\gamma} - \hat{\beta} \bar{\gamma}_m, \quad \hat{\alpha} \sim \mathcal{N}_N \left(\alpha; \frac{1}{T} \left[1 + \frac{\bar{\gamma}_m^2}{s_m^2} \right] \Sigma \right) \quad (19)$$

$$\hat{\beta} = \frac{\frac{1}{T} \sum_{t=1}^{t=T} (\gamma_t - \bar{\gamma})(\gamma_{mt} - \bar{\gamma}_m)}{s_m^2}, \quad \hat{\beta} \sim \mathcal{N}_N \left(\beta; \frac{1}{T} \left[\frac{1}{s_m^2} \right] \Sigma \right) \quad (20)$$

$$\hat{\Sigma} = \frac{1}{T} \sum_{t=1}^{t=T} (\gamma_t - \hat{\alpha} - \hat{\beta} \gamma_{mt})^T (\gamma_t - \hat{\alpha} - \hat{\beta} \gamma_{mt}), \quad T\hat{\Sigma} \sim \mathcal{W}_N(T-2; \Sigma) \quad (21)$$

kde
$$\bar{\gamma} = \frac{1}{T} \sum_{t=1}^{t=T} \gamma_t, \quad \bar{\gamma}_m = \frac{1}{T} \sum_{t=1}^{t=T} \gamma_{mt}, \quad s_m^2 = \frac{1}{T} \sum_{t=1}^{t=T} (\gamma_{mt} - \bar{\gamma}_m)^2. \quad (22)$$

Pre testovanie nulovej hypotézy $H_0: \alpha = 0$ oproti alternative $H_1: \alpha \neq 0$ použijeme Waldovu testovaciu štatistiku a MacKinley testovaciu štatistiku.

Waldova testovacia štatistika

$$W = \hat{\alpha}^T [\text{cov} \hat{\alpha}]^{-1} \hat{\alpha} = \frac{T s_m^2}{s_m^2 + \bar{\gamma}_m^2} \hat{\alpha}^T \Sigma^{-1} \hat{\alpha} \quad (23)$$

má za predpokladu platnosti nulovej hypotézy asymptoticky chí-kvadrát rozdelenie s N stupňami voľnosti.

MacKinlay-Gibbons-Ross-Shanken (1989) odvodili testovaciu štatistiku pre konečný výber

$$MK = \frac{T - N - 1}{N} \frac{s_m^2}{s_m^2 + \bar{\gamma}_m^2} \hat{\alpha}^T \hat{\Sigma}^{-1} \hat{\alpha}, \quad (24)$$

ktorá ma za predpokladu platnosti nulovej hypotézy F-rozdelenie s počtom stupňov voľnosti N ; $(T-N-1)$.

Gibbons, Ross a Shanken dokázali, že kvadratická forma $\alpha^T \Sigma^{-1} \alpha$ môže byť vyjadrená pomocou Sharpovej miery portfólia

$$\alpha^T \Sigma^{-1} \alpha = \frac{\bar{\gamma}_q^2}{s_q^2} - \frac{\bar{\gamma}_m^2}{s_m^2}, \quad (25)$$

kde q je *ex post* dotyčnicové portfólio a m je trhové portfólio. Potom testovacia štatistika sa dá napísať v tvare

$$MK = \frac{T - N - 1}{N} \left(\frac{\frac{\bar{\gamma}_q^2}{s_q^2} - \frac{\bar{\gamma}_m^2}{s_m^2}}{1 + \frac{\bar{\gamma}_m^2}{s_m^2}} \right). \quad (26)$$

Podľa (26) portfólio s maximálnou hodnotou Sharpovej miery zo všetkých prípustných portfólií je dotyčnicové portfólio. To znamená, že ak *ex post* trhové portfólio je dotyčnicové portfólio, potom hodnota MK bude rovná nule, čo vedie k záveru, že trhové portfólio je efektívne portfólio.

6. Dáta

Uvedené testy sme aplikovali na skupine akcií 10 amerických spoločností, ktoré sú súčasťou indexu S&P 500. K dispozícii sme mali denné ceny akcií a denné hodnoty indexu S&P 500, z ktorých sme vypočítali mesačné logaritmické výnosy a analizovali sme ich. Trhové výnosy boli nahradené výnosmi akciového indexu S&P 500. Bezriziková úroková miera bola reprezentovaná časovým radom 6-mesačných úrokových sadzieb LIBOR USD p. a. Časový rad výnosov zahŕňal 240 mesačných pozorovaní v priebehu 20 rokov od 29. januára 1988 do 31. decembra 2007.

Časový rad 240 pozorovaní sme rozdelili na štyri neprekrývajúce sa čiastkové obdobia po 5-rokoch (60 mesačných pozorovaní v každom období) a na dve 10-ročné čiastkové obdobia (120 mesačných pozorovaní). Odhady parametrov α , β a použitie vyššie popísaných testov sme aplikovali osobitne pre 5-ročné obdobia, potom pre 10-ročné obdobia a pre celé 20-ročné obdobie. Výsledky testov sú v tabuľke 1.

Výsledky testov sú pre jednotlivé čiastkové obdobia rozdielne. Vidíme, že v 5 ročnom období 1/1998 – 12/2002 na základe oboch testov nezamietame nulovú hypotézu a vzhľadom na použité dáta sa potvrdila platnosť Sharpovej-Lintnerovej verzie CAPM modelu. Keďže dané 5 ročné obdobie je súčasťou desaťročného obdobia 1/1998- 12/2007, tento výsledok sa potvrdil aj v tomto čiastkovom období. V ostatných obdobiach sme na 5 %-nej hladine významnosti zamietli nulovú hypotézu, čo vedie k záveru, že sa nepotvrdila platnosť Sharpovej-Lintnerovej verzie CAPM a vzhľadom na dané dáta trhové portfólio nie je efektívne portfólio.

Tabuľka 1: Výsledky testovania validity CAPM

Obdobie	počet pozorovaní	Wald test	p-value	MacKinlay test	p-value
5 ročné obdobie					
1/1988-12/1992	60	24,203	0,007	1,977	0,057
1/1993-12/1997	60	20,346	0,026	1,662	0,117
1/1998-12/2002	60	4,548	0,919	0,371	0,953
1/2003-12/2007	60	24,158	0,007	1,973	0,057
10 ročné obdobie					
1/1988-12/1997	120	30,42	0,0007	2,763	0,004
1/1998-12/2007	120	13,812	0,182	1,255	0,265
20 ročné obdobie					
1/1988-12/2007	240	22,361	0,013	2,133	0,023

7. Záver

Testy, ktoré sme použili na overenie platnosti Sharpe-Lintnerovej verzie CAPM modelu sú založené na predpoklade, že výnosy aktív majú združené N -rozmerné rozdelenie. Pri testovaní združenej normality našich dát sme použili testy založené na viacrozmerných charakteristikách šikmosti a špicatosti (Mardinove testy z roku 1970). Na základe týchto testov sa ukázalo, že rozdelenie výnosov je symetrické, ale špicatejšie ako viacrozmerné normálne rozdelenie. Takéto rozdelenie je typické pre rozdelenie trhových výnosov a je v súlade s inými empirickými štúdiami. Porušenie predpokladu o združenej normalite výnosov môže byť jednou z príčin, prečo sa nepotvrdila platnosť CAPM.

Literatúra

- ANDĚL, Jiří 2005. *Základy matematické statistiky*. Praha : Matfyzpress, 2005. 358 s. ISBN 80-86732-40-1.
- BENNINGA, Simon 2000. *Financial modeling*. With a section on Visual Basic for Applications by Benajmin Czaczkes. Second Edition. Cambridge / London : The MIT Press, 2000. 622 s. ISBN 0-262-02482-9.
- CAMPBELL, John Y., LO, Andrew W., MACKINLAY, A. Craig 1997. *The econometrics of financial markets*. Second printing, with corrections. Princeton [New Jersey, USA] : Princeton University Press, 1997. 611 s. ISBN 0-691-04301-9.
- ELTON, Edwin J., GRUBER, Martin J. et al. *Modern portfolio theory and investment analysis*. Seventh edition. Hoboken [New Jersey, USA] : Wiley, 2007. 728 pp. ISBN 0-470-05082-9.
- MLYNAROVÍČ, Vladimír 2001. *Finančné investovanie. Teórie a aplikácie*. Bratislava: Iura Edition, 2001. 293 s. ISBN 80-89047-16-5.
- POTOCKÝ, Rastislav, LAMOŠ, František 1989. *Pravdepodobnosť a matematická štatistika*. Bratislava : Alfa, 1989. 344 s. ISBN 80-7184-767-4.
- RAO, C. Radhakrishna, 2002. *Linear statistical inference and its applications*. 1973 (second) edition. New York : Wiley, 2002. 625 s. ISBN 0-471-21875-8.

Adresa autorov

Martin Bod'a, Ing. Et Bc., Mária Kanderová, Ing., PhD.
 Katedra kvantitatívnych metód a informatiky
 Ekonomická fakulta, UMB
 Tajovského 10
 975 90 Banská Bystrica
 martin.boda.@umb.sk; maria.kanderova@umb.sk

Branch structure of employment in particular regions of the Czech Republic

Hana Boháčová, Pavla Jindrová, Jana Heckenbergerová

Abstract: The aim of this paper is to analyse the branch structure of employment in the Czech Republic regions according to the Industrial Classification of Economics Activities. The branch structure resemblance between the regions during 1993 - 2006 was studied and an algorithm of the future resemblance forecast was proposed.

Key words: Industrial classification of economic activities, Czech Republic regions, cluster analysis.

1. Industrial classification of economic activities

The industrial classification of economic activities focuses on all the working activities performed by economic subjects in a given region.

There was the OKEČ classification used in Czech Republic during the period 1994 to 2007. Since January 2008 CZ-NACE classification is being used. The similarities and differences in the branch structure of employment in particular regions of the Czech Republic in years 1993 – 2006 are investigated in this paper in so far that the OKEČ classification is adequate for our purpose. Following OKEČ categories are involved here:

A 01	Agriculture, hunting and related activities,
A 02, B	Forestry, fishing, fish farming and related activities,
C	Mineral raw materials mining,
D	Manufacturing industry,
E	Production and distribution of electricity, gas and water,
F	Building industries,
G	Business, car and consumer items repair,
H	Accommodation and boarding,
I	Transfer, stocking and communication,
J	Factoring,
K	Real properties and lease, related entrepreneurial activities,
L	Public administration and defence, obligatory social security,
M	Education,
N	Health care and social-service work, veterinary activities,
O	Other public, social and personal services.

2. Resemblance between the regions of the Czech Republic in view of the branch structure of employment

The task now is to find out the resemblance and discrepancies in the branch structure of employment between particular regions of the Czech Republic in the period of 1993 to 2006.

The time series of the numbers of employees in separate categories listed above per a NUTS 3 regions in mentioned period published on the web of the Czech Statistical Office (cf. [3]) are used. Let us choose the time series from the beginning, the middle and the end of the period under consideration – it means from 1993, 2000 and 2006. Cluster analysis is applied on these data - average distance method and Euclidean norm are used, the clustering indicators are the numbers of employees and classified objects are particular regions (cf. [1] or [2]). Unistat 5.6 was used.

Let us have a look at the dendrogram in figure 1. It is the output of cluster analysis for the year 1993. Particular regions are denoted as follows:

- 1 Královéhradecký,
- 2 Jihočeský,
- 3 Jihomoravský,
- 4 Karlovarský,
- 5 Liberecký,
- 6 Moravskoslezský,
- 7 Olomoucký,
- 8 Pardubický,
- 9 hl. m. Praha,
- 10 Plzeňský,
- 11 Středočeský,
- 12 Ústecký,
- 13 Vysočina,
- 14 Zlínský.

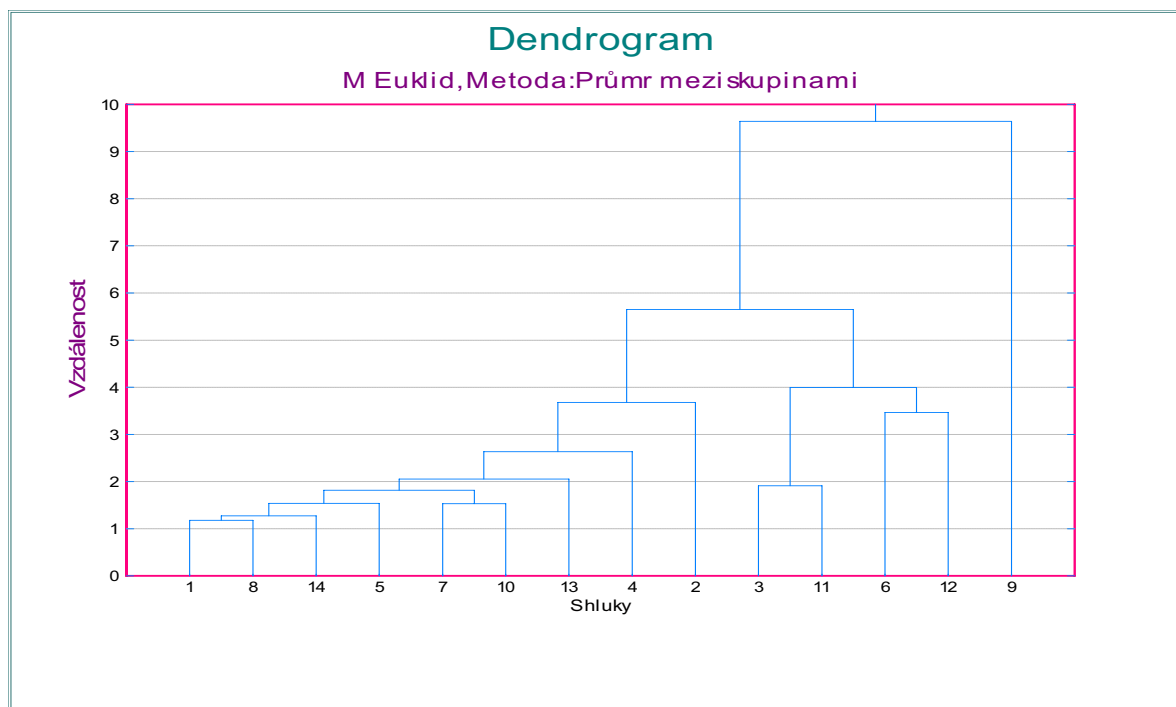


Figure 1: Dendrogram – cluster analysis, branch structure of employment in the Czech Republic, 1993.

We can see the biggest resemblance between Pardubický and Královéhradecký region and resemblance between Olomoucký and Plzeňský region at the first sight. Next we can notice that the regions are splitted into three groups – the first group is formed by Pardubický, Královéhradecký, Zlínský, Liberecký, Olomoucký, Plzeňský, Karlovarský, Jihočeský and Vysočina region. The second group consists of Jihomoravský, Středočeský, Moravskoslezský and Ústecký region. The capital city of Prague is in the third group separately.

We will try to explain the reason why these three groups are formed a little bit later. Let us now see the result of cluster analysis for year 2000 in figure 2.

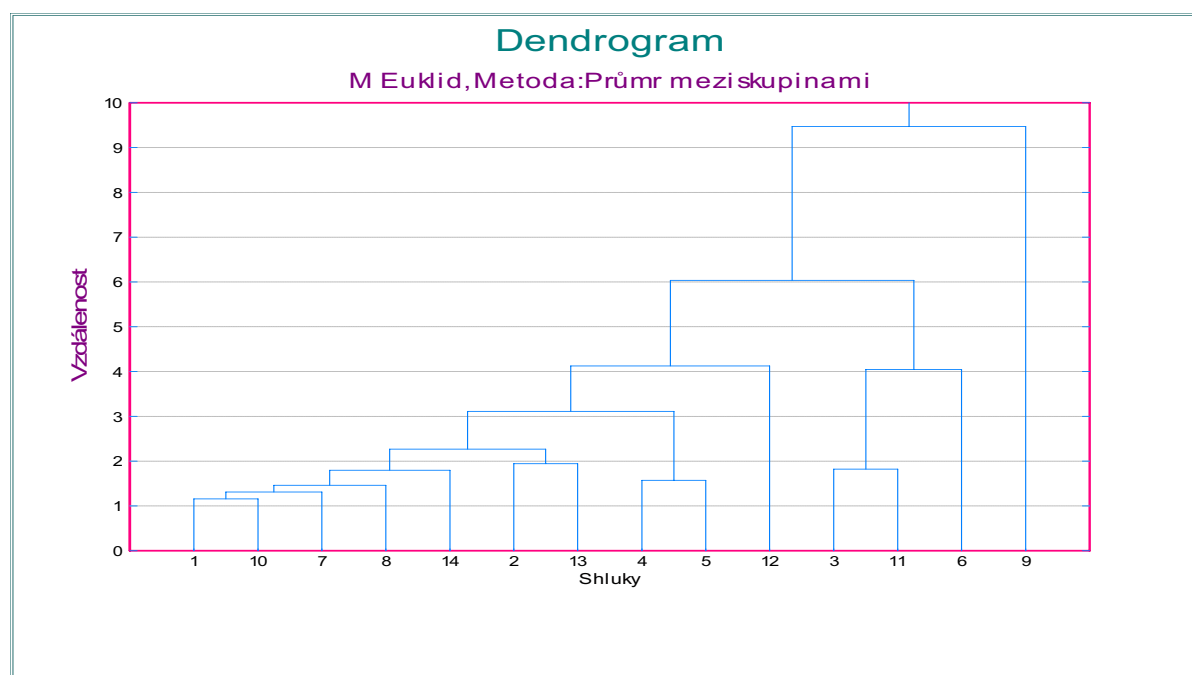


Figure 2: Dendrogram – cluster analysis, branch structure of employment in the Czech Republic, 2000.

The most telling resemblance is between Královéhradecký and Plzeňský region in this case, then Olomoucký and Pardubický region come in within the next two steps. Figure 2 shows three groups of regions as well. The first group is formed by the above mentioned regions, Zlínský, Jihočeský, Karlovarský, Liberecký, Ústecký and Vysočina region. Jihomoravský, Středočeský and Moravskoslezský region are in the second group. The last group is the capital city of Prague on its own again.

In figure 3 we can find the dendrogram for the last year of the considered period – for the year 2006. The situation is quite similar to that in figure 1. The most significant resemblance can be seen between Královéhradecký and Pardubický region and then between Olomoucký and Plzeňský region. This figure has the three groups too. The first group consists of Královéhradecký, Pardubický, Zlínský, Vysočina, Olomoucký, Plzeňský, Jihočeský, Karlovarský and Liberecký region. Jihomoravský, Středočeský, Moravskoslezský and Ústecký region are in the second group. The third group is the capital city of Prague.

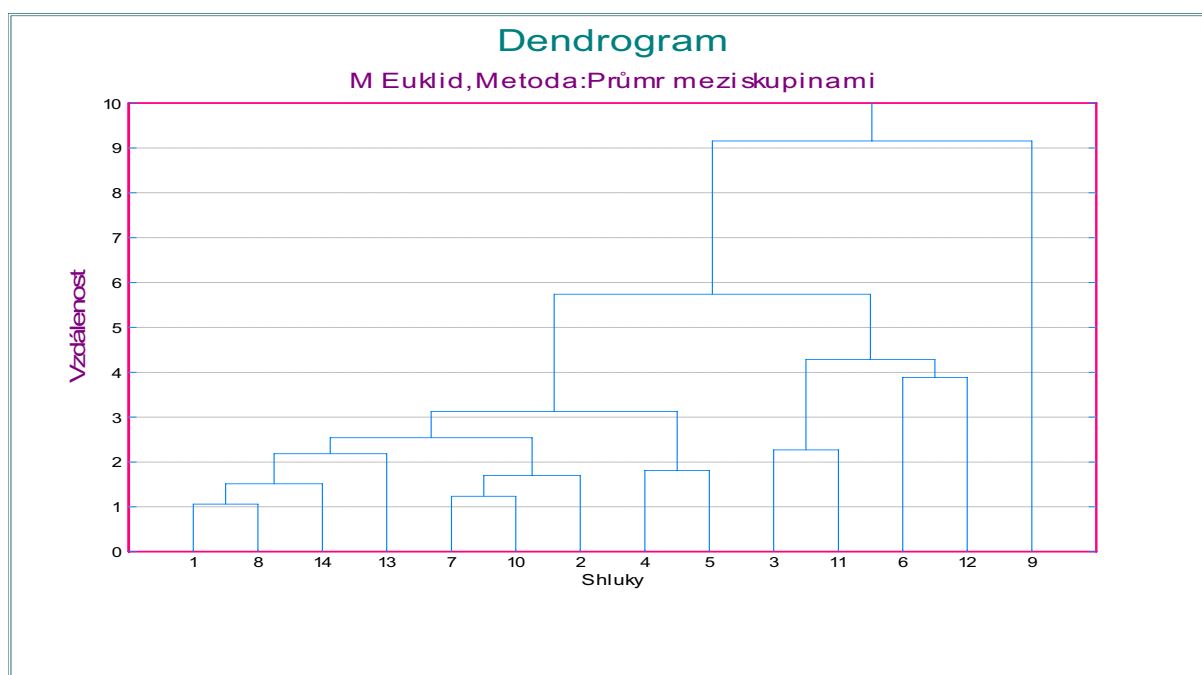


Figure 3: Dendrogram – cluster analysis, branch structure of employment in the Czech Republic, 2006.

When we look for the reason that causes the splitting of the regions in the dendrogram into the three groups we have to focus on the percentual distribution of the branches in each region. We can find that the main difference between the first and the second group is in the distribution of the D category (manufacturing industry). E.g. in 2006 30,5% of the economically active people in the whole Czech Republic were employed in the manufacturing industry. In the regions of the first group in figure 3 the percentual distribution of category D – Manufacturing industry was slightly higher then the republic one. In the regions from the second group it was a little lower. Next we can see that the capital city of Prague is much different from the other groups (which is of course exactly what we could expect). The main difference is that the capital city has markedly lower distribution of categories D – Manufacturing industry, A 01 – Agriculture, hunting and related activities, A 02 and B - Forestry, fishing fish farming and related activities and higher distribution of K - Real properties and lease, related entrepreneurial activities and O - Other public, social and personal services on the opposite side.

3. Resemblance estimate algorithm proposal

We have no objective reason to expect some expressive changes in the resemblance and/or differences in the branch structure of employment between the regions of the Czech Republic. But when we decide to estimate the future situation (or when we need to estimate the future cluster splitting for some other problem) we can fit the time series data with a suitable trend curve, make the prediction of the future values and then use these predicted values as the input of the cluster analysis.

4. References

- [1]ANDĚL, J. 1978. Mathematical Statistics. (in Czech), Praha: SNTL, 1978.
[2]KUBANOVÁ, J. 2004. Statistical Methods for Economical and Technical Practices.
(in Czech), Bratislava: Statis, 2004. ISBN 80-85659-37-9.
[3]www.czso.cz

Authors addresses:

Mgr. Hana Boháčová
Institute of Mathematics
Faculty of Economics and Administration
Studentská 84
532 10 Pardubice
hana.bohacova@upce.cz

Mgr. Pavla Jindrová
Institute of Mathematics
Faculty of Economics and Administration
Studentská 84
532 10 Pardubice
pavla.jindrova@upce.cz

Mgr. Jana Heckenbergerová
Institute of Information Technologies
Faculty of Electrical Engineering and Informatics
Studentská 95
532 10 Pardubice
jana.heckenbergerova@upce.cz

Vývoj rozvodovosti na Slovensku 1997-2007

Divorce Rate Progress in Slovakia from 1997 to 2007

Rastislav Čief

Abstract: Divorce is one of the forms of the marriage disintegration, based on a legal way to end the marriage. The divorce rate negatively affects the population reproduction and has a very negative effects on children's upbringing. In Slovakia, from 1997 to 2005, the number of divorces was growing, only in last two years it has recorded some decrease. The new trends have been noticed - increasing of the average duration of divorced marriages, increasing proportion of men suing for divorce and increasing proportion of divorces on divorce proceedings' results.

key words: Divorce, divorce rate, reason for divorce, duration of divorced marriages

kľúčové slová: rozvod, rozvodovosť, príčiny rozvratu manželstiev, dĺžka trvania rozvedených manželstiev

1. Úvod

Rodina je základnou bunkou spoločnosti. Na vytvorenie rodiny má kľúčový vplyv inštitúcia manželstva, ktorá je známa vo všetkých civilizáciách na svete. Vo vyspelých krajinách v 20-tom storočí nastala kríza rodiny, ktorá pokračuje aj v súčasnosti. Dôsledkom je masový rozpad manželstiev, prípadne v niektorých oblastiach aj minimálny vznik nových manželstiev, ktoré sú nahradzované partnerstvom bez sobáša. Tieto trendy zasahujú aj Slovensko a v tomto sa príspevku venujeme najnovší vývoj.

Rozvod je jedna z foriem rozpadu manželstva založená na zákonnom spôsobe ukončenia manželstva. Rozvodovosť negatívne ovplyvňuje reprodukciu obyvateľstva a má výrazne nepriaznivé dôsledky na výchovu detí (Mládek, J.,1992).

2. Vývoj rozvodovosti na Slovensku 1997-2007

V období 1997-2007 počet ukončených konaní neustále rástol z 11 838 v roku 1997 až na 14 346 v roku 2005, no v posledných dvoch rokoch zaznamenal menší pokles. Podobne aj počet rozvodov stúpал z 9 138 v roku 1997 až na 12 716 v roku 2006. Ostatné výsledky konania zaznamenali pokles: zamietnuté návrhy klesli zo 187 v roku 1997 na 87 v roku 2007, vzaté návrhy späť klesli z 1 946 na 653 a iné rozhodnutia z 566 na 133 za rovnaké obdobie.

Stúpajúci trend mali rozvedené manželstva s maloletými deťmi, kým v roku 1997 ich bolo 6 670 v roku 2006 až 8 474, popritom, ale poklesol priemer počtu maloletých detí v rozvedenom manželstve na 1,5 a potvrdilo sa, že manželstva s 3 a viac deťmi sa ukázali stabilnejšie, keďže ich počty v danom období stagnovali (ročne okolo 540 rozvodov).

Veľmi výrazné zmeny vidieť v dĺžke trvania rozvedeného manželstva, kde vidieť trend každoročne sa zvyšujúceho priemeru z 11,8 rokov v 1997 na 14,2 rokov v 2007 a zdá sa, že tento trend ešte môže pokračovať. Znamená to, že sa čoraz častejšie rozvádajú aj staršie manželstvá, ktoré by už mali mať prekonané kritické obdobia. U krátko trvajúcich manželstiev vidno pokles, prípadne stagnáciu počtu rozvodov (s výnimkou dvojročných manželstiev). Naproti tomu v kategórii 9 a viacročných manželstiev stúpol počet rozvodov s 5 212 v roku 1997 na 8 344 v roku 2007 (pritom v roku 2006 bol ešte vyšší 8 621 rozvodov!).

Zaujímavú zmenu vidno v podávaní návrhu dlhé desaťročia výrazne dominujú v podávaní návrhov ženy (69% v roku 1997), ale v súčasnosti sa ich dominancia začína

znižovať (65% v roku 2006). Zdá sa, že kým v minulosti boli so svojim manželstvom oveľa viac nespokojné ženy v súčasnosti stúpa počet mužov ktorí sú nespokojný v danom manželstve a dávajú návrh na rozvod.

Od roku 1963 súd pri rozvode povinne určuje príčinu rozvratu manželstva. Momentálne klesá kvalita tohto určovania, čo je vidieť v tom, že sa výrazne zvyšuje kategória 6-rozdielnosť pováh, názorov a záujmov, za ktorú sa zrejme skrývajú aj iné príčiny, ktoré je pred súdom neprijemné rozoberať (z 4 640 v roku 1997 stúpila na 7 796! v roku 2006). Stúpili aj kategórie 3-nevera (u mužov o 38% u žien len o 18%) a kategória 9-ostatné príčiny (u mužov o polovicu a u žien takmer na dvojnásobok). Zaujímavo sa vyvíjala kategória 2-alkoholizmus, kde u mužov mierne poklesla (z 1109 v roku 1997 na 1064 v roku 2007) a naopak u žien zaznamenala nárast (z 102 v roku 1997 na 125 v roku 2007).

Ostatné kategórie stagnovali alebo dlhodobo klesali, logický je pokles kategórie 1-neuvážené uzavretie manželstva jednak vzhľadom na nižšiu sobášnosť a jednak vzhľadom na zvyšovanie veku rozvádzaných manželstiev. Najvýraznejší pokles zaznamenala kategória 8-sexuálne nezhody, ktorá klesla 104 v roku 1997 na 47 v roku 2007.

Dlhodobo vidieť pokles „nevinných“ žien, čiže pokles v kategórii 0-nezistené zavinenie, ktoré v roku 1985 bolo až 46,5% v roku 1997 bolo 22% a v roku 2007 bolo už len necelých 15%. Muži v tejto kategórii stagnovali na okolo 300 prípadov (necelé 3%).

Tab.č.1: Rozvody na Slovensku v rokoch 1997-2007

Územie	Počet ukončených konaní	z nich podal návrh		Výsledok konania				
		muž	žena	manžel. rozvedené	návrh zamietnutý	návrh vzatý späť	iné rozhodnutie	manželstvo vyhlásené za neplat.
1997	11 838	3 675	8 152	9 138	187	1 946	566	1
1998	12 116	3 858	8 244	9 312	211	2 075	518	0
1999	12 457	3 858	8 587	9 664	186	1 937	669	1
2000	12 027	3 777	8 224	9 273	125	1 931	696	2
2001	12 443	4 066	8 360	9 817	147	1 787	692	0
2002	13 752	4 544	9 182	10 960	116	1 842	834	0
2003	13 606	4 411	9 167	10 716	120	1 812	957	1
2004	13 857	4 663	9 176	10 889	132	1 822	1 008	6
2005	14 346	4 914	9 416	11 553	128	1 817	843	5
2006	14 007	4 816	9 164	12 716	97	981	207	6
2007	13 048	4 454	8 593	12 174	87	653	133	1
2001-2007	13 580	4 553	9 008	11 261	118	1 531	668	3

zdroj: Základné demografické údaje Prešovského kraja. ŠÚ SR, Prešov, 1997-2007

Tab.č.2: Rozvody manželstiev s maloletými deťmi

Z rozvedených manželstiev					
manželstvá s maloletými deťmi					priemerná dĺžka trvania manž.
spolu	počet maloletých detí				
	1	2	3+	priemer	
6670	3611	2498	561	1,6	11,8
6 755	3 730	2 454	571	1,6	12,1
6 836	3 855	2 415	566	1,5	12,3
6 514	3 771	2 235	508	1,5	12,7

6 880	3 965	2 375	540	1,5	13,1
7 691	4 464	2 618	609	1,5	13,1
7 470	4 391	2 543	536	1,5	13,3
7 329	4 297	2 536	496	1,5	13,6
7 609	4 469	2 651	489	1,5	13,9
8 474	4 968	2 945	561	1,5	14,2
7 994	4 684	2 738	572	1,5	14,2
7 635	4 463	2 629	543	2	14

zdroj: Základné demografické údaje Prešovského kraja. ŠÚ SR, Prešov, 1997-2007

Tab.č.3: Rozvody podľa dĺžky trvania manželstva

roky	Dĺžka trvania rozvedeného manželstva (v rokoch)								
	do 1	2	3	4	5	6	7	8	9 a viac
1997	354	424	502	544	563	535	547	457	5212
1998	346	391	483	492	580	502	536	499	5483
1999	349	449	475	520	522	528	498	507	5816
2000	319	392	471	472	479	461	483	492	5 704
2001	134	450	526	465	515	457	459	478	6333
2002	332	455	528	584	549	511	500	467	7 034
2003	323	382	496	494	549	498	512	452	7 010
2004	310	390	459	504	597	518	514	446	7 151
2005	332	391	448	534	546	506	521	488	7 787
2006	311	454	502	548	538	559	610	573	8 621
2007	318	428	539	524	502	480	534	505	8 344
2001-2007	294	421	500	522	542	504	521	487	7 469

zdroj: Základné demografické údaje Prešovského kraja. ŠÚ SR, Prešov, 1997-2007

Tab.č.4: Rozvody podľa príčiny rozvratu manželstva a trvalého pobytu manželov v SR

rok	príčina rozvratu manželstva (kód)									
	na strane muža									
	1	2	3	4	5	6	7	8	9	0
1997	357	1109	971	858	178	4640	60	104	550	311
1998	376	1069	987	847	191	4962	58	91	445	286
1999	379	1159	1052	1006	202	4947	50	87	473	309
2000	331	1100	946	893	179	4883	41	84	493	323
2001	296	1043	998	865	154	5586	49	68	472	286
2002	367	1159	1125	884	170	6235	40	81	581	318
2003	375	1117	1118	708	191	6187	24	62	647	287
2004	267	1168	1251	617	202	6334	40	63	647	300
2005	245	1200	1373	628	194	6844	28	71	696	274
2006	281	1255	1292	653	225	7796	19	57	833	305
2007	232	1064	1341	532	198	7519	36	47	883	322
1997-2007	3506	12443	12454	8491	2084	65933	445	815	6720	3321

rok	príčina rozvratu manželstva (kód)									
	na strane ženy									
	1	2	3	4	5	6	7	8	9	0
1997	357	102	662	385	19	4640	59	104	780	2030
1998	376	108	661	358	19	4962	59	91	653	2025
1999	379	86	591	429	28	4947	73	87	685	2359
2000	331	112	625	433	11	4883	49	84	650	2095
2001	296	82	531	369	4	5586	50	68	735	2096
2002	367	86	604	369	14	6235	51	81	859	2294
2003	375	112	660	278	20	6187	46	62	1095	1881
2004	267	104	747	235	16	6334	42	63	1149	1932
2005	245	128	720	251	24	6844	46	71	1218	2006
2006	281	149	767	268	19	7796	38	57	1391	1950
2007	232	125	781	249	19	7519	50	47	1344	1808
1997-2007	3506	1194	7349	3624	193	65933	563	815	10559	22476

zdroj: Základné demografické údaje Prešovského kraja. ŠÚ SR, Prešov, 1997-2007

kód- príčina rozvratu:

- 1-neuvážené uzavretie manželstva
- 2-alkoholizmus
- 3-nevera
- 4-nezáujem o rodinu, vrátane ukončenia spolužitia
- 5-zle zaobchádzanie, odsúdenie pre trestný čin
- 6-rozdielnosť pováh, názorov a záujmov
- 7-zdravotné dôvody, vrátane neplodnosti
- 8-sexuálne nezhody
- 9-ostatné príčiny
- 0-súd nezistil zavinenie

3. Záver

Na Slovensku v období 1997 až 2005 počet rozvodov rástol až v posledných dvoch rokoch zaznamenal menší pokles. Z nových trendov sú najmä neustále zvyšovanie priemernej dĺžky trvania rozvedených manželstiev a zvýšenie podielu rozvodov z výsledkov konania, keďže sa výrazne znížili zamietnuté a stiahnuté návrhy.

Kým v minulosti výrazným modelom bolo, že nespokojná manželka dávala, návrh na rozvod a súd ju v polovici prípadoch uznal nevinou, v súčasnosti sa tento model narúša zvyšujúcim sa počtom podaní návrhov mužov a klesajúcim podielom žien u ktorých súd nezistil zavinenie. Zdá sa, že tento trend je dôsledkom zvyšujúcej sa emancipácie žien. Ďalším novým trendom je hľadanie si mladších partneriek u mužov, čo môže mať za dôsledok nárast rozvodov dlhotrvajúcich manželstiev.

Celkovo možno povedať, že Slovensko napodobňuje v rozvodovosti celosvetové trendy a v budúcnosti zrejme nemožno očakávať výrazné zlepšenie.

4. Literatúra

- [1] MLADEK, J. 1992. Základy geografie obyvateľstva. SPN, Bratislava, 1992.
- [2] Základné demografické údaje Prešovského kraja 1997. ŠÚ SR, KS Prešov, 1998.
- [3] Základné demografické údaje Prešovského kraja 1998. ŠÚ SR, KS Prešov, 1999.
- [4] Základné demografické údaje Prešovského kraja 1999. ŠÚ SR, KS Prešov, 2000.

- [5] Základné demografické údaje Prešovského kraja 2000. ŠÚ SR, KS Prešov, 2001.
- [6] Základné demografické údaje Prešovského kraja 2001. ŠÚ SR, KS Prešov, 2002.
- [7] Základné demografické údaje Prešovského kraja 2002. ŠÚ SR, KS Prešov, 2003.
- [8] Základné demografické údaje Prešovského kraja 2003. ŠÚ SR, KS Prešov, 2004.
- [9] Základné demografické údaje Prešovského kraja 2004. ŠÚ SR, KS Prešov, 2005.
- [10] Základné demografické údaje Prešovského kraja 2005. ŠÚ SR, KS Prešov, 2006.
- [11] Základné demografické údaje Prešovského kraja 2006. ŠÚ SR, KS Prešov, 2007.
- [12] Základné demografické údaje Prešovského kraja 2007. ŠÚ SR, KS Prešov, 2008.

Adresa Autora:

Rastislav Čief
Dobrianského 45
066 01 Humenné
cief@post.sk

Regionální demografie v ČSÚ - prezentace a analytické využívání

Eduard Durník, Jiří Kamenický

Regional Demography in CSO – presentations and analytical applications

Eduard Durník, Jiří Kamenický

Abstrakt: Český statistický úřad přechází stále častěji na využívání administrativních dat z externích zdrojů. Demografická statistika vychází z tradičních hlášení o jednotlivých demografických jevech, ale od roku 2005 se opírá i o údaje z centrální evidence obyvatel Ministerstva vnitra. Při zpracování výsledků se potýká s problematikou územních změn, ale také s častou změnou legislativy týkající se dat za cizince. Pro účely regionální statistiky se připravují přepočty několika set demografických ukazatelů podle aktuální příslušnosti do územních celků, čímž vzniká důležitý zdroj pro analýzy regionálních časových řad.

Prezentace demografických dat je zajišťována vydáváním klasických publikací, ale perspektivu má především webová prezentace demografických dat. K tomu slouží nejen specializované aplikace (umožňují mimo jiné prezentovat i územní změny), ale také univerzální aplikace (veřejná databáze) umožňující ad-hoc specifikaci uživatelů včetně možnosti přepínat zobrazení tabulka – graf – mapa.

Vedle popisných výstupů ročenkového typu nebo časových řad se věnuje velký prostor i analytickým publikacím. V nabídce ČSÚ najdeme jak komplexní analýzy, jejichž součástí jsou i kapitoly demografické analýzy (např. demografický pilíř v publikacích Demografický, sociální a ekonomický vývoj v kraji, analýza demografických jevů v publikaci Život cizinců v ČR) nebo specializované monotematické analýzy zaměřené výhradně na oblast demografie (např. Vysokoškoláci z demografického pohledu).

Klíčová slova: Organizace regionální a demografické statistiky – Datové zdroje – Standardní publikační výstupy – Webové prezentace – Demografické analýzy – Multikriteriální hodnocení

Key words: organisation of regional and demographic statistics – Data sources – Standard publications – Web presentations – Demographical analyses – multicriterial evaluations

1. Úvod

Demografická statistika zůstává i v současné době rodinným stříbrem většiny statistických úřadů, stejně tak tomu je i v Českém statistickém úřadě. Poměrně stabilizovaný systém zjišťování dat, vysoká spolehlivost údajů, atraktivnost výstupů, schopnost vyrovnat se i s požadavky lokální statistiky, to patří k výrazným pozitivům.

Tématem našeho příspěvku je poskytnutí základního obrázku, jak je VČSÚ organizovaná demografická a regionální statistika. Načtneme, jak zacházíme s výstupy demografické statistiky, abychom uspokojili dlouhodobě vysoký zájem uživatelů o relevantní data. Po vstupu Česka do Evropské unie se razantně zvyšuje význam výstupů regionální statistiky. Při přípravě programů rozvoje regionů, typování zaostalých regionů z různých hledisek, patří mezi významné pilíře data demografické statistiky. Chceme nastínit i přístupy regionálních analytiků ČSÚ, kteří se touto tematikou často zabývají, představit některé analytické práce.

Rovněž chceme dokumentovat stěžejní roli webu při prezentaci demografických a regionálních dat, ale také se zmíníme o standardních publikacích demografické statistiky.

2. Organizace regionální a demografické statistiky v ČSÚ, základní metodické přístupy

V ČSÚ neexistuje specializovaný útvar, který by měl na starosti regionální statistikou. Regionální statistikou se na ČSÚ zabývají:

- oddělení regionálních databází (správa a rozvoj regionálních databází, regionální analýzy, mezirezortní vazby zejména v pracovních skupinách pro udržitelný rozvoj, webové prezentace regionálních údajů)
- oddělení informačních služeb a regionálních analýz (14 krajských pracovišť – tvorba a správa krajských webových stránek ČSÚ, produkce tabulkových publikací, ale i monotematických či komplexnějších analýz (zpracovaných obvykle podle jednotné osnovy vytvořené v rámci pracovních skupin)
- specialisté na regionální účty – zabývají se aplikací různých metod odhadů regionálního hrubého domácího produktu a dalších ukazatelů

Demografickou statistikou se na ČSÚ zabývají:

- oddělení demografické statistiky (zpracování hlášení o demografických událostech a zpracování údajů centrální evidence obyvatel, produkce dat běžné demografické statistiky, analýzy populačního vývoje, demografické projekce)
- referát regionální demografie v Olomouci (funguje zhruba od roku 2003, zabývá se především přípravou tabulkových výstupů – časových řad za různé územní celky (od úrovně LAU1 ale i ad-hoc mikroregiony), územními přepočty, ale i demografickými analýzami, vede severomoravsko pobočku České demografické společnosti, která je živou diskusní platformou pro aktuální otázky populačního vývoje)
- samostatné oddělení specifických statistik obyvatelstva (gender statistika, cizinci)
- oddělení sčítání lidu (příprava SLDB 2011, metodika, organizace, ale i diseminace)

Základem zpracování demografických údajů zůstává v ČSÚ zpracování statistik jednotlivých hlášení o demografických událostech (narození, úmrtí, sňatek, rozvod). ČSÚ čerpá rovněž z registru o potratech ministerstva zdravotnictví. Významnou změnou byl v roce 2005 nový způsob zpracování migrace. Začaly se využívat data z centrální evidence obyvatel ministerstva vnitra (vnitřní a v současné době i zahraniční stěhování), bylo zrušeno hlášení o stěhování.

Data demografické statistiky významně doplňují populační censy, vláda ČR již schválila návrh zákona o SLDB 2011, nyní se jím bude zabývat Parlament. Sčítání by mělo proběhnout 26.3.2011. Metodickou otázkou zatím v řešení zůstává vztah standardní bilance stavu obyvatel a výsledků SLDB.

Česká demografická data za jednotlivá území jsou přepočítávána podle aktuálního územního začlenění obcí do vyšších územních celků. Sice již klesá počet změn v počtu obcí (rozdělování a slučování), ale významně se mění příslušnost obcí do okresů v souvislosti se změnou hranic okresů i krajů. V ČR stál ještě není schválen zákon o územním uspořádání země, mění se zejména struktura územních celků na úrovni LAU 1. Přepočty demografických údajů vycházejí z přesné evidence vazeb územní příslušnosti – tedy mezi číselníkem obcí a dalšími číselníky různých územních typů.

Do databází jsou ukládány informace jak poplatné územní příslušnosti v referenčním období, tak zpětné přepočty starších údajů s respektováním aktuální územní příslušnosti.

Regionální údaje včetně demografických se snažíme prezentovat na webu, v uživatelsky přátelském prostředí. Cílem je, aby uživatel mohl výběrem parametrů získat co nejrychleji data za požadované území. Snažíme se k tomu využívat i prostředků GIS technologií. Zatím nabízíme jen generování předdefinovaných tabulek v podobě jednoduchých kartogramů. V jednotlivých publikacích však využíváme funkcionality produktů firmy ESRI (ArcMap) a vyrábíme pestré mapy.

Připravuje se projekt, který by měl pomoci dokončit vybudování GIS databáze v úřadě a jejím propojení s budovaným centrálním datovým skladem. Rovněž je naším záměrem využít zkušeností vysokých škol, testujeme např. využívání mapových služeb (MASHUPS) při prezentaci dat demografické statistiky.

3. Regionální publikace s demografickými údaji

Standardní publikace regionální statistiky obsahující nejdůležitější demografické údaje za kraje a nižší územní celky (do úrovně obcí, v Praze – 57 městských částí). Jejich výhodou je možnost prezentace reprezentativního vzorku demografických dat v kontextu ostatních socioekonomických charakteristik zemních celků. V ústředí jsou každoročně vydávány publikace: *Okresy, Kraje, Malý lexikon obcí ČR*. Jednotlivá pracoviště ČSÚ v regionech prezentují za svá území ve stejné periodicitě demografické údaje v *krajských ročenkách*, v menším rozsahu ve čtvrtletní periodicitě v *bulletinech* (jejich součástí bývají 1-2x ročně i ad-hoc analýzy, které se často zabývají aktuálním demografickými problémy-migrace obyvatel, bytová výstavba).

K podrobnějším regionálním analýzám jsou určeny specializované demografické publikace. Základem je tzv. pramenné dílo – *Demografická ročenka ČR*. Jeho výhodou je dlouhodobá tradice, možnost komparace základních regionálně členěných (oblasti, kraje, okresy, velikostní skupiny obcí, od r. 2006 rozšířeny o pohled města-venkov) demografických údajů s celorepublikovými tendencemi, v územním členění jsou prezentovány i specializované analytické výstupy (podrobné úmrtnostní tabulky).

Specializované pracoviště ČSÚ v Olomouci připravuje každoročně obdobná *pramenná díla* zaměřené na *kraje, správní obvody obcí s rozšířenou působností a města* (s desetiletými časovými řadami, všechny údaje jsou při územním změnách zpětně přepočítány). Uživateli se zde otevírá možnost využití i analytičtějších ukazatelů populačního vývoje (průměrný věk při sňatku, rozvodu, podíl předmanželských koncepcí, úhrnná potratovost, naděje dožití ve vybraných letech aj.). Vedle toho se věnuje i rekonstrukci historických časových řad (*Přirozená měna obyvatelstva v zemích Koruny české v letech 1. světové války 1914 až 1918, Zemřelí podle podrobného seznamu příčin smrti, pohlaví a věku v ČR (1919 až 2006)*).

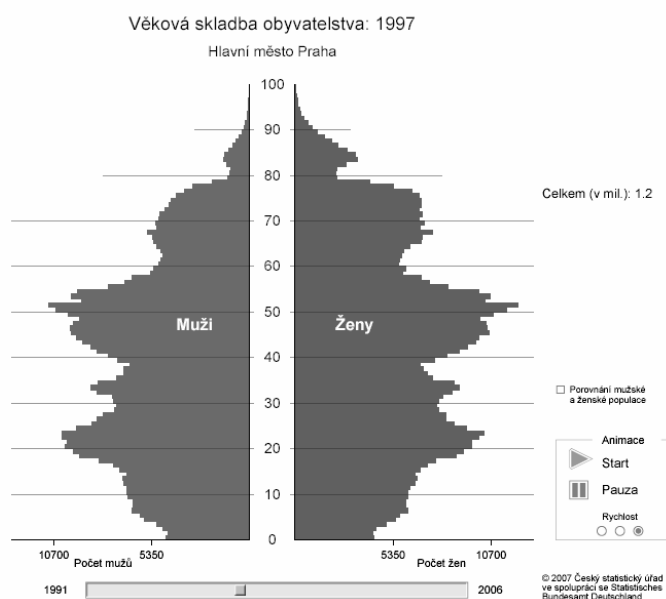
Samostatnou kapitolu představují výstupy ze SLDB (jejich popis je nad rámec tohoto příspěvku, pro připravovaný census v roce 2011 se počítá se zachování jejich podrobnosti v intencích roku 2001, podstatně by však měly být rozšířeny formy diseminace – veřejná databáze, specializované mapové výstupy). Pro regionální statistiku bylo povzbuzujícím počinem zpracování *Historického lexikonu obcí ČR (1869-2001)*, jehož význam je potvrzen faktem, že byl zpracován v době probíhající reformy veřejné správy. Změny v územní struktuře státu z toho plynoucí (mj. vznik nových územně-správních celků-205 obvodů obcí s rozšířenou působností) byly do připravovaného lexikonu promítnuty (lexikon je navíc vydáván každý rok v elektronické podobě – s přepočtem cenzálních dat na aktuální územní strukturu). Historický lexikon byl zpracován v některých krajských mutacích, obohacený i detailnější místopisné zajímavosti (dobové fotografie obcí), např. *za Liberecký kraj*. ČSÚ reagoval na legislativní změny (umožnily vybraným obcím navrácení historického statutu *města a městyse*) a připravil rozsáhlou datovou publikaci za tato území (vč. současných měst)

obsahující vedle dlouhodobé řady počtu obyvatel také aktuální sociodemografické údaje z městské a obecní statistiky, které mají vztah k „střediskovosti“ obcí.

4. Webová prezentace demografických dat

Demografické údaje na webu ČSÚ lze nalézt v několika sekcích, například:

- ve specializované sekci Nejžádanější pod odkazem *Obyvatelstvo* – zde můžete nalézt nejaktuálnější údaje (tzv. Rychlé informace), pramenné dílo, specializované aplikace či odkazy (databáze demografických údajů za obce ČR včetně územních přesunů), časové řady, zajímavé grafy a kartogramy, ale i třeba atraktivní animované grafy (např. s vývojem obyvatel v jednotlivých krajích do roku 2006 – viz **Obrázek 1**), ale také například příspěvky z demografické konference



Obrázek 1 Animovaný graf Populační vývoj v hlavním městě Praha

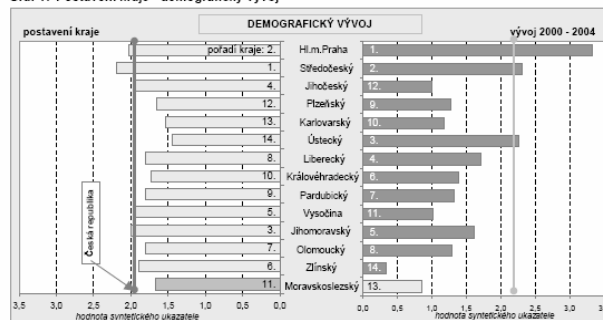
- v sekci Data v publikacích po odkazem *Obyvatelstvo, volby* – souhrnný přehled všech vydaných publikací v elektronické podobě, s možností zobrazení a ukládání jednotlivých tabulek v XLS nebo PDF formátu, přehledem metodik apod. Je zde k dispozici i rozsáhlý archiv.

- v sekci Regiony, města, obce pod odkazem *Stránky krajských pracovišť* - každý kraj zde prezentuje publikace (viz např. **Obrázek 2**), časové řady, animované grafy, ale třeba i krátké demografické analýzy – rychlé informace.

- v sekci Databáze, registry, IČO pod odkazem *Veřejná databáze* – s odkazem na specializovanou zobrazovací aplikaci – umožňuje přístup do centrální veřejné databáze včetně kapitoly Demografie, uživatelé mohou uživatelsky přívětivou formou vybírat i požadované parametry pro jednotlivé předdefinované tabulky a volit alternativní způsoby zobrazení – Tabulka-Graf-Mapa. Jsou navíc okamžitě dostupné definice ukazatelů a další metainformace vztahované ke každému údaji (viz **Obrázek 3**).

Ustecký a Liberecký. Naproti tomu v Karlovarském kraji, který má z hlediska demografie obdobné postavení jako Ústecký, nebyl vývoj v posledních letech tak příznivý – vinou přetrvávající vysoké úmrtnosti a rychlejší dynamiky populačního stárnutí. Nejméně příznivý demografický vývoj v čase zaznamenal kraj Zlínský a Moravskoslezský, podprůměrné hodnoty však nacházíme ve všech migračně uzavřenějších moravských krajích. Postavení Jihočeského kraje se postupně zhoršovalo zejména vlivem vývoje standardizované úmrtnosti.

Graf 17 Postavení kraje - demografický vývoj



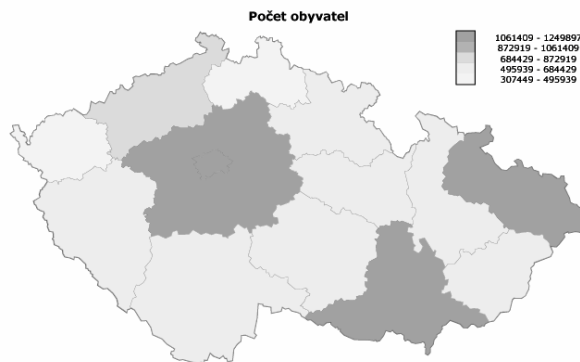
Sociální vývoj

V sociální oblasti se potvrdilo dominantní postavení Hlavního města Prahy; příznivých hodnot syntetického ukazatele dosáhly dále kraje Plzeňský, Středočeský a Jihočeský. Naproti tomu ze sociálního hlediska byla situace nejproblematičtější v Moravskoslezském, Ústeckém a Zlínském kraji, nepříznivě je možná hodnotit

Obrázek 2 Ukázka z analytické publikace publikace krajského pracoviště v Ostravě

MOS B04KR_M Počet obyvatel v krajích
stav k 31.12.

Období: 31.12.2007



Obrázek 3 Ukázka zobrazení mapy ve veřejné databázi ČSÚ

5. Regionální analýzy v demografii

Regionální analýzy se v oblasti demografie mohou opřít o tradičně širokou datovou základnu (roční bilance demografických událostí, populační census). Mezi uživateli statistických údajů patří takové produkty k nejžádanějším.

Každoroční rozbor populačního vývoje je možno nalézt ve zprávě *Vývoj obyvatelstva ČR*, která má od roku 2002 i kapitulu věnující se analýze populačního vývoje v krajích. Především pracoviště v Olomouci připravuje tematicky zaměřené ad-hoc analýzy, většina z nich pojednává o zvoleném demografickém problému i v regionální dimenzi, byť jejich hlavní zaměření je převážně republikové (*Sebevraždy*, *Vysokoškoláci v demografickém pohledu*, *Vnitřní stěhování v ČR*).

Ke speciálním výstupům můžeme zařadit Populační projekce. K jejich zpracování je třeba kvalifikovaného posouzení tendencí všech složek populačního vývoje, jde tedy o práci analytického zamření. Poslední takový odhad zpracoval ČSÚ v roce 2003 (Populační prognóza ČR do r. 2050) Byla to již pátá projekce zpracovaná ČSÚ po vzniku samostatné ČR a zároveň první, která obsahovala výhled populačního vývoje pro jednotlivé kraje (zpracovaný bez vlivu migrace). Na rok 2009 připravuje ČSÚ jejich aktualizaci.

Regionální analýzy s demografickou problematikou zpracovávají také pracovníci informačních služeb ČSÚ v jednotlivých krajích. Každý rok je pro daný kraj zpracovávána 1 „velká“ analýza. Její jednotnou metodologickou přípravu má na starosti pracovní skupina složená z regionálních analytiků (z ústředí i jednotlivých krajských pracovišť). V letech 2005-2007 byly touto formou zpracovány 3 významné analýzy, ve kterých se mj. uplatnilo víceaspektní hodnocení demografických ukazatelů:

- Demografický, sociální a ekonomický vývoj kraje v letech 2000 – 2004
- Vývoj lidských zdrojů v kraji v letech 2000 – 2005
- Regionální rozdíly v demografickém, sociálním a ekonomickém vývoji kraje v letech 2000 – 2005

V současnosti probíhající v jednotlivých krajských pracovištích ČSÚ práce na dvou významných analýzách: vývoje venkova a bytová výstavba, v obou z nich se ve výrazné míře uplatňují aktuální demografické otázky (např. suburbanizační tendence, vzestup porodnosti).

6. Závěr

Příspěvkem jsme se pokusili ilustrovat současnou podobu prezentace demografických dat v Českém statistickém úřadě a zachytit i vazbu na regionální statistiku. Její klíčový význam spočívá v poskytování objektivních a vysoce kvalitních údajů, které by měly nejen poskytnout souhrnný obrázek o demografickém vývoji v republice a jeho částech, ale také pomoci při formování strategií udržitelného rozvoje se zajištěním harmonického rozvoje v jednotlivých regionech.

7. Literatura

- [1] ČSÚ 2005. Demografický, sociální a ekonomický vývoj kraje v letech 2000 až 2004 (mutace za jednotlivé kraje).
- [2] ČSÚ 2006. Cizinci v regionech ČR (celorepubliková publikace), kap E (Analýza regionálních rozdílů), str. 135-139.
- [3] ČSÚ 2007. Regionální rozdíly v demografickém, sociálním a ekonomickém vývoji ČR v letech 2000 až 2005 (celorepubliková publikace), 90 s.
- [4] ČSÚ 2007. Analýza regionálních rozdílů v ČR. (celorepubliková publikace), 82 s.
- [5] KUPROVÁ, L. – KAMENICKÝ, J. 2004. Multikriteriální hodnocení postavení krajů v rámci ČR v letech 2000-2004 In: STATISTIKA. 2006, č. 4. Praha : český statistický úřad, 2006, s. 303 – 315.

Adresa autorů:

Eduard Durník, Ing.
Český statistický úřad
Na padesátém 81, Praha 10, 100 82
eduard.durnik@czso.cz

Jiří Kamenický, Bc.
Český statistický úřad
Na padesátém 81, Praha 10, 100 82
jiri.kamenicky@czso.cz

Analýza růstu střední délky života v ČR metodou klouzavé lineární regrese **The analysis of the growths of the life expectation in the Czech Republic** **by the method of moving linear regression**

Tomáš Fiala

Abstract: The paper contains the analysis of the development of the life expectancy in the Czech Republic after the World War II by the method of moving linear regression computed for 10 years' period. The initial rapid growth of the life expectancy was followed by the stagnation or even drop in the sixties and seventies. Relatively rapid growth in the nineties was slowed down at present time. The present annual average increment of life expectancy (measured by regression coefficient) is about 0.27 years for males and 0.20 years for females. In 2050 the estimated life expectancy for males is about 85 years, for females about 88 years.

Key words: mortality, life expectation, linear regression.

Klíčová slova: úmrtnost, střední délka života, lineární regrese.

1. Úvod

Jednou z nejčastěji používaných souhrnných charakteristik úmrtnosti je střední délka života novorozence. Její hodnoty v ČR po druhé světové válce měly rostoucí trend přerušeny obdobím stagnace či dokonce mírného poklesu. Jedná se o trend po částech lineární, proto byl vývoj střední délky života modelován lineární regresní funkcí, ovšem pouze za 10letá období vzájemně se překrývající. Byl analyzován vývoj střední délky života v letech 19472007, první „desetiletá“ regrese byla proto vypočtena za období 194756, druhá za období 194857, atd., poslední za období 19982007. Jedná se o určitou analogii výpočtu klouzavých průměrů, používám proto pojmu „klouzavá lineární regrese“. Parametr u lineárního členu regresní funkce (regresní koeficient) lze pak interpretovat jako průměrný roční nárůst střední délky života v příslušném desetiletém období, za které byla regrese vypočtena. Na základě regresních modelů byl rovněž proveden extrapolovaný odhad střední délky života v roce 2050. Analýza byla prováděna zvlášť pro muže a pro ženy.

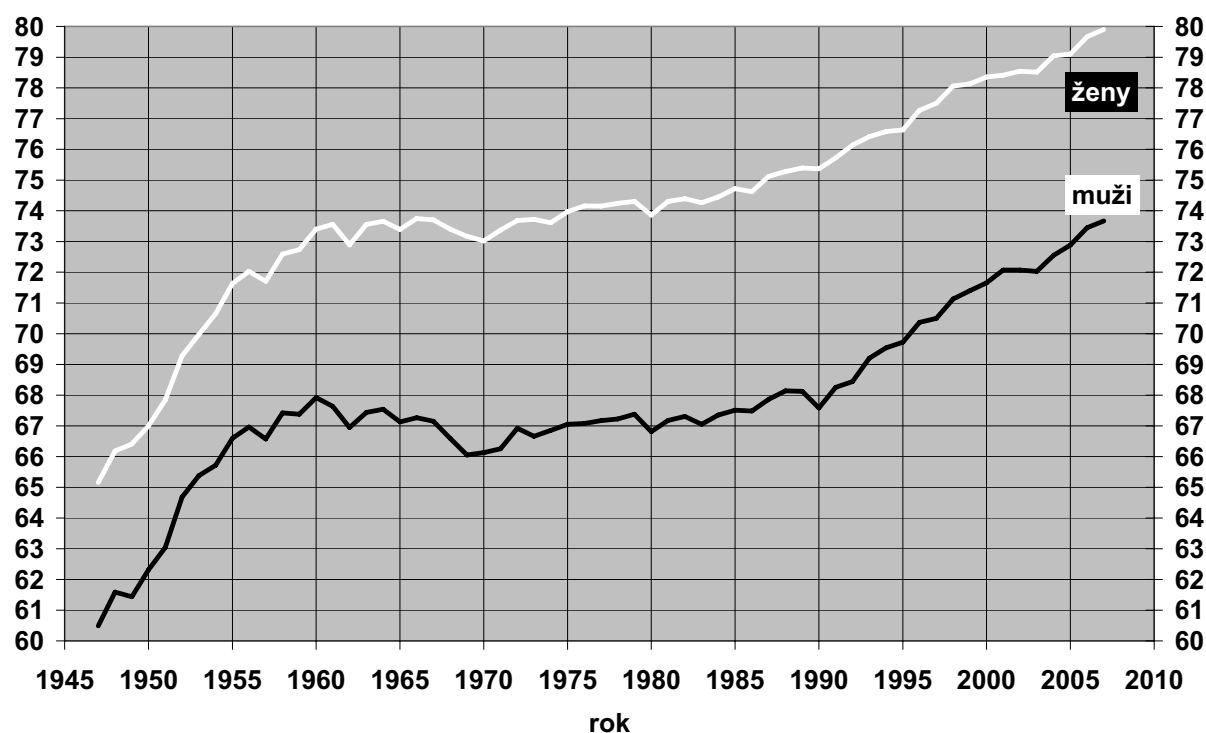
2. Vývoj střední délky života

Empirické hodnoty střední délky života mužů a žen v ČR od roku 1947 zobrazuje Graf 1. Základní trend vývoje je stejný pro obě pohlaví: poměrně prudký růst do roku 1955 a jeho postupné zpomalení do roku 1960. V 60. letech pozorujeme u žen stagnaci, u mužů dokonce pokles střední délky života. Teprve v 70. letech začíná opět růst, ovšem poměrně mírný trvající do konce 80. let. Od roku 1991 se růst střední délky života u obou pohlaví výrazně zvýšil a trvá do současné doby.

3. Hodnoty regresních koeficientů

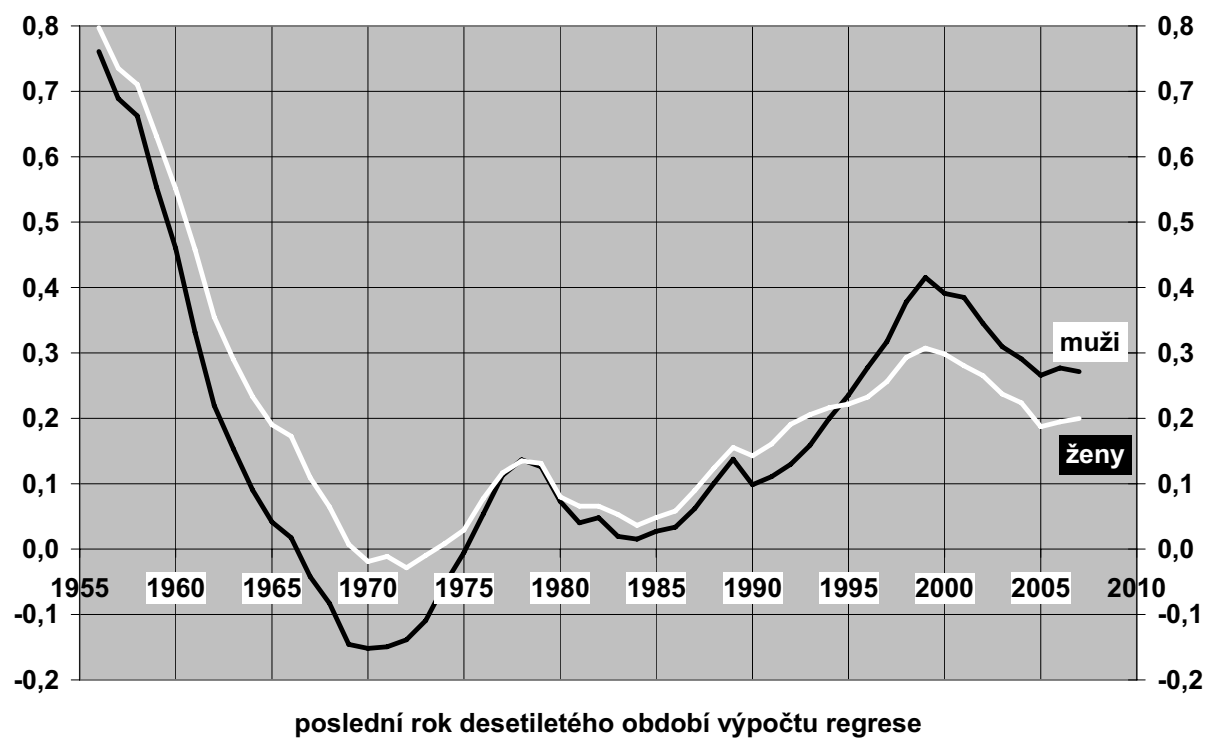
Podívejme se nyní na vývoj regresních koeficientů klouzavé lineární regrese počítané vždy pro 10leté období (viz Graf 2). Letopočet na vodorovné ose v grafech označuje vždy konec příslušného období. Např. hodnota pro rok 2000 znamená hodnotu regresního koeficientu lineární regrese vypočtené pro období 19912000 atd.

Hodnoty regresních koeficientů se během sledovaného období poměrně výrazně měnily. Vysoké počáteční hodnoty poměrně rychle klesly. Zatímco za období 194756 byla hodnota regresního koeficientu pro muže i pro ženy blízká 0,8, za období 195261 již byla pouze zhruba poloviční.



Zdroj: Vlastní graf na základě dat CSÚ

Graf 1: Vývoj střední délky života mužů a žen v ČR



Zdroj: Vlastní výpočet

Graf 2: Hodnoty regresních koeficientů klouzavé lineární regrese pro desetiletá období

Počínaje obdobím 195867 a konče obdobím 196675 jsou hodnoty regresního koeficientu pro muže záporné, což svědčí o poklesu střední délky života. Pro ženy se v uvedených obdobích pohybují regresní koeficienty kolem nulových hodnot, střední délky života tedy stagnovala.

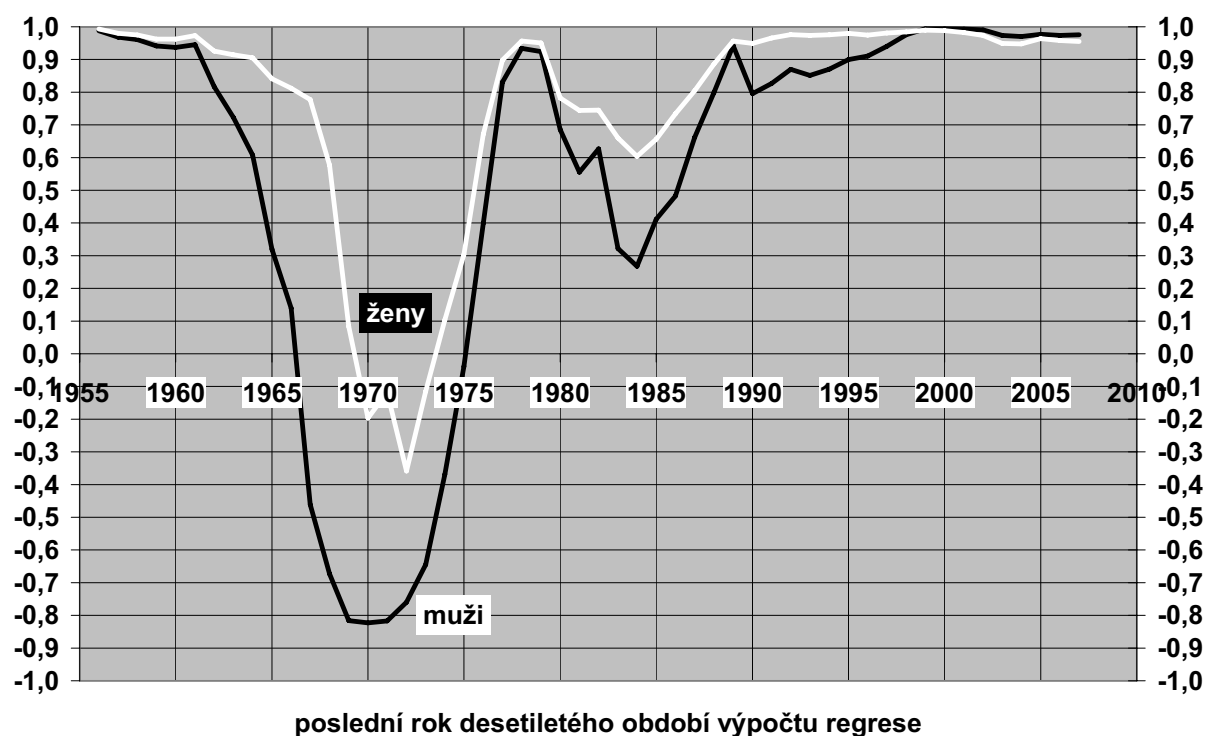
Počínaje obdobím 196776 jsou hodnoty regresních koeficientů pro muže i pro ženy trvale kladné, což svědčí o obnovení trendu růstu střední délky života. Tempo růstu však kolísá. Zatímco za období 196978 dosahují regresní koeficienty hodnot téměř 0,15, za období 1975-84 jsou jen o málo vyšší než 0. Teprve pak pozorujeme dlouhodobé zvyšování tempa růstu střední délky života vrcholící v období 199099, kdy je hodnota regresního koeficientu pro muže vyšší než 0,4 a pro ženy vyšší než 0,3.

V posledních letech došlo k určitému zpomalení růstu střední délky života. Hodnota regresního koeficientů pro muže za období 19982007 činila 0,27, pro ženy pak 0,20.

Poznamenejme ještě, že v minulých obdobích byl regresní koeficient pro ženy vyšší než pro muže (nebo zhruba stejný), střední délky života žen rostla rychleji než u mužů nebo alespoň stejně rychle. Teprve počínaje obdobím 198695 dochází ke změně; střední délka života mužů roste rychleji než střední délka života žen.

4. Hodnoty koeficientů korelace

Jednoduchou mírou těsnosti lineární závislosti je koeficient korelace. Jeho hodnoty zachycuje Graf 3. Vidíme, že zatímco v minulosti může být použití lineární regrese poněkud diskutabilní, počínaje obdobím zhruba 197988 je lineární závislost poměrně těsná (koeficient korelace je vyšší než 0,8, později dokonce vyšší než 0,9).



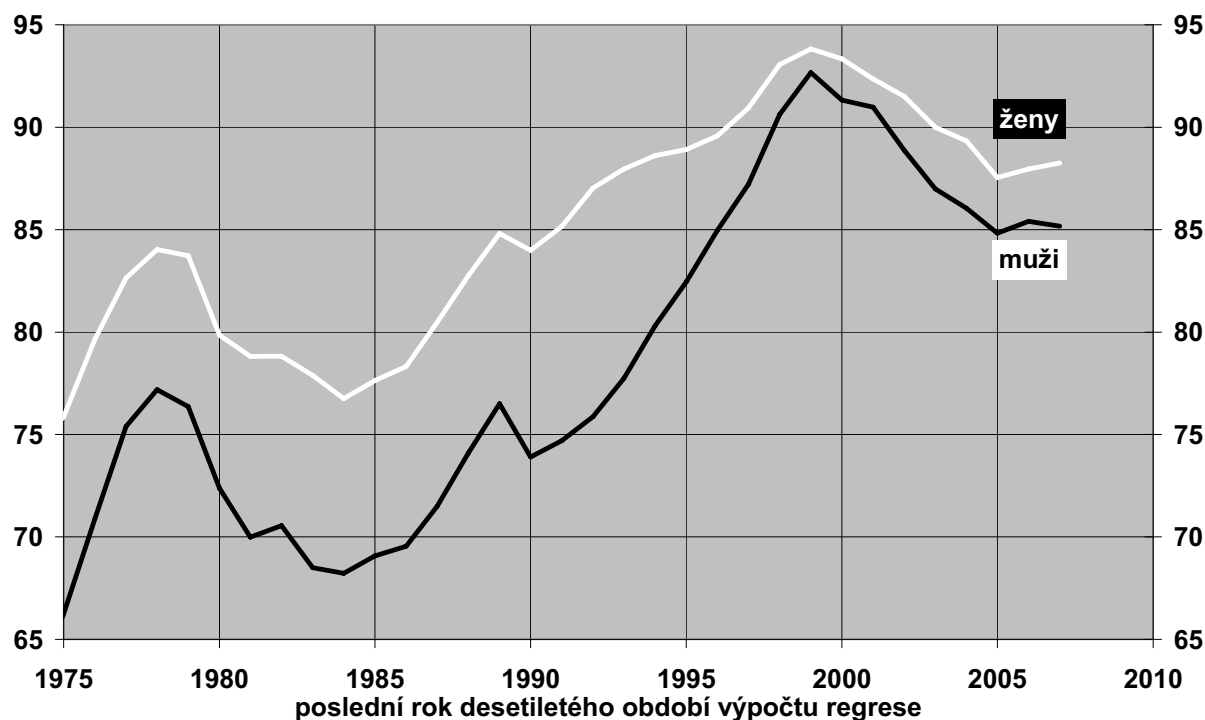
Zdroj: Vlastní výpočet

Graf 3: Hodnoty koeficientů korelace pro desetiletá období

5. Extrapolovaná střední délka života v roce 2050

Na základě regresního modelu lze provádět extrapolaci budoucího vývoje. Následující graf zachycuje hodnoty střední délky života v roce 2050 vypočtené extrapolací na základě re-

gresních modelů vycházejících z jednotlivých desetiletých období. Byly použity pouze modely počínaje obdobím 196675, kdy regresní model opět předpokládal růst střední délky života. Pokud by i v tomto tisíciletí pokračoval růst střední délky života stejným tempem jako v období 199099, mohla by být v roce 2050 střední délky života mužů i žen vyšší než 90 let. Při zachování současného pomalejšího tempa růstu však lze předpokládat v roce 2050 střední délku života mužů pouze kolem 85 let, u žen zhruba o 3 roky vyšší.



Zdroj: Vlastní výpočet

Graf 4: Hodnoty extrapolované střední délky života v roce 2050 na základě klouzavé lineární regrese pro desetiletá období

6. Závěr

Střední délka života v ČR po druhé světové válce do roku 1960 poměrně rychle rostla. Pak však následuje dvacetileté období stagnace, u mužů dokonce mírného poklesu střední délky života. K výraznému růstu střední délky života dochází od druhé poloviny 80. let, především pak po roce 1990.

Měříme-li nárůst střední délky života regresním koeficientem počítaným za 10letá období, k nejvyššímu nárůstu v poslední době došlo v období 199099 (roční nárůst zhruba 0,4 let u mužů a 0,3 roku u žen). Za poslední desetileté období (19982007) byl průměrný roční nárůst zhruba 0,27 let u mužů a 0,2 roku u žen. Při zachování tohoto tempa růstu by se střední délka života mužů v roce 2050 pohybovala kolem 85 let, střední délka života žen kolem 88 let.

7. Literatura

- [1] FIALA, T. 2006. Průřezová a kohortní analýza úmrtnosti v ČR pomocí životních potenciálů. Bratislava 07.12.2006 – 08.12.2006. In: Forum Statisticum Slovacum. Bratislava : Slovenská statistická a demografická spoločnosť, 2006, s. 54–59. ISSN 1336-7420.
- [2] KOSCHIN, F. 2006. Základní demografické trendy v České republice a v Evropě. Praha 25.10.2006 – 26.10.2006. In: BEST. [online] Berlín : European Academy of the Urban Environment, 2006, s. 1–21. URL: <http://www.eaue.de/best/best-Prag-Koschin.pdf>.

- [3]LANGHAMROVÁ, J. – FIALA, T. 2007. Střední délka života zemí EU. Praha : ČSÚ, 2007. 4 s. ISSN 0011-8265. Tabulkový soubor – časopis Demografie, roč. 49, č. 3.

Adresa autora:

Tomáš Fiala, RNDr., CSc.
katedra demografie
fakulta informatiky a statistiky VŠE Praha
nám. W. Churchilla 4
130 67 Praha 3
Česká republika
fiala@vse.cz

Článek vznikl v rámci dlouhodobého výzkumného projektu 2D06026 „Reprodukce lidského kapitálu“ financovaného MŠMT v rámci Národního programu výzkumu II.

Aplikácia metódy QFD v hutníckej výrobe **Applikation of method QFD in metallurgical production**

Zuzana Hajduová, Cyril Závadský, Erika Liptáková

Abstract: The presented work paper deals with the systems of quality control, specially with systems and methods of measurements of quality costs. An implementation of such a measurement in practice is presented. The QFD method uses data and rigorous statistical analysis to identify defects in a process or a product. In this paper we present basic informations about this method.

Key words: quality, statistical method, QFD, copper wire.

Kľúčové slová: kvalita, štatistická metóda, QFD, medený drôt.

1. Úvod

Zlá kvalita vyvoláva celý rad ekonomických problémov u producentov i užívateľov. Naopak, špičková schopnosť uspokojovať požiadavky zákazníkov prináša všetkým účastníkom trhu značné efekty. Preto sa vo fungujúcich systémoch kvality musí rozvíjať aj prvok, ktorý sa najčastejšie označuje ako zlepšenie kvality. Štatistické metódy tvoria jeho neoddeliteľnú súčasť.

Cieľom našej práce bolo navrhnúť postupy pre vykonávanie meraní, monitorovania a zlepšovania procesov zabezpečujúcich výrobu medeného drôtu v konkrétnej hutníckej výrobe. Bolo poukázané na opodstatnené využívanie metódy QFD (Quality Function Deployment) na identifikáciu faktorov, ktoré majú vplyv na kvalitu vyrábaného medeného drôtu.

2. Charakteristika metódy QFD

QFD je metódou formálne založenou na systéme matíc, ktorý umožňuje aktuálne reagovať na priania zákazníka a rešpektovať ich vo všetkých fázach tvorby výrobku .

Dôvody použitia QFD:

- ❖ Sústreďenie úsilia na začiatok procesu tvorby výrobku znižuje počet konštrukčných a technologických zmien.
- ❖ Skracuje dobu vývoja, čo znamená, že zákazník dostane výrobok, ktorý práve požaduje a prostredníctvom takto získanej flexibility sa zrýchľuje reakčná doba.
- ❖ Zlepšuje komunikáciu vo vývojovom štádiu a špecifikuje požiadavky do výrobného procesu, čo má za následok menej problémov pri rozbehu výroby.
- ❖ Znižuje náklady na výrobu nových výrobkov.
- ❖ Redukuje problémy v distribučnej sieti.
- ❖ Orientácia na zákazníka zabezpečí predajnosť výrobkov a stálu pozíciu na trhu.

Metóda je charakterizovaná systematickým poznávaním zákazníckych prianí a ich transformáciou do technických špecifikácií výrobku, riadením k stanoveným cieľom.

1. krok: Identifikovať potreby zákazníkov

Potreby zákazníkov sú ich požiadavky na produkty alebo služby vyjadrené slovami a termínmi, ktoré oni sami používajú. Dôležitú úlohu v tejto etape zohrávajú marketingoví pracovníci. Jednotlivé oblasti definovaných potrieb (obsluha, ekonomická stránka) sa členia na ich hlavné komponenty (napríklad v obsluhu je to individuálny prístup, včasnosť, jednoduchosť, vhodné správanie). Tie sa ďalej členia na jednoduchšie konkrétne zvládnuteľné prvky (v správaní je to priateľské vystupovanie, nijaké čakanie v radoch, atď.).

2. krok: Zostaviť zoznam charakteristík procesu (produktu) zodpovedajúcich definovaným potrebám

Charakteristiky produktu či procesu predstavujú technické parametre, ktoré sú vyvíjané za účelom čo najlepšieho uspokojovania potrieb. V domčeku kvality sa umiestňujú pod „striešku“.

Jednotlivé charakteristiky alebo interné procesy často navzájom súvisia. Tieto vzťahy (korelácie) sa znázorňujú v časti predstavujúcej „striešku“ pomocou symbolov napríklad:

- krúžok s bodkou uprostred znázorňuje veľmi silný vzťah,
- prázdny krúžok predstavuje silný vzťah,
- △ trojuholník znamená slabú vzájomnú súvislosť.

3. krok: Zostrojiť maticu vzťahov medzi atribútmi zákazníka a zodpovedajúcimi charakteristikami produktu

V tretej fáze budujeme základnú maticu (pod „strieškou“). Exaktne vypočítavame, alebo podľa možnosti čo najviac objektívne odhadujeme vzťahy medzi atribútmi zákazníka (potrebami) a zodpovedajúcimi charakteristikami produktu alebo služby. Na zaznačenie vzťahu sa spravidla používajú tie isté symboly, ktoré boli definované v 2. kroku.

4. krok: Zhodnotiť situáciu na trhu

V tejto etape sa každá požiadavka zákazníka ohodnotí podľa jeho vyjadrenia známku (prioritou) od 1 (menej významné) po 5 (extrémne významné). V časti matice pre hodnotenie konkurencie sa zvýraznia silné a slabé stránky konkurencie v dosahovaní uspokojenia daných potrieb zákazníkov. Toto hodnotenie napomáha pri rozhodovaní o prioritách pre nové produkty, respektíve určuje stratégiu rozvoja kvality. Tam, kde napríklad zákazník požaduje uspokojenie na úrovni 5 a pôsobenie konkurencie je len na úrovni 1 alebo 2, existuje spravidla reálna možnosť uplatnenia sa na tomto trhu.

5. krok: Zhodnotiť zodpovedajúce charakteristiky procesov (produktov) konkurencie a stanoviť ciele

V tomto bode zohráva dôležitú úlohu prieskum trhu. V dolnej časti domčeka kvality sa zaznačujú miery (od 1 po 5), akými jednotlivé konkrétne charakteristiky procesu (produktu) konkurencie uspokojujú zákazníka. Ak konkurenčný produkt alebo služba uspokojuje potreby a požiadavky zákazníkov na dobrej úrovni (matica vpravo), ale hodnotenie charakteristík procesov (produktov) nezodpovedá vyjadreniu tohto uspokojenia, buď hodnotenie príslušných charakteristík nie je pravdivé, alebo existuje niečo iné, čo ovplyvňuje vnímanie zákazníka (nepostrehnutá dôležitá potreba zákazníka).

6. krok: Vytipovať zodpovedajúce charakteristiky pre následný rozvoj

Podstatu tvorí identifikácia charakteristík, ktoré predstavujú významnú súvislosť s definovanými potrebami zákazníkov, alebo je v danej oblasti chabé uplatnenie konkurencie. Vybrané charakteristiky sa stanú cieľom následných opatrení na zlepšenie uspokojenia zákazníkov.

3. Implementácia metódy QFD

Výhodou metódy QFD (Quality Function Deployment), ktorej teoretický základ už bol opísaný, je to, že v sebe efektívnym spôsobom zahŕňa niekoľko dôležitých funkcií potrebných na definovanie hlavných smerov a problémov zlepšovania.

Ide o:

- ❖ formuláciu najdôležitejších požiadaviek zákazníka,
- ❖ ich hierarchizáciu, teda určenie významu pre zákazníka,

- ❖ vyhodnotenie konkurencieschopnosti z pohľadu vnímania zákazníkom,
- ❖ vyhodnotenie konkurencieschopnosti z pohľadu dosahovania konkrétnych technických parametrov,
- ❖ transformáciu požiadaviek zákazníka do technických parametrov výroby,
- ❖ určenie kritických parametrov výroby na základe odhadu tesnosti väzby.

Jedným z najdôležitejších atribútov kvality medeného drôtu je jeho odolnosť proti pretrhnutiu. Problém spočíva v tom, že opätovne ťahaný drôt sa pretrhne skôr, než výrobcovia dosiahnu potrebný rozmer (od $\varnothing = 0,2$ mm až po $\varnothing = 1,85$ mm). Veľký vplyv na to majú materiálové vlastnosti (kompaktnosť materiálu), obsah kyslíka v medi a iné dôležité faktory.

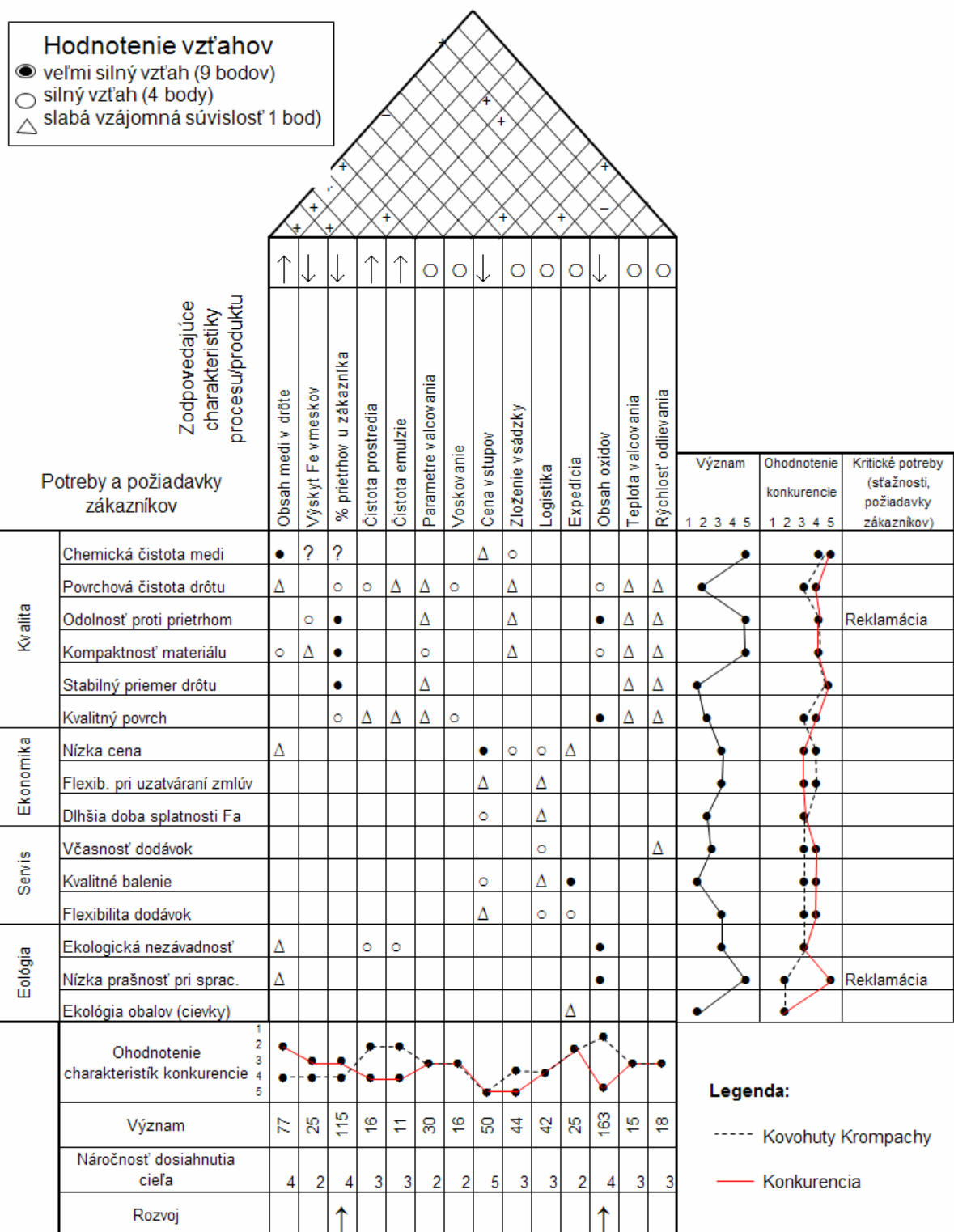
Úlohou je dosiahnuť stav, aby sa drôty s požadovanými rozmermi nepretrhli pri výrobe a mali požadované vlastnosti. Aby sme splnili tento cieľ, musíme si zvoliť vhodné kvalitatívne znaky a kritéria, podľa ktorých vieme zhodnotiť, či sme splnili vytýčený cieľ alebo nie. Za hlavnú odozvu sme si zvolili pretrhnutie drôtu, atribútovú náhodnú veličinu, ktorá je merateľná u zákazníka. Problém s takto definovanou odozvou spočíva v relatívne komplikovanom spôsobe identifikácie drôtu u zákazníka a veľkým časovým oneskorením.

Podklady pre vykonanie metódy QFD boli získané na základe reklamácií, ako aj ohlasov zákazníkov danej hutníckej firmy. Ďalším zdrojom informácií bol brainstorming, vykonaný so skupinou odborníkov, bývalých zamestnancov prevádzky Cu drôt. Na základe týchto informácií boli vyhodnotené vstupné charakteristiky metódy QFD. Požiadavky zákazníkov na vstupe boli rozdelené do štyroch základných častí – kvalita, ekonomika, servis a ekológia.

<i>Kvalita</i>	<i>Ekonomika</i>	<i>Servis</i>	<i>Ekológia</i>
Chemická čistota medi	Nízka cena	Včasnosť dodávok	Ekologická nezávadnosť
Kompaktnosť materiálu	Flexibilita pri uzatváraní zmlúv	Flexibilita dodávok	Nízka prašnosť pri spracovaní
Stabilný priemer drôtu	Dlhšia doba splatnosti faktúr	Kvalitné balenie	Ekológia obalov (cievky)
Kvalitný povrch			
Povrchová čistota drôtu			

Tabuľka 1: Požiadavky zákazníkov

Formulár, na základe ktorého je možné identifikovať priebeh vypracovania analýzy QFD, je na Obr. 1.



Obrázok 1: QFD pre skúmanú hutnícku výrobu

Z pohľadu vnímania zákazníkom najväčší význam mali požiadavky na chemickú čistotu medi, odolnosť voči prietrhom a požiadavka nízkej prašnosti pri ťahaní drôtu na tandemových ťahačkách u zákazníka. Pri hodnotení konkurencieschopnosti sa práve posledne uvádzaná požiadavka javí ako konkurenčná nevýhoda sledovanej spoločnosti

Konkurenčné podniky totiž vplyvom inej technológie výroby z pohľadu zákazníka dosahovali podstatne menšie percento prašnosti. Jednotlivé súvzťažnosti, na základe ktorých možno transformovať požiadavky zákazníka do technických parametrov, sú zaznamenané

v štruktúre matice QFD. Vo vypracovanom formulári QFD (Obrázok 1) sú tiež uvedené číselné hodnoty charakterizujúce význam jednotlivých technických parametrov, pri ktorých ako hranicu kritičnosti sme uvažovali hodnotu 100 bodov. Na základe takéhoto kritéria nám z metódy QFD ako kritické faktory výroby jednoznačne vyšli dva parametre – percento prietrhov u zákazníka s bodovou hodnotou 115 a obsah oxidov s bodovou hodnotou 163.

Práve pomocou skúmanej metódy sme došli k záveru, že týmto spomínaným dvom parametrom, ktoré tvoria jednoznačne základné kritické prvky prípadného zlepšovania, musia zamestnanci podniku venovať zvýšenú pozornosť, ak chcú by na trhu firmou, ktorá prosperuje a je konkurencie schopná.

4. Záver

Využívanie štatistických metód vo výrobe už v súčasnosti nie je ničím výnimočným. Tento príspevok v skrátenej forme poukázal na možnosť, opodstatnenosť, ale aj vykonateľnosť implementácie sofistikovaných štatistických metód ako je QFD v rámci systému zlepšovania v praktických podmienkach výrobného procesu. Okrem teoretického zdôvodnenia použitia tejto metódy, bola prezentovaná aplikácia QFD na konkrétnych údajoch získaných z praxe. Prostredníctvom nej boli určené kritické aspekty merania výkonnosti a efektívnosti jednotlivých podprocesov.

5. Literatúra

- [1] TKÁČ, M.: Matematická štatistika a teória pravdepodobnosti. In: PROJEKT TEMPUS – PHARE IB_JEP – 13406-98, SJF TU Košice, 1999-2000. ISBN 80-7099-536-X.
- [2] WENDLING, R., LESEUX, B.: Ekonomické aspekty prevencie. In: Projekt TEMPUS-PHARE IB_JEP_13406-98, SJF, Technická univerzita v Košiciach, 1999-2000, ISBN 80-7099-583-1.
- [3] PETER S. PANDE, ROBERT P. NEUMAN, ROLAND R. CAVANAGH: The Six Sigma Way: How GE, Motorola, and other companies are honing their performance. New York 2000.
- [4] TIŽOVÁ, M. - TURISOVÁ, R. - HAJDUOVÁ, Z.: Environmentálna a ekonomická analýza nákladov a úžitkov. In: Trendy v systémoch riadenia podnikov: 7. medzinárodná vedecká konferencia, Herľany, 9.-10. november 2004: Zborník príspevkov. Košice: TU, 2004. 5 s.. ISBN 80-8073-202-7
- [5] MIXTAJ, L.: Economics of quality in an aviation company. In ACTA AVIONICA 2007, roč. XI, č.13, s. 83-87, ISSN 1335-9479.
- [6] TKÁČ, M.: Štatistické riadenie kvality v praxi. In: EKONOMICKÁ UNIVERZITA BRATISLAVA, Ekonóm, Bratislava, 2001, s. 50-54, ISBN 80-225-0145-X.

Adresa autora (-ov):

Zuzana Hajduová, RNDr., PhD.
Tajovského 11
041 30 Košice
Zuzana.Hajduova@tuke.sk

Cyril Závadský, Ing.
Tajovského 11
041 30 Košice
cyril.zavadsky@euke.sk

Erika Liptáková, RNDr.
Tajovského 11
041 30 Košice
erika.liptakova@euke.sk

Asymptotické vlastnosti odhadů s minimální Kolmogorovskou vzdáleností

Asymptotic properties of Minimum Kolmogorov distance density estimators

Bc. Jitka Hanousková a Ing. Václav Kůs, PhD.

Abstrakt: This paper focuses on the minimum Kolmogorov distance density estimate $\tilde{f}_n$ of probability density f on the real line. The rate of consistency for $n \rightarrow \infty$ is known if the degree of variations is finite. The rate of consistency is studied under more general conditions. The generalization of degree of variation - the partial degree of variation is defined for density g of nonparametric family D containing the unknown density f . If the partial degree is finite and some additional assumptions are fulfilled then the Kolmogorov distance estimate is consistent of the order of $n^{-1/2}$ in L_1 -norm and also in the expected L_1 -norm. The examples confront this method with the method based on Vapnik-Chervonenkis dimension.

Key words: Kolmogorov estimators, degree of variation, Vapnik-Chervonenkis dimension

Klíčová slova: Kolmogorovské odhady, stupeň variace, Vapnik-Chervonenkissova dimenze

1. Úvod

Zavedme značení, které budeme v dalším textu používat. Necht' λ je σ – finitní míra na $(\mathbb{R}, \mathcal{B})$, kde $\mathcal{B}$ je borelovská σ – algebra na $\mathbb{R}$. Buď $\mathcal{F}(\mathbb{R})$ množina všech distribučních funkcí na $(\mathbb{R}, \mathcal{B})$. Označme $\mathcal{F}_\lambda$ podmnožinu všech rozdělení absolutně spojitých vzhledem k míře λ na $(\mathbb{R}, \mathcal{B})$. Označme $\mathcal{D}_\lambda$ množinu v Banachově prostoru $L_1(\mathbb{R}, d\lambda)$ obsahující hustoty odpovídající distribučním funkcím z $\mathcal{F}_\lambda$, $\mathcal{D}$ její libovolnou neprázdnou podmnožinu a $\mathcal{F}$ podmnožinu $\mathcal{F}_\lambda$, obsahující distribuční funkce odpovídající hustotám z podmnožiny $\mathcal{D}$.

V dalším textu bude $\tilde{F}_n$ označovat distribuční funkce odpovídající odhadu hustoty $\tilde{f}_n$. Dále buďte $X_1, \dots, X_n$ stejně a nezávisle rozdělená pozorování, symbolem F_n budeme označovat empirickou distribuční funkci na založenou na $(X_1, \dots, X_n)$ a symbolem $\nu_n(A)$ budeme označovat empirickou distribuci založenou na $(X_1, \dots, X_n)$

$$F_n(x) = \frac{1}{n} \sum_{j=1}^n I_{\{X_j \leq x\}}, \quad \nu_n(A) = \frac{1}{n} \sum_{j=1}^n I_{\{X_j \in A\}}, \quad A \subset \mathbb{R},$$

kde $I_{\{X_j \leq x\}}$ je indikátor jevu $X_j \in (-\infty, x]$ a $I_{\{X_j \in A\}}$ je indikátor jevu $X_j \in A$.

Definice 1. Řekneme, že odhad $\tilde{f}_n$ hustoty $f \in \mathcal{D}$ je Kolmogorovským odhadem právě tehdy, když odpovídající distribuční funkce $\tilde{F}_n \in \mathcal{F}$ vyhovuje podmínce:

$$K(\tilde{F}_n, F_n) = \inf_{F \in \mathcal{F}} K(F, F_n) \quad \text{s.j.,}$$

kde K je Kolmogorovská vzdálenost na $\mathcal{F}(\mathbb{R})$ definována níže.

Každá metrická vzdálenost D na $\mathcal{F}(\mathbb{R})$ definuje pseudometriku ρ_D na $\mathcal{D}_\lambda$ tímto způsobem: $\rho_D(f, g) = D(F, G)$, kde $F, G \in \mathcal{F}_\lambda$ jsou distribuční funkce odpovídající hustotám $f, g \in \mathcal{D}_\lambda$.

Definice 2. Říkáme, že odhad $\tilde{f}_n$ hustoty $f \in \mathcal{D}$ je konzistentní v dané ρ_D vzdálenosti, resp. v její střední hodnotě, právě když $\rho_D(\tilde{f}_n, f) \rightarrow 0$ skoro jistě, resp. když $E\rho_D(\tilde{f}_n, f) \rightarrow 0$. Říkáme, že odhad $\tilde{f}_n$ je konzistentní řádu $r_n \rightarrow 0$ ve vzdálenosti ρ_D , resp. v její střední hodnotě, právě tehdy, když $\rho_D(\tilde{f}_n, f) = O_p(r_n)$, resp. když $E\rho_D(\tilde{f}_n, f) = O(r_n)$.

Nebudeme se zabývat obecnými vzdálenostmi D a ρ_D , ale pouze Kolmogorovskou vzdáleností na $\mathcal{D}(\rho_K)$ a vzdáleností v totální variaci na $\mathcal{D}(\rho_V)$ definovanými jako:

$$\rho_K(f, g) = K(F, G) = \sup_{x \in \mathbb{R}} |F(x) - G(x)|, \quad \rho_V(f, g) = V(F, G) = \int_{\mathbb{R}} |f - g| d\lambda,$$

kde $F, G \in \mathcal{F}_\lambda$ a f, g jsou jim odpovídající hustoty.

2. Konzistence v L_1 -normě

Definice 3. Řekneme, že ρ_K dominuje ρ_V na $\mathcal{D}$ (ozn. $\rho_K \succ \rho_V$) právě, když pro každou posloupnost $f, f_1, f_2, \dots \in \mathcal{D}$ konvergence $f_n \rightarrow f$ v ρ_K pro $n \rightarrow \infty$ implikuje konvergenci $f_n \rightarrow f$ v ρ_V . A řekneme, že ρ_K stejnoměrně dominuje ρ_V lokálně vzhledem k ρ_K na $\mathcal{D}$ (ozn. $\rho_K \succ^u \rho_V/\rho_K$) právě, když pro každou hustotu $g \in \mathcal{D}$ existuje $c > 0$ a Kolmogorovské okolí hustoty g , $B_K(g) \subset \mathcal{D}$ takové, že $\rho_K(f, g) \geq c\rho_V(f, g)$ pro všechny $f \in B_K(g)$.

Věta 1. Necht' $\rho_K \succ \rho_V$ na $\mathcal{D}$, potom každý Kolmogorovský odhad hustoty z $\mathcal{D}$ je konzistentní v L_1 -normě. Pokud $\rho_K \succ^u \rho_V/\rho_K$ na $\mathcal{D}$, potom každý Kolmogorovský odhad hustoty z $\mathcal{D}$ je konzistentní řádu $n^{-1/2}$ v L_1 -normě i střední hodnotě L_1 -normy.

Je známo, že každá dvojice hustot $f, g \in \mathcal{D}_\lambda$ definuje finitní míru ν s hustotou

$$\frac{d\nu}{d\lambda} = f - g. \text{ Tato míra je rozdílem dvou finitních měr na } (\mathbb{R}, \mathcal{B}), \text{ horní variace } \nu^+ \text{ a dolní}$$

variace ν^- s hustotami $\frac{d\nu^+}{d\lambda} = (f - g)^+ = \max\{0, f - g\}$ a ν^- s

$$\frac{d\nu^-}{d\lambda} = (f - g)^- = \max\{0, g - f\}.$$

Definice 4. Řekneme, že $A \in \mathcal{B}$ separuje ν^+ a ν^- , právě když platí buď $\nu^+(A) = \nu^+(\mathbb{R})$ a $\nu^-(\mathbb{R} - A) = \nu^-(\mathbb{R})$, a nebo $\nu^+(\mathbb{R} - A) = \nu^+(\mathbb{R})$ a $\nu^-(A) = \nu^-(\mathbb{R})$.

Definice 5. Buďte $f, g \in \mathcal{D}_\lambda$. Potom stupeň variace $DV(f, g) \in [0, +\infty]$ je definován takto: $DV(f, g) = 0$, když A separuje ν^+ a ν^- a $\lambda(A) = 0$. Jinak

$$DV(f, g) = \inf \left\{ m \in \mathbb{N} : A = \bigcup_{j=1}^m J_j, A \text{ separuje } \nu^+, \nu^- \right\},$$

kde $J_1, \dots, J_m$ jsou neprázdné intervaly v $\mathbb{R}$. Je-li minimalizovaná množina prázdná, tj. neexistuje-li žádné m s požadovanými vlastnostmi, pokládáme $DV(f, g) = +\infty$.

Definice 6. Pro danou $f \in \mathcal{D} \subset \mathcal{D}_\lambda$ a $\delta > 0$ definujeme lokální stupeň variace $LDV_\delta(f)$

hustoty f vzhledem ke Kolmogorovské vzdálenosti v $\mathcal{D}$ vztahem

$$LDV_{\delta}(f) = \sup\{DV(f, g) : g \in B_{K, \delta}(f) \cap \mathcal{D}\},$$

kde $B_{K, \delta}(f)$ je Kolmogorovská koule v $\mathcal{D}$ o poloměru δ se středem v f . Dále stupněm variace $DV(\mathcal{D})$ rodiny $\mathcal{D} \subset \mathcal{D}_{\lambda}$ nazveme: $DV(\mathcal{D}) = \sup\{DV(f, g) : f, g \in \mathcal{D}\}$.

Věta 2. *Bud' $\mathcal{D} \subset \mathcal{D}_{\lambda}$, necht' pro každou hustotu $f \in \mathcal{D}$ existuje $\delta > 0$ tak, že $LDV_{\delta}(f) < +\infty$. Potom ρ_K stejnoměrně dominuje ρ_V lokálně vzhledem k ρ_K ($\rho_K \succ^u \rho_V/\rho_K$) na $\mathcal{D}$.*

3. Konzistence v L_1 -normě za obecnějších předpokladů

Nyní zobecníme teorii ze Sekce 2 a ukážeme, že i Kolmogorovské odhady, pro které $LDV_{\delta}(f)$ nemusí být konečný a které splňují jisté dodatečné předpoklady, jsou konzistentní v L_1 -normě a její střední hodnotě. Za tímto účelem zavedeme nové typy dominancí.

Definice 7. Řekneme, že ρ_K asymptoticky dominuje ρ_V řádu a_n lokálně s ohledem na ρ_K na $\mathcal{D}$ (označme $\rho_K \succeq \rho_V/\rho_K (a_n \rightarrow 0)$) právě tehdy, když $(\forall f \in \mathcal{D}) (\exists B_K(f))$ takové, že $(\forall (f_n)_1^{\infty} \in B_K(f), f_n \rightarrow f \text{ v } \rho_K) (\exists c > 0)$ tak, že $\rho_K(f_n, f) \geq c\rho_V(f_n, f) - a_n$, kde $B_K(f)$ je Kolmogorovské okolí hustoty f a a_n je nezáporná posloupnost splňující $\lim_{n \rightarrow +\infty} a_n = 0$.

Nově definovaná dominance je skutečně zobecněním původní Definice 3 (viz příklad v [1]). Důkazy dále uvedených vět lze nalézt v [7].

Věta 3. *Necht' $\rho_K \succeq \rho_V/\rho_K (a_n \rightarrow 0)$ na $\mathcal{D}$. Potom každý Kolmogorovský odhad $\tilde{f}_n$ hustoty f z $\mathcal{D}$ je konzistentní v L_1 -normě a ve střední hodnotě L_1 -normy. Pokud navíc $a_n = o(n^{-1/2})$, potom každý Kolmogorovský odhad hustoty f z $\mathcal{D}$ je konzistentní řádu $n^{-1/2}$ v L_1 -normě a ve střední hodnotě L_1 -normy.*

Nyní budeme formulovat podmínky, za kterých rodina $\mathcal{D} \subset \mathcal{D}_{\lambda}$ vyhoví podmínkám dominance předpokládaným ve větě 3. Definujme zobecněnou podobu stupně variace a podívejme se na jeho vztah k dříve definovanému stupni variace z Definice 6.

Definice 8. Bud' $f, g \in \mathcal{D}_{\lambda}$ a necht' $a \in [0, +\infty]$. Potom parciální stupeň variace $DV_a(f, g) \in [0, +\infty]$ je definován takto:

$$DV(f, g) = \inf \left\{ m \in \mathbb{N} \cup \{0\} : A = \bigcup_{j=1}^m J_j + I, A \text{ separuje } \nu^+, \nu^- \right\},$$

kde $J_1, \dots, J_m$ jsou neprázdné disjunktní intervaly v $\mathbb{R}$ a $I \subset \mathbb{R}$ množina taková, že $\nu^+(I) \leq a$ a $\nu^-(I) \leq a$ přičemž $\bigcup_{j=1}^m J_j \cap I = \emptyset$. Je-li minimalizovaná množina prázdná, tj. neexistuje-li žádné m a množina I s požadovanými vlastnostmi, potom pokládáme $DV_a(f, g) = +\infty$.

Pokud $a > b > 0$, potom $DV_a \leq DV_b \leq DV_0 = DV$. Původní stupeň variace DV nás informuje o počtu znaménkových změn rozdílu $f - g$. Z parciálního stupně variace nezjistíme počet

znaménkových změn, víme jen, že pokud $DV_a(f, g) < +\infty$, potom až na konečný počet výjimek se všechny znaménkové změny odehrávají na množině I .

Věta 4. *Bud' $\mathcal{D} \subset \mathcal{D}_\lambda$. Necht' pro každou hustotu $f \in \mathcal{D}$ existuje Kolmogorovské okolí $B_K(f)$ a konstanta $K \in [0, \infty)$ a dále nezáporná posloupnost a_n s $\lim_{n \rightarrow \infty} a_n = 0$ tak, že $(\forall (f_n)_1^\infty \in B_K(f), f_n \rightarrow f \nu_{\rho_K})$ platí, že $\forall n \in \mathbb{N}$ je $DV_{a_n}(f_n, f) < K$. Potom ρ_K asymptoticky stejnoměrně dominuje ρ_V řádu a_n lokálně s ohledem na ρ_K na $\mathcal{D}$.*

4. Vapnik-Chervonenkisova dimenze

Krátce zmiňme jiný přístup pro ověřování konzistence Kolmogorovských odhadů [2].

Definice 9. Bud' $\mathcal{A}$ třída měřitelných množin. Pro $(z_1, \dots, z_n) \in \{\mathbb{R}^d\}^n$ označme $N_{\mathcal{A}}(z_1, \dots, z_n)$ počet různých množin ve třídě $\{\{z_1, \dots, z_n\} \cap A; A \in \mathcal{A}\}$. Dále n -tý shatter koeficient definujeme jako $s(\mathcal{A}, n) = \max_{(z_1, \dots, z_n) \in \{\mathbb{R}^d\}^n} N_{\mathcal{A}}(z_1, \dots, z_n)$. Tedy shatter koeficient je maximální počet různých podmnožin z n bodů, které mohou být vybrány pomocí třídy množin $\mathcal{A}$.

Definice 10. Bud' $\mathcal{A}$ třída množin $|\mathcal{A}| \geq 2$ ($|\mathcal{A}|$ je počet prvků množiny $\mathcal{A}$). Největší přirozené číslo $k \geq 1$, pro které platí $s(\mathcal{A}, k) = 2^k$ nazýváme Vapnik-Chervonenkisovou (VC) dimenzí třídy $\mathcal{A}$, označme $V_{\mathcal{A}}$. Pokud $s(\mathcal{A}, n) = 2^n$ pro každé n , položme $V_{\mathcal{A}} = \infty$.

Definice 11. Bud' $\mathcal{F}_{\Theta} = \{f_{\theta} : \theta \in \Theta\}$ parametrická třída hustot v $\mathbb{R}^d$ ($\Theta \subset \mathbb{R}^k$) a $X_1, \dots, X_n$ stejně nezávisle rozdělená pozorování na $f_{\theta} \in \mathcal{F}_{\Theta}$. Definujme třídu množin

$$\mathcal{A} = \{\{x \in \mathbb{R}^d : f_{\theta_1} > f_{\theta_2}\} \mid \theta_1, \theta_2 \in \Theta\}$$

a následně odhad parametru θ s minimální $D_{\mathcal{A}}$ vzdáleností vztahem $\tilde{\theta}_n = \arg \min_{\theta \in \Theta} D_{\mathcal{A}}(P_{\theta}, \nu_n)$,

kde P_{θ} je distribuce odpovídající hustotě f_{θ} a $D_{\mathcal{A}}$ je zobecněná Kolmogorov-Smirnovova vzdálenost

$$D_{\mathcal{A}}(P, Q) = \sup_{A \in \mathcal{A}} |P(A) - Q(A)|.$$

Věta 5. *Pokud $\mathcal{A}$ z Definice 11 má konečnou VC dimenzi, pak odhad hustoty f_{θ} pro zobecněnou Kolmogorovskou vzdálenost je konzistentní řádu $(n^{-1/2})$ ve střední hodnotě L_1 -normy.*

5. Vapnik-Chervonenkisova dimenze a stupeň variace

Na příkladech porovnejme oba zmíněné přístupy a podívejme se na vzájemný vztah jim odpovídajících charakteristik: stupně variace, parciálního stupně variace a VC dimenze. Protože VC dimenze je citlivá na změny hustot na množinách nulové míry, zatímco stupeň variace ne, uvažujeme dále rodiny $\mathcal{D}$ obsahující hustoty lišící se pouze na množinách nenulové míry.

Na jednoduchém příkladě snadno ověříme, že v obecném případě z konečnosti Vapnik-Chervonenkisovy dimenze neplyne konečnost stupně variace, tedy je možné použít Vapnik-Chervonenkisovu teorii, ale teorii založenou na stupni variace nikoli.

Příklad 1. Necht' rodina hustot $\mathcal{D}$ obsahuje rovnoměrnou hustotu g intervalu $(0,1)$ a hustoty f_n definované takto

$$f_n = \begin{cases} \frac{n}{n+1} & x \in (\frac{1}{2^{2i}}, \frac{1}{2^{2i-1}}), i=1,2,.. \\ \frac{n+3}{n+1} & x \in (\frac{1}{2^{2i+1}}, \frac{1}{2^{2i}}), i=1,2,.. \\ 1 & x \in (\frac{1}{2}, 1) \end{cases}$$

Potom $DV(\mathcal{D}) = \infty$, protože například $DV(g, f_1) = \infty$, zatímco $V_{\mathcal{A}} < \infty$, protože třída množin $\mathcal{A}$ obsahuje pouze dvě různé množiny. Avšak je možné použít teorii parciálního stupně variace vybudovanou v části 3. Snadno zjistíme, že totiž existuje posloupnost $a_n, \lim_{n \rightarrow \infty} a_n = 0$, taková, že $DV_{a_n}(f_n, g) < K < \infty$ pro $\forall n \in \mathbb{N}$.

V následujících dvou příkladech zkonstruujeme rodinu hustot, na které bude možné použít obě teorie (tj. $DV(\mathcal{D}) < \infty \wedge V_{\mathcal{A}} < \infty$) a dále rodinu hustot, na které obě teorie selžou (tj. $DV(\mathcal{D}) = \infty \wedge V_{\mathcal{A}} = \infty$).

Příklad 2. Uvažujme rodinu $\mathcal{D}$ všech hustot Gaussova normálního rozdělení. Potom zřejmě $DV(\mathcal{D}) = 1$ a $V_{\mathcal{A}} = 3$.

Příklad 3. Zvolme $k \in \mathbb{N}$ libovolně a vyberme k různých bodů $z_1, \dots, z_k$ v intervalu $(2k, 2k+1)$ a označme $z_{(1)}, \dots, z_{(k)}$ jejich vzestupné přerovnání. Pro $j = 1, \dots, k$ definujme intervaly

$$U_j = (u_{j-1}, u_j) \subset (2k, 2k+1), \text{ kde } u_j = \frac{z_{(j+1)} + z_{(j)}}{2} \text{ pro } j = 1, \dots, k-1 \\ u_0 = 2k, u_k = 2k+1.$$

Existuje 2^k různých podmnožin množiny $\{1, \dots, k\}$, označme je $M_i^k, i = 1, \dots, 2^k$ a definujme

$$f_{M_i^k} = \begin{cases} \lambda \left(\bigcup_{j \in M_i^k} U_j \right)^{-1} & \text{na } \bigcup_{j \in M_i^k} U_j, \text{ kde } \lambda \text{ je Lebesgueova míra na } \mathbb{R}. \\ 0 & \text{jinde na } \mathbb{R}, \end{cases}$$

Definujme rodinu hustot $\mathcal{D} = \{f_{M_i^k} : i = 1, \dots, 2^k, k \in \mathbb{N}\}$. Z konstrukce je zřejmé, že pro každé $k \in \mathbb{N}$ existuje množina $\{z_1, \dots, z_k\}$, která je roztržena třídou množin $\mathcal{A} = \{x \in \mathbb{R} : f_{M_i^k} > f_{M_j^k}, i, j \in \{1, \dots, 2^k\}, k \in \mathbb{N}\}$, tzn. $V_{\mathcal{A}} = \infty$. V tomto případě také $DV(\mathcal{D}) = \infty$.

Následující příklad ukáže, že z konečnosti stupně variace neplyne konečnost Vapnik-Chervonenkisovy dimenze.

Příklad 4. (viz [1]) Zvolme $k \in \mathbb{N}$ libovolně a vyberme k různých bodů $z_1, \dots, z_k$ v intervalu $(2k-1, 2k)$ a označme $z_{(1)}, \dots, z_{(k)}$ jejich vzestupné přerovnání a $Z_k = \{z_1, \dots, z_k\}$. Pro $j = 1, \dots, k$ definujme intervaly $U_i = (u_{i-1}, u_i)$, kde $u_i = \frac{z_{(j+1)} + z_{(j)}}{2}$ pro $i = 1, \dots, k-1$ $u_0 = 2k-1, u_k = 2k$. Existuje 2^k různých podmnožin $S_j \subset Z_k, j = 1, \dots, 2^k$. Pro $j = 1, \dots, 2^k$ a $i = 1, \dots, k$ definujme hustoty

$$g_j^k(x) = \begin{cases} 1 - \frac{1}{2^j} & \text{prox} \in U_i, \text{ když } U_i \cap S_j = \emptyset \\ 1 - \frac{1}{2^{j+1}} & \text{prox} \in U_i, \text{ když } U_i \cap S_j \neq \emptyset, \\ 1 - \sum_{i=1}^k \alpha_j^i d_i & \text{prox} \in (2k-2, \dots, 2k-1) \end{cases}$$

kde $d_i = |u_i - u_{i-1}|$ je délka intervalu U_i a

$$\alpha_j^i = \begin{cases} 1 - \frac{1}{2^j} & \text{ když } U_i \cap S_j = \emptyset \\ 1 - \frac{1}{2^{j+1}} & \text{ když } U_i \cap S_j \neq \emptyset \end{cases}.$$

Sestrojíme rodinu hustot $\mathcal{D} = \{g_j^k : j = 1, \dots, 2^k, k \in \mathbb{N}\} \cup \{f_k, k \in \mathbb{N}\}$, kde f_k je rovnoměrná hustota na intervalu $(2k-1, 2k)$. Vidíme, že $DV(\mathcal{D}) = 1$, zatímco VC dimenze třídy množin $\mathcal{A} = \{x \in \mathbb{R} : f(x) > g(x), f, g \in \mathcal{D}\}$ je $V_{\mathcal{A}} = \infty$, protože z konstrukce je patrné, že pro každé $k \in \mathbb{N}$ existuje množina $Z_k = \{z_1, \dots, z_k\}$, která je roztržena třídou množin $\mathcal{A}$.

Je tedy vidět, že podmínka $DV(\mathcal{D}) < \infty$ pro Kolmogorovské odhady je přímo neporovnatelná s podmínkou $V_{\mathcal{A}} < \infty$ pro zobecněné Kolmogorovské odhady bez dalších omezujících požadavků na rodiny $\mathcal{D}$. Nicméně zobecněný Kolmogorovský odhad je výpočetně značně náročnější než Kolmogorovský odhad, neboť minimalizace se provádí přes mnohem větší třídu množin. Také ověření podmínky $DV(\mathcal{D}) < \infty$ je snazší než ověřování podmínky $V_{\mathcal{A}} < \infty$.

6. Literatura

- [1] KŮS, V. 2004. Nonparametric density estimates consistent of the order of $n^{-1/2}$ in the L_1 -norm. In: Metrika, 2004, s. 1-14.
- [2] DEVROYE, L. - GYÖRFI, L. - LUGOSI, G.: A Probabilistic Theory of Pattern Recognition, New York: SPRINGER, 1996
- [3] IZENMAN, A. J. 1991. Recent Developments in Nonparametric Density Estimation. In: Journal of the American Statistical Association, 1991, s.205-224.
- [4] YATRACOS, Y. G. 1985. Rates of Convergence of Minimum Distance Estimators and Kolmogorov's Entropy. In: Annals of Statistics, č. 2, 1985, s. 768-774
- [5] DEVROYE, L. - GYÖRFI, L.: Nonparametric density Estimate, the L_1 - view, New York: WILEY, 1985
- [6] GYÖRFI, L. - VAJDA I. - VAN DER MEULEN, E. 1996.: Minimum Kolmogorov Distance Estimates of Parameters and Parametrized Distributions. In: Metrika, 1996, s.237-255
- [7] HANOUSKOVÁ, J. 2008. Asymptotické vlastnosti Kolmogorovských odhadů hustot pravděpodobnosti, Katedra matematiky FJFI ČVUT Praha, 2008.

*Jitka Hanousková, Bc. a Václav Kůs, Ing., PhD.
KM FJFI ČVUT Praha, Trojanova 13, 120 00 Praha 2.*

Štatistické riadenie kvality – Taguchiho návrhy experimentov v SAS-e Statistical Quality Control – Taguchi Design of Experiments within SAS

Adriana Horníková

Abstract: Taguchi experimental designs consider the economical qualities when judging the feature of a design. These designs are strongly related to quality control theory. In this contribution come up principles of Taguchi designs and what distinguishes them from other kinds of design. Loss function is the primary difference in understanding of quality of products followed by orthogonal arrays (inner and outer arrays), signal to noise ratio characteristic and others. All important steps of creating a Taguchi design, e.g. orthogonal arrays or orthogonal matrix, inner and outer arrays, SN ratio variable, etc., are illustrated with an example 'manufacturing of nylon tubes' (after Byrne and Taguchi, 1986) while utilizing the statistical module SAS JMP. Confidence intervals for average factor levels, particular treatment or prediction intervals are detailed according to Taguchi.

Keywords: design of experiments, Taguchi designs, statistical software SAS, SAS Slovakia

Kľúčové slová: navrhovanie experimentov, Taguchiho návrhy, štatistický softvér SAS.

1. Úvod

Metodika vytvárania návrhov experimentov nesie pomenovanie japonského profesora Genichiho Taguchiho, ktorý sa zaslúžil o vytvorenie nového prístupu kontroly kvality, pričom použil teóriu navrhovania experimentov. Taguchi navrhol a použil iný model pre charakteristiku nákladov výrobku, aký je typický v klasickom ponímaní. V ekonomickom ponímaní znižovanie nákladov je spojené so znižovaním variability. Definoval stratu (stratovú funkciu), ktorá zohľadňuje odberateľove očakávania, aby kúpil kvalitnejší výrobok a tiež potreby výrobcu produkovať nízko nákladové výrobky. Princípom Taguchiho návrhov je maximalizácia 'strát', kde stratu matematicky vyjadruje stratová (penalizačná) funkcia, $L(y)$. Unikátne je používanie pomeru signálu k šumu, ozn. SN (z angl. 'signal to noise ratio').

Tabuľka 1: Zhrnutie - typy stratovej funkcie a priemerná strata na výrobok.

Typ charakteristiky	Strata jedného výrobku	Priemerná strata na jeden výrobok	Pomer signálu k šumu (SN)
čím väčšie, tým lepšie	$\frac{k}{Y^2}$	$\frac{k}{\bar{Y}^2} \left[1 + \frac{3\sigma_Y^2}{\bar{Y}^2} \right]$	$-10 \cdot \log \left(\frac{1}{n} \sum_{i=1}^n \frac{1}{Y_i^2} \right)$
čím bližšie k T , tým lepšie	$k \cdot (Y - T)^2$	$k \cdot \sigma_Y^2 + k(\bar{Y} - T)^2$	$10 \cdot \log \left(\frac{\bar{Y}^2}{s^2} \right)$
čím menšie, tým lepšie	kY^2	$k \cdot \sigma_Y^2 + k \cdot \bar{Y}^2$	$-10 \cdot \log \left(\frac{1}{n} \sum_{i=1}^n Y_i^2 \right)$

V tabuľke 1 je $\bar{Y} = \frac{\sum_{i=1}^n Y_i}{n}$ a $s^2 = \frac{\sum_{i=1}^n (Y_i - \bar{Y})^2}{n-1}$, k je koeficient úmernosti, T je cieľová hodnota, Y_i je skutočná hodnota výkonovej charakteristiky a $L(Y)$ je strata spôsobená odchýlkou charakteristiky od cieľovej hodnoty T .

2. Inžinierska stratégia

V inžinierskej stratégii sa podľa G. Taguchiho postup navrhovania výrobku člení do troch základných metodických postupov:

- 1) **návrh systému (tiež primárny návrh):** Ide o krok, kedy sa navrhne nový výrobok použitím nových metód, koncepcií, nápadov, prístupov, atď. Jednou z možností ako si zachovať konkurencieschopnosť na trhu je, stať sa lídrom v aplikovaní moderných technológií. Avšak v dôsledku kopírovania tento krok dlho nevydrží.
- 2) **návrh parametrov (tiež sekundárny návrh):** V tejto fáze je cieľom zlepšenie uniformnosti výrobku, teda jeho vlastností a to s veľmi nízkymi nákladmi alebo ušetrením. Niektoré parametre výrobku sa upravujú tak, že správanie sa výrobku je menej citlivé k zdrojom variability. Toto je uskutočnené analýzou podielu variability vzhľadom na ozvu, teda pomer signálu k šumu. Z pohľadu získania cieľového nastavenia faktora je z inžinierskeho hľadiska menej náročné prestavenie strednej hodnoty faktora na cieľovú hodnotu, ako zníženie variability. Návrh parametrov je návrh, ktorého poslaním je ekonomické zvýhodnenie produkcie výrobku.
- 3) **návrh tolerancií (tiež terciálny návrh):** Ide o fázu zlepšovania kvality s minimálnymi nákladmi a to pomocou zužovania tolerančných hraníc výrobku alebo výrobného procesu s cieľom znížiť vlastnú variabilitu výrobku. Tento návrh sa v praxi používa hlavne vtedy, keď zníženie variability z návrhu parametrov nebolo postačujúce.

Ide o akúsi trojkrovovú procedúru, kedy sa pre kontrolné faktory vnútorného poľa nájdu optimálne nastavenia. Potom sa v rámci jednotlivých pokusov vonkajšieho poľa určí, ako na ozvu vplývajú šumové faktory. Toto sa realizuje pre všetky kombinácie pokusov vnútorného a vonkajšieho poľa. V prípade, že je potrebné optimalizovať nie len variabilitu, ale aj polohu premennej, je vhodné použiť novú premennú, ktorá je vytvorená ako pomer medzi signálom k šumu. Pri hľadaní optimálneho riešenia maximalizujeme pomer signálu k šumu pre potrebné charakteristiky a strednú hodnotu udržiavame napríklad na cieľovej hodnote (viď tab. 1).

Ešte k štatistickej podstate Taguchiho návrhov. Teoretické základy ortogonálnych matic (Taguchiho návrh sa označuje aj ortogonálnym poľom, ktorý sa dá matematicky vyjadriť formou ortogonálnej matice) siahajú už do roku 1897, kedy Jacques Hadamard, francúzsky matematik popisoval polia a matice. O znovuobjavenie sa až po druhej svetovej vojne postarali Blackett a Burman, štatistickí z Veľkej Británie, ktorí vytvorili tiež po nich pomenované návrhy, a to hlavne saturované návrhy. Vo všeobecnosti sú Hadamardové matice identické s maticami Taguchiho, líšia sa len v poradí medzi riadkami a stĺpcami.

Výstupy z analýzy experimentov sa zobrazujú aj graficky, pričom sa znázorňujú úrovne významných faktorov, ich kombinácie na posúdenie interakcií, poradie pokusov alebo sa vytvárajú normálne diagramy, polovičné normálne diagramy alebo Danielové diagramy.

Takýto prístup podporuje aj systém JMP SAS. Modul obsahuje rovnomennú vlastnú platformu pomenovanú Taguchiho polia (Taguchi Arrays). Softvér v súlade s teóriou navrhovania Taguchiho návrhov umožňuje rozlíšiť medzi kontrolovanými a šumovými faktormi, vytvoriť vnútorné a vonkajšie polia, ako aj vyšpecifikovať pomer signálu k šumu. S cieľom vytvoriť Taguchiho návrh sa zvolí v hlavnom menu rovnomenná platforma. Vo vstupnom okne ponuky sa zvolí cieľ pre ozvu (Response), pomenujú sa faktory a priradí sa im úloha, signálny faktor (signal factor) alebo šumový faktor (noise factor). Následne sa z ponúkaných vnútorných a vonkajších polí zvolia požadované návrhy. Následne program vytvorí návrh tabuľky údajov, aby sa mohlo začať s experimentovaním.

3. Interval spoľahlivosti

Ak je návrh konečný, je možné získať výsledok aj v podobe odhadu strednej hodnoty, ktorá predstavuje odhad skutočnej hodnoty zisťovanej veličiny. Rovnako Taguchi definoval tri rôzne intervaly spoľahlivosti okolo strednej hodnoty:

- interval spoľahlivosti pre priemernú hodnotu úrovne faktora
- predikčný interval okolo odhadu priemernej hodnoty pokusu
- interval okolo odhadu priemernej hodnoty pokusu, ktorý sa použije v potvrdzujúcom experimente na overenie predikcie.

Prvý z troch intervalov spoľahlivosti okolo strednej hodnoty úrovne faktora sa vypočíta:

$$\mu_{A_1} = \bar{A}_1 \pm CI_1$$

alebo

$$\bar{A}_1 - CI_1 \leq \mu_{A_1} \leq \bar{A}_1 + CI_1$$

pričom $\bar{A}_1$ je priemerná hodnota faktora A pri úrovni 1 a

$$CI_1 = \sqrt{\frac{F_\alpha(1, \nu_2) \cdot s_e^2}{n}},$$

kde s_e^2 je rozptyl náhodnej chyby, n je počet opakovaní v danom pokuse a $F_\alpha(1, \nu_2)$ je príslušný α kvantyl Fisherovho rozdelenia pravdepodobnosti so stupňami voľnosti 1 a ν_2 , ν_2 sú stupne voľnosti asociované so stupňami voľnosti pre chybu ν_e .

Interval spoľahlivosti pre strednú hodnotu pokusu uvažuje s nasledovným intervalom spoľahlivosti:

$$CI_2 = \sqrt{\frac{F_\alpha(1, \nu_e) \cdot s_e^2}{n_{eff}}},$$

kde s_e^2 je rozptyl náhodnej chyby, n_{eff} je efektívny počet opakovaní v danom pokuse

$$n_{eff} = \frac{N}{1 + [\nu_A + \nu_B + \dots + \nu_F]},$$

pričom v menovateli zlomku v zátvorke je celkový počet stupňov voľnosti v spojitosti s odhadom strednej hodnoty, konkrétne ν_A sú stupne voľnosti faktora A , ktoré sa vypočítajú ako počet úrovní faktora A mínus jedna a $F_\alpha(1, \nu_e)$ je α kvantyl Fisherovho rozdelenia pravdepodobnosti so stupňami voľnosti 1 a ν_e .

Predikčný interval pre nový pokus použiteľný na overenie správnosti nastavení sa dá určiť podľa nasledujúceho vzťahu:

$$CI_3 = \sqrt{F_\alpha(1, \nu_e) \cdot s_e^2 \left[\frac{1}{n_{eff}} + \frac{1}{r} \right]},$$

kde s_e^2 je rozptyl náhodnej chyby, n_{eff} je efektívny počet opakovaní v danom pokuse a r predstavuje veľkosť výberu v prognózovanom pokuse.

4. Príklad aplikácie Taguchiho polí do systému SAS JMP

Príklad popísal Byrne a Taguchi v roku 1986 na štyridsiatom výročnom kongrese kontroly kvality v Milwaukee, WI v USA. Cieľom návrhu experimentu je nájsť také

nastavenia kontrolovaných faktorov, aby sa docielilo maximalizácie priľnavosti a minimalizácie nákladov na montáž nylonových výrobkov. Experiment definuje sedem faktorov, tri šumové (s dvomi úrovňami) a štyri kontrolované (každý s tromi úrovňami).

DOE- Taguchi Arrays

Taguchi Design

Response

Response Name	Goal	Lower Limit	Upper Limit	Importance
Y	Larger Is Better	.	.	.

Factors

Name	Role	Values
Interfer	Signal	1 2 3
Wall	Signal	1 2 3
Depth	Signal	1 2 3
Adhesive	Signal	1 2 3
Time	Noise	L1 L2
Temperature	Noise	L1 L2
Humidity	Noise	L1 L2

Taguchi Design
7Factors

Choose Inner and Outer Array Designs

Inner Array		Outer Array
Number	Design Name	Design Name
9	L9 - Taguchi	L4
27	L27 - Taguchi	L8
81	Full Factorial	

Continue
Back

Obrázok 1: Zadanie faktorov, cieľa a vytvorenie Taguchiho polí.
Ide o návrh so siedmimi faktormi L9 - vnútorné pole a L8 - vonkajšie pole.

Byrne Taguchi Data

	Interfer	Wall	Depth	Adhesive	Pattern	Y---	Y--+	Y+--	Y+--	Y+--	Y+--	Y+--	Y+--	Mean Y	St Ratio Y
1	1	1	1	1	----	15,6	9,5	16,9	19,9	19,6	19,6	20	19,1	17,525	24,02534
2	1	2	2	2	-000	15	16,2	19,4	19,6	19,7	19,8	24,2	21,9	19,475	25,52164
3	1	3	3	3	+++	16,3	16,7	19,1	15,6	22,6	18,2	23,3	20,4	19,025	25,33476
4	2	1	2	3	0-0+	18,3	17,4	18,9	18,6	21	18,9	23,2	24,7	20,125	25,90425
5	2	2	3	1	00+-	19,7	18,6	19,4	25,1	25,6	21,4	27,5	25,3	22,825	26,90753
6	2	3	1	2	0+-0	16,2	16,3	20	19,8	14,7	19,6	22,5	24,7	19,225	25,32574
7	3	1	3	2	++0	16,4	19,1	18,4	23,6	16,8	18,6	24,3	21,6	19,85	25,71081
8	3	2	1	3	+0-+	14,2	15,6	15,1	16,8	17,8	19,6	23,2	24,4	18,3375	24,83231
9	3	3	2	1	++0-	16,1	19,9	19,3	17,3	23,1	22,7	22,6	28,6	21,2	26,15198

Obrázok 2: Tabuľka návrhu experimentu aj s údajmi.

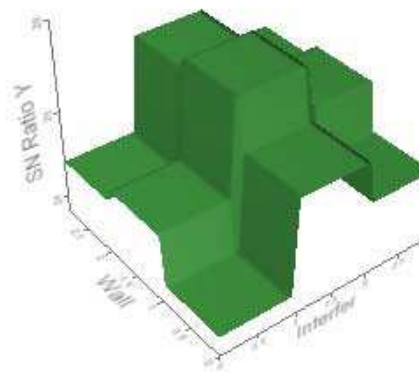
G. Taguchi navrhol procedúru, ktorú je možné všeobecne aplikovať pri Taguchiho návrhoch na identifikovanie nastavení parametrov, ktoré maximalizujú výkonovú charakteristiku:

- určenie najvhodnejších a iných porovnateľne výhodných nastavení parametrov návrhu a identifikovanie významných šumových faktorov

- skonštruovanie vnútorného a vonkajšieho poľa a vytvorenie návrhu experimentu parametrov
- uskutočnenie návrhu experimentu parametrov
- vyjadrenie nového nastavenia parametrov podľa výkonovej charakteristiky
- potvrdenie, že nové nastavenia naozaj zlepšujú výkonovú štatistiku (najčastejšie experimentálne). Postup je aplikovaný s menšími obmedzeniami aj v tomto príklade.

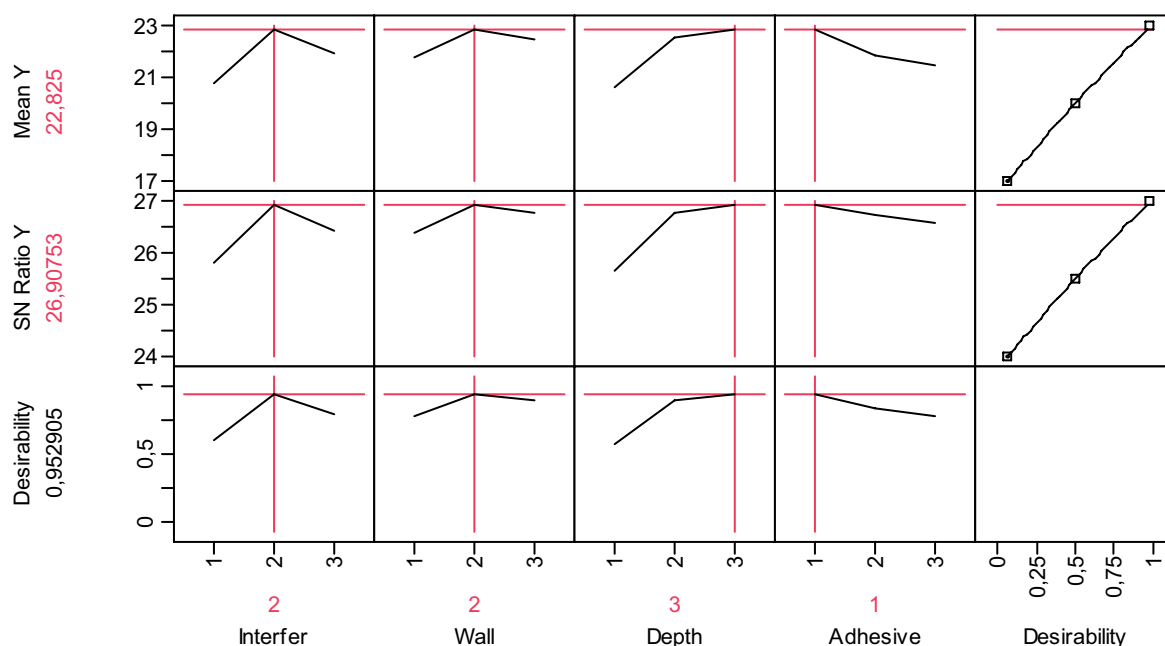
Tabuľka 2: Analýza rozptylu pre SN pomer.

Zdroj	Df	Súčet štvorcov	Priemerné štvorce	F štatistika
Model	8	5,2875264	0,660941	.
Chyba	0	0,0000000	.	Prob > F
C. Total	8	5,2875264		.



Obr. 3: Vľavo 3D vyobrazenie premennej SN pomer.

Štatistický modul SAS JMP poskytol výsledky metódou najmenších štvorcov, vo forme výpočtov aj grafických výstupov. Vo funkcii 'Prediction profiler' (profil predikcií) možnosť zvoliť funkciu 'Maximalize desirability' (maximalizácia funkcie vhodnosti nastavení úrovni faktorov), ktorá umožňuje nájsť také nastavenia faktorov (predikciu), aby získané výsledky boli, čo možno najžiadanejšie a to vzhľadom na 'Desirability function' (funkcia vhodnosti). Pomer SN sa touto maximalizáciou zmenil z hodnoty 24,02534 na 26,90753. Maximalizácia zmenila nastavenie faktora interferencia a stena z úrovne L1 na L2, nastavenie faktora hĺbka z L1 na L3. Posledný štvrtý kontrolovaný faktor svoje nastavenie nezmenil. Takto sa získalo požadované nastavenie kontrolovaných faktorov tak, aby boli minimalizované náklady a maximalizovaná priľnavosť nylonových výrobkov.



Obrázok 4: Funkcia Prediction profiler okrem účinku faktorov umožňuje nájsť aj najvhodnejšie nastavenie, resp. vo voľbe Desirability functions (funkcia vhodnosti nastavení) získať požadované maximum tejto funkcie.

5. Záver

Navrhovanie experimentov je metodika, ktorá sa používa pred samotným uskutočnením experimentu, pred tým ako sa vôbec začne niečo testovať, merať, skúšať alebo sledovať. Experimentovanie sa používa za účelom pochopenia a zlepšenia systému. Príspevok prináša základy tvorby Taguchiho návrhov a všetky významné kroky, ako vytvorenie ortogonálneho poľa a pod. vysvetľuje aj pomocou štatistického softvéru SAS JMP.

Najvhodnejšie nastavenie z aplikačného príkladu je v pokuse č. 5 podľa obr. 2, kedy je hodnota ozvy najvyššia (vyše 22 a SN pomer vyše 26). Pre danú hodnotu je možné vypočítať podľa podkapitoly 3 interval spoľahlivosti pre jeden pokus v 5 riadku, ktorý sa rovná:

- na $\alpha = 0,05$ pre ozvu $22,03 \pm 3,60$ m.j.,
- na $\alpha = 0,01$ pre ozvu $22,03 \pm 5,24$ m.j.

V závere príspevku je treba poznamenať, že Taguchi robustný návrh je možné aplikovať na statické systémy. Rovnako však je aplikovateľný aj na dynamické systémy. Dynamickým označujeme taký systém, kedy výstupná premenná závisí od vstupnej premennej (signálu). Príkladom je brzdomý systém automobilu, kedy vodič generuje stlačením brzdy vstupný signál, ktorý sa lineárne prenáša ďalej v podobe brzdného účinku (výstupný signál), a pod.

Podakovanie:

Príspevok vznikol s príspevom grantovej agentúry VEGA v rámci projektov 1/0437/08 Kvantitatívne metódy v stratégii šesť sigma a 1/3182/06 Zlepšovanie kvality produkcie strojárskych výrobkov pomocou moderných štatistických metód.

5. Literatúra

- [1] Byrne, D. M., Taguchi, G.: ASQC 40th Anniversary Quality Control Congress Transactions, Milwaukee, WI: American Society of Quality Control, 168-177.
- [2] Dehnad, K.: Quality Control, Robust Design, and the Taguchi Method (edited by), Wadsworth and Brooks/Cole 1989. (ISBN 0-534-09048-6)
- [3] Janiga, I., Garaj, I. One-sided tolerance factors of normal distributions with unknown mean and variability. In: Measurement Science Review. ISSN 1335-8871. Vol. 6, No. 2 (2006), s. 12-16.
- [4] Horníková, A.: Navrhovanie experimentov so štatistickým softvérom SAS. 14. Slovenská štatistická konferencia Regionálna štatistika, Strečno 17.-19. 9. 2008.
- [5] Müllerová, A.: Taguchiho prístup k štatistickému riadeniu kvality (dizertačná práca) Ekonomická univerzita v Bratislave, 2002.
- [6] Ross, F. J.: Taguchi Techniques for Quality Engineering (Loss Function, Orthogonal Experiments, Parameter and Tolerance Design, McGraw-Hill Book Company 1988.
- [7] Taguchi, G.: Introduction to Quality Engineering (Designing Quality into Products and Processes), Asian Productivity Organization, 1986. (ISBN 92-833-1083-7)
- [8] TEREK, M., HRNČIAROVÁ, L.: Štatistické riadenie kvality, Edícia ekonómia, 2004.
- [9] TEREK, M., HRNČIAROVÁ, L.: Výberové skúmanie. Vydavateľstvo EKONÓM, 2008.
- [10] Tošenovský J., Noskovičová, D.: Štatistické metódy pro zlepšování jakosti. Montanex a.s. 2000. (ISBN 80-7225-040-X)

Adresa autorky:

Adriana Horníková, Ing., Mgr., PhD.,
Fakulta hospodárskej informatiky, Katedra štatistiky
Ekonomická univerzita v Bratislave
Dolnozemska cesta 1/B, 853 25 Bratislava
E-mail: ahornik@euba.sk

Možnosti využitia MS Excel pri operatívnom sledovaní kvality výrobného procesu

Possibilities of using MS Excel to operative control of manufacturing process quality

Stella Hrehová

Abstract: In this paper are presented the possibilities of using MS Excel such as a tool of operative process control. There is shown a connection between MS Excel and the internet explorer in time. The principle of this method is using the empty cells to prepare the graph of manufacturing process which will be gradually completed.

Key words: Quality, MS Excel, Process, Control

Kľúčové slová: kvalita, MS Excel, proces, riadenie.

1. Úvod

Štatistické riadenie procesov aplikujeme na kritické znaky, ktoré sú dôležité či už z hľadiska funkčnosti výrobku, alebo sú požadované priamo odberateľom. Dôležité je stanoviť si jasné pravidlá a postupy, ako sa zachovať v prípade prekročenia povolených hraníc. Čoraz častejším štandardom je stanovenie dvoch regulačných medzí pre rôzne sigma úrovne. Podľa prekročenia týchto regulačných medzí sa rozhoduje o tom, či je potrebné danú dávku aj triediť, alebo je potrebné iba regulačný zásah.

Pri hodnotení kvality výrobného procesu sa vytvára regulačný graf, ktorý v grafickej podobe popisuje schopnosť procesu dodržiavať požadované tolerančné hranice. Používanie regulačných diagramov sa člení na dve etapy.

V prvej etape je úlohou zostrojiť regulačný graf, určiť jeho :

- dolnú tolerančnú medzu,
- hornú tolerančnú medzu.

V rámci týchto tolerančných pásiem sa môže pohybovať skúmaný parameter.

Druhou etapou je vlastné monitorovanie priebehu výrobného procesu. V prípade, že hodnota regulovaného parametra padne mimo regulačného pásma, výrobca zisťuje príčinu tohto javu a prijíma opatrenia na návrat parametra medzi regulačné medze. V priebehu monitorovania možno tiež sledovať priebeh regulovaných hodnôt a podľa ich vývoja možno prijať opatrenia na korekciu tohto smeru vývoja skôr, než nastane mimoriadna situácia. Je predpoklad, že zabehnutá výroba už má vytvorený regulačný graf. Pre takýto prípad, je možné použitie programového prostriedku MS Excel.

2. Spôsoby hodnotenia regulovaných hodnôt

Hodnotenie regulovaných hodnôt je možné vykonávať dvoma možnosťami.

- off-line techniky. Tieto techniky sú založené na hodnotení parametrov až po ich vyhodnotení, teda až po určitom čase. V mnohých prípadoch však tieto techniky nie sú vhodné pre plynulé, kontinuálne systémy [6].

- on-line techniky. Pri takomto prístupe nie sú definované len regulačné hranice, ale sú tiež definované pravidlá ako regulovať proces tak, aby boli regulované parametre v rámci regulačných medzí. Tieto techniky sa ukazujú byť efektívne pre regulovanie procesu, avšak ich použitie si vyžaduje špecifický softvér a hardvér a samozrejme sú aj finančné náročné. Pre použitie u malých a stredných firiem by to zrejme nebolo efektívne.

V takýchto prípadoch je možné použitie konvenčného prostriedku, ktorý je schopný vytvárať ilúziu on-line hodnotenia kvality výrobného procesu. Predpokladom je, že regulačný graf je vytváraný pomocou tabuľkového prostriedku MS Excel.

Pre potreby využitia jeho prostredia je dôležitá jeho vlastnosť prepojenia zdrojovej tabuľky pre graf a samotného grafického zobrazenia. V prípade takéhoto použitia je potrebné dopredu myslieť na tento fakt a označiť oblasť pre vytvorenie grafu aj z prázdnych buniek, ktoré sa potom priebežne pri kontrolách budú dopĺňať. Takto zostrojený graf sa bude sám aktualizovať po doplnení hodnoty do zdrojovej tabuľky buď pri každom otvorení, alebo pri každom stlačení tlačidla "Obnoviť" (Refresh) na paneli nástrojov internetového prehliadača.

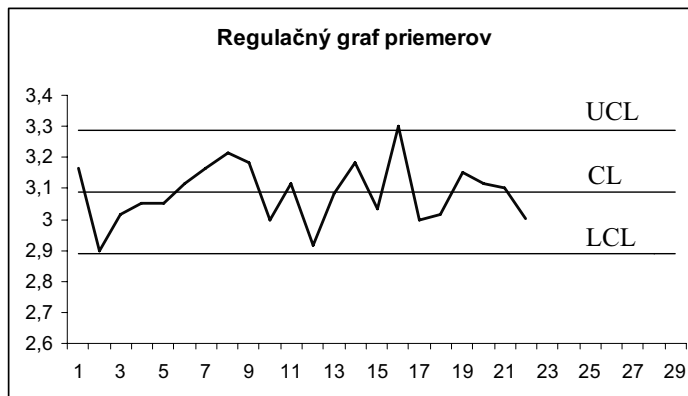
Nasledujúci obrázok ukazuje zostrojený regulačný graf s použitím prázdnych buniek.

Vytvoríme graf priemerov z nasledujúcich hodnôt.

Tabuľka 1 : Zdrojová tabuľka

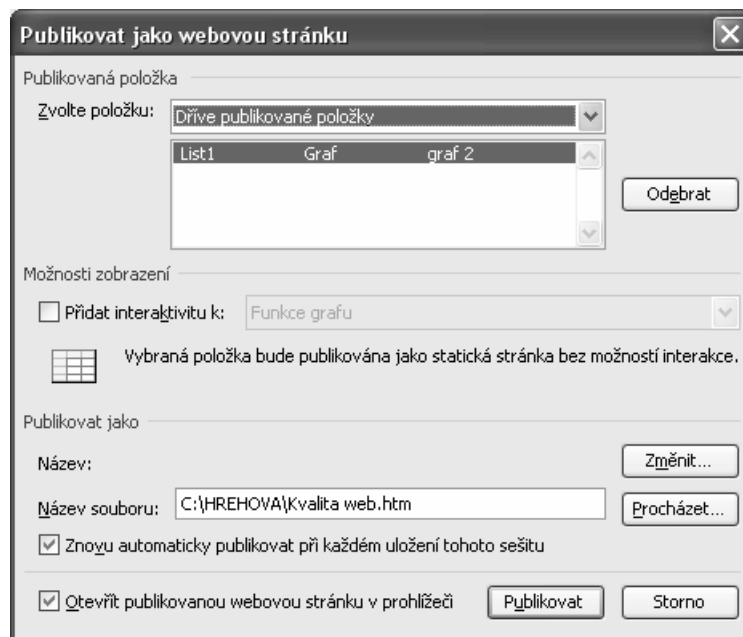
P.č.	Priemery						Rozpätie	CL	UCL	LCL
1	2,8	3,4	3,1	3	2,8	2,7	2,966667	0,7	3,091867	2,890969
2	3,4	3,3	3,3	3,2	3,1	3,3	3,266667	0,3	3,091867	2,890969
3	3,1	2,9	3,2	3,2	3,3	3,3	3,166667	0,4	3,091867	2,890969
4	3	2,8	2,7	2,7	3,2	3	2,9	0,5	3,091867	2,890969
5	3	3	3,1	2,9	3,1	3	3,016667	0,2	3,091867	2,890969
6	2,9	2,9	3,2	2,9	3,3	3,1	3,05	0,4	3,091867	2,890969
7	2,8	3	3,2	2,9	3,2	3,2	3,05	0,4	3,091867	2,890969
8	3	3,2	3,4	3	3	3,1	3,116667	0,4	3,091867	2,890969
9	3,2	3,4	3,2	2,9	3,1	3,2	3,166667	0,5	3,091867	2,890969
10	3,5	3,4	3,2	3,1	3	3,1	3,216667	0,5	3,091867	2,890969
11	3	3,2	3,3	3,3	3,2	3,1	3,183333	0,3	3,091867	2,890969
12	3,2	3,1	3,1	2,9	2,9	2,8	3	0,4	3,091867	2,890969
13	3	3,1	3,2	3,2	3,1	3,1	3,116667	0,2	3,091867	2,890969
14	3	3	2,5	3,1	3	2,9	2,916667	0,6	3,091867	2,890969
15	3	3,3	3,1	3	3	3,1	3,083333	0,3	3,091867	2,890969
16	3,4	3,2	3,2	3,1	3,1	3,1	3,183333	0,3	3,091867	2,890969
17	3,2	3,2	3,1	3	2,9	2,8	3,033333	0,4	3,091867	2,890969
18	3,2	3,2	3,2	3,4	3,4	3,4	3,3	0,2	3,091867	2,890969
19	3,2	3,2	2,7	3	2,9	3	3	0,5	3,091867	2,890969
20	3	2,9	3,2	3	3	3	3,016667	0,3	3,091867	2,890969
21	3,3	3,3	3	3	3,1	3,2	3,15	0,3	3,091867	2,890969
22	3,1	2,8	3,4	3	3,3	3,1	3,116667	0,6	3,091867	2,890969
23	2,9	2,9	3,1	3,2	3,4	3,1	3,1	0,5	3,091867	2,890969
24	3,1	3	2,7	2,8	3,3	3,1	3	0,6	3,091867	2,890969
									3,091867	2,890969
									3,091867	2,890969
									3,091867	2,890969
									3,091867	2,890969
									3,091867	2,890969

Na základe týchto získaných hodnôt bol vytvorený regulačný graf priemerov v prostredí MS Excel, pričom boli použité prázdne bunky (sivo vyplnené), ktoré sa rezervovali na ďalšie dopĺňanie hodnôt podľa toho, ako sa bude proces vyvíjať. Hodnoty sú určené automaticky z nameraných hodnôt pomocou vzorca, ktorý je skopírovaný aj do týchto buniek. Podobne sú určené hodnoty tolerančných hraníc. Je potrebné však dôsledne dodržiavať odkazy na bunky, aby sa ich hodnoty menili podľa aktuálneho stavu.



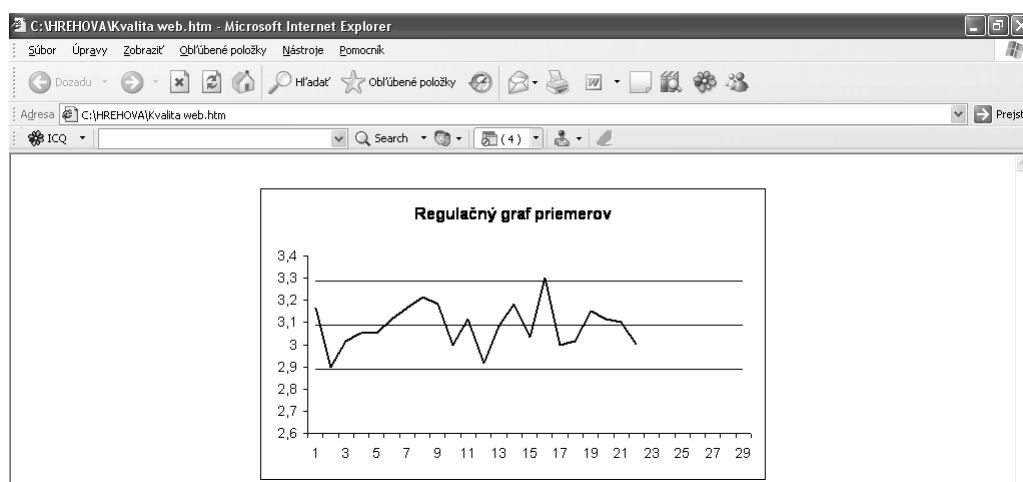
Obrázok 1: Zostrojený regulačný graf

Pre vytvorenie webovskej stránky použijeme voľbu MS Excel a to uloženie súboru v tvare web stránky. Z možnosti si vyberieme len položku grafu, ktorý má pre nás najvyššiu vypovedaciu hodnotu a prípadne si dáme zobrazit' aj vytvorenú stránku, ktorú si uložíme.



Obrázok 2: Okno pre aktiváciu grafu v internetovom prehliadači

Pretože predpokladáme, že graf sa bude aktualizovať po nejakom čase, je potrebné zaškrtnúť **Znovu automaticky publikovať pri každém uložení tohoto sešitu**. Po otvorení tejto stránky budú zobrazené hodnoty už podľa novej situácie a je možné rozhodnúť či proces je stabilný, alebo nie.



Obrázok 3: Zobrazenie webovskej stránky

3. Záver

V príspevku bol uvedený možný spôsob využitia MS Excel, ktorý umožňuje ukladanie súborov aj ako web súbory a teda je ich možné zobraziť v internetovom prehliadači. Spojením tejto vlastnosti a vlastnosti prepojenia oblasti z ktorej vytvárame graf sa z MS Excelu stáva použiteľný prostriedok pre operatívne riadenie kvality výrobného procesu. Možnosťou tvorby web stránky, ktorá by zobrazovala regulačný graf má zodpovedný pracovník možnosť operatívne riadiť tento proces. Okrem tejto možnosti využitia sa táto informácia môže poskytovať aj vyššiemu manažmentu pri potrebe hodnotenia výrobného procesu.

4. Literatúra

- [1]FLOREKOVÁ, L. 1998. Metódy štatistického hodnotenia kvality SPC, Acta Montanistica Slovaca, roč. 3, č.1, Edičné stredisko na fakulte BERG Košice 1998, ISSN 1335-1788
- [2]CHAJDIK, J. 1998: Štatistické riadenie kvality, Statis Bratislava, ISBN 80-85659-12-3
- [3]NENÁDAL, J. A KOL. 2007: Moderní systémy řízení jakosti. Quality management., Management Press, Praha 2007, ISBN 978-80-7261-071-6
- [4]ŠESTÁK, M. 2007 : LEAN Six Sigma a štatistická regulácia procesov, Kvalita, MASM Žilina 2007, roč. XV, č.2-2007, ISSN 1335-9231
- [5]TEREK, M.; HRNČIAROVÁ, L. 2004: Štatistické riadenie kvality, IURA EDITION spol s r.o., Bratislava 2004, ISBN 80-89047-97-1
- [6][HTTPS://WWW.GOLDPRACTICES.COM//PRACTICES/SPC/INDEX.PHP](https://www.goldpractices.com/practices/spc/index.php)

Adresa autora (-ov):

Stella Hrehová, Ing. PhD.
FVT TU v Košiciach
Katedra matematiky, informatiky
a kybernetiky
Bayerova 1
080 01 Prešov
stella.hrehova@tuke.sk

Príspevok bol vypracovaný v rámci projektu : "Asymptotické a oscilatorické vlastnosti riešení nelineárnych diferenciálnych systémov", č. 1/4004/07

Štatistické riadenie kvality. Analýza spôsobilosti procesu. Statistical Quality Control. Process Capability Analysis.

Ľubica Hrnčiarová, Milan Terek

Abstract: Statistical techniques are able to be much useful before process activities, to modelling of process variation, to analysis of this variation in respect of demands or specifications and also they can help to its reduction. This common activity is known as process capability analysis. The content of this contribution is short description of the techniques, which are possible to use in process capability analysis.

Key words: Statistical Quality Control, Process Quality Analysis, Histograms, Control Charts, Process Capability Indices, Experimental Design.

Kľúčové slová: štatistické riadenie kvality, analýza spôsobilosti procesu, histogramy, regulačné diagramy, indexy spôsobilosti procesu, navrhovanie experimentov.

1. Úvod

Na dnešných svetových trhoch, s množstvom konkurencieschopných firiem, otázka kvality má významnejšie postavenie, ako v minulosti. Efektívne *programy zlepšovania kvality* sú nástrojom udržiavania, resp. zlepšovania pozície podnikov na trhu. Integrálnou súčasťou týchto programov je *analýza spôsobilosti procesu*.

Obsahom príspevku je stručný prehľad techník, ktoré možno použiť v *analýze spôsobilosti procesu*.

Získané *informácie o spôsobilosti procesu* sú:

1. *pre výrobcov* veľmi dôležitým podkladom pre kvalifikované rozhodnutia pri plánovaní kvality.
2. *pre zákazníka* dôkazom o tom, či výrobok bol vyrobený v stabilných výrobných podmienkach a boli dodržané predpísané kritéria kvality.

2. Analýza spôsobilosti procesu pomocou histogramu

Histogram sa používa na grafické znázornenie rozdelenia početností hodnôt znaku. V riadení kvality možno túto grafickú metódu použiť pri analýze spôsobilosti procesu.

Ak budeme mieru spôsobilosti definovať ako 6σ a možno prijať predpoklad o pomernej stabilite rozdelenia početností, potom 6σ rozpätie rozdelenia pravdepodobnosti ukazovateľa kvality centrované na strednej hodnote μ možno odhadnúť pomocou výberového priemeru $\bar{X}$ a výberovej smerodajnej odchýlky S ako $\bar{X} \pm 3S$.

Použitie histogramu na odhad spôsobilosti procesu umožňuje získať okamžitú vizuálnu predstavu o spôsobilosti procesu, ale nie vždy o potenciálnej spôsobilosti. Príslušné rozdelenie môže byť napríklad ovplyvnené vymedziteľnými príčinami, elimináciou ktorých sa môže variabilita ukazovateľa kvality redukovať. Na redukciu variability eliminovaním vymedziteľných príčin variability je vhodné použiť regulačné diagramy.

3. Analýza spôsobilosti procesu pomocou regulačných diagramov

Regulačné diagramy sa považujú za primárnu techniku analýzy spôsobilosti [6].

Analýza spôsobilosti pomocou regulačného diagramu môže indikovať, že proces sa dostal do štatisticky nezvládnutého stavu. Vtedy nie je vhodné odhadovať spôsobilosť procesu. Stabilita procesu je nevyhnutným predpokladom na získanie spoľahlivého odhadu

spôsobilosti procesu. Keď je proces nestabilný, prvou úlohou je samozrejme nájsť a eliminovať vymedziteľné príčiny variability.

V analýze spôsobilosti možno využiť 3σ - regulačné diagramy na reguláciu meraní aj 3σ - regulačné diagramy na reguláciu porovnávaním.

Pri aplikácii regulačných diagramov je nevyhnutné stanoviť rozsah výberu n , frekvenciu výberov a parametre regulačného diagramu (regulačné hranice UCL a LCL, centrálna priamka CL).

Vzťahy na výpočet hornej regulačnej hranice UCL a dolnej regulačnej hranice LCL v 3σ - regulačnom diagrame sú nasledovné:

$$UCL = \mu_V + 3\sigma_V$$

$$LCL = \mu_V - 3\sigma_V$$

V uvedených vzťahoch V je výberová charakteristika, pomocou ktorej odhadujeme nejaký parameter rozdelenia pravdepodobnosti regulovaného ukazovateľa kvality a predpokladáme, že V má strednú hodnotu μ_V a smerodajnú odchýlku σ_V .

Konštrukcie základných typov Shewhartových regulačných diagramov na reguláciu meraní a na reguláciu porovnávaním možno nájsť napríklad v [1] [6], [8] [9]. Podľa [6], by sa mali vždy keď je to možné uprednostniť regulačné diagramy priemerov a regulačné diagramy rozpätí, pretože poskytujú lepšiu informáciu v porovnaní s regulačnými diagramami na reguláciu porovnávaním, aj keď aj tieto môžu byť v analýze spôsobilosti užitočné.

Pomocou regulačných diagramov priemerov a regulačných diagramov rozpätí možno analyzovať proces bez ohľadu na tolerančné hranice. Naopak, napríklad regulačné diagramy podielu nezhodných jednotiek umožňujú analyzovať proces len s ohľadom na tolerančné hranice [8].

4. Analýza spôsobilosti procesu pomocou indexov spôsobilosti

Na hodnotenie spôsobilosti procesu sa najprv používali histogramy. Požiadavkou manažérov z praxe bolo, aby to, čo prezrádza histogram, bolo možné vyjadriť nejakým číslom (ukazovateľom). To bol základ pre vznik indexov spôsobilosti procesu (*Process Capability Indices*). V čitateli indexov je požadovaná (predpísaná) presnosť procesu, v menovateli inherentná (prirodzená) variabilita procesu.

Pri výpočte indexov spôsobilosti sa všeobecne predpokladá že proces je stabilný, pozorovania sú štatisticky nezávislé a majú normálne rozdelenie.

J. M. Juran, so svojimi spolupracovníkmi v roku 1974 navrhol index spôsobilosti procesu C_p :

$$C_p = \frac{USL - LSL}{6\sigma} \quad (1)$$

kde USL je horná tolerančná hranica, LSL je dolná tolerančná hranica a σ je smerodajná odchýlka procesu.

Index spôsobilosti prvej generácie C_p je iba potencionálnou mierou spôsobilosti procesu, pretože nevyčísluje, kde je proces centrováný. Pri používaní indexu C_p iba

predpokladáme, že stredná hodnota procesu μ sa rovná cieľovej hodnote T sledovaného ukazovateľa kvality.

Od začiatku používania indexu spôsobilosti C_p sa používa hodnota 1,33 za minimálne požadovanú hodnotu spôsobilosti procesu.

Ďalším *indexom spôsobilosti prvej generácie* je C_{pk} . Tento index už reaguje na vychýlenie strednej hodnoty procesu μ od stredu tolerančného intervalu procesu MSL, resp. od cieľovej hodnoty T sledovaného ukazovateľa kvality, keď $T = \text{MSL}$.

Je ukazovateľom aktuálnej spôsobilosti procesu. Index C_{pk} kvantifikuje spôsobilosť „horšej polovice“ údajov procesu. Vyplýva to z toho, že

$$C_{pk} = \min(C_{pu}, C_{pl}), \quad (2)$$

kde $C_{pu} = \frac{\text{USL} - \mu}{3\sigma}$ a $C_{pl} = \frac{\mu - \text{LSL}}{3\sigma}$. C_{pu} je horný index spôsobilosti a C_{pl} je dolný index spôsobilosti pri zadaní jednostranného tolerančného intervalu (teda buď $\text{LSL} = -\infty$ (resp. 0) alebo $\text{USL} = \infty$) a minimálne C_{pu} , resp. C_{pl} obsahuje vyšší podiel hodnôt sledovaného ukazovateľa kvality mimo požadovaných tolerančných hraníc USL, resp. LSL.

Kde je proces centrováný, možno zistiť napríklad z nasledujúcich vzťahov:

keď $C_p = C_{pk}$, proces je centrováný v strede tolerančného intervalu,

keď $C_p > C_{pk}$, proces nie je ideálne centrováný,

keď $C_{pk} = 0$, proces je centrováný na jednej z tolerančných hraníc,

keď $C_{pk} < 0$, proces je centrováný mimo tolerančných hraníc.

V prípade, keď sa stredná hodnota procesu μ približuje k niektorej z tolerančných hraníc procesu (LSL, resp. USL) a súčasne sa znižuje smerodajná odchýlka procesu σ , nereaguje ani C_{pk} na vychýlenie strednej hodnoty procesu μ od stredu tolerančného intervalu procesu MSL, resp. keď $T = \text{MSL}$ od cieľovej hodnoty T sledovaného ukazovateľa kvality.

Keď $T \neq \text{MSL}$, C_{pk} vôbec nevystihuje odchýlku μ od cieľovej hodnoty T .

Kde je proces centrováný možno zistiť aj pomocou *indexu prvej generácie* K . Tento index sa niekedy nazýva ukazovateľ presnosti nastavenia a je definovaný vzťahom:

$$K = \frac{|T - \mu|}{\frac{1}{2}(\text{USL} - \text{LSL})} \quad (3)$$

Keď $T = \text{MSL}$, index K zachytáva odchýlku strednej hodnoty od stredu tolerančného intervalu. Keď $T \neq \text{MSL}$, index K zachytáva odchýlku strednej hodnoty od požadovanej hodnoty, čo je veľmi užitočná informácia.

Obor hodnôt indexu K je $[0, \infty)$. Keď:

4. $K = 0$, potom $\mu = T$,
5. $0 < K < 1$, potom μ sa nachádza medzi LSL a USL,
6. $K > 1$, potom μ sa nachádza mimo tolerančných hraníc.

Pretože index K nevyjadruje vzťah veľkosti prirodzeného tolerančného rozpätia k predpísanému tolerančnému rozpätiu, je vhodné ho používať spolu s niektorým zo skôr uvedených indexov.

Konštrukcia indexov prvej generácie vychádza z klasického prístupu k riadeniu kvality, podľa ktorého sa výrobky vyrobené v rozsahu požadovanej tolerancie považujú za zhodné a výrobky vyrobené mimo požadovanej tolerancie za nezhodné. Klasické štatistické riadenie kvality je však podriadené programom zlepšovania kvality zameraným na výrobu výrobkov, ktorých ukazovatele kvality by sa rovnali cieľovej hodnote.

Použitím Taguchiho stratovej funkcie boli indexy spôsobilosti procesu prvej generácie modifikované na *indexy spôsobilosti procesu druhej generácie*.

Index C_{pm} je definovaný takto

$$C_{pm} = \frac{USL - LSL}{6\psi} \quad (4)$$

kde ψ je druhá odmocnina zo ψ^2 .

ψ^2 vyjadruje očakávanú mieru kolísania ukazovateľa kvality X okolo cieľovej hodnoty T a je definovaná:

$$\psi^2 = E[(X - T)^2] = \sigma^2 + (\mu - T)^2$$

Keď sa rozptyl σ^2 zväčšuje a (alebo) sa stredná hodnota procesu μ vzdďaľuje od cieľovej hodnoty T , menovateľ indexu rastie a C_{pm} klesá. Obyčajne sa predpokladá, že cieľová hodnota je v strede tolerančného rozpätia ($T = MSL$). Keď $T \neq MSL$, index C_{pm} nie je vhodnou mierou spôsobilosti procesu. Keď $T = \mu$, potom $C_{pm} = C_p$.

Keď proces nie je centrovaný na cieľovú hodnotu ($T \neq MSL$), index C_{pm} nie je vhodnou mierou spôsobilosti procesu. V tomto prípade je vhodné použiť *index druhej generácie* C_{pm}^* .

Je definovaný takto

$$C_{pm}^* = \frac{\min(USL - T; T - LSL)}{3\sqrt{\sigma^2 + (\mu - T)^2}} \quad \text{alebo} \quad C_{pm}^* = \min(C_{pml}, C_{pmu}), \quad (5)$$

$$\text{kde} \quad C_{pml} = \frac{T - LSL}{3\sqrt{\sigma^2 + (\mu - T)^2}} \quad \text{a} \quad C_{pmu} = \frac{USL - T}{3\sqrt{\sigma^2 + (\mu - T)^2}}$$

Keď $T = MSL$, potom $C_{pm}^* = C_{pm}$. Keď sa okrem toho $\mu = MSL$, potom $C_{pm} = C_{pm}^* = C_{pk}$.

Nevýhodou indexu C_{pm}^* , podobne ako C_{pk} je, že hoci poskytuje informáciu, že proces je aktuálne nespôsobilý, neposkytuje informáciu o jej príčinách. Preto je vhodné spolu s indexom C_{pm}^* sledovať aj index K .

Medzi *indexy tretej generácie* patria napríklad indexy C_{pmk} , C_{jpk} .

Index C_{pmk} , je kombináciou indexov C_{pk} a C_{pm} a je definovaný takto

$$C_{pmk} = \frac{\min(USL - \mu, \mu - LSL)}{3\sqrt{\sigma^2 + (\mu - T)^2}} \quad (6)$$

Výhodou tohto indexu je, že je citlivejší na variabilitu okolo cieľovej hodnoty ako indexy C_{pk} a C_{pm}

Ďalší index spôsobilosti procesu tretej generácie meria variabilitu ukazovateľa kvality X pred aj za cieľovou hodnotou. Index C_{jpk} možno použiť aj v prípade, keď rozdelenie pravdepodobnosti X je zošikmené. Index C_{jpk} je definovaný takto

$$C_{jpk} = \min(CU_{jpk}, CL_{jpk}). \quad (7)$$

Indexy CU_{jpk} , CL_{jpk} definujú spôsobilosť procesu na základe hornej a dolnej polovice údajov:

$$CU_{jpk} = \frac{1}{3\sqrt{2}} \left[\frac{USL - T}{\sqrt{E_{X>T}[(X - T)^2]}} \right] \quad CL_{jpk} = \frac{1}{3\sqrt{2}} \left[\frac{T - LSL}{\sqrt{E_{X\leq T}[(X - T)^2]}} \right], \text{ kde}$$

$$E_{X>T}[(X - T)^2] = E[(X - T)^2 | X > T]P(X > T) \quad E_{X\leq T}[(X - T)^2] = E[(X - T)^2 | X \leq T]P(X \leq T)$$

Podrobnejšie informácie o indexoch spôsobilosti možno nájsť napríklad v [3] [4], [5], [6], [8], [9].

5. Analýza spôsobilosti procesu pomocou navrhnutých experimentov

Zlepšovanie kvality možno chápať jednoducho ako redukciu variability procesu a produktov[6].

V procese hľadání možností redukcie tejto variability sú veľmi účinným nástrojom metódy *navrhovania (plánovania) experimentov* (*Experimental Design* alebo *Design of Experiments*). Základné poznatky o navrhovaní experimentov s jedným faktorom a viacfaktorových experimentov (dvojúrovňových) a analýze účinkov faktora (faktorov) na výstupnú premennú možno nájsť napríklad v [7], [9].

6. Záver

Úsilie o zlepšovanie kvality bude naďalej pokračovať a vyvíjať sa. Všeobecne sa zlepšovaním kvality chápe časť manažérstva kvality zameraná na zvyšovanie spôsobilosti plniť požiadavky na kvalitu [10].

Literatúra

- [1] JANIGA, I. – FILIP, I.: *Výpočet ARL pre regulačné diagramy aritmetických priemerov s výstražnými medzami pomocou programu Microsoft Excel*. . Forum Statisticum Slovaca 3/2006, s. 85 – 95. ISSN 1336-7420
- [2] JANIGA, I.: *Normalizácia metód štatistického riadenia kvality*. Forum Statisticum Slovaca 2/2007. ISSN 1336-7420
- [3] KOTZ, S.-JOHNSON, N. L.: *Process Capability Indices*. London, Chapman & Hall 1993
- [4] KUREKOVÁ, EVA - PALENČÁR, RUDOLF: *Posudzovanie spôsobilosti meracích procesov*. In: Strojnícky časopis = Journal of Mechanical engineering. - ISSN 0039-2472. - Roč. 53, č. 3 (2002), s. 141-151
- [5] KUREKOVÁ, EVA: *Analýza metód výpočtu indexov spôsobilosti*. In: Strojné inžinierstvo 2004. - Bratislava : STU v Bratislave, 2004. - ISBN 80-227-2105-0. - S1 94-98
- [6] MONTGOMERY, D. C.: *Introduction to Statistical Quality Control*. New York, J. Wiley 1991.
- [7] MONTGOMERY, D. C.: *Design and Analysis of Experiments*. New York, J. Wiley 1984.
- [8] TEREK, M. – HRNČIAROVÁ, Ľ.: *Analýza spôsobilosti procesu*. Bratislava: Ekonóm 2001, ISBN 80-225-1443-8.
- [9] TEREK M., HRNČIAROVÁ Ľ.: *Štatistické riadenie kvality*. Bratislava: Vydavateľstvo Iura Edition., 2004, ISBN 80-89047-97-1.
- [10] STN EN ISO 9000. *Systémy Manažérstva kvality. Základy a slovník*. Úrad pre normalizáciu, metrológiu a skúšobníctvo SR, 2001.

Tento príspevok vznikla s príspevom grantovej agentúry VEGA v rámci projektu číslo 1/0437/08 Kvantitatívne metódy v stratégii šesť sigma a projektu číslo 1/3182/06 Zlepšovanie kvality produkcie strojárnských výrobkov pomocou štatistických metód.

Adresa autorov:

Ľubica Hrnčiarová, doc.Ing. PhD.
Ekonomická univerzita,
Dolnozemska cesta 1
852 35 Bratislava
hrnciario@euba.sk

Milan Terek, prof.Ing. PhD.
Ekonomická univerzita,
Dolnozemska cesta 1
852 35 Bratislava
terek@euba.sk

Štatistické porovnanie zhody analyzovaného a referenčného súboru Statistical comparison deuces to analyse and reference file

Jozef Chajdiak, Maroš Suchánek

Abstract: Príspevok obsahuje štatistické porovnanie súboru pacientiek s indikáciou abruptio placentae praecox a referenčného súboru rodiacich žien bez tejto indikácie.

Paper includes statistical comparison file patient with indication abruptio placentae praecox and reference file bearing woman without this indication.

Key words: abruptio placentae praecox, F test, t test, Chi square test.

Kľúčové slová: abruptio placentae praecox, F test, t test, Chi kvadrát test.

1. Úvod

Analyzuje sa súbor pacientiek s indikáciou abruptio placentae praecox. Na kontrolu slúži referenčný súbor rodiacich žien bez tejto indikácie. Našou úlohou je porovnať zhodu súborov podľa vybraných znakov.

2. Použité testy

Pri kardinálnych znakoch a nominálnych znakoch s alternatívnymi hodnotami kódovanými číslami 0 a 1 resp. 1 a 2 sme zhodu hodnôt príslušných premenných overovali prostredníctvom zhody stredných hodnôt a zhody rozptylov. Testovali sme hypotézy zhody rozptylov σ^2 (homoskedasticitu resp. heteroskedasticitu):

$$H_0 : \sigma_1^2 = \sigma_2^2$$

$$H_1 : \sigma_1^2 \neq \sigma_2^2$$

kde dolný index 1 predstavuje rozptyl hodnôt v súbore PLACE (súbor pacientiek s indikáciou abruptio placentae praecox) a index 2 rozptyl v referenčnom súbore REFER (súbor rodičiek bez indikácie abruptio placentae praecox).

Podľa výsledku vyjadrujúcom zhodu rozptylov (homoskedasticitu) alebo rôznosť rozptylov (heteroskedasticitu) sme použili príslušnú verziu t testu zhody stredných hodnôt μ v súboroch PLACE a REFER:

$$H_0 : \mu_1 = \mu_2$$

$$H_1 : \mu_1 \neq \mu_2$$

Pri nominálnych znakoch s viac ako dvoma hodnotami sme použili χ^2 test (chí kvadrát test) testujúcom zhodu napozorovaných hodnôt $f(x)$ (Actual) a teoretických hodnôt $g(x)$ (Expected):

$$H_0 : f(x) = g(x)$$

$$H_1 : f(x) \neq g(x)$$

kde funkcie f a g predstavujú funkcie rozdelenia hodnôt premennej x .

K numerickým výpočtom boli použité funkcie FTEST, TTEST a CHITEST systému Excel verzie 2007. Výstupom funkcií je p-hodnota, ktorú porovnávame s zvolenou hladinou významnosti α (v našom prípade sme zvolili $\alpha=0,05$). V prípade, že p-hodnota je väčšia ako α nie je dôvod zamietnuť hypotézu H_0 (prijmame hypotézu H_0). V prípade, že p-hodnota je menšia ako α je dôvod zamietnuť hypotézu H_0 (prijmame hypotézu H_1).

3. Overovanie vstupnej zhody

Súbor PLACE obsahuje údaje o 106 pacientkach a referenčný súbor REFER údaje o 110 rodičkách.

Testovali sme zhodu stredných hodnôt a rozptylov za vek, hmotnosť pri pôrode, prírastok hmotnosti za tehotenstvo, hmotnosť pred pôrodom, výšku a týždeň ukončenia tehotenstva. Základný štatistický rozbor premenných a výsledky testov sú v tab.1.

Tabuľka 1: Výsledky za číselné premenné

	vek	hmot	Kg+	Kg pred	vyška	týždeň
PLACE						
N	106	106	106	106	106	106
MIN	19	53	0	46	147	26
MAX	56	97	21	95	181	42
PRIEMER	30,80	71,28	11,89	59,48	165,25	36,28
MEDIAN	30	70	11	57,5	165	37
STDEV	5,67	9,15	4,20	8,46	6,68	3,66
REFER						
N	110	110	110	110	110	110
MIN	19	56	0	48	154	37
MAX	57	124	35	110	180	42
PRIEMER	33,87	77,03	14,47	62,52	167,76	39,71
MEDIAN	33	75	14	60	168	40
STDEV	8,02	11,80	5,54	10,52	5,40	1,03
TESTY						
F	0,00041	0,009397	0	0,02533	0,028	2E-32
Výsledok	≠	≠	≠	≠	≠	≠
T	0,00131	8,61E-05	0	0,02017	0,003	7E-16
Výsledok	≠	≠	≠	≠	≠	≠

Referenčný súbor obsahuje variabilnejšie hodnoty ako analyzovaný súbor - ženy sú trochu staršie, hmotnejšie, vyššie a porodili neskôr. Všetky testy zhody naznačujú rôznosť údajov.

U premených počet plodov, DVT (deep venous thrombosis - hlboká žilová trombóza), DM (diabetes mellitus - pravá netehotenská cukrovka) a IUD (intrauterine device - vnútromaternicové antikoncepčné teliesko pred tehotnosťou) boli zistené len hodnoty 0 (nie) - nevyskytujú sa.

4. Overovanie výsledkov experimentu

Postatnými výsledkami experimentov je preukázanie/nepreukázanie štatisticky významného vplyvu členenia žien do experimentálneho súboru PLACE a do referenčného súboru REFER pri premenných FVL (faktor V Leiden – dedičná mutácia 5. faktora krvnej zrážanlivosti - zvyšuje tendenciu ku krvnej zrážanlivosti), PL (mutácia génu protrombínu – dedičná mutácia zvyšujúca krvnú zrážanlivosť) a MTHFR (mutácia génu enzýmu metyléntetrahydrofolátreduktázy – zvyšuje krvnú zrážanlivosť). Premenné FVL, PL a MTHFR majú možnosti:

mutácia neprítomná (0),

heterozygotná mutácia (1) - 1 gén zdravý a 1 chorobný,

homozygotná mutácia (2) oba chorobné gény (od oboch rodičov).

K testu výsledkov experimentu sme použili χ^2 test (chí kvadrát test) testujúcom zhodu napozorovaných hodnôt $f(x)$ (Actual) a teoretických hodnôt $g(x)$ (Expected).

Výsledky testov sú v tab.2, 3 a 4.

Tabuľka 2: Výsledky významnosti FVL

FVL(Actual)	PLACE	REFER	Spolu	FVL(Actual)	PLACE	REFER	Spolu
0	87	103	190	0	87	103	190
1	17	7	24	1+2	19	7	2624
2	2	0	2	Spolu	106	110	216
Spolu	106	110	216				
FVL(Expected)	PLACE	REFER	Spolu	FVL(Expected)	PLACE	REFER	Spolu
0	93,24	96,76	190	0	93,24	96,76	190
1	11,78	12,22	24	1+2	12,76	13,24	26
2	0,98	1,02	2	Spolu	106	110	216
Spolu	106	110	216				
CHI TEST	0,024204			CHI TEST	0,009044		
Zhoda	≠			Zhoda	≠		

Poznámka: Vzhľadom k tomu, že početnosť tried s výskytom FVL=2 je menšia ako 5, triedy FVL=1 a FVL=2 sme zlúčili do spoločnej triedy (1+2).

Tabuľka 3: Výsledky významnosti PL

PL(Actual)	PLACE	REFER	Spolu
0	102	109	211
1	4	1	5
2	0	0	0
Spolu	106	110	216
PL(Expected)	PLACE	REFER	Spolu
0	103,55	107,45	211
1	2,45	2,55	5
2	0	0	0
Spolu	106	110	216
CHI TEST	.		
Zhoda	.		

Tabuľka 4: Výsledky významnosti MTHFR

MTHFR(Actual)	PLACE	REFER	Spolu
0	41	41	82
1	40	61	101
2	25	8	33
Spolu	106	110	216
MTHFR(Expected)	PLACE	REFER	Spolu
0	40,24	41,76	82
1	49,56	51,44	101
2	16,19	16,81	33
Spolu	106	110	216
CHI TEST	0,001463		
Zhoda	≠		

Experiment preukázal štatisticky významný vplyv faktorov FVL a MTHFR (p-hodnota v riadku CHI TEST v tabuľkách 2 a 4 je menšia ako hladina významnosti $\alpha=0,05$). Pri faktore PT vzhľadom k rozdeleniu hodnôt možno len slovné konštatovať, že výskyt PT v súbore PLACE je 4/106 čo je zhruba štyri krát viac ako výskyt v súbore REFER rovný 1/110.

5. Záver

Z výsledkov experimentálneho súboru a referenčného súboru vyplýva, že vplyv faktora FVL a MTHFR je štatisticky významný, pri PT faktore je len vyšší vplyv. Ženy v referenčnom súbore sú v priemere biologicky mohutnejšie, t.j. staršie, hmotnejšie (pred a po pôrode ako aj prírastkom hmotnosti počas tehotenstva), sú vyššie a rodia neskôr (nerodia predčasne) ako ženy s experimentálneho súboru s indikáciou abruptio placentae praecox.

6. Literatúra

- [1]CHAJDIK, J. 2003. Štatistika jednoducho. Bratislava: Statis 2003. 194 s. ISBN 80-85659-28-X.
- [2]CHAJDIK, J. 2003. Štatistické úlohy a ich riešenie v Exceli. Bratislava: Statis 2005. 262 s. ISBN 80-85659-39-5.

Adresa autorov:

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

Maroš Suchánek, MUDr.
FN Ružinov
920 00 Bratislava
marossuchanek@gmail.com

Štatistická regulácia v plniacom procese Statistical control in filling process

Ivan Janiga

Abstract: In the paper we deal with the filling process. We want to find out if the filling process can be controlled.

Key words: filling process, statistical process control, control chart, normal distribution, statistical hypothesis, goodness-of-fit test.

Kľúčové slová: plniaci proces, štatistická regulácia procesu, regulačný diagram, normálne rozdelenie, štatistická hypotéza, test dobrej zhody.

1. Úvod

V príspevku sa zaoberáme možnosťou zavedenia štatistickej regulácie do plniaceho procesu. Uvažujeme objemové plničky, ktoré tvoria najrozsiahlejší používaný systém plnenia homogénnych tekutých látok. Charakteristickým predstaviteľom objemovej plničky je univerzálna štvorpolohová objemová kruhová kroková plnička, ktorá sa používa na plnenie rastlinných stolových olejov, ale aj na plnenie motorových i ďalších syntetických olejov. Linka na plnenie oleja obsahuje plniacu jednotku, ktorú tvoria dva rotory spojené dopravným pásom, na ktorom sú uchytené nosiče. Nosiče vykonávajú plynulý pohyb z jednej polohy do druhej. Keď sa nosiče zastavia v danej polohe, nasleduje proces plnenia.

Plnička má dva odmerné valce (OV1, OV2) a štyri hubice (H1, H2, H3, H4). Objemové plnenie sa vykonáva v štyroch fázach, ktoré sa cyklicky opakujú. Pri štarte je nultá fáza pri, ktorej sa naplňa OV1, potom nasledujú jednotlivé fázy plnenia:

1. fáza. OV1 sa vyprázdňuje cez H1 do fľaše č. 1 & OV2 sa naplňa,
2. fáza. OV2 sa vyprázdňuje cez H2 do fľaše č. 2 & OV1 sa naplňa,
3. fáza. OV1 sa vyprázdňuje cez H3 do fľaše č. 3 & OV2 sa naplňa,
4. fáza. OV2 sa vyprázdňuje cez H4 do fľaše č. 4 & OV1 sa naplňa.

2. Realizácia experimentu plniaceho procesu

Boli vykonané tri etapy merania pre H1, H2, H3 a H4. V druhej etape vypadlo jedno meranie pri H 2 aj H 4. Teda v etapách 1 a 3 bolo naplnených a odmeraných po 40 fliaš, v etape 2 len 38 fliaš. Fľaše o objeme 1 liter sa plnili vodou. Označené boli tak, aby sa dalo identifikovať, na ktorej z plniacich hubíc bola fľaša naplnená. Dbali sme pritom aj na dodržanie časovej následnosti plnenia.

Objem plnenia bol nastavený na 1000 ml, avšak namerané hodnoty sa pohybovali v intervale od 989 ml do 1006 ml. Merania ukazujú na závislosť medzi plniacimi hubicami H1 a H3, ktoré sú plnené plniacim valcom OV1, a hubicami H2 a H4, ktoré sú plnené odmerným valcom OV2. Pri H1 a H3 sa zdajú odchýlky od 1000 ml nevýrazné, avšak H2 a H4 plnili s výnimkou 5 hodnôt pod 1000 ml.

3. Model vyhodnotenia experimentálnych dát

Predpokladajme, že náhodná premenná $X_i^{(k)}$, ktorá predstavuje meranie hodnôt vo fľašiach naplnených na i -tej hubici v k -tej etape, má rozdelenie s distribučnou funkciou $F_i^{(k)}(x)$, pričom parametre rozdelenia sú $\mu_i^{(k)}$ – stredná hodnota – a $\sigma_i^{(k)}$ – smerodajná

odchýlka. Celkove je to dvanásť náhodných premenných, pretože v experimente sa použijú štyri plniace hubice ($i = 1, 2, 3, 4$) a tri etapy plnenia ($k = 1, 2, 3$).

Ďalej uvažujeme ďalšie štyri náhodné premenné:

X_i s distribučnou funkciou $F_i(x)$, s parametrami μ_i a σ_i , ktorá predstavuje meranie hodnôt vo fľašiach naplnených na i -tej hubici, pričom $i = 1, 2, 3, 4$,

X_{13} s distribučnou funkciou $F_{13}(x)$, s parametrami μ_{13} a σ_{13} , ktorá predstavuje meranie hodnôt vo fľašiach naplnených hubicami H1 a H3,

X_{24} s distribučnou funkciou $F_{24}(x)$, s parametrami μ_{24} a σ_{24} , ktorá predstavuje meranie hodnôt vo fľašiach naplnených hubicami H2 a H4,

X_{1-4} s distribučnou funkciou $F_{1-4}(x)$, s parametrami μ_{1-4} a σ_{1-4} , ktorá predstavuje meranie hodnôt vo fľašiach naplnených hubicami H1 až H4.

Predpokladáme, že rozdelenia všetkých 19 náhodných premenných sú spojité.

Definujme a označme náhodné výbery takto:

$V_i^{(k)}$ je množina hodnôt náhodného výberu z rozdelenia s distribučnou funkciou (ďalej d. f.) $F_i^{(k)}(x)$, ktoré predstavujú merania vo fľašiach naplnených na i -tej hubici v k -tej etape. Pretože v experimente sa použijú štyri plniace hubice ($i = 1, 2, 3, 4$) a tri etapy plnenia ($k = 1, 2, 3$), predstavuje to celkove hodnoty dvanástich náhodných výberov.

V_i je množina hodnôt náhodného výberu z rozdelenia s d. f. $F_i(x)$, ktoré predstavujú všetky hodnoty namerané vo fľašiach naplnených i -tou hubicou, pričom $i = 1, 2, 3, 4$. Matematicky to vyjadríme ako $V_i = V_i^{(1)} \cup V_i^{(2)} \cup V_i^{(3)}$.

V_{13} je množina hodnôt náhodného výberu z rozdelenia s d. f. $F_{13}(x)$, ktoré predstavujú všetky hodnoty namerané vo fľašiach naplnených hubicami H1 a H3, čo vyjadríme matematicky $V_{13} = V_1 \cup V_3$.

V_{24} je množina hodnôt náhodného výberu z rozdelenia s d. f. $F_{24}(x)$, ktoré predstavujú všetky hodnoty namerané vo fľašiach naplnených hubicami H2 a H4, pričom matematický zápis je $V_{24} = V_2 \cup V_4$.

V_{1-4} je množina hodnôt náhodného výberu z rozdelenia s d. f. $F_{1-4}(x)$, čo sú hodnoty namerané vo fľašiach naplnených všetkými hubicami. Matematicky to zapíšeme $V_{1-4} = V_1 \cup V_2 \cup V_3 \cup V_4$.

Overenie predpokladu normality

Pri implementácii štatistickej regulácie do procesu predpokladáme, že merania (výberové dáta) pochádzajú z normálneho rozdelenia. Tento predpoklad treba overiť na základe výberových dát. Overovanie sa robí pomocou testov zhody. Testuje sa pritom nulová hypotéza H_0 , že daná náhodná premenná má normálne rozdelenie oproti alternatívnej hypotéze H_1 , že nemá normálne rozdelenie.

V našom prípade testujeme tieto nulové hypotézy:

$H_0: X_i^{(k)}$ má normálne rozdelenie s d. f. $F_i^{(k)}(x)$ pre $i = 1, 2, 3, 4$ a $k = 1, 2, 3$,

$H_0: X_i$ má normálne rozdelenie s d. f. $F_i(x)$ pre $i = 1, 2, 3, 4$,

$H_0 : X_{13}$ má normálne rozdelenie s d. f. $F_{13}(x)$,

$H_0 : X_{24}$ má normálne rozdelenie s d. f. $F_{24}(x)$ a

$H_0 : X_{1-4}$ má normálne rozdelenie s d. f. $F_{1-4}(x)$.

Chceme zistiť či sa dá proces regulovať tak, že sa reguluje proces plnenia každej hubica samostatne, alebo sa reguluje proces naplňovania každého odmerného valca samostatne, alebo sa reguluje celý proces plnenia.

Pri štatistickej regulácii procesu musí byť splnený aj predpoklad, že merania pochádzajú z toho istého rozdelenia, z čoho vyplýva, že musia mať rovnaké rozptyly aj stredné hodnoty. Tento predpoklad treba overiť na základe výberových dát. Overenie predpokladov sa robí pomocou dvoch testov v uvedenom poradí:

1. Testuje sa nulová hypotéza H_0 , že všetky rozptyly sú rovnaké oproti alternatívnej hypotéze H_1 , že aspoň jeden sa líši od ostatných.
2. Testuje sa nulová hypotéza H_0 , že všetky stredné hodnoty sú rovnaké oproti alternatívnej hypotéze H_1 , že aspoň jeden sa líši od ostatných.

V ďalšom vyhodnocovaní porovnávame tri etapy plnenia pre každú zo štyroch hubíc samostatne. Ďalej porovnávame plnenie hubicou H1 s plnením hubicou H3 a plnenie hubicou H2 s plnením hubicou H4. Nakoniec porovnávame plnenie hubicami H1 a H3 s plnením hubicami H2 a H4.

Najprv testujeme nulové hypotézy o rovnosti rozptylov:

$$H_0 : (\sigma_i^{(1)})^2 = (\sigma_i^{(2)})^2 = (\sigma_i^{(3)})^2, \quad i = 1, 2, 3, 4,$$

$$H_0 : \sigma_1^2 = \sigma_3^2, \quad H_0 : \sigma_2^2 = \sigma_4^2, \quad H_0 : \sigma_{13}^2 = \sigma_{24}^2, \quad H_0 : \sigma_1^2 = \sigma_2^2 = \sigma_3^2 = \sigma_4^2.$$

Ak danú hypotézu rovnosti rozptylov nezamietneme, považujeme rozptyly za rovnaké.

V ďalšom kroku testujeme nulové hypotézy o rovnosti stredných hodnôt:

$$H_0 : \mu_i^{(1)} = \mu_i^{(2)} = \mu_i^{(3)}, \quad i = 1, 2, 3, 4,$$

$$H_0 : \mu_1 = \mu_3, \quad H_0 : \mu_2 = \mu_4, \quad H_0 : \mu_{13} = \mu_{24}, \quad H_0 : \mu_1 = \mu_2 = \mu_3 = \mu_4.$$

V ďalšom postupe na testovanie rovnosti príslušných stredných hodnôt resp. mediánov použijeme

- analýzu rozptylu (ANOVA), ak ide o porovnávanie normálnych rozdelení alebo
- Kruskalov-Wallisov test (K-W test), ak porovnávame iné ako normálne rozdelenia.

V prípade, že hypotézu rovnosti rozptylov zamietneme, nemáme momentálne k dispozícii vhodnú metódu na testovanie hypotézy o rovnosti príslušných stredných hodnôt.

V každom prípade však môžeme použiť metódy opisnej štatistiky (MOŠ), keď sa nedajú použiť exaktnejšie metódy.

4. Vyhodnotenie experimentálnych dát

Postupujeme podľa časti 3. Na testovanie normality sme použili päť testov (Kolmogorovov-Smirnovov D, Kuiperov V, Cramerov-Von Misesov W^2 , Watsonov U^2 a Andersonov-Darlingov A^2). V tabuľke 1 sú uvedené P-hodnoty pre jednotlivé testy.

Tabuľka 1: Testy normality výberov (P-hodnoty)

Výber	Rozsah	Test D	Test V	Test W^2	Test U^2	Test A^2	Normálne rozdelenie	1 - α
$V_1^{(1)}$	10	$\geq 0,10$	$\geq 0,10$	0,3476	0,3097	0,4408	má	$\geq 90\%$
$V_1^{(2)}$	10	$\geq 0,10$	$\geq 0,10$	0,3436	0,3099	0,4498	má	$\geq 90\%$
$V_1^{(3)}$	10	$\geq 0,10$	$\geq 0,10$	0,5738	0,5289	0,5412	má	$\geq 90\%$
V_1	30	$< 0,10$	$< 0,01$	0,0203	0,0142	0,0104	nemá	99%
$V_2^{(1)}$	10	$\geq 0,10$	$\geq 0,10$	0,4606	0,4899	0,5030	má	$\geq 90\%$
$V_2^{(2)}$	9	$< 0,05$	$< 0,05$	0,0098	0,0069	0,0153	nemá	95%
$V_2^{(3)}$	10	$\geq 0,10$	$\geq 0,10$	0,1725	0,1710	0,1871	má	$\geq 90\%$
V_2	29	$< 0,01$	$< 0,01$	0,0016	0,0010	0,0024	nemá	99%
$V_3^{(1)}$	10	$\geq 0,10$	$\geq 0,10$	0,8359	0,8060	0,8630	má	$\geq 90\%$
$V_3^{(2)}$	10	$< 0,10$	$\geq 0,10$	0,1748	0,1619	0,1618	nemá	90%
$V_3^{(3)}$	10	$< 0,10$	$< 0,05$	0,0512	0,0511	0,0696	nemá	95%
V_3	30	$\geq 0,10$	$\geq 0,10$	0,4831	0,4399	0,5083	má	90%
$V_4^{(1)}$	10	$\geq 0,10$	$\geq 0,10$	0,8396	0,8404	0,8345	má	$\geq 90\%$
$V_4^{(2)}$	9	$\geq 0,10$	$\geq 0,10$	0,2796	0,2559	0,2179	má	$\geq 90\%$
$V_4^{(3)}$	10	$\geq 0,10$	$\geq 0,10$	0,8014	0,7965	0,7999	má	$\geq 90\%$
V_4	30	$\geq 0,10$	$\geq 0,10$	0,1703	0,2159	0,1378	má	$\geq 90\%$
V_{13}	60	$\geq 0,10$	$< 0,10$	0,2232	0,2045	0,1741	nemá	90%
V_{24}	58	$< 0,05$	$< 0,01$	0,0164	0,0107	0,0150	nemá	99%
V_{1-4}	118	$< 0,01$	$< 0,01$	0,0000	0,0000	0,0000	nemá	99%

Z tabuľky 1 vyplýva, že normálne rozdelenie nemajú náhodné premenné X_1 , X_2 , $X_2^{(2)}$, $X_3^{(2)}$, $X_3^{(3)}$, X_{13} , X_{24} a X_{1-4} . Naopak náhodné premenné $X_1^{(1)}$, $X_1^{(2)}$, $X_1^{(3)}$, $X_2^{(1)}$, $X_2^{(3)}$, $X_3^{(1)}$, X_3 , $X_4^{(1)}$, $X_4^{(2)}$, $X_4^{(3)}$ a X_4 majú normálne rozdelenia.

Pri testovaní rovnosti rozptylov sme použili Cochranov test, Bartlettov a Leveneov test. Výsledky testov o rozptyloch sú uvedené v tabuľke 2. Bartlettov test sa používa vtedy, keď je splnený predpoklad normality. Leveneov možno použiť vtedy, keď nie je splnený predpoklad normality. Výsledky testov sú v tabuľke 2.

Tabuľka 2: Testy rovnosti rozptylov (P-hodnoty)

Nezáv. výbery	Hypotéza H_0	C-test	B-test	L-test	Rovnosť	1 - α
$V_1^{(1)}, V_1^{(2)}, V_1^{(3)}$	$(\sigma_1^{(1)})^2 = (\sigma_1^{(2)})^2 = (\sigma_1^{(3)})^2$	0,00	0,00	0,03	nie	95%
$V_2^{(1)}, V_2^{(2)}, V_2^{(3)}$	$(\sigma_2^{(1)})^2 = (\sigma_2^{(2)})^2 = (\sigma_2^{(3)})^2$	0,36	0,55	0,99	áno	95%
$V_3^{(1)}, V_3^{(2)}, V_3^{(3)}$	$(\sigma_3^{(1)})^2 = (\sigma_3^{(2)})^2 = (\sigma_3^{(3)})^2$	0,37	0,55	0,82	áno	95%

$V_4^{(1)}, V_4^{(2)}, V_4^{(3)}$	$(\sigma_4^{(1)})^2 = (\sigma_4^{(2)})^2 = (\sigma_4^{(3)})^2$	0,41	0,06	0,16	áno	95%
V_1, V_3	$\sigma_1^2 = \sigma_3^2$	0,63	0,63	0,59	áno	95%
V_2, V_4	$\sigma_2^2 = \sigma_4^2$	0,19	0,19	0,28	áno	95%
V_{13}, V_{24}	$\sigma_{13}^2 = \sigma_{24}^2$	0,00	0,00	0,00	nie	95%
V_1, V_2, V_3, V_4	$\sigma_1^2 = \sigma_2^2 = \sigma_3^2 = \sigma_4^2$	0,00	0,00	0,00	nie	95%

Z tabuľky 2 vidieť, že sme zamietli 3 hypotézy o rovnosti rozptylov

$$H_0 : (\sigma_1^{(1)})^2 = (\sigma_1^{(2)})^2 = (\sigma_1^{(3)})^2, \quad H_0 : \sigma_{13}^2 = \sigma_{24}^2, \quad H_0 : \sigma_1^2 = \sigma_2^2 = \sigma_3^2 = \sigma_4^2.$$

Na sprehľadnenie korektnosti použitých metód a praktického významu testovaných štatistických hypotéz sme vytvorili tabuľku 3, z ktorej možno nasledujúce porovnania.

Tabuľka 3: Použitie korektných metód a praktický význam hypotéz

Rozdelenia normálne?	Rozptyly rovnaké?	Test hypotézy H_0	Metódy	Praktický význam hypotézy. Plnenie fliaš. Porovnanie
áno	nie	$\mu_1^{(1)} = \mu_1^{(2)} = \mu_1^{(3)}$	MOŠ	etáp plnenia 1. hubicou
nie	áno	$\mu_2^{(1)} = \mu_2^{(2)} = \mu_2^{(3)}$	K-W test	etáp plnenia 2. hubicou
nie	áno	$\mu_3^{(1)} = \mu_3^{(2)} = \mu_3^{(3)}$	K-W test	etáp plnenia 3. hubicou
áno	áno	$\mu_4^{(1)} = \mu_4^{(2)} = \mu_4^{(3)}$	ANOVA	etáp plnenia 4. hubicou
nie	áno	$\mu_1 = \mu_3$	K-W test	plnenia 1. hubicou s plnením 3. hubicou
nie	áno	$\mu_2 = \mu_4$	K-W test	plnenia 2. hubicou s plnením 4. hubicou
nie	nie	$\mu_{13} = \mu_{24}$	MOŠ	plnenia 1. a 3. hubicou s plnením 2. a 4. hubicou
nie	nie	$\mu_1 = \mu_2 = \mu_3 = \mu_4$	MOŠ	plnení štyroch hubíc medzi sebou

Porovnanie etáp plnenia na jednotlivých hubiciach

H1. Dátové súbory v jednotlivých etapách pochádzajú z normálnych rozdelení s rozdielnymi rozptylmi, nespĺňajú teda požiadavky pre použitie vyšších štatistických metód. Použili sme nástroje opisnej štatistiky (MOŠ), konkrétne číselné charakteristiky (priemer, medián, smerodajné odchyľku, minimálnu a maximálnu hodnotu, dolný a horný kvartil) potrebné na konštrukciu Boxovho-Whiskerovho. Rozpätie hodnôt v etape 1 je 3,6 krát väčšie ako v etapách 2 a 3. Najmenší priemer hodnôt je v etape 1, najväčší v etape 3. Medián je naopak najmenší v etape 2, najväčší v etape 1. Rozdielnosť medzi plnením fliaš v etape 1 a v ďalších evidentne existuje, nevieme ho však rozumne vysvetliť.

H2. Dáta z etapy 2 a 3 nespĺňajú požiadavku normality, rozptyly v jednotlivých etapách sú však rovnaké. Použili sme preto Kruskalov-Wallisov test, ktorý ukázal štatisticky významný rozdiel medzi mediánmi jednotlivých etáp na úrovni spoľahlivosti 95%. Z Boxovho-Whiskerovho grafu sme zistili, že medián v etape 3 sa významne líši od ostatných.

Teda v etape 2 plnenia fliaš hubicou H2 boli výrazne nižšie hodnoty ako v ostatných dvoch etapách. Rozdielnosť plnenia v jednotlivých etapách nevieme rozumne vysvetliť.

H3. Dáta z etapy 2 nespĺňajú požiadavku normality, rozptyly v jednotlivých etapách sú však rovnaké. Použili sme preto Kruskalov-Wallisov test, ktorý ukázal štatisticky významný rozdiel medzi mediánmi jednotlivých etáp na úrovni spoľahlivosti 95%. Z Boxovho-Wiskerovho grafu sme zistili, že mediány sa medzi sebou výrazne odlišujú. V etape 1 je medián najnižší (998,5 ml), vyšší je v etape 2 (1000,5 ml), avšak najvyšší je v etape 3 (1002,25 ml). Toto rozdielne plnenie v jednotlivých etapách nevieme rozumne vysvetliť.

H4. Dáta vo všetkých troch etapách plnenia spĺňajú obidve požiadavky (normalitu, rovnosť rozptylov) pre použitie metódy ANOVA. Pretože P-hodnota (0,0012) F-testu je menšia ako 0,05, existuje štatisticky významný rozdiel medzi strednými hodnotami jednotlivých etáp plnení na úrovni spoľahlivosti 95%. Na zistenie rozdielov medzi strednými hodnotami sme použili Fisherovu LSD (Fisher's least significant difference) metódu, na základe ktorej sme zistili, že existuje štatisticky významný rozdiel medzi strednými hodnotami, stredná hodnota etapy 3 je štatisticky významne menšia ako stredné hodnoty etáp 1 a 2.

Porovnanie dvoch hubíc (spoločný odmerný valec)

H1 s H3. Namerané hodnoty vo fľašiach naplnených hubicou 1 a hubicou 2 (dáta zo všetkých troch etáp) nespĺňajú predpoklad normality, spĺňajú však predpoklad rovnosti rozptylov. Použili sme preto Kruskalov-Wallisov test, ktorý nepreukázal štatisticky významný rozdiel medzi dvomi mediánmi na úrovni spoľahlivosti 95%. Z Boxovho-Wiskerovho grafu sme zistili, že mediány sa rovnajú. Môžeme teda konštatovať, že obidve hubice plnia rovnako správne.

H2 s H4. Dáta v tomto prípade nespĺňajú predpoklad normality, spĺňajú však predpoklad rovnosti rozptylov. Použili sme preto Kruskalov-Wallisov test, ktorý síce nepreukázal štatisticky významný rozdiel medzi mediánmi dvoch základných súborov. Keď sa však porovnáme výberové hodnoty priemerov, mediánov, rozptylov a pozrieme na Boxov-Wiskerov graf zistíme, že ich nemôžeme považovať za rovnaké s 95% spoľahlivosťou. Nemôžeme teda tvrdiť, že obidve hubice plnia rovnako.

Porovnanie všetkých štyroch hubíc

Na porovnanie môžeme korektne použiť iba metódy MOŠ. Na základe Boxovho-Wiskerovho grafu môžeme tvrdiť, že H1 a H2 plnia rovnaké objemové množstvo, ktoré je väčšie ako 1000 ml, kým H2 plní priemerne o 5 ml menej a H4 priemerne o 6 ml menej. V mediánoch je medzi H2 a H4 ešte väčší rozdiel, teda plnenie týmito dvoma hubicami nemôžeme považovať za rovnaké.

Porovnanie dvoch odmerných valcov (V_{13} , V_{24})

Vieme, že V_{13} predstavuje všetky hodnoty namerané vo fľašiach naplnených hubicami H1 a H3, ktoré sú napĺňané z odmerného valca OV1. V_{24} predstavujú všetky hodnoty namerané vo fľašiach naplnených hubicami H2 a H4, ktoré sú napĺňané z odmerného valca OV2. Dáta nespĺňajú ani jeden z predpokladov (normalita, rovnosť rozptylov). Môžeme použiť iba metódy MOŠ. Boxov-Wiskerov graf ukazuje pomerne veľký rozdiel medzi mediánmi (5,5 ml) aj medzi priermi (5,4 ml). Môžeme teda tvrdiť, že dva odmerné valce neplnia do príslušných hubíc rovnaký objem.

5. Záver

Na základe výsledkov analýzy dát sme dospeli k názoru, že všetky uvažované procesy – proces plnenia každej hubice samostatne, proces napĺňania každého valca samostatne a celý

proces plnenia – vykazujú nestabilitu. Nesplňajú základne predpoklady ako je normalita, homoskedasticita a pod. Na záver možno konštatovať, že prezentované procesy sa v danom stave nedajú regulovať ani pomocou jednorozmerných ani pomocou viacrozmerých regulačných diagramov. Štatistickými metódami v riadení kvality sa zaoberajú citované publikácie [1], [2], [3], [5], [6], [7], [8], [9] [10].

6. Literatúra

1. GARAJ, I., JANIGA, I. *Dvojstranné tolerančné medze pre neznámu strednú hodnotu a rozptyl normálneho rozdelenia*. Bratislava: STU, 2002. 147 s. ISBN 80-227-1779-7.
2. GARAJ, I., JANIGA, I. *Dvojstranné tolerančné medze normálnych rozdelení s neznámymi strednými hodnotami a s neznámym spoločným rozptylom. Two sided tolerance limits of normal distributions with unknown means and unknown common variability*. Bratislava: STU, 2004. 218 s. ISBN 80-227-2019-4.
3. GARAJ, I., JANIGA, I. *Jednostranné tolerančné medze normálneho rozdelenia s neznámou strednou hodnotou a rozptylom. One sided tolerance limits of normal distribution with unknown mean and variability*. Bratislava: STU, 2005. 214 s. ISBN 80-227-2218-9.
4. MICHALKOVÁ, M. Štatistická kontrola presnosti plnenia tekutých látok vo firme AMT, s. r. o. Nové Mesto nad Váhom. Diplomová práca. Vedúci DP: doc. RNDr. Ivan Janiga, PhD.
5. TEREK, M., HRNČIAROVÁ, Ľ. *Štatistické riadenie kvality*. Vydavateľstvo IURA EDITION, 2004, 234 s. ISBN 80-89047-97-1.
6. TEREK, M.: *Analýza rozhodovania*. Bratislava: IURA Edition 2007. ISBN 978-80-8078-131-6.
7. TEREK, M., HRNČIAROVÁ, Ľ. *Výberové skúmanie*. Bratislava: Ekonóm 2008. ISBN 978-80-225-2440-7.
8. GARAJ, I. Požiadavky na rozsah náhodného výberu jednostranných preberacích plánov meraním. In *FORUM STATISTICUM SLOVACUM*. ISSN 1336-7420, 1/2006, s. 38-43.
9. HORNÍKOVÁ, A.: Navrhovanie experimentov so štatistickým softvérom SAS. In *FORUM STATISTICUM SLOVACUM*. ISSN 1336-7420, 5/2008 (v tlači).
10. HORNÍKOVÁ, A.: Normalizácia pri aplikácii štatistických metód. Štatistické metódy v ekonómii: 1. vedecký seminár, Horný Smokovec 10.-16. september 2007 : zborník príspevkov. - S. 43-47. - Bratislava : Vydavateľstvo EKONÓM, 2007.

Podakovanie:

Tento príspevok vznikol s podporou grantových projektov VEGA č. 1/3182/06 Zlepšovanie kvality produkcie strojárskych výrobkov pomocou moderných štatistických metód a VEGA č. 1/1247/08 Kvantitatívne metódy v stratégii šesť sigma.

Adresa autora:

Ivan Janiga, doc. RNDr. PhD.
Strojnícka fakulta STU, Nám. slobody 17,
812 31 Bratislava, ivan.janiga@stuba.sk

Nezaměstnanost cizinců v České republice podle vzdělání Unemployment of Foreigners in the Czech Republic

Eva Kačerová

Abstract: The number of long-term or permanently residing foreigners in the CR exceeded in the 2008 the amount of 410 000. More than 50 % of them are in the age of economic activity. This paper is focused on unemployment of foreigners in the Czech Republic according to education level. This article came into being within the framework of the long-term research project 2D06026, "Reproduction of Human Capital", financed by the Ministry of Education, Youth and Sport within the framework of National Research Program II.

Key words: Unemployment, education, Czech Republic.

Klíčové slová: Nezaměstnanost, vzdělání, Česká republika

1. Úvod

Nezaměstnanost cizinců je v České republice relativně novým jevem, který může v budoucnu nabývat na významu. Dlouhodobý trend nárůstu počtu cizinců na území České republiky, růst počtu cizinců s trvalým pobytem, přičemž většina cizinců je ve věku ekonomické aktivity, a tedy pohybují se aktivně na trhu práce, zvyšují riziko nárůstu počtu nezaměstnaných cizinců na území České republiky. Rizikovým faktorem pro nezaměstnanost cizinců je zejména neznalost jazyka a nedostatečná integrace do společnosti. Tyto faktory se však netýkají zaměstnanců-cizinců v nadnárodních společnostech, v nichž dorozumívacím jazykem není jenom čeština. Cizinci jsou na trhu práce velmi zranitelným segmentem, jehož (ne)zaměstnanost je podstatně ovlivněna výkyvy hospodářského cyklu, a proto je na místě uvažovat o aktivní politice zaměstnanosti cílené právě na tuto rizikovou skupinu. V neposlední řadě zůstává též nezodpovězenou otázkou, zda-li by se v případě zvyšující se nezaměstnanosti cizinců na území České republiky (zejména pak těch, kteří by měli nárok na podporu v nezaměstnanosti a jiné dávky ze systému státní sociální péče a podpory) nezačaly projevovat nacionalistické tendence v majoritní společnosti, jak tomu bylo v nedávné době ve Francii a v Nizozemsku.

2. Nezaměstnanost cizinců

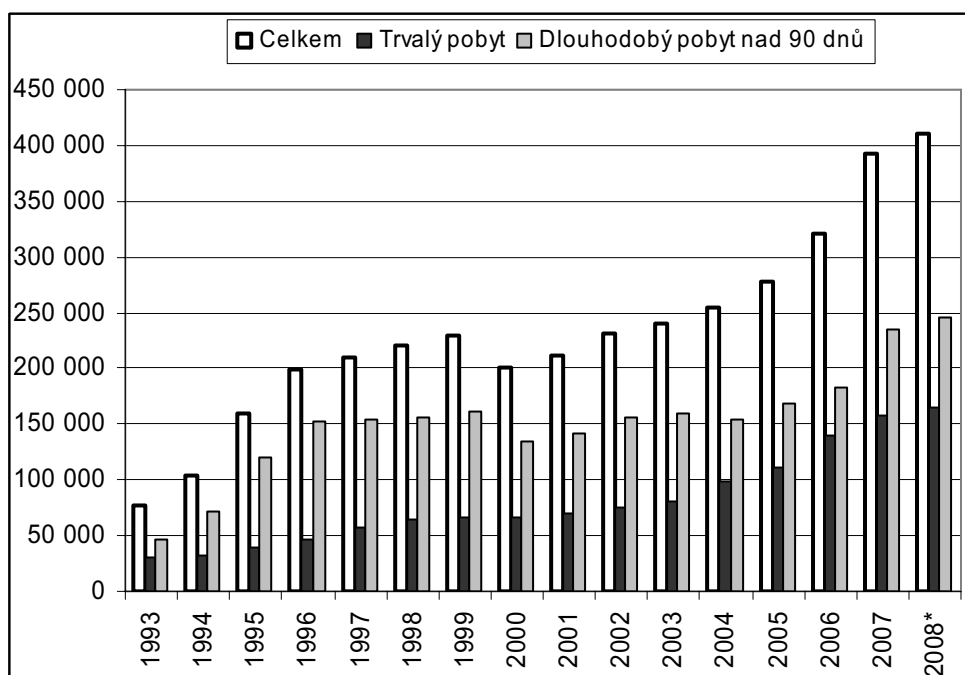
Kvantitativní data o zaměstnanosti cizinců poskytují jednotlivé okresní úřady práce, dále Ministerstvo práce a sociálních věcí a Český statistický úřad. Avšak údaje jimi poskytované mnohdy postrádají časovou kontinuitu. Tak je tomu například při zjišťování vlivu struktury vzdělání cizinců na jejich nezaměstnanost.

Počet nezaměstnaných meziročně (30.6.2007-30.6.2008) klesl z 371 tisíce na 298 tisíc, přičemž klesající tendence vývoje (s výjimkou sezónního kolísání) je patrná již od poloviny roku 2004 kdy došlo ke změně metodiky v měření nezaměstnanosti. U uchazečů o zaměstnání vedených v evidenci úřadu práce se kromě mnoha jiných údajů zjišťuje také údaj o nejvyšším dosaženém vzdělání. Dosažené vzdělání je tříděno do 14 skupin (tab. 1). Počet nezaměstnaných v jednotlivých vzdělanostních skupinách samozřejmě kromě „rizika být nezaměstnaný s určitou úrovní vzdělání“ je ovlivněn celkovým počtem osob s dosaženým vzděláním v celé populaci, resp. v některé její zkoumané části. Úroveň dosaženého vzdělání se zjišťuje při sčítání lidu, poslední bylo na území ČR k 1.3.2001. Tato struktura se samozřejmě od uvedeného data změnila o nové absolventy jednotlivých typů škol a také se mění úmrtním. Bohužel nejsou k dispozici údaje o hledajících zaměstnání podle všech

občanství, a tak jsou počty nezaměstnaných v tabulkách 1 a 2 tříděny z hlediska občanství dle dostupných kritérií. „EU“ v tomto případě znamená občany nejen zemí Evropské unie ale i Islandu, Norska, Lichtenštejnska a Švýcarska. Časová řada za občany „třetích zemí“ nebo občany ČR podle vzdělání není k dispozici, tak data „celkem bez EU“ zahrnují jak občany ČR tak i uchazeče o zaměstnání ze třetích zemí, kterých bylo evidováno okolo 3 tisíc. Ženy představují mezi nezaměstnanými většinu (56 %). V celkové populaci je poměr žen a mužů 51:49, takže lze konstatovat, že v majoritní populaci jsou nezaměstnaností více ohroženy ženy. Mezi nezaměstnanými cizinci z „EU“ tvoří ženy také 56 %, přitom však v populaci cizinců z těchto zemí ženy tvoří pouze 40 %.

Z hlediska vzdělání jsou nejvíce nezaměstnaností ohroženy osoby se základním vzděláním, osoby s výučním listem a osoby se středoškolským vzděláním, kterému nepředcházelo vyučení. Tyto kategorie jsou zcela jasně vysvětlitelné i četností výskytu těchto vzdělanostních skupin v celkové populaci. Podobná situace je i u občanů „EU“, avšak zde více než polovinu nezaměstnaných představují osoby se základním vzděláním.

V roce 2006 provedlo Ministerstvo práce a sociálních věcí ČR analýzu míry využívání výuky českého jazyka v rámci rekvalifikačních kurzů pro cizince vedené v evidenci úřadů práce a její efektivity. Údaje byly zjišťovány za rok 2005 a za období leden až říjen 2006. Z analýzy vyplývá, že v roce 2005 byla z celkového počtu 6 778 cizinců v evidenci úřadů práce poskytnuta výuka českého jazyka jako rekvalifikace 66 cizincům (tj. 1,0 %). V období leden až říjen 2006 to bylo z 7 642 cizinců 115 cizinců (1,5 %). Podle uvedené analýzy je tento nízký podíl dán zejména malým zájmem cizinců o účast na rekvalifikaci. V roce 2005 projevilo zájem o výuku 89 cizinců, v období leden až říjen roku 2006 to bylo 157 cizinců. V některých případech však nebylo možné zájemce do rekvalifikace zařadit, protože se jednalo o jednotlivce a nebylo možné otevřít kurz a individuální kurzy úřady práce nehradí. Co se týče efektivity kurzů zhruba 58 % cizinců, kteří se v roce 2005 zúčastnili výuky českého jazyka jako rekvalifikace, našlo do 6 měsíců práci.



Obr. 1: Počet cizinců s trvalým a dlouhodobým pobytem nad 90 dnů v ČR v letech 1993-2007 k 31.12. a v roce 2008 k 31.5

Zdroj: MPSV

Tabulka 1: Počet nezaměstnaných v ČR k 30.6.2008

VI.08	Celkem bez "EU"		"EU"	
	ženy	muži	ženy	muži
bez vzdělání	312	197	13	3
neúplné základ. vzdělání	727	874	9	24
základní vzdělání	53 062	39 377	813	605
nižší střední vzdělání	169	74	0	0
nižší střední odborné vzdělání	3 054	2 821	10	4
střední odborné vzdělání s výuč.listem	56 584	54 473	299	308
stř.nebo stř.odb. bez mat.i výuč.listu	2 370	452	4	10
ÚSV	6 034	3 054	35	21
ÚSO s vyučením i maturitou	6 208	5 434	36	36
ÚSO s maturitou (bez vyučení)	29 007	15 798	189	123
vyšší odborné vzdělání	1 037	493	4	6
bakalářské vzdělání	1 091	798	12	4
vysokoškolské vzdělání	5 218	6 059	77	66
doktorské vzdělání	161	227	3	1
	165 034	130 131	1 504	1 211

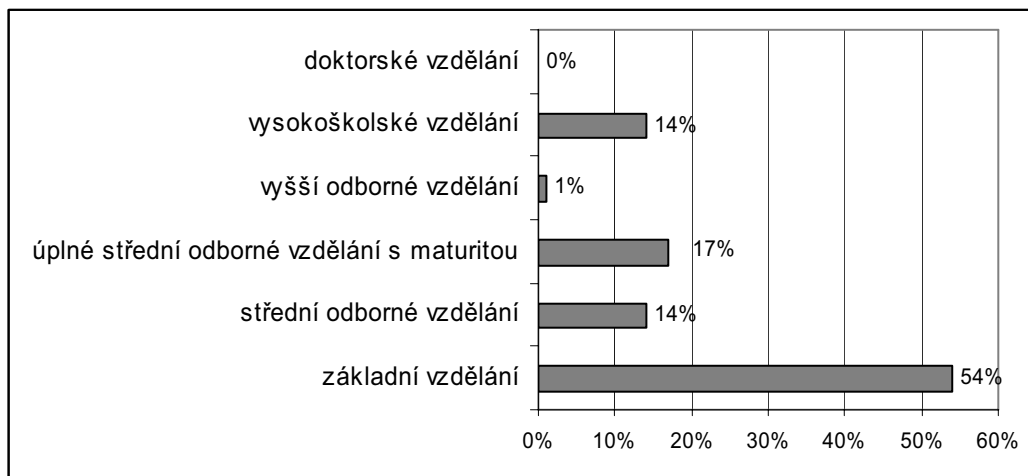
Tabulka 2: Počet nezaměstnaných v ČR k 30.6.2007

VI.07	Celkem bez "EU"		"EU"	
	ženy	muži	ženy	muži
Bez vzdělání	389	271	11	6
Neúplné základ. vzdělání	833	1 011	10	28
Základní vzdělání	66 370	51 531	876	833
Nižší střední vzdělání	250	98	0	1
Nižší střední odborné vzdělání	3 774	3 683	17	10
Střední odborné vzdělání s výuč.listem	72 085	69 727	333	342
Stř.nebo stř.odb. bez mat.i výuč.listu	2 954	586	7	11
ÚSV	7 427	3 799	46	25
ÚSO s vyučením i maturitou	7 322	6 302	43	30
ÚSO s maturitou (bez vyučení)	34 595	18 621	184	114
Vyšší odborné vzdělání	1 152	620	5	2
Bakalářské vzdělání	1 040	774	7	3
Vysokoškolské vzdělání	5 490	6 561	89	72
Doktorské vzdělání	170	248	2	1
Celkem	203 851	163 832	1 630	1 478

3. Nezaměstnanost cizinců ze třetích zemí

Podle údajů k 31.3.2008 bylo vedeno v evidenci úřadů práce celkem 2 710 občanů třetích zemí, z čehož 58 % tvořily ženy. Nárok na dávky v nezaměstnanosti mělo 45 % těchto cizinců. Vzdělanostní strukturu občanů třetích zemí vedených v evidenci uchazečů o zaměstnání bylo možné jistit pouze k 30.6.2007 (obr. 2). Potíže při opětovném začlenění na trhu práce lze očekávat u uchazečů ze základním vzděláním, kterých je více než polovina. Protože občané třetích zemí často vykonávají zaměstnání s nižšími požadavky na vzdělání může být ukazatel o nejvyšším dosaženém vzdělání cizinců zavádějící. U nezaměstnaných z řad cizinců ze třetích zemí průměrná mzda v posledním zaměstnání jen o málo překračovala minimální mzdu, která 1.1.2007 byla 8 000,- Kč. Pouze 34 (1,2 %) osob z uvedeného

celkového počtu nezaměstnaných občanů třetích zemí bylo zařazeno do některého z opatření aktivní politiky zaměstnanosti.



Obr. 2: Vzdelanostní struktura občanů třetích zemí vedených v evidenci uchazečů o zaměstnání (stav k 30.6.2007)

Zdroj: MPSV

4. Závěr

Pro podrobnější zkoumání integrace cizinců na trhu práce bude do budoucna třeba vytvořit systém pravidelného monitoringu nezaměstnanosti cizinců, a to jak občanů třetích zemí, tak také občanů EU/EHP. Jedná se o potřebu pravidelného sběru dat týkajících se zejména délky nezaměstnanosti cizinců a jejího opakování, výše dávek v nezaměstnanosti a posledních výdělků, tříd zaměstnání a opatření aktivní politiky zaměstnanosti, do nichž jsou cizinci zařazováni. Tyto údaje by mohly pomoci odhalit případnou marginalizaci cizinců na českém pracovním trhu. Údaje o cílenosti opatření aktivní politiky zaměstnanosti na cizince pak mohou odhalit možné institucionální vyloučení cizinců v České republice.

5. Literatura

- [1]KOSCHIN, F. A KOL. 2007. Prognóza lidského kapitálu obyvatelstva České republiky do roku 2050, VŠE, Praha 2007.
- [2]BULLETIN Č. 20 – MEZINÁRODNÍ PRACOVNÍ MIGRACE V ČR. Výzkumný ústav práce a sociálních věcí, Praha 2008.
- [3]www.portal.mpsv.cz
- [4]www.vupsv.cz
- [5]www.czso.cz
- [6]www.uiv.cz

Adresa autora

Eva Kačerová, RNDr.
Vysoká škola ekonomická v Praze
Fakulta informatiky a statistiky
Katedra demografie
nám. W. Churchilla 4
130 67 Praha 3
kacerova@vse.cz

Robust GMM estimation

Jan Kalina

Abstract: The robustification of the GMM (generalized method of moments) estimation has obtained an attention in robust econometrics only recently. This work discusses the weighted generalized method of moments (WGMM) estimator, which is a robust analogy of the GMM estimator based on downweighting less reliable observations. Our aim is to describe possible applications and two special cases of the WGMM estimator for the linear regression and their computational aspects. The first case is the least weighted squares regression modified to be efficient under heteroscedasticity. Another application is the instrumental weighted variables (I WV) estimator, which is a robust analogy of the instrumental variables. We present the computational aspects of the least weighted squares regression, which is crucial for understanding the computational aspects of the robust GMM estimator.

Key words: robust econometrics, computational aspects, generalized method of moments, least weighted squares, instrumental variables.

Acknowledgement: This work is supported by the grant 402/06/408 (Robustification of the generalized method of moments) of the Grant Agency of the Czech Republic.

1. GMM estimation

The generalized method of moments (GMM) was proposed by [4] as a general tool for statistical estimation. The estimator is given by orthogonality conditions and is defined in a very abstract way for a general parametric situation, involving instrumental variables. This chapter starts with a short general description of the estimator and proceeds to two special cases which are popular in applied econometrics.

[12] showed the connection to the classical method of moments, so that the estimator can be defined by means of moment conditions allowing for an overidentification. In other words there can be more conditions than parameters of the model. While the classical method of moments estimator is given as the solution of the system of equations $\mathbf{g}=0$, where $\mathbf{g}$ is the vector of particular moment conditions, there is no such solution if the overidentification occurs. There [4] proposed the estimator in the form

$$\min \mathbf{g}^T \mathbf{W} \mathbf{g}$$

and computed the optimal weight matrix $\mathbf{W}$ to be the inverse of the variance matrix of the equations $\mathbf{g}$. In practice a consistent estimator is however used instead of $\mathbf{W}$ itself. [12] explains that this estimator can be useful as a combination of several available consistent estimators. A numerical example is presented for example in [3] for estimating two parameters of the gamma distribution combining four different moment conditions. The computation involves optimizing a nonlinear function and iterative procedures must be applied, usually based on a linear approximation to the nonlinear function $\mathbf{g}^T \mathbf{W} \mathbf{g}$. Further paragraphs are devoted to two special cases for the linear regression context.

One special case of the GMM estimator is the well known instrumental variables estimator, which has become a standard tool for estimation in econometrics (see [3]).

Another special case of the GMM estimator proposed by John Cragg in [2] is less known. It is a modification of the least squares regression, which is efficient also under heteroscedasticity. Let us consider the linear regression model in the form

$$Y_i = b_0 + b_1 x_{i1} + \dots + b_p x_{ip} + e_i, i = 1, 2, \dots, n, \quad (1)$$

which can be rewritten in the usual matrix notation as $Y = X\beta + e$. [3] or [6] warn that under heteroscedasticity of the random errors e the least squares estimator b of the regression parameters β is not efficient, and further the estimator of $\text{var } b$ is biased. It follows that the confidence intervals and hypothesis tests concerning β are not valid. Therefore it is very useful that [2] proposed such transformation, which allows to obtain a more reliable estimator of β itself and mainly of $\text{var } b$ even without testing if the heteroscedasticity is present in the model (1) or not.

The idea of [2] is to use some auxiliary variables which could contribute to explaining the variability of the errors e . The usual choice contains squares of all independent variables from (1) and also products of always two different independent variables. This corresponds to using a matrix Q consisting of all columns of the original design matrix X and the auxiliary variables as additional columns. The model (1) is transformed to

$$Q^T Y = Q^T X\beta + Q^T e. \quad (2)$$

The parameters β can be estimated by Aitken's estimator. The variance of matrix of the errors e can be estimated by diagonal matrix S containing squares of residuals

$$u_i = Y_i - b_0 - b_1 x_{i1} - \dots - b_p x_{ip}.$$

Now $\text{var } b$ can be estimated by

$$X^T Q (Q^T S Q)^{-1} Q^T X.$$

2. Least weighted squares regression

The least weighted squares (LWS) regression is a robust regression method with a high breakdown point proposed by [9]. There must be nonnegative weights $w_1, w_2, \dots, w_n$ specified before the computation of the estimator. While the classical weighted regression assigns a fixed and known weight to each observation, in the context of least weighted squares only the magnitudes of the weights are known a priori. These are assigned to the data after a permutation, which is determined automatically only during the computation based on the residuals. It is reasonable to choose such weights so that the sequence $w_1, w_2, \dots, w_n$ is decreasing (non-increasing), so that the most reliable observations obtain the largest weights, while outliers with large values of the residuals get small (or zero) weights.

Let us denote the i^{th} order value among the squared residuals for a particular value of the estimate b of the parameter β by $u_i^2(b)$. The least weighted squares estimator b_{LWS} for model (1) is defined as

$$b_{LWS} = \operatorname{argmin} \sum_{i=1}^n w_i u_{(i)}^2(b).$$

The least trimmed squares (LTS) regression proposed by [8] represents a special case of least weighted squares with weights equal to zero or one only. The computation of the LWS estimator is intensive and an approximative algorithm is proposed in [6].

The least weighted squares estimator has interesting applications, which follow from its robustness and at the same time efficiency for normal data. Theoretical properties including the breakdown point of the estimator are studied by [1]. It is especially suitable to use the LWS estimator rather than other robust regression estimators, because diagnostic tools (such as tests of heteroscedasticity and autocorrelation of the errors e) can be computed directly using the weighted residuals and again are not affected by outliers. Such diagnostic tests are equivalent with those computed for least squares regression, see [5]. Another advantage of the estimator is that no detection of outliers is actually needed to compute it, because outlying data are downweighted automatically. [11] conjectures that the LWS estimator is a reasonable compromise between the least squares and least trimmed squares, namely the estimator combines the efficiency of the least squares with the robustness of the least trimmed squares.

3. Robust GMM estimation

This section discusses a robust version of the GMM estimator together with robustifications for the two special cases from Section 1, namely the instrumental variables and the regression efficient under heteroscedasticity. All these approaches are based on downweighting less reliable observations similarly with the idea of the LWS regression. While some methods are proposed already by [10] or [11], we discuss other possible ways how to propose robust estimators.

[10] proposes the weighted generalized method of moments (WGMM) estimator as a robustification of the GMM estimator in a very general context. There are weights in the system of equations g assigned to particular observations in an implicit way based on the values of the residuals. The equations can be interpreted as weighted sums of residuals. While an algorithm is not available, it must combine the iterative search for the optimal permutation of weights with the iterative nonlinear optimization. Such computation is very intensive. We believe that only the two special cases in the following paragraphs are useful for practical applications. Recalling that [12] recommends the GMM estimator in this general form mainly as a tool to combine several estimators, we suggest to use always robust estimators (robust equations g) and to use the GMM estimator itself, which does not need to be robustified any more.

A robust version of the instrumental variables was proposed by [11] and called instrumental weighted variables (IWV) estimator. [11] studies the consistency and asymptotic normality of the estimator and [7] derives asymptotic diagnostic tools based on its residuals, namely tests of heteroscedasticity and autocorrelation of the errors e . The definition is rather technical. The usual instrumental variables is computed as the two-stage least squares estimator, using the regression of the regressor against the instruments in the first step and using the regression of the response against the fitted values from the first step (see [3]). Therefore we propose to compute the robust instrumental variables estimator by applying the least weighted squares in both steps of this two-stage procedure.

The other special case of the GMM estimator described in Section 1 was Cragg's modification of least squares ensuring efficiency under heteroscedasticity. The robustification was proposed by [10], who introduces weights to the model (2). We stress that there are as many rows in the model (2) as there are columns in the matrix Q . This is the number of variables in the original model (1) together with the number of auxiliary variables contained

in Q . Therefore is an error in the definition of the estimator in [10] and we propose a different approach. Namely our idea is to introduce the weights directly in the model (2) in the form

$$Q^T WY = Q^T WX\beta + Q^T We, \quad (3)$$

where W is a weight matrix containing weights determined by the least weighted squares in the original model (1). Therefore the whole procedure starts by the LWS and then uses the (classical) weighted regression to estimate β in the transformed model (3).

The computation of the robust instruments and robust regression efficient under heteroscedasticity is based on computing the least weighted squares regression. Therefore it is crucial to study its computational aspects now.

4. Computational aspects of the least weighted squares regression

A fast and reliable algorithm for computing the least weighted squares regression is proposed by [6]. Now we study the computational aspects of both the least weighted squares (LWS) and the least trimmed squares (LTS) estimators. We use the function `ltsReg()` from package `VR` of the software package `R` to compute the LTS estimator and our own subroutine programmed in `C++` to compute the LWS estimator, which is not implemented in any commercial software.

The following example studies the time series of real gross domestic product (GDP) and real gross private domestic investment (INVESTMENT) in the United States in the years from 1980 to 2001. Both variables are expressed as multiples of 10^9 dollars. The data are summarized in Table 1 and come from the website of the U.S. Department of Commerce. The data are adjusted for deflation of money rather than nominal. Therefore it has a reasonable economic interpretation to model INVESTMENT as a linear response of GDP. [3] presents a similar investment equation.

Year	1980	1981	1982	1983	1984	1985
GDP	4900.9	5021	4949.3	5132.3	5505.2	5717.1
Investment	655.3	715.6	615.2	673.7	871.5	863.4
Year	1986	1987	1988	1989	1990	1991
GDP	5912.4	6113.3	6368.4	6591.8	6707.9	6676.4
Investment	857.7	879.3	902.8	936.5	907.3	829.5
Year	1992	1993	1994	1995	1996	1997
GDP	6880	7062.6	7347.7	7543.8	7813.2	8159.5
Investment	899.8	977.9	1107	1140.6	1242.7	1393.3
Year	1998	1999	2000	2001		
GDP		8508.9	8856.5	9224	9333.8	
Investment		1558	1660.1	1772.9	1630.8	

Table 1: Investment data.

The least squares estimate for the slope b_1 is highly significant, the t statistic equals $t=15.28$ and the corresponding p -value is less than 0.0001. Therefore the linear model appears to be reasonable. The coefficient of determination equals $R^2=0.921$.

	LS	LTS; $h=19$	LWS
Intercept	-582.0	-371.4	-465.4
GDP	0.239	0.203	0.221
$\sum_{i=1}^{22} u_i^2$	198 827	272 146	269 335
$\sum_{i=1}^{19} u_{(i)}^2(b)$	115 891	107 262	120 656
$\sum_{i=1}^{19} w_i u_{(i)}^2(b)$	48 538	48 597	44 977

Table 2: Results of the example with investment data.

Table 2 starts with results of the least squares. A subjective search for a proper value of the trimming constant for the least trimmed squares suggests to choose $h=19$. The years 1998, 1999 and 2000 are detected as outliers. The weights determined by the least weighted squares are values of a linear decreasing function with values 1, $1-1/n$, ..., $1/n$ standardized so that the sum of weights equals 1.

The least weighted squares regression searches for the minimal weighted sum of squared residuals over all permutations of the weights. In our example, these optimal weights are equal to

$$(15,9,22,20,6, 10,18,21,17,13, 8,2,4,7,14, 12,19,11,5,3, 1,16)^T$$

standardized so that the sum of weights equals 1. The estimates of the intercept and slope in Table 2 suggest the least weighted squares to be a reasonable compromise between the least squares and least trimmed squares.

Now we verify that our algorithm for computing the least weighted squares estimate gives a reliable result. Our argument resembles that of [9] who studies the tightness of his least trimmed squares algorithm. We compare the trimmed sum of squared residuals and the weighted sum of squared residuals of the least trimmed squares (again with $h=19$) and the least weighted squares with linear weights from Table 2.

The least trimmed squares minimize the value of

$$\sum_{i=1}^{19} u_{(i)}^2(b)$$

and namely they find the minimal value of this loss function to be 107 262 (see again Table 2). While the least weighted squares yield a much larger value 120 656, one would deduce that also

$$\sum_{i=1}^{19} w_i u_{(i)}^2(b)$$

is larger compared to least trimmed squares. However the LTS yield the value 48 597 and the LWS designed to minimized this weighted sum of squared residuals find a much smaller value 44 977. Therefore the approximation to least weighted sum of squares minimizes the loss function reliably, which gives evidence in favour of the approximative algorithm.

5. References

- [1] ČÍŽEK, P. 2008. Efficient robust estimation of time-series regression models. In: Applications of Mathematics 53, Nr. 3, 267-279.
- [2] CRAGG, J.G. 1983. More efficient estimation in the presence of heteroscedasticity of unknown form. In: Econometrica 51, No. 3, 751-763.
- [3] GREENE, W.H. 2002. Econometric analysis. Fifth edition. Macmillan, New York.
- [4] HANSEN, L.P. 1982. Large samples properties of generalized method of moments estimators. In: Econometrica 50, Nr. 4, 1029-1054.
- [5] KALINA, J. 2007. Autocorrelated errors of least weighted squares. In: Forum Statisticum Slovacum 3, Nr. 2, 136-141.
- [6] KALINA, J. 2008. Robust regression and diagnostic tools. In Kupka K. (ed.): Data analysis 2007/II, Progressive methods of statistical data analysis and modelling for research and technical practice, Trilobyte, Pardubice. In print.
- [7] KALINA, J. 2008. Computing robust GMM estimators. Submitted to Computational statistics and data analysis.
- [8] ROUSSEEUW P.J. – LEROY A.M. 1987. Robust regression and outlier detection. Wiley, New York.
- [9] VÍŠEK, J.Á. 2001. Regression with high breakdown point. In Antoch J., Dohnal G. (eds.): Proceedings of ROBUST 2000, Summer School of JČMF, JČMF and Czech statistical society, 324-356.
- [10] VÍŠEK, J.Á. 2005. Robustifying generalized method of moments. In Kupka K. (ed.): Data analysis 2004/II, Progressive methods of statistical data analysis and modelling for research and technical practice, Trilobyte, Pardubice, 171-193.
- [11] VÍŠEK, J.Á. 2006. Instrumental weighted variables. In: Austrian Journal of Statistics 35, Nr. 2&3, 379-387.
- [12] WOOLDRIDGE, J.M. 2001. Applications of generalized method of moments estimation. In: Journal of Economic Perspectives 15, Nr. 4, 87-100.

Address of the author:

Jan Kalina, Dr.rer.nat.
KPMS MFF UK
Sokolovská 83
186 75 Praha 8
Czech Republic
kalina@karlin.mff.cuni.cz

Stabilita poptávky po penězích v České republice Money Demand Stability in the Czech republic

Svatopluk Kapounek

Abstrakt

The effectiveness and success of a monetary policy depends on a stable money demand function. The stable money demand function ensures that the money supply would have predictable impact on the macroeconomic variables such as inflation and real economic growth. This article deals with the money demand definition under the Keynes theoretical approach and selected monetary aggregates M1, M2 and M3 in the Czech republic for the year 2003-2008.

The target of the article is empirical analysis of the money demand stability with the CUSUM and Hansen's stability tests.

Key words

stability test CUSUM, Hansen's stability test, monetary aggregates, optimum currency area, asymmetric shock

Klíčové slová

test stability CUSUM, Hansenův test stability, peněžní agregáty, efektivita hospodářské politiky

1. Úvod

Stabilita poptávky po penězích je spolu s její úrokovou elasticitou hlavní podmínkou efektivnosti měnové a fiskální politiky, neboť umožňuje zajistit přímou vazbu mezi relevantními měnovými agregáty a nominálním důchodem. Její důležitost roste v prostředí inflačního cílení spolu s potřebou predikce dopadů změn charakteru monetární politiky na inflaci a ekonomický růst. Kvantitativní rovnice peněz definuje vztah mezi množstvím peněz v ekonomice, reálnou produkcí a inflací prostřednictvím následujícího vztahu:

$$M \times V = P \times Y, \quad (1)$$

kde M vyjadřuje měnovou zásobu, V rychlost oběhu peněz v ekonomice, P agregátní cenovou hladinou a Y reálný produkt. Z pohledu keynesiánců lze za předpokladu konstantní rychlosti oběhu peněz v ekonomice považovat za jediný spojovací článek (mezi změnou množstvím peněz v ekonomice a reálným důchodem) úrokovou sazbu.

Předpokládejme, že pro určitý objem celkového bohatství drží ekonomické subjekty jednu jeho část v peněžní formě a zbývající část bohatství v jiných formách aktiv. V cenných papírech, nemovitostech, dlouhodobých vkladech, cennostech apod. Keynes vymezuje poptávku po penězích prostřednictvím funkce preference likvidity a motivů držby peněz (transakční, opatrnostní a spekulativní motiv). Poptávka po penězích z transakčního a opatrnostního motivu je přitom pozitivně závislá zejména na úrovni reálného důchodu a negativně na úrokové sazbě. Spekulativní motiv představuje poptávku po alternativních aktivech, než jsou peníze. Tato alternativní aktiva se ovšem vyznačují různými výnosy i riziky. Reakce poptávky po penězích na výši úrokové sazby v případě spekulativního motivu je proto více či méně diskutabilní. Při očekávání zvýšení úrokových sazeb tento motiv může pozitivně ovlivnit poptávku investorů po penězích, neboť se ti mohou přesunovat svá aktiva z alternativních, rizikovějších aktiv zpět do peněžních zůstatků.

Keynesovu funkci poptávky po penězích (funkci preference likvidity) lze zapsat vztahem:

$$\frac{M_D}{P} = f(Y, IR), \quad (2)$$

kde M_D je poptávka po nominálních peněžních zůstatcích, P agregátní cenová hladina, Y reálný důchod a IR úroková sazba.

Keynes dále předpokládá, že agregátní cenová hladina je v krátkém období nepružná z důvodu rigidit mezd a cen (Mach, 2002). Pro další empirickou analýzu proto faktor změny cenové hladiny nebude uvažován. „Při empirických odhadech různě specifikovaných modelů poptávky po penězích k těmto (výše uvedeným) základním vysvětlujícím faktorům přistupují i exogenní vlivy, které se však v řadě případů nahrazují z důvodu obtížné kvantifikace pomocí proxy proměnných (Hušek a Pelikán, 2003, s.114).

Zatímco teoretické vymezení poptávky po penězích je založeno na přístupech jednotlivých ekonomických škol, její empirická definice vychází z předpokladu rovnosti skutečné peněžní zásoby a poptávky po penězích. S rozvojem finančních trhů lze do empirické poptávky po penězích zahrnout nejen poptávku po oběživu jako takovém ale také bezhotovostní peníze (vklady u bankovních či jiných institucí). Závislost různých druhů peněz, odlišujících se různým stupněm jejich likvidity, na úrokové sazbě a důchodu z důvodu různých motivů poptávky po penězích je diskutabilní. Jak uvádí Revenda (Revenda a kol, 1999, s.26): „nejlepším vymezením peněz je to, které nejlépe předpovídá vývoj těch proměnných, jež by peníze mely vysvětlit.“ V empirické části tohoto článku bude proto uvažováno s více způsoby vymezení peněz prostřednictvím peněžních agregátů M1, M2 a M3.

Peněžní agregát M1 tvoří hotovostní oběživo a vklady na běžných účtech v bankách. Peněžní agregát M2 se skládá z peněžního agregátu M1, z úsporných vkladů v bankách a termínovaných vkladů v bankách. Peněžní agregát M3 je pak oproti M2 rozšířen o euroměnové vklady u domácích bank. V některých případech lze použít také agregáty M4 i M5, od širšího vymezení peněz než je vymezuje peněžní agregát M3 bude ovšem vzhledem k teoretickým úvahám pramenících z keynesiánské poptávky po penězích v tomto příspěvku abstrahováno.

Příspěvek se zabývá nejen výše uvedenou definicí poptávky po penězích z pohledu keynesiánské ekonomické teorie a jejím empirickým odvozením pro Českou republiku, ale především testováním její stability. Cílem příspěvku je ověřit hypotézu, zda je poptávka po penězích měřená různými peněžními agregáty v České republice stabilní či nikoliv.

2. Metodika

Thomas (1993) popisuje stabilitu funkce poptávky po penězích jako stálost vztahu mezi poptávkou po penězích a několika málo proměnnými. Předpokládá, že více než dvě či tři významné vysvětlující proměnné identifikované vícerozměrnou regresní analýzou negativně ovlivňují stabilitu poptávky po penězích. Hušek a Pelikán (2003) hovoří o stabilitě poptávky po penězích v souvislosti se stálostí parametrů vícerozměrné regresní funkce v čase a zároveň v souvislosti s relativně malým rozptylem náhodné složky.

Nejčastěji používaným testem stability poptávky po penězích je CUSUM test založený na kumulovaných součtech reziduí (Greene, 2003). Testovou statistikou je:

$$S_t = \frac{\sum_{t=k+1}^n w_t^2}{\sum_{t=k+1}^N w_t^2} \quad (3)$$

kde w_t jsou rekursivní rezidua modelu. Za předpokladu hypotézy o konstantnosti parametrů

$$\langle E|S_t \rangle = \frac{(t-k)}{(T-k)} \quad (4)$$

pro t jdoucí od nuly za předpokladu $t=k$ k $t=T$. Test je prezentován prostřednictvím grafu proměnné S_t vůči t , kde horní a dolní mez tvoří hranice 5% hladiny významnosti. (Brown, Durbin a Evans, 1975)

Pomocným nástrojem pro identifikaci stability/nestability modelu slouží také graf rekurzivních reziduí w_t

$$w_t = \frac{y_t - x_t' b_{t-1}}{(1 + x_t' (X_t' X_{t-1})^{-1} x_t)^{1/2}} \quad \text{pro } t=k+1, \dots, T. \quad (5)$$

Podmínkou je nezávislé a normální rozdělení s nulovou střední hodnotou a konstantním rozptylem rekurzivních reziduí. Pokud rekurzivní rezidua v grafu překročí horní či dolní mez, model lze považovat za nestabilní.

Velmi elegantním nástrojem pro identifikaci nestability parametrů v lineárním modelu je Hansenův test. (Hansen, 1992). Mějme standardní lineární regresní model, kde:

$$y_t = \beta_1 x_{1,t} + \beta_2 x_{2,t} + \beta_m x_{m,t} + \varepsilon_t, \quad (6)$$
$$E\langle e_t | x_t \rangle = 0, \quad E\langle e_t^2 \rangle = \sigma^2.$$

Prostřednictvím OLS odhadu získáváme odhad parametrů $\hat{\beta}$ a $\hat{\sigma}$. Za podmínky

$$0 = \sum_{t=1}^n x_{i,t} \hat{e}_t \quad \text{pro } i = 1, \dots, m \quad (7)$$
$$0 = \sum_{t=1}^n (\hat{e}_t^2 - \hat{\sigma}^2),$$

kde $\hat{e}_t = y_t - x_t' \hat{\beta}$. Vztah (7) lze pro účely definice testového kritéria zapsat jako

$$f_{i,t} \begin{cases} x_{i,t} \hat{e}_t, i = 1, \dots, m \\ \hat{e}_t^2 - \hat{\sigma}^2, i = m+1 \end{cases} \quad (8)$$

Za předpokladu, že je kumulativní součet reziduí $S_{i,t}$ roven nule, kde

$$S_{i,t} = \sum_{t=1}^n f_{i,t}, \quad (9)$$

je testové kritérium pro stabilitu jednotlivých parametrů modelu definován vztahem

$$L_i = \frac{1}{nV_i} \sum_{t=1}^n S_{i,t}^2, \quad (10)$$

kde

$$V_i = \sum_{t=1}^n f_{i,t}^2. \quad (11)$$

Testové kritérium pro stabilitu všech parametrů modelu současně pak

$$L = \frac{1}{n} \sum_{t=1}^n S_t' V^{-1} S_t, \quad (12)$$

kde

$$V = \sum_{t=1}^n f_t f_t'. \quad (13)$$

Hypotéza přitom předpokládá nulovou střední hodnotu podmínky $f_{i,t}$ a kumulativnímu součtu blížícímu se nule, proti hypotéze alternativní, že parametry modelu jsou nestabilní.

3. Výsledky empirické analýzy

Pro empirickou analýzu byly zvoleny meziroční změny čtvrtletních absolutních hodnot měnové zásoby definované peněžním agregátem M1, M2 a M3, dle meziroční změny čtvrtletních absolutních hodnot HDP v konstantních cenách a meziroční změny čtvrtletních průměrů krátkodobé jednodenní úrokové sazby na mezibankovním trhu a to v období 2003/1 – 2008/1.

Výsledky vícerozměrné regresní analýzy v Tabulce 1 prezentují 3 statisticky významné modely na 1% hladině významnosti. Druhý a třetí model odhaduje poptávku po penězích definovanou prostřednictvím peněžních agregátů M2 a M3. Oba tyto modely sice vykazují významnost a to včetně parametrů, nicméně pozitivní kauzalitu mezi úrokovou sazbou a poptávkou po penězích lze předpokládat pouze v případě spekulativního motivu, kdy je poptávka po penězích determinována očekáváním ekonomických subjektů, tedy budoucími hodnotami úrokové sazby. Výsledky odhadů modelů předpokládají zpoždění řádu t-1, což nekoresponduje s předpoklady ekonomické teorie.

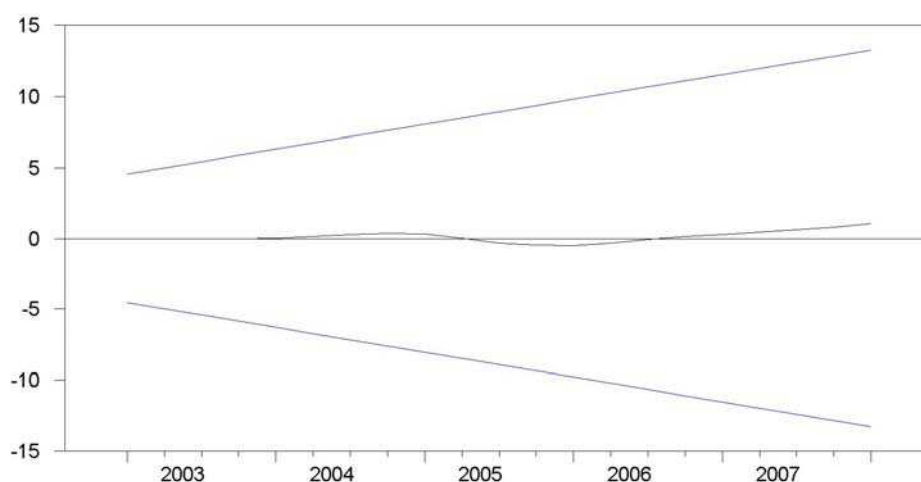
Tabulka 1: Výsledky vícerozměrné regresní analýzy

	odhad parametrů	SE	t-statistika	F-test	p-value	DW	n
Model M1				53,930	0,000	0,266	21
Y_t	2,317	0,223	10,380		0,000		
IR_t	-0,094	0,047	-1,999		0,060		
Model M2				148,880	0,000	0,378	20
Y_t	1,701	0,107	15,940		0,000		
IR_{t-1}	0,058	0,023	2,486		0,023		
Model M3				153,560	0,000	0,550	20
Y_t	1,754	0,108	16,196		0,000		
IR_{t-1}	0,059	0,024	2,506		0,022		

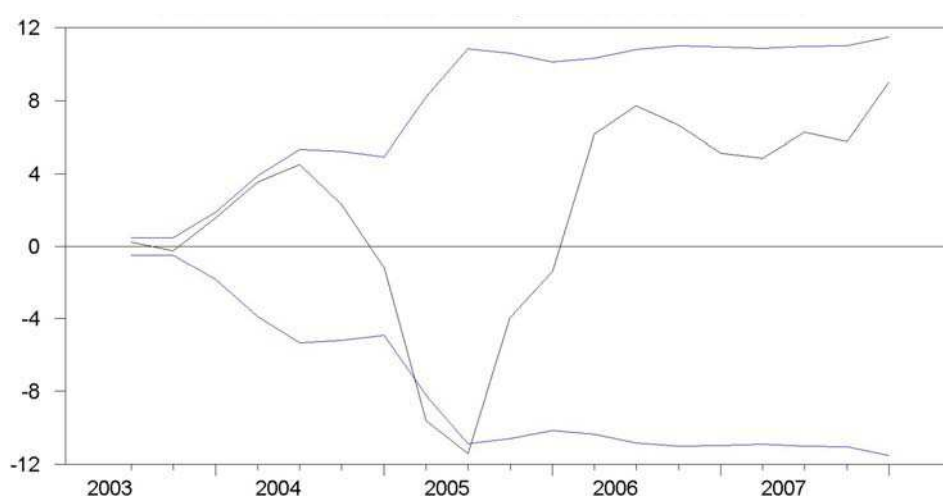
Naopak, model poptávky po penězích definovaný prostřednictvím peněžního agregátu M1 naplňuje keynesiánské předpoklady motivů držby peněz v souvislosti s transakčním i opatrnostním motivem. Nízký rozptyl náhodné složky daný významností modelu (F-test) a pouze dvě proměnné mající vliv na poptávku po penězích potvrzují volbu prvního modelu pro testování stability jeho parametrů.

Obrázek 1 prezentuje výsledky CUSUM testu. Kumulativní součet reziduí modelu poptávky po penězích neprotíná horní či dolní mez danou 5% hladinou významnosti a tedy lze předpokládat, že funkce poptávky po penězích odvozená v tomto příspěvku je stabilní. Graf rekurzivních reziduí (obrázek 2) již tak přesvědčivě o stabilitě poptávky po penězích definované prostřednictvím peněžního agregátu M1 nehovoří.

Výsledky Hansenova testu stability parametrů modelu prezentuje Tabulka 2. Na 5% hladině významnosti nezamítnul Hansenův test hypotézu o stabilitě parametrů pouze v případě parametru reálného HDP. Společný test parametrů i parametr úrokové sazby zamítl hypotézu o stabilitě.



Obrázek 1: Kumulativní součet reziduí (CUSUM test)



Obrázek 2: Rekurzivní rezidua

Tabulka2: Výsledky Hansenova testu stability parametrů

Test parametrů	t-statistika	p-value
Společný	1,9141	0,0000
Y_t	0,4466	0,0600
IR_t	1,2057	0,0000

4. Závěr

Jak již bylo řečeno v úvodu tohoto příspěvku, stabilita poptávky po penězích je základní podmínkou efektivní aplikace nástrojů hospodářské politiky, zejména pak monetární politiky ve smyslu podpory trvalého a udržitelného ekonomického růstu a stabilizace cenové hladiny ze strany centrální banky. Poptávka po penězích byla dle keynesiánského konceptu definována prostřednictvím transakčního, opatnostního a spekulativního motivu a determinována reálným ekonomickým růstem a krátkodobou úrokovou sazbou na mezibankovním trhu. Z empirického hlediska bylo pro odvození keynesiánské funkce preference likvidity použito peněžního agregátu M1, zahrnujícího hotovostní oběživo a vklady na běžných účtech v bankách. Pro agregáty M2 a M3 nebyl v tomto příspěvku nalezen významný regresní model.

Hypotéza o stabilitě poptávky po penězích byla podpořena jednoduchostí modelu (pouze dvě proměnné), jeho významností i významností jeho parametrů. Stabilitu podpořil také test založený na kumulovaných součtech reziduí, CUSUM test. Test rekurzivních reziduí a Hansenův test stability parametrů modelu hovoří ovšem proti stabilitě poptávky po penězích definované v tomto příspěvku.

Základní problém odhadu poptávky po penězích je možné vidět v její determinaci. Pokud je poptávka po penězích určována exogenním způsobem ze strany centrální banky, lze jednorovnicový model odhadnutý v tomto příspěvku považovat za dostatečný. Za předpokladu endogenního charakteru poptávky po penězích je nutné uvažovat např. soustavu simultánních rovnic či kointegrační analýzu zahrnující error correction model.

Z výsledků empirické analýzy v tomto příspěvku lze říci, že poptávka po penězích v České republice v letech 2003-2008 je významně determinována reálným ekonomickým růstem a krátkodobou tržní úrokovou sazbou a je tedy možné ji efektivně ovlivňovat nástroji monetární politiky centrální banky (operace na mezibankovním trhu určující výši tamní krátkodobé úrokové sazby). Česká republika se ovšem stále řadí mezi posttransformační ekonomiky s ne plně rozvinutým kapitálovým trhem. Stejně tak řada strukturálních změn související s procesem evropské integrace významně omezují stabilitu poptávky po penězích testovanou metodami vyvinutými především pro rozvinuté tržní ekonomiky a dlouhé časové řady.

5. Literatura

- [1]BROWN, R.L., DURBIN, J. AND EVANS, J.M. (1975): TECHNIQUES FOR TESTING THE CONSTANCY OF REGRESSION RELATIONSHIPS OVER TIME. JOURNAL OF THE ROYAL STATISTICAL SOCIETY B37, pp.149-163;
- [2]GREENE, W.H. (2003): ECONOMETRIC ANALYSIS. FIFTH EDITION. NEW JERSEY: PEARSON EDUCATION. 2003. pp.134-143. ISBN: 0-13-066189-9;
- [3]HANSEN, B.E. (1992): TESTING FOR PARAMETER INSTABILITY IN LINEAR MODELS. JOURNAL OF POLICY MODELING. VOL 14 PP. 517-533;
- [4]HUŠEK, R., PELIKÁN, J. (2003): APLIKOVANÁ EKONOMETRIE. TEORIE A PRAXE. PRAHA: PROFESSIONAL PUBLISHING. 2003. PP 113-138, ISBN: 80-86419-29-0;
- [5]MACH, M. (2002): MAKROEKONOMIE, POKROČILEJŠÍ ANALÝZA. PRAHA: MELANDRIUM. 2002. PP 7-74, ISBN: 80-86175-22-7;
- [6]REVENDA, Z. A KOL. (1999): PENĚŽNÍ EKONOMIE A BANKOVNICTVÍ. PRAHA: MANAGEMENT PRESS. 1999. PP 389-413. ISBN: 80-85943-49-2;
- [7]THOMAS, R.L. (1993): INTRODUCTORY ECONOMETRICS: THEORY AND APPLICATIONS. LONDON: LONGMAN ECONOMICS. 1993. ISBN: -0-582-07378-2

Adresa autora:

Ing. Svatopluk Kapounek Ph.D.

Ústav financí

Mendelova zemědělská a lesnická univerzita
v Brně

Zemědělská 1

Česká republika

613 00 Brno

skapounek@mendelu.cz

Diferenciácia slovenských domácností vyplývajúca z ekonomickej aktivity

Differentiation of Slovak Households Followed from the Activity Status

Alena Kaščáková, Gabriela Nedelová

Abstract: The aim of this article is to detect the difference in demographic and income-expenditure characteristics among the Slovak households separated into groups according to activity status using the logistic regression.

Key words: *household income, consumption expenditures, Household Budget Surveys, activity status, logistic regression.*

Kľúčové slová: *příjem domácnosti, spotrebné výdavky, štatistika rodinných účtov, ekonomická aktivita, logistická regresia.*

1. Úvod

Cieľom príspevku je popísať rozdiely v demografických a príjmo-výdavkových charakteristikách medzi skupinami slovenských domácností v členení podľa ekonomickej aktivity s využitím nominálnej (multinomickej) logistickej regresie.

Pre analýzu sme použili najnovšie údaje štatistiky rodinných účtov, ktoré sme mali k dispozícii a to za rok 2006, pre ich spracovanie sme využili štatistický softvér SPSS.

2. Vybrané ukazovatele domácností v členení podľa ekonomickej aktivity

Do roku 2003 sa na Slovensku realizoval kvótový zámerný výber spravodajských jednotiek do výberovej vzorky. Na základe odporúčania Eurostatu sa od roku 2004 zmenil spôsob výberu a od 1.1.2004 sa výber uskutočňuje náhodne. Zmenil sa aj spôsob klasifikácie domácností podľa ekonomickej aktivity.

Tabuľka 1. Klasifikácia domácností podľa ekonomickej aktivity

Do roku 2003		Od roku 2004	
1	manuálne pracujúci mimo poľnohospodárstva	1	pracujúci (na plný aj čiastočný úväzok)
2	samostatne zárobkovo činný	2	zamestnaný, ale dočasne mimo prácu
3	nemanuálne pracujúci	3	nezamestnaný
5	manuálne pracujúci v poľnohospodárstve	4	nepracujúci starobný dôchodca
7	nepracujúci starobný dôchodca	5	študent, učeň, vojak základnej služby
		6	ekonomicky neaktívny, žena v domácnosti
		7	neschopný práce
		9	neaplikovateľné (dieťa do 16 rokov, ak nezaradené, do kódu 5)

Z rodinných účtov je možné poznať informácie o počte členov domácností, ich štruktúre podľa ekonomickej aktivity, zamestnania, postavení referenta domácnosti. Rodinné účty podávajú široký a podrobný prehľad o príjmoch domácností aj o výdavkovej štruktúre podľa medzinárodnej klasifikácie COICOP, podľa ktorej je možné všetky výdavky domácností rozčleniť do 12 výdavkových skupín. Z nasledujúcich tabuliek je zrejma diferencovanosť sledovaných ukazovateľov v domácnostiach členených podľa ekonomickej aktivity.

Tabuľka 2. Demografické a príjmové charakteristiky domácností podľa skupín ekonomickej aktivity (v „novom“ členení platnom od roku 2004) v priemere za mesiac v roku 2006

EA	Počet členov domácnosti	Počet nezaopatrených detí	Hrubý príjem	Sociálny príjem	Príjmy z majetku	Iné príjmy	Čistý príjem
1	3,29	1,13	34 125,39	3 901,90	90,57	1 799,51	28 998,42
2	3,09	0,82	25 632,27	9 314,55	0,00	1 133,55	22 993,73
3	2,91	0,94	14 756,86	8 092,84	15,77	2 182,86	14 029,91
4	1,82	0,08	15 975,56	12 776,86	42,87	391,59	15 509,65
5	2,00	0,00	13 680,00	5 430,00	0,00	7 250,00	13 680,00
6	2,77	0,91	20 476,10	8 010,57	17,61	2 746,33	19 330,92
7	2,46	0,33	20 546,67	11 512,37	65,04	3 509,11	19 677,50
9	1,00	0,00	7 317,50	7 317,50	0,00	0,00	7 317,50

Zdroj: Vlastné spracovanie. Údaje rodinných účtov, ŠÚ SR, 2007.

Tabuľka 3. Objem priemerných mesačných výdavkov domácností podľa skupín ekonomickej aktivity v roku 2006

EA	Potraviny a nealkoholické nápoje	Alkohol a tabak	Odievanie a obuv	Bývanie, voda, energia, palivá	Nábytok a zariadenie domácnosti	Zdravie
1	6 017,49	738,63	1 616,96	5 546,47	1 201,08	589,24
2	5 703,02	395,42	1 780,35	4 838,05	394,85	676,82
3	4 361,97	590,59	684,43	3 921,64	329,16	258,80
4	4 253,33	406,34	621,25	4 690,73	741,86	714,06
5	2 420,25	116,50	327,00	3 333,00	22,60	630,30
6	4 898,60	612,45	528,24	3 905,39	391,11	409,51
7	4 884,37	594,50	735,70	4 652,82	787,88	691,35
9	1 824,25	54,00	205,00	3 925,00	161,30	368,50

pokračovanie tabuľky

EA	Doprava	Pošta a telekomunikácie	Rekreácia a kultúra	Vzdelanie	Hotely a reštaurácie	Ostatné tovary a služby
1	2 469,46	1 426,05	1 918,60	226,06	1 535,64	2 326,94
2	1 920,83	972,46	971,44	90,91	986,98	1 484,12
3	701,84	836,08	738,80	30,29	487,71	949,58
4	723,94	621,46	787,81	40,83	283,62	771,70
5	1 252,00	840,50	308,45	212,50	394,00	843,40
6	797,47	784,35	1 087,58	12,50	518,01	1 135,29
7	823,00	879,12	1 070,41	41,06	449,23	956,08
9	70,00	429,45	145,25	0,00	0,00	282,40

Zdroj: Vlastné spracovanie. Údaje rodinných účtov, ŠÚ SR, 2007.

3. Použitie logistickej regresie

Pre analýzu domácností za rok 2006 bolo k dispozícii 4701 pozorovaní o hospodárení domácností a ich demografických charakteristikách. Štruktúra sledovaných domácností podľa ekonomickej aktivity je uvedená v tabuľke 4.

Pre veľmi malú početnosť sme z analýzy vylúčili skupiny 2, 5 a 9, čiže celkovo 15 pozorovaní. Najpočetnejšiu skupinu 1- domácnosti pracujúcich - sme ponechali bezo zmien,

skupiny 3, 6 a 7 sme pre analýzu zlúčili, ďalej ich nazveme skupinou, kde referent domácnosti je nepracujúci, ale nie dôchodca (nepracujúce domácnosti) a vo výstupe v prílohe je táto skupina označená ako skupina 3. Poslednou sledovanou skupinou domácností sú dôchodcovské domácnosti (skupina 4). Referenčnou skupinou oproti ktorej sa predchádzajúce dve skupiny domácností porovnávali bola skupina 1 – domácnosti pracujúcich.

Tabuľka 4. Štruktúra sledovaného súboru podľa ekonomickej aktivity

Skupina EA	Počet pozorovaní	Podiel (v %)	Skupina EA	Počet pozorovaní	Podiel (v %)
1	3 033	64,5	5	2	0,0
2	11	0,2	6	44	0,9
3	128	2,7	7	123	2,6
4	1 358	28,9	9	2	0,3
Celkom				4 701	100,0

Do modelu boli najprv vybrané všetky demografické, príjmové a výdavkové charakteristiky, postupne však boli vylúčené nevýznamné charakteristiky. Vysvetľovanou premennou bola ekonomická aktivita domácnosti, štatisticky významnými vysvetľujúcimi premennými demografické charakteristiky - počet členov domácností a počet nezaopatrených detí, charakteristiky príjmu – čistý peňažný príjem a sociálny príjem a výdavkové položky – náklady na bývanie, dopravu, vzdelanie a zdravie. Na hladine významnosti 0,05 je model ako celok aj všetky uvedené vysvetľujúce premenné štatisticky významné. Nagelkerkovo $R^2 = 0,735$ vyjadruje dobrú kvalitu modelu, diskriminačná sila modelu je vysoká, model dobre klasifikoval 89,3 % prípadov z celkového počtu domácností. V prípade pracujúcich domácností bolo dobre klasifikovaných 95,7 %, z dôchodcovských domácností bolo správne klasifikovaných 92,6 % prípadov a zostávajúca skupina domácností nepracujúcich osôb bola správne klasifikovaná len v 7,8 % prípadov (ale v tejto skupine domácností je aj najnižší počet pozorovaní – len 6,3 % z celkového počtu sledovaných domácností).

Všetky výsledky využitia multinomickej logistickej regresie sú uvedené v prílohe a popísané v závere.

4. Záver

Interpretáciou získaných parametrov modelu zisťujeme, že v roku 2006 pre slovenské domácnosti pracujúcich, nepracujúcich a dôchodcov platili nasledujúce vzťahy:

- ak sa zvýšil počet členov domácnosti o 1 osobu, šanca, že domácnosť je
 - o nepracujúca vzrástla oproti tomu, že je pracujúca 1,54 krát
 - o dôchodcovská oproti tomu, že je pracujúca poklesla 0,73 krát,
- ak sa zvýšil počet nezaopatrených detí, pri nezmenených ostatných vysvetľujúcich charakteristikách, šance, že ide o pracujúcu domácnosť v oboch skupinách (nepracujúca i dôchodcovská) domácností poklesli (viac v skupine dôchodcovských domácností),
- to isté platí aj pre čistý príjem domácnosti,
- ak sa zvýšil sociálny príjem domácnosti o 1 tis. Sk, šanca, že ide o nepracujúcu domácnosť sa zvýšila 1,31 krát, že ide o dôchodcovskú domácnosť 1,62 krát oproti pracujúcej domácnosti,
- ak sa zvýšili výdavky na bývanie, dopravu aj na zdravie, šance, že domácnosť je dôchodcovská alebo nepracujúca klesli oproti pracujúcej domácnosti, avšak v dôchodcovských domácnostiach bol pokles šance len veľmi mierny,
- zaujímavý je výsledok týkajúci sa výdavkov na vzdelanie – pri zvýšení výdavkov na vzdelanie o 1000 Sk sa mení šanca, že domácnosť je nepracujúca 0,653 krát (klesá)

ale šanca, že domácnosť je dôchodcovská rastie oproti tomu, že domácnosť je pracujúca 1,15 krát.

V slovenských domácnostiach v roku 2006 je možné považovať za najvýznamnejšie premenné identifikujúce príslušnosť ku skupine domácností podľa ekonomickej aktivity (pracujúcich, nepracujúcich a dôchodcovských domácností) nasledujúce charakteristiky: počet členov domácnosti, počet nezaopatrených detí, čistý a sociálny príjem a z výdavkových položiek hlavne náklady na vzdelanie a na zdravie, čiastočne aj na bývanie a dopravu.

5. Literatúra

- [1] HEBÁK, P. A KOLEKTÍV. 2005. Vícerozměrné statistické metody [3]. Praha: Informatorium, 2005. ISBN 80-7333-039-3.
- [2] PECÁKOVÁ, I. 2007. Logistická regrese s vícekategoriální vysvětlovanou proměnnou. In: Acta Oeconomica Pragensia. 2007, roč. 15, č. 1. Praha: Vysoká škola ekonomická, 2007, s. 86 – 96. ISSN 0572-3043.
- [3] ŘEHÁKOVÁ, B. 2000. Nebojte se logistické regrese. In: Sociologický časopis. 2000, roč. 36, č. 4. Praha: Sociologický časopis, 2000 AV ČR, s. 475 – 492. ISSN 0038-0288.
- [4] STANKOVIČOVÁ, I., VOJTKOVÁ, M. 2007. Viacrozmerné štatistické metódy s aplikáciami. Bratislava: Iura Edition, 2007. ISBN 978-80-8078-152-1.
- [5] STANKOVIČOVÁ, I. 2007. Logistická regresia a jej využitie v ekonomickej praxi. In: Forum Statisticum Slovaca. 2007, roč. III, č. 1. Bratislava: Slovenská štatistická a demografická spoločnosť, 2007, s. 42– 54. ISSN 1336-7420.
- [6] STANKOVIČOVÁ, I. 2008. Logistická regresia a jej využitie v ekonomickej praxi. In: Forum Statisticum Slovaca. 2008, roč. IV, č. 2. Bratislava: Slovenská štatistická a demografická spoločnosť, 2008, s. 117 – 128. ISSN 1336-7420.

Adresa autora:

Alena Kaščáková, Ing., PhD.
KKMI EF UMB
Tajovského 10
975 90 Banská Bystrica
alena.kascakova@umb.sk

Gabriela Nedelová, RNDr., PhD.
KKMI EF UMB
Tajovského 10
975 90 Banská Bystrica
gabriela.nedelova@umb.sk

Príspevok vznikol s podporou grantu KEGA 3/5214/07.

Prílohy:

Model Fitting Information

Model	Model Fitting Criteria	Likelihood Ratio Tests		
	-2 Log Likelihood	Chi-Square	df	Sig.
Intercept Only	7634,383			
Final	3307,791	4326,592	16	,000

Goodness-of-Fit

	Chi-Square	df	Sig.
Pearson	1,74 E + 021	9354	,000
Deviance	3307,791	9354	1,000

Likelihood Ratio Tests

Effect	Model Fitting Criteria	Likelihood Ratio Tests		
	-2 Log Likelihood of Reduced Model	Chi-Square	df	Sig.
Intercept	3344,281	36,490	2	,000
pocclen	3348,098	40,307	2	,000
nezdet	3366,358	58,568	2	,000
cptis	3698,213	390,422	2	,000
socprtis	5384,942	2077,152	2	,000
byvtis	3317,599	9,808	2	,007
doprtis	3314,637	6,846	2	,033
vzdtis	3315,841	8,050	2	,018
zdrtis	3317,859	10,069	2	,007

Classification

Observed	Predicted			
	1	3	4	Percent Correct
1	2903	6	124	95,7%
3	118	23	154	7,8%
4	94	6	1258	92,6%
Overall Percentage	66,5%	,7%	32,8%	89,3%

Pseudo R-Square

Cox and Snell	,608
Nagelkerke	,735
McFadden	,533

Parameter Estimates

EA ^a		B	Std. Error	Wald	df	Sig.	Exp(B)	95% Confidence Interval for Exp(B)	
								Lower Bound	Upper Bound
3	Intercept	-1,175	,195	36,227	1	,000			
	pocclen	,430	,109	15,691	1	,000	1,538	1,243	1,903
	nezdet	-,552	,132	17,613	1	,000	,576	,445	,745
	cptis	-,138	,012	127,943	1	,000	,871	,851	,892
	socprtis	,272	,015	313,165	1	,000	1,312	1,273	1,352
	byvtis	-,087	,032	7,407	1	,006	,916	,861	,976
	doprtis	-,129	,066	3,848	1	,050	,879	,773	1,000
	vzdtis	-,427	,330	1,671	1	,196	,653	,342	1,246
	zdrtis	-,345	,123	7,866	1	,005	,708	,557	,901
4	Intercept	-,266	,163	2,650	1	,104			
	pocclen	-,319	,108	8,661	1	,003	,727	,588	,899
	nezdet	-1,149	,160	51,501	1	,000	,317	,232	,434
	cptis	-,174	,011	249,002	1	,000	,840	,822	,858
	socprtis	,482	,017	828,309	1	,000	1,619	1,567	1,673
	byvtis	-,006	,019	,089	1	,766	,994	,958	1,032
	doprtis	-,012	,011	1,151	1	,283	,988	,967	1,010
	vzdtis	,137	,054	6,503	1	,011	1,147	1,032	1,275
	zdrtis	-,102	,056	3,330	1	,068	,903	,809	1,008

a. The reference category is: 1.

Grilichesove konzistentné hranice parametrov Cobbovej-Douglasovej produkčnej funkcie stavebných podnikov

Griliches' consistent bounds of Cobb-Douglas production function parameters of construction companies

Samuel Koróny¹

Abstract: The paper deals with application of Griliches' consistent bounds of Cobb-Douglas production function parameters. Examined production function was estimated from absolute economic production indicators of Slovak construction joint stock and ltd. companies from the year 2001 (value added, sum of assets and average number of employees).

Keywords: Regression analysis, Errors in variables, Construction sector

Kľúčové slová: Regresná analýza, Chyby v premenných, Stavebníctvo

1. Úvod

Príspevok voľne nadväzuje na Klepperove ohraničenia regresných parametrov (Klepper 1984) aplikované na Cobbovu-Douglasovu produkčnú funkciu (ďalej CDPF) slovenských stavebných podnikov (Koróny 2008). Grilichesov postup pre zistenie konzistentných ohraničení regresných parametrov je starší (Maddala 1992 podľa Griliches 1971) a vychádza z predpokladu znalosti podielu rozptylu chybovej zložky voči príslušnej premennej.

2. Dáta

Pre zistenie konzistentných hraníc parametrov CDPF boli opäť použité vybrané ukazovatele - pridaná hodnota (v mil. Sk) ako výstup, priemerný evidenčný počet zamestnancov vo fyzických osobách a celkový kapitál (v mil. Sk) ako vstupy súkromných slovenských stavebných podnikov právnej formy a. s. a s. r. o. Údaje sú z výkazu ŠÚ SR Prod 3-04 za rok 2001. Regresná analýza dát bola urobená v štatistickom systéme SPSS verzia 13.

3. Vzťahy pre Grilichesove konzistentné hranice regresných parametrov

Majme lineárny regresný model s dvomi nezávislými premennými v normalizovanom tvare

$$y = \beta_1 x_1 + \beta_2 x_2 + e,$$

v ktorom všetky premenné sú merané s chybou. Nech pozorované premenné sú

$$Y = y + v, X_1 = x_1 + u_1, X_2 = x_2 + u_2.$$

Za navzájom nekorelované považujeme chybové zložky u_1, u_2, v a tiež nekorelované s premennými x_1, x_2, y . Chybové zložky e a v sa dajú zahrnúť do jednej zloženej chyby v premennej Y , preto $\sigma^2 = \text{var}(e + v)$. Nech $\text{cov}(X_1, X_2) = \rho$. Rozptyly relatívnych chýb nezávislých premenných označíme

¹ Samuel Koróny, Ústav vedy a výskumu UMB, Banská Bystrica

$$\lambda_1 = \frac{\text{var}(u_1)}{\text{var}(X_1)} = \text{var}(u_1), \quad \lambda_2 = \frac{\text{var}(u_2)}{\text{var}(X_2)} = \text{var}(u_2).$$

Regresná rovnica v pozorovaných premenných je

$$Y = \beta_1 X_1 + \beta_2 X_2 + w,$$

kde

$$w = e + v - \beta_1 u_1 - \beta_2 u_2.$$

Pre rozptyl závislej premennej potom platí

$$\text{var}(Y) = \beta_1^2 + \beta_2^2 + 2\rho\beta_1\beta_2 + \sigma^2 - \beta_1^2\lambda_1 - \beta_2^2\lambda_2 = 1. \quad (1)$$

Po úpravách dostaneme pravdepodobnostné limity $\hat{\beta}_1, \hat{\beta}_2$ pre regresné parametre β_1, β_2 získané metódou najmenších štvorcov

$$\begin{aligned} \text{plim} (\hat{\beta}_1 - \beta_1) &= -\frac{\beta_1\lambda_1 - \rho\beta_2\lambda_2}{1 - \rho^2}, \\ \text{plim} (\hat{\beta}_2 - \beta_2) &= -\frac{\beta_2\lambda_2 - \rho\beta_1\lambda_1}{1 - \rho^2}. \end{aligned} \quad (2)$$

Na základe vzťahu (2) a rozsahu predpokladaných hodnôt pre parametre relatívnej chyby λ_1, λ_2 je potom možné dostať intervaly konzistentných odhadov β_1, β_2 . Na základe toho sa dá urobiť analýza citlivosti regresných parametrov na relatívnu veľkosť rozptylu chýb. Vychýlenie odhadu je navyše zväčšené faktorom $1/(1 - \rho^2)$, pričom pre väčšie ρ , je väčšie aj vychýlenie. Multikolinearita teda zhoršuje odhady aj v tomto prípade.

Pre hodnoty parametrov λ_1, λ_2 sú však isté ohraňovania, ktoré vyplývajú z toho, že musí platiť podmienka

$$(1 - \lambda_1)(1 - \lambda_2) - \rho^2 \geq 0. \quad (3)$$

Z nej vyplýva, že pre vyššie hodnoty ρ^2 (vysokú multikolinearitu medzi premennými X_1 a X_2) nemôžeme vôbec uvažovať niektoré hodnoty parametrov λ_1, λ_2 .

Pre získanie správnych hodnôt regresných parametrov musí byť splnená ešte druhá podmienka kladnosti rozptylu celkovej chybovej zložky

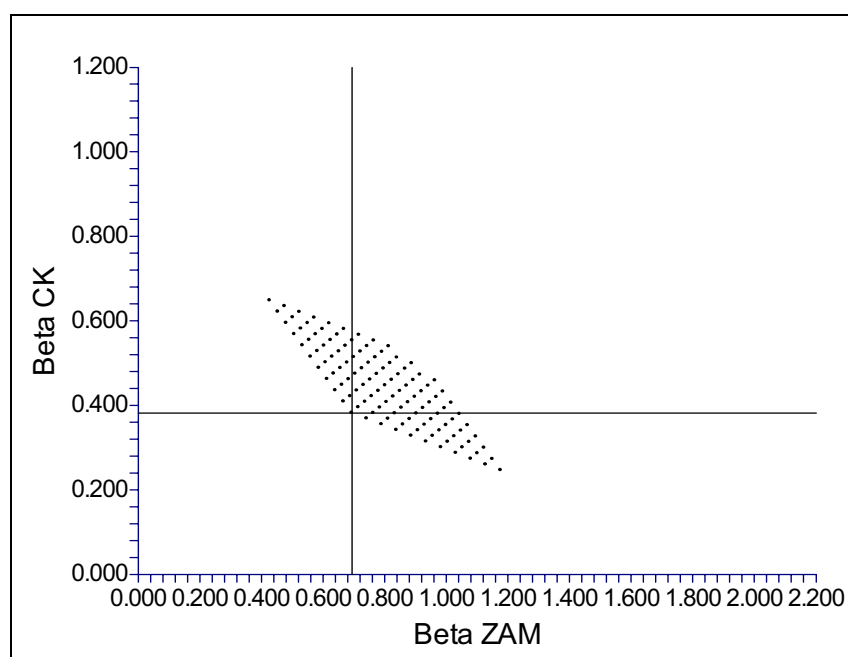
$$\sigma^2 = 1 - (\beta_1^2 + \beta_2^2 + 2\rho\beta_1\beta_2 - \lambda_1\beta_1^2 + \lambda_2\beta_2^2) > 0. \quad (4)$$

Grilichesov postup je teda vhodný, ak máme určitú predstavu o veľkosti chybovej zložky v premenných. Klepperov postup to nevyžaduje, ale jeho výsledkom sú väčšie oblasti konzistencie.

4. Výsledky pre Grilichesove konzistentné hranice regresných parametrov

Na základe Grilichesovho postupu je možné získať intervaly pre konzistentné hodnoty regresných parametrov CDPF pri uvažovaní rozsahov konkrétnych hodnôt relatívnej chyby v premenných. V tabuľke 1 a grafe 1 sú uvedené výsledky. Uvažovali sme o hodnotách relatívneho rozptylu chýb v intervale od 0 po 0,5 s krokom 0,05.

Pre stavebné podniky právnej formy s. r. o. pôvodné hodnoty parametrov CDPF bez uvažovania chýb v premenných sú $\beta_{ZAM} = 0,694$ a $\beta_{CK} = 0,382$. Na grafe 1 je to bod označený priesečníkom priamok. Parameter účinnosti pracovnej sily má minimálnu hodnotu $\beta_{ZAM} = 0,427$ pre $\lambda_{ZAM} = 0$, $\lambda_{CK} = 0,5$. Vtedy je parameter účinnosti celkového kapitálu maximálny a rovná sa $\beta_{CK} = 0,648$. Maximálnu hodnotu má parameter účinnosti pracovnej sily $\beta_{ZAM} = 1,176$ pre $\lambda_{ZAM} = 0,5$; $\lambda_{CK} = 0$. Účinnosť celkového kapitálu je naopak minimálna a rovná sa $\beta_{CK} = 0,246$.



Graf 1: Zobrazenie Grilichesových konzistentných oblastí regresných parametrov CDPF podnikov formy s. r. o. pre relatívne chyby premenných do 0,5

V tabuľke 1 sú uvedené súradnice vrcholov obalu konvexnej množiny konzistentných parametrov CDPF stavebných podnikov právnej formy s. r. o.

Tabuľka 1: Hraničné hodnoty Grilichesových konzistentných regresných parametrov CDPF stavebných podnikov právnej formy s. r. o.

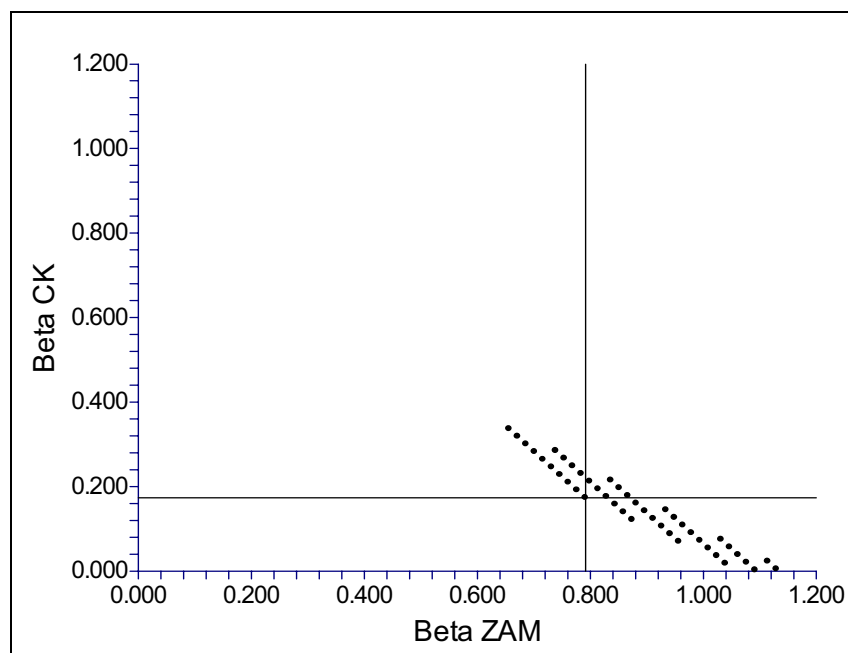
λ_{ZAM}	λ_{CK}	σ^2	β_{ZAM}	β_{CK}
0.00	0.0	0.378	0.694	0.382
0.00	0.5	0.397	0.427	0.648
0.40	0.5	0.271	0.813	0.539
0.50	0.0	0.395	1.176	0.246
0.50	0.4	0.269	0.963	0.459

V tabuľke 2 a grafe 2 sú uvedené výsledky pre stavebné podniky právnej formy a. s. Znova sme uvažovali o hodnotách relatívneho rozptylu chýb v intervale od 0 po 0,5 s krokom 0,05, ale mnoho ich pre väčšie hodnoty chybovej zložky vypadlo, lebo nedávali konzistentné hodnoty regresných parametrov CDPF (nesplňali podmienku 3 alebo 4).

Pre stavebné podniky právnej formy a. s. pôvodné hodnoty parametrov účinnosti CDPF bez uvažovania chýb sú $\beta_{ZAM} = 0,792$ a $\beta_{CK} = 0,174$.

Minimálnu hodnotu má účinnosť pracovnej sily $\beta_{ZAM} = 0,657$ pre $\lambda_{ZAM} = 0$, $\lambda_{CK} = 0,45$. Vtedy je účinnosť celkového kapitálu maximálna a rovná sa $\beta_{CK} = 0,337$.

Maximálnu hodnotu má účinnosť pracovnej sily $\beta_{ZAM} = 1,039$ pre $\lambda_{ZAM} = 0,15$; $\lambda_{CK} = 0$. Vtedy je účinnosť celkového kapitálu naopak minimálna a rovná sa $\beta_{CK} = 0,018$. Hodnota účinnosti celkového kapitálu je prakticky rovná nule, čo ukazuje na nezrovnalosti v dostupných dátach. Možných miním je viac, vybrali sme daný prípad, pretože je v zhode s Klepperovým ohraničením.



Graf 2: Zobrazenie Grilichesových konzistentných oblastí regresných parametrov CDPF podnikov formy a. s. pre relatívne chyby premenných do 0,5

V tabuľke 2 sú prípady, ktoré ležia na hranici a tvoria vrcholy obal konvexnej množiny konzistentných parametrov CDPF stavebných podnikov právnej formy a. s.

Tabuľka 2: Hraničné hodnoty Grilichesových konzistentných regresných parametrov CDPF stavebných podnikov právnej formy a. s.

λ_{ZAM}	λ_{CK}	σ^2	β_{ZAM}	β_{CK}
0.00	0.00	0.317	0.792	0.174
0.00	0.45	0.329	0.657	0.337
0.15	0.00	0.275	1.039	0.018

Na grafoch 1 a 2 sú zobrazené oblasti konzistencie regresných parametrov účinnosti vstupných produkčných faktorov priemerného evidenčného počtu zamestnancov a celkového úhrnu aktív CDPF stavebných podnikov. Pre orientáciu v nich je užitočné vedieť, že ak rastie relatívna chyba danej premennej (produkčného faktora), tak rastie aj konzistentná hodnota jej regresného parametra. Relatívna chyba priemerného evidenčného počtu zamestnancov CDPF rastie smerom zľava doprava. Relatívna chyba celkového úhrnu aktív CDPF rastie smerom zdola nahor.

Z grafu 1, ktorý zobrazuje situáciu v prípade stavebných podnikov právnej formy s. r. o. vyplýva určitá symetria hodnôt voči hodnote získanej metódou najmenších štvorcov.

Pre stavebné podniky právnej formy a. s. (graf 2) je zreteľná nesúmernosť v hodnotách parametrov účinnosti Cobbovej-Douglasovej produkčnej funkcie. Účinnosť celkového kapitálu (celkového úhrnu aktív) má nižšie hodnoty pre všetky kombinácie relatívnych chýb produkčných faktorov. Výsledkom toho je, že mnoho hodnôt parametra účinnosti celkového úhrnu aktív Cobbovej-Douglasovej produkčnej funkcie je prakticky rovné nule. Pravdepodobne to ukazuje na ďalej nezisťované nezrovnalosti vo vstupných údajoch.

5. Záver

Grilichesov postup pre zistenie konzistentných ohraničení regresných parametrov je vhodný vtedy, ak máme aspoň určité konkrétne informácie o veľkosti chýb v nezávislých premenných. Jeho ďalšou výhodou je názorná geometrická interpretácia, ktorá sa samozrejme stráca pri použití viac ako dvoch nezávislých premenných. Autor dúfa, že takto upozorní na určité možnosti riešenia regresie s chybami v premenných.

6. Literatúra

- [1] GRILICHES, Z. – RINGSTAD, V. 1971. Economies of Scale and the Form of the Production Function. Amstredam : North-Holland, 1971. ISBN 9780720431728
- [2] KLEPPER, S. – LEAMER E.E. 1984. Consistent Sets for Regressions with Errors in All Variables. In: Econometrica, vol.52, No. 1, pp. 163-183
- [3] KORÓNY, S. 2008. Klepperove konzistentné hranice pre parametre Cobbovej-Douglasovej produkčnej funkcie stavebných podnikov. In: Forum Statisticum Slovacum. Roč.3, č.1, 2007, s. 55-66. ISSN 1336-7420
- [4] MADDALA, G. S. 1992. Introduction to Econometrics. London: Prentice-Hall, 1992. ISBN 0-13-880352-8

Adresa autora:

RNDr. Samuel Koróny
Ústav vedy a výskumu UMB
Cesta na amfiteáter 1
974 01 Banská Bystrica
Email: samuel.korony@umb.sk

Situation in University Educational System in the Czech Republic

Bohdan Linda, Jana Kubanová

Abstract:

The main target of the paper is the university educational system analysis. Authors believe, that the contemporary system of both three-stage education, and system of financing negative influence level of education.

Key words: number of graduates of high schools, number of students at universities, quality of graduates, students enrolled to universities

This paper deals with quality of the university education and it's an attempt to document the tendency of its development. In advance we want to draw attention to reality, that none tests of the level of university education were provided, but speculations are performed both on the basis of personal experiences, and some experience transmitted by colleagues from other universities. Is generally known that the pedagogues systematically appeal decreasing quality of students. This is a problem of all kind of schools. We try to analyze, from statistical point of view, the reasons of this situation in the time horizon of post November period.

Generally dominates the opinion, that if we want to grant certain education to more people, the lower is final level of this education. To be able to express this opinion, it is necessary to clarify at first the term „to grant certain education". As long as we understand this term „give education" (accurately to give any relevant certificate about graduation of studies), then this statement is true. However as far as we understand this term „to give a possibility to obtain this education in the case, that the applicant is interested in knowledge", then the mentioned statement is insignificant. It is possible to reply that this isn't problem, if we during studies affect the unthrifty students. It would not be really problem, if wouldn't have existed a contemporary system of universities financing. These universities are funding on the basis of a number of students, who are registered in certain academic year. This financing system could be accepted in case, when “surplus of applicants”. Have a look at the situation in the Czech Republic. It is evident, that the fundamental determining element for the number of applicants for university education is the number of inhabitants and their age structure. The figure number 1 shows the time series of a number of inhabitants from the year 1989.

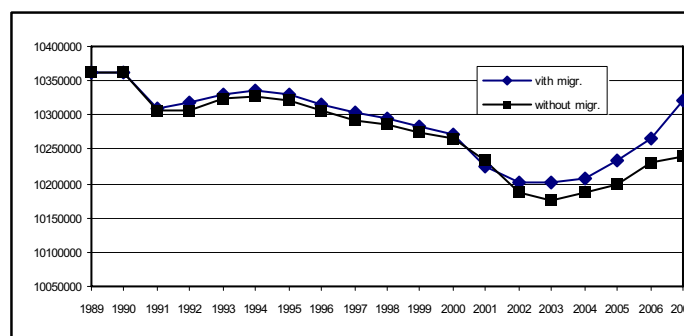


Figure 1: Development of mid-year population with and without migration

The upper line presents the absolute number of inhabitants, that systematically dropped till the year 2004 and after this year the trend reversed. Unfortunately the optimistic growth, above all in the last years, is caused above all by migration. According to the statistical methodology all foreigners either with permanent residence or with visa over 90 days are included with into population. It is evident, that the great deals of foreigners come into the Czech Republic only to work, they haven't family members here and they can't be counted in the reproduction able population.

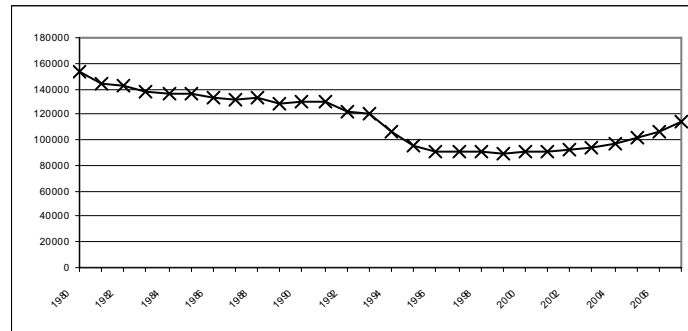


Figure 2. Development of number of births

The lower line in the figure 1 presents the number of population purified from the migration. It hasn't so optimistic trend, above all when we remember, that from the year 2004, when the fertility period of the "Husaks' children" finishes and postponed parenting are probably realized in this period.

The students who were born in the year 1989 enter at university now. The number of born children toboggans from the year 1991 till the year 1996 and then following five years keep approximately equivalently (see figure number 2). It means that minimal by other 6 years the number of young people, that could attend the university, will naturally decline, and minimal by other five years will this number stagnate.

From the other hand, let's look the situation of the number of students at the high schools. The figure number 3 presents the time series of the development of the number of public and private universities and the total number of faculties at both kinds of schools. The intense growth of faculties was registered in the first two past November years. At this time the number increased to 21 which presents approximately 30% growth.

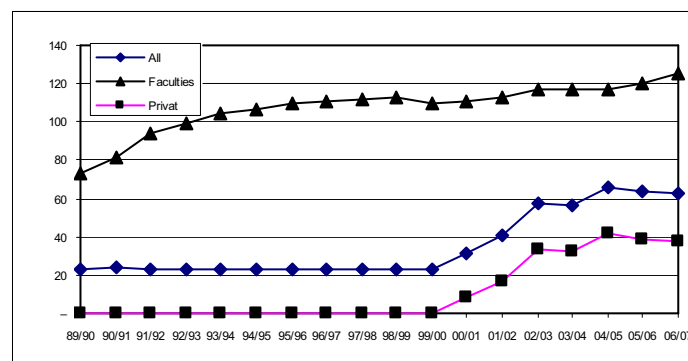


Figure 3. Number of public and private high schools and all faculties

The number of universities stagnated at the number 23 till the year 2000, when started to raise the private universities. The whole growth can be ascribed to private universities, number 38 of them already predominate the number of public ones. More important is the number of university students. The development of this from the year 1993 is demonstrated in the figure 4. The numbers of students in bachelors, masters and doctoral study programs are also stated in the figure. The number of private universities students is stated summary, the number of master studies students was insignificant in past years. The educational programs, in which get on the high school alumni, are important for our further thinking. They are, except some few exceptions mainly bachelor educational programs, because majority of universities passed to the three-stage educational system from the year 2005.

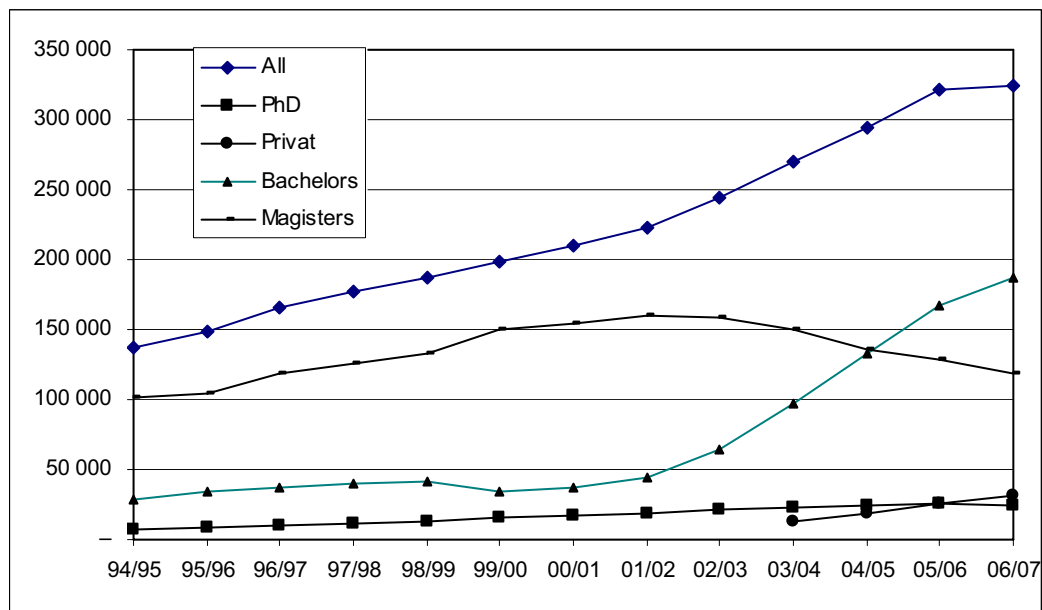


Figure 4 *The development of the number of students at universities totally and according to the study programmes.*

The potential of students, who enter for the first time university, create alumni of the high schools. There're stated the numbers of graduate students at the high schools with GCE examination and for confrontation even numbers of students, entering university in the figure 5. The graph of ingoing students to university provides the relevant information from the year 2004 when finished the five-year master study at majority of schools. The number of entering bachelor students are determining for the numbers of students, who for the first time register at university.

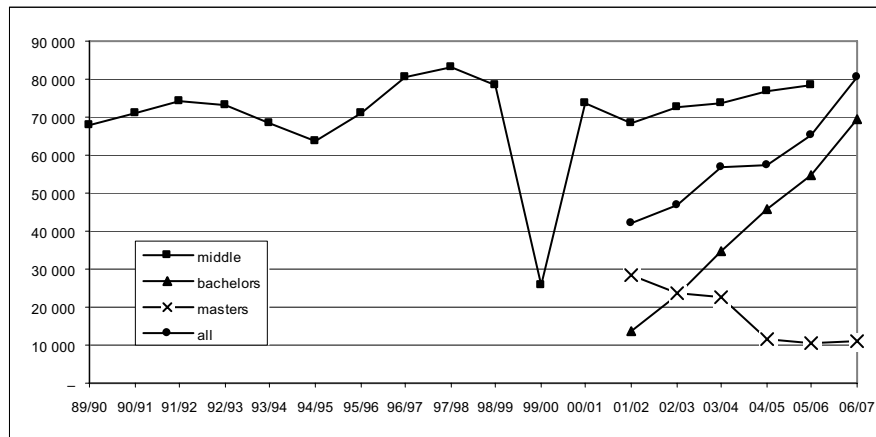


Figure 5. Numbers of graduate students at high schools and number of students entering the universities

The considerable disproportion among number of graduate students at the high schools and number of student entering university is evident in the figure 6. To sketch the situation we introduce even the figure 6. The trend of the number of pupil at basic school and high school is shown.

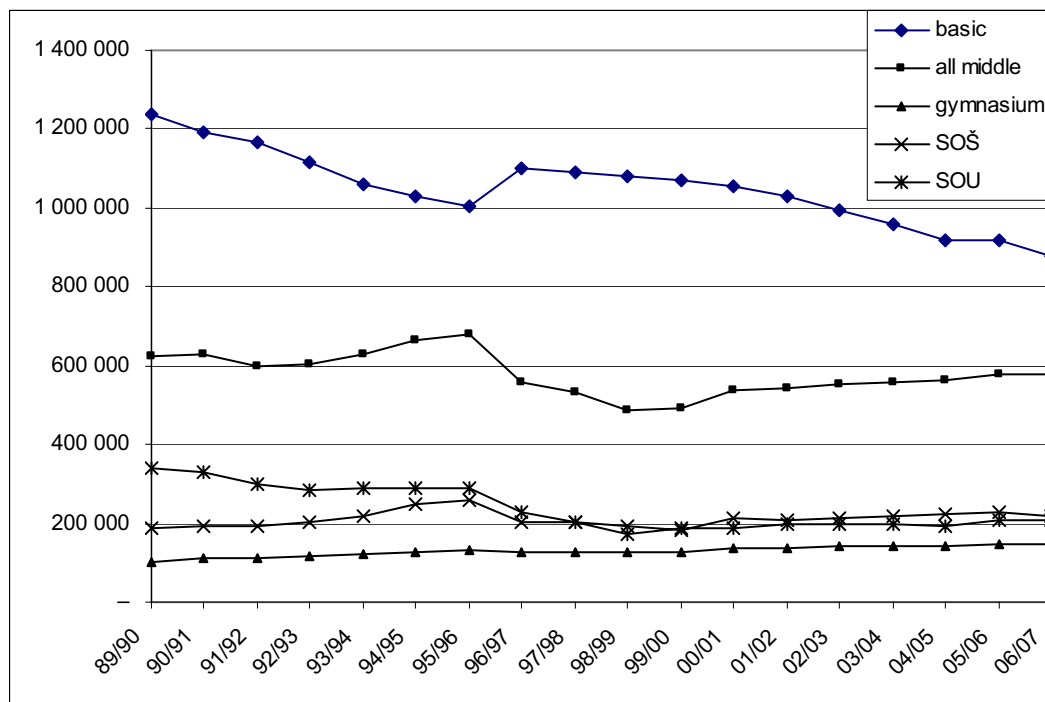


Figure 6. Numbers of pupils of basic and high schools

We can see again any disproportion among these time series. At systematic shortage of pupils of the basic schools grows the number of high schools students.

Finally it is possible to claim, that if declines the number of pupils of the basic schools and grows the number of students at high schools and if the number of university students grows faster than at high schools, this all must express in the quality of graduates. The struggle for students started, unfortunately no resulted in quality production, but in demand reduction. The figure 7 compares fruitfulness of the university study applicants at entrance examinations at public and private universities.

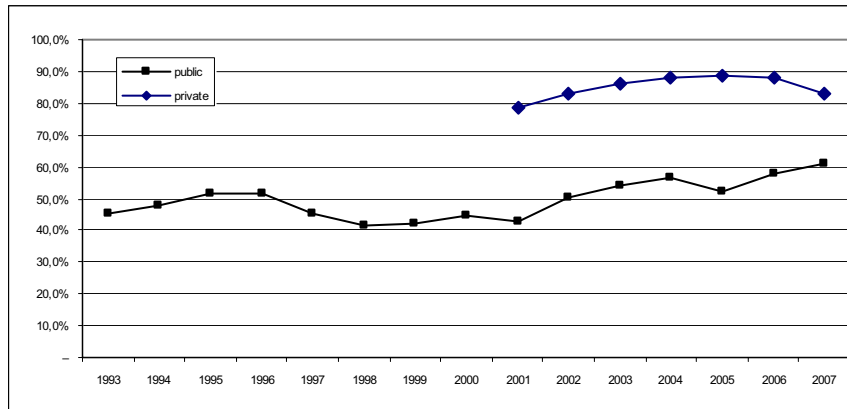


Figure 7. Fruitfulness of the university study applicants at entrance examinations at public and private universities.

Very interesting is improvement of fruitfulness from the moment, when started to work the private schools. Fruitfulness increased at public universities from 45% to more than 60%. Fruitfulness at private schools was over 80%. We can believe, that the ministry should re-value the system of universities financing with regard to presented numbers. The system of three-stage education should be revaluated as well. It is acceptable only for some types of study programs, however certainly not for technical and economical programs. It is good idea to provide the university education to widest population; however it shouldn't decrease quality of education. Some kind of solution might be comeback to the parallel system of the three-year bachelor and five-year master study, mutually impenetrable.

References

[1]<http://www.czso.cz>

[2]<http://www.uiv.cz>

Addresses of authors:

Jana Kubanová, doc., PaedDr, CSc.
Studentská 95
532 10 Pardubice
jana.kubanova@upce.cz

Bohdan Linda, doc., RNDr., CSc.
Studentská 95
532 10 Pardubice
bohdan.linda@upce.cz

Korelácia komplementárnych a disjunktných javov Correlation of complementary and disjoint events

Ján Luha

Abstract: Article dealt with specific model correlation of complementary and disjoint events. Application is provided by results from population research with questions with nominal scales.

Key words: complementary events, disjoint events, qualitative attribute, nominal scale, correlation of disjoint and complementary events, indicator variables.

Kľúčové slová: komplementárne javy, disjunktné javy, kvalitatívny znak, nominálna škála, korelácia disjunktných a komplementárnych javov, indikátorové premenné.

1. Úvod

V príspevku sa zaoberáme špecifickou situáciou korelácie elementárnych javov definovaných kvalitatívnym znakom. Pri skúmaní ich závislostí môžeme využiť vyjadrenie pomocou indikátorových (nula/jedna) premenných. Na výpočet korelácie používame vtedy Pearsonov korelačný koeficient. V príspevku ukážeme zjednodušené vzorce na výpočet korelácií pre komplementárne a disjunktné javy a dôkaz o ich zápornej korelácií.

Príspevok nadväzuje na prácu [9]Luha J.(2008): Korelácie vnorených javov. FORUM STATISTICUM SLOVACUM 5/2008. SŠDS Bratislava 2008. ISSN 1336-7420.

Indikátorová premenná Z je definovaná:

$$Z = \begin{cases} 1, & \text{ak skúmaný jav nastal,} \\ 0, & \text{ak skúmaný jav nenastal.} \end{cases}$$

Ak $E(Z)=p$, tak $D(Z)=E(Z - E(Z))=p(1-p)$.

Označme Ω množinu všetkých skúmaných prvkov. Potom zrejme $1 - Z$ je indikátorová premenná komplementárneho javu $\Omega - Z$.

Disjunktné javy definujeme jednoducho pomocou príslušných podmnožín elementárnych javov skúmaného kvalitatívneho znaku. Uvažujme kvalitatívny znak Z s množinou hodnôt $\{1, 2, \dots, r\}$. Tento znak reprezentujeme pomocou r indikátorových premenných $Z_1, Z_2, \dots, Z_r$, odpovedajúcich množinám rozkladu, ktoré kvôli zjednodušeniu označme rovnakými symbolmi $Z_1, Z_2, \dots, Z_r$, znaku Z , tým sú určené aj **disjunktné javy určené kvalitatívnym znakom** Z . Špeciálnym prípadom sú **komplementárne javy** v prípade keď $r=2$.

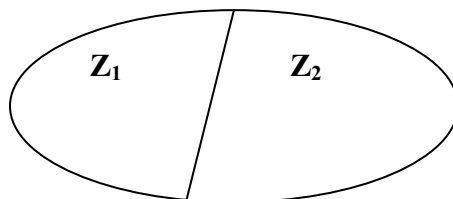
2. Korelačný koeficient komplementárnych javov

Majme kvalitatívny znak Z s dvojprvkovou množinou hodnôt. Tento môžeme vyjadriť pomocou jednej indikátorovej premennej alebo pomocou dvoch **komplementárnych** indikátorových premenných. Najprv uvažujme množinové vyjadrenie:

Z_1 a Z_2 sú komplementárne, ak pre množinové vyjadrenie platí: $Z_1 \cap Z_2 = \emptyset$, a $Z_1 \cup Z_2 = \Omega$,

kde $\emptyset$ symbol reprezentuje prázdnu množinu a Ω množinu všetkých skúmaných prvkov. Potom $Z_2 = \Omega - Z_1$.

Graficky zobrazíme komplementárne javy definované nominálnym znakom s dvojprvkovou množinou hodnôt pomocou odpovedajúcich množín:



Zrejme, pre indikátorové premenné dvojice komplementárnych javov, platí:

$$Z_2 = 1 - Z_1.$$

Využijeme vzťah uverejnený v Luha J.(2008). Možno ho nájsť aj v Řehák J., Řeháková B. (1986). Kvôli úplnosti uvádzame **korelačný koeficient dvoch elementárnych javov**:

$$\begin{aligned} \rho(Z_1, Z_2) &= E((Z_1 - EZ_1)(Z_2 - EZ_2)) / \sqrt{D(Z_1)D(Z_2)} = E(Z_1 Z_2) - E(Z_1)E(Z_2) / \sqrt{D(Z_1)D(Z_2)} = \\ &= (p_{11} - p_{1.}p_{.1}) / \sqrt{(p_{1.}p_{.1}p_{2.}p_{.2})} = \\ &= (p_{11}p_{22} - p_{12}p_{21}) / \sqrt{(p_{1.}p_{.1}p_{2.}p_{.2})} = \\ &= (p_{11} - p_{1.}p_{.1}) / \sqrt{(p_{1.}p_{.1}(1-p_{1.})(1-p_{.1}))}. \end{aligned}$$

kde:

p_{11}	$p_{12} = p_{1.} - p_{11}$	$p_{1.}$
$p_{21} = p_{.1} - p_{11}$	$p_{22} = p - p_{1.} - p_{.1} + p_{11}$	$1 - p_{1.}$
$p_{.1}$	$1 - p_{.1}$	1

Uvažujeme dvojicu komplementárnych javov Z_1, Z_2 . Pre ich indikátorové premenné platí:

$Z_1 Z_2 = 0$, a preto ďalej platí:

$$Z_1(1 - Z_2) = Z_1 - Z_1 Z_2 = Z_1,$$

$$(1 - Z_1)Z_2 = Z_2 - Z_1 Z_2 = Z_2 \text{ a}$$

$$(1 - Z_1)(1 - Z_2) = 1 - Z_1 - Z_2 + Z_1 Z_2 = 1 - Z_1 - Z_2 = 0, \text{ pretože}$$

pre komplementárne javy platí: $Z_2 = 1 - Z_1$.

Využijeme vyjadrenie kontingenčnej tabuľky pomocou indikátorových premenných:

Z_1/Z_2	Z_2	$\Omega - Z_2$	suma
Z_1	$Z_1 Z_2$	$Z_1(1 - Z_2)$	Z_1
$\Omega - Z_1$	$(1 - Z_1)Z_2$	$(1 - Z_1)(1 - Z_2)$	$1 - Z_1$
suma	Z_2	$1 - Z_2$	1

Potom pravdepodobnosti v políčkach 2x2 tabuľke dostaneme ako stredné hodnoty hore uvedených vzťahov pre indikátorové premenné. Výsledok je v redukovanej tabuľke:

0	$p_{1.}$	$p_{1.}$
$p_{.1}$	0	$1 - p_{1.}$
$p_{.1}$	$1 - p_{.1}$	1

Pre komplementárne javy ale platí $p_{1.} = 1 - p_{1.}$. Potom môžeme využiť vyjadrenie:

0	$p_{1.}$	$p_{1.}$
$p_{1.}$	0	$p_{1.}$
$p_{1.}$	$p_{1.}$	1

Hore zistené zjednodušenie vyplývajúce z vlastnosti komplementárnosti dvojice javov aplikujeme na všeobecný vzorec a dostávame koeficient korelácie komplementárnych javov:

$$\begin{aligned}\rho(Z_1, Z_2) &= E(Z_1 Z_2) - E(Z_1)E(Z_2) / \sqrt{D(Z_1)D(Z_2)} = \\ &= (p_{11} - p_{1.}p_{.1}) / \sqrt{(p_{1.}p_{.1}(1-p_{1.})(1-p_{.1}))} = \\ &= - p_{1.}p_{.1} / \sqrt{(p_{1.}p_{.1}p_{1.}p_{.1})} = -1.\end{aligned}$$

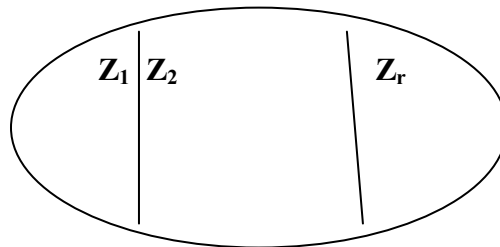
Korelačný koeficient komplementárnych javov je rovný – 1. Uvedený výsledok je prirodzeným dôsledkom vzájomnej komplementárnosti dvoch javov.

3. Korelácie disjunktných javov určených kvalitatívnym znakom

Pre množinové vyjadrenie *disjunktných javov určených kvalitatívnym znakom* s r prvkovou množinou hodnôt platí:

$$Z_i \cap Z_j = \emptyset \text{ pre } i \neq j \text{ a } \bigcup_{i=1}^r Z_i = \Omega.$$

Resp. odpovedajúce grafické vyjadrenie:



Uvažujeme r disjunktných javov $Z_1, Z_2, \dots, Z_r$ s pravdepodobnosťami ich nastatia $p_1, p_2, \dots, p_r$. Zrejme platí $p_i \geq 0$ a $\sum_{i=1}^r p_i = 1$.

Pre indikátorové premenné uvažovaných disjunktných javov platí:

$Z_i Z_j = 0$, pre $i \neq j$, a preto ďalej platí:

$$Z_i(1 - Z_j) = Z_i - Z_i Z_j = Z_i,$$

$$(1 - Z_i)Z_j = Z_j - Z_i Z_j = Z_j \text{ a}$$

$$(1 - Z_i)(1 - Z_j) = 1 - Z_i - Z_j + Z_i Z_j = 1 - Z_i - Z_j.$$

Využijeme vyjadrenie kontingenčnej tabuľky pomocou indikátorových premenných:

Z_i/Z_j	Z_j	$\Omega - Z_j$	suma
Z_i	$Z_i Z_j$	$Z_i(1 - Z_j)$	Z_i
$\Omega - Z_i$	$(1 - Z_i)Z_j$	$(1 - Z_i)(1 - Z_j)$	$1 - Z_i$

Z_i/Z_j	Z_j	$\Omega - Z_j$	suma
Z_i	0	Z_i	Z_i
$\Omega - Z_i$	Z_j	$1 - Z_i - Z_j$	$1 - Z_i$

suma	Z_j	$1-Z_j$	1
------	-------	---------	---

suma	Z_j	$1-Z_j$	1
------	-------	---------	---

Potom pravdepodobnosti (alebo ich odhady) v políčkach 2x2 tabuľke dostaneme ako stredné hodnoty hore uvedených vzťahov pre indikátorové premenné. Výsledok je v redukovanej tabuľke:

0	p_i	p_i
$p_{21}=p_j$	$p_{22}=1-p_i-p_j$	$1-p_i$
p_j	$1-p_j$	1

Pomocou zjednodušenia vyplývajúceho z vlastnosti disjunktnosti javov a dostaneme **koeficient korelácie medzi disjunktnými javmi i a j**:

$$\begin{aligned}\rho(Z_i, Z_j) &= E(Z_i Z_j) - E(Z_i)E(Z_j) / \sqrt{D(Z_i)D(Z_j)} = \\ &= (p_{ij} - p_i p_j) / \sqrt{(p_i p_j (1-p_i)(1-p_j))} = \\ &= - p_i p_j / \sqrt{(p_i p_j (1-p_i)(1-p_j))} = \\ &= - \sqrt{(p_i p_j / ((1-p_i)(1-p_j)))}.\end{aligned}$$

Ak označíme **šancu** javu Z_i oproti komplementárnemu javu $\Omega - Z_i$ $s_i = p_i / (1-p_i)$ a odpovedajúco aj pre jav Z_j a $\omega_{ij} = s_i / s_j$ **pomer šancí** (cross product ratio), tak

$$\rho(Z_i, Z_j) = - \sqrt{s_i s_j} = - s_j \sqrt{\omega_{ij}}.$$

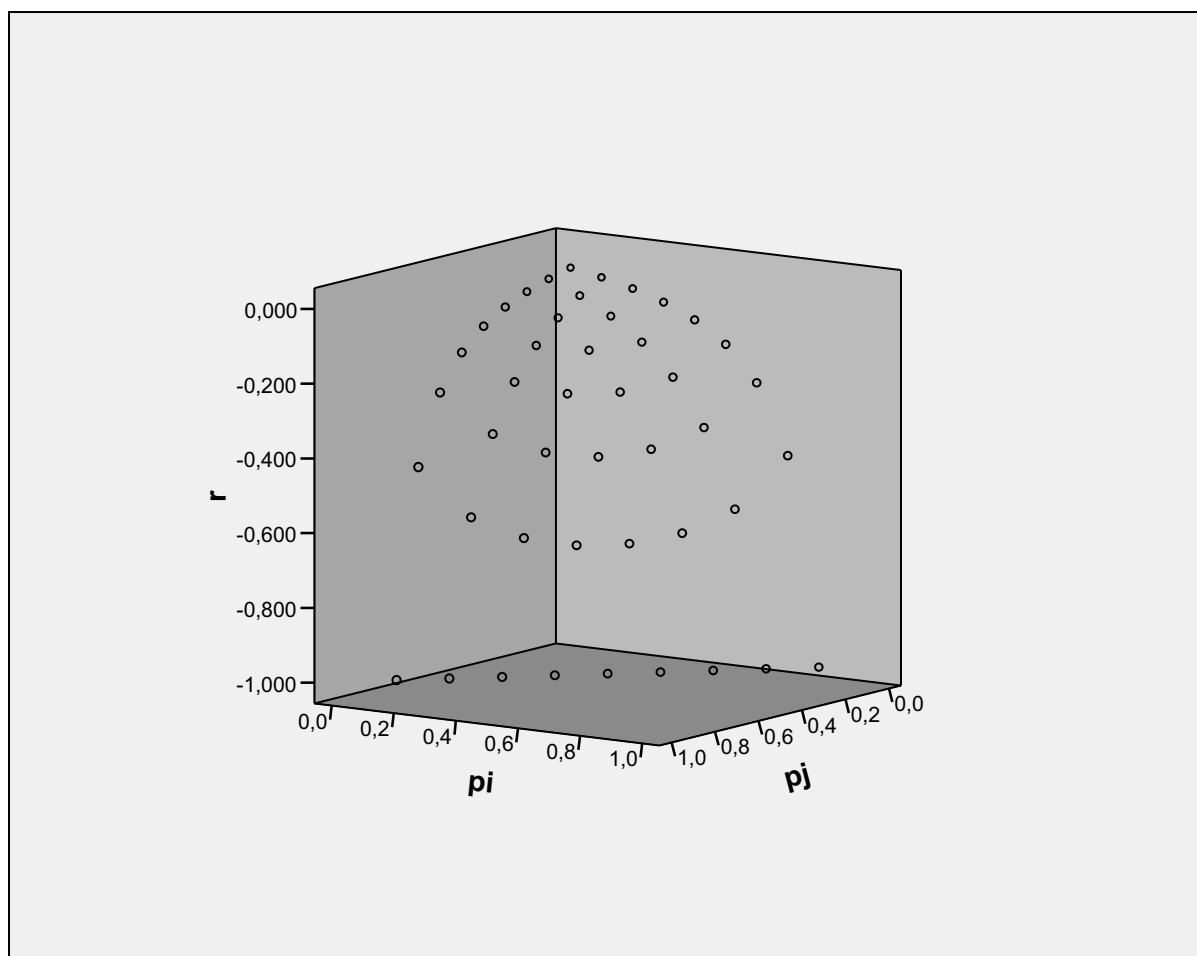
Korelačný koeficient disjunktných javov je záporný.

V priestore $0 < p_i < p_j \leq 1$, $p_i + p_j \leq 1$ môžeme vypočítať priebeh hodnôt korelácie dvoch disjunktných javov. V tabuľke uvádzame okrem vybraných hodnôt podielov aj príslušné šance, pomer šancí a korelačný koeficient. Pričom hodnotu korelácie -1 dostávame v prípade $p_i + p_j = 1$, teda keď ide o dvojicu komplementárných javov.

p_i	p_j	s_i	s_j	ω_{ij}	ρ
0,1	0,1	0,111	0,111	1,000	-0,012
0,1	0,2	0,111	0,250	2,250	-0,028
0,1	0,3	0,111	0,429	3,857	-0,048
0,1	0,4	0,111	0,667	6,000	-0,074
0,1	0,5	0,111	1,000	9,000	-0,111
0,1	0,6	0,111	1,500	13,500	-0,167
0,1	0,7	0,111	2,333	21,000	-0,259
0,1	0,8	0,111	4,000	36,000	-0,444
0,1	0,9	0,111	9,000	81,000	-1,000
0,2	0,1	0,250	0,111	0,444	-0,028
0,2	0,2	0,250	0,250	1,000	-0,063
0,2	0,3	0,250	0,429	1,714	-0,107
0,2	0,4	0,250	0,667	2,667	-0,167
0,2	0,5	0,250	1,000	4,000	-0,250
0,2	0,6	0,250	1,500	6,000	-0,375
0,2	0,7	0,250	2,333	9,333	-0,583
0,2	0,8	0,250	4,000	16,000	-1,000
0,3	0,1	0,429	0,111	0,259	-0,048
0,3	0,2	0,429	0,250	0,583	-0,107
0,3	0,3	0,429	0,429	1,000	-0,184
0,3	0,4	0,429	0,667	1,556	-0,286
0,3	0,5	0,429	1,000	2,333	-0,429

p_i	p_j	s_i	s_j	ω_{ij}	ρ
0,3	0,6	0,429	1,500	3,500	-0,643
0,3	0,7	0,429	2,333	5,444	-1,000
0,4	0,1	0,667	0,111	0,167	-0,074
0,4	0,2	0,667	0,250	0,375	-0,167
0,4	0,3	0,667	0,429	0,643	-0,286
0,4	0,4	0,667	0,667	1,000	-0,444
0,4	0,5	0,667	1,000	1,500	-0,667
0,4	0,6	0,667	1,500	2,250	-1,000
0,5	0,1	1,000	0,111	0,111	-0,111
0,5	0,2	1,000	0,250	0,250	-0,250
0,5	0,3	1,000	0,429	0,429	-0,429
0,5	0,4	1,000	0,667	0,667	-0,667
0,5	0,5	1,000	1,000	1,000	-1,000
0,6	0,1	1,500	0,111	0,074	-0,167
0,6	0,2	1,500	0,250	0,167	-0,375
0,6	0,3	1,500	0,429	0,286	-0,643
0,6	0,4	1,500	0,667	0,444	-1,000
0,7	0,1	2,333	0,111	0,048	-0,259
0,7	0,2	2,333	0,250	0,107	-0,583
0,7	0,3	2,333	0,429	0,184	-1,000
0,8	0,1	4,000	0,111	0,028	-0,444
0,8	0,2	4,000	0,250	0,063	-1,000
0,9	0,1	9,000	0,111	0,012	-1,000

Priebeh korelácie disjunktých javov možno ilustrovať v grafe:



Uvedené ilustrácie nezachytávajú všetky obmedzenia kladené na disjunktné javy. Na komplexnejší pohľad musíme uvažovať disjunktné javy pri $r > 2$ simultánne.

Priebeh korelácie r disjunktých v priestore $p_i > 0$, $\sum_{i=1}^r p_i = 1$ môžeme reprezentovať $r \times r$ korelačnou maticou. Kvôli úplnosti uvádzame aj korelácie dvojíc javov Z_i, Z_i , ktorá je zrejme rovná 1.

Z_i/Z_j	Z_1	Z_2	...	Z_j	...	Z_{r-1}	Z_r
Z_1	1	$\rho(Z_1, Z_{r-1})$	...	$\rho(Z_1, Z_j)$	...	$\rho(Z_1, Z_{r-1})$	1
Z_2	$\rho(Z_2, Z_1)$	1		$\rho(Z_2, Z_j)$		1	$\rho(Z_2, Z_r)$
...							
Z_i	$\rho(Z_i, Z_1)$			$\rho(Z_i, Z_j)$		$\rho(Z_i, Z_{r-1})$	$\rho(Z_i, Z_r)$
...							
Z_{r-1}	$\rho(Z_{r-1}, Z_1)$			$\rho(Z_{r-1}, Z_j)$		1	$\rho(Z_{r-1}, Z_r)$
Z_r	$\rho(Z_r, Z_1)$	$\rho(Z_r, Z_2)$		$\rho(Z_r, Z_j)$		$\rho(Z_r, Z_{r-1})$	1

V políčkach korelačnej matice sú korelácie počítané pomocou vzťahu $\rho(Z_i, Z_j) = -\sqrt{s_i s_j}$.

Na ilustráciu uvádzame príklad priebehu korelácií, šancí a pomeru šancí pre vybrané hodnoty podielov pre $r=3$.

p ₁	p ₂	p ₃	suma	s ₁	s ₂	s ₃	ω ₁₂	ω ₁₃	ω ₂₃	r ₁₂	r ₁₃	r ₂₃
0,10	0,10	0,80	1	0,111	0,111	4,000	1,000	36,000	36,000	-0,111	-0,667	-0,667
0,10	0,20	0,70	1	0,111	0,250	2,333	2,250	21,000	9,333	-0,167	-0,509	-0,764
0,10	0,30	0,60	1	0,111	0,429	1,500	3,857	13,500	3,500	-0,218	-0,408	-0,802
0,10	0,40	0,50	1	0,111	0,667	1,000	6,000	9,000	1,500	-0,272	-0,333	-0,816
0,10	0,50	0,40	1	0,111	1,000	0,667	9,000	6,000	0,667	-0,333	-0,272	-0,816
0,10	0,60	0,30	1	0,111	1,500	0,429	13,500	3,857	0,286	-0,408	-0,218	-0,802
0,10	0,70	0,20	1	0,111	2,333	0,250	21,000	2,250	0,107	-0,509	-0,167	-0,764
0,10	0,80	0,10	1	0,111	4,000	0,111	36,000	1,000	0,028	-0,667	-0,111	-0,667
0,20	0,10	0,70	1	0,250	0,111	2,333	0,444	9,333	21,000	-0,167	-0,764	-0,509
0,20	0,20	0,60	1	0,250	0,250	1,500	1,000	6,000	6,000	-0,250	-0,612	-0,612
0,20	0,30	0,50	1	0,250	0,429	1,000	1,714	4,000	2,333	-0,327	-0,500	-0,655
0,20	0,40	0,40	1	0,250	0,667	0,667	2,667	2,667	1,000	-0,408	-0,408	-0,667
0,20	0,50	0,30	1	0,250	1,000	0,429	4,000	1,714	0,429	-0,500	-0,327	-0,655
0,20	0,60	0,20	1	0,250	1,500	0,250	6,000	1,000	0,167	-0,612	-0,250	-0,612
0,20	0,70	0,10	1	0,250	2,333	0,111	9,333	0,444	0,048	-0,764	-0,167	-0,509
0,30	0,10	0,60	1	0,429	0,111	1,500	0,259	3,500	13,500	-0,218	-0,802	-0,408
0,30	0,20	0,50	1	0,429	0,250	1,000	0,583	2,333	4,000	-0,327	-0,655	-0,500
0,30	0,30	0,40	1	0,429	0,429	0,667	1,000	1,556	1,556	-0,429	-0,535	-0,535
0,30	0,40	0,30	1	0,429	0,667	0,429	1,556	1,000	0,643	-0,535	-0,429	-0,535
0,30	0,50	0,20	1	0,429	1,000	0,250	2,333	0,583	0,250	-0,655	-0,327	-0,500
0,30	0,60	0,10	1	0,429	1,500	0,111	3,500	0,259	0,074	-0,802	-0,218	-0,408
0,40	0,10	0,50	1	0,667	0,111	1,000	0,167	1,500	9,000	-0,272	-0,816	-0,333
0,40	0,20	0,40	1	0,667	0,250	0,667	0,375	1,000	2,667	-0,408	-0,667	-0,408
0,40	0,30	0,30	1	0,667	0,429	0,429	0,643	0,643	1,000	-0,535	-0,535	-0,429
0,40	0,40	0,20	1	0,667	0,667	0,250	1,000	0,375	0,375	-0,667	-0,408	-0,408
0,40	0,50	0,10	1	0,667	1,000	0,111	1,500	0,167	0,111	-0,816	-0,272	-0,333
0,50	0,10	0,40	1	1,000	0,111	0,667	0,111	0,667	6,000	-0,333	-0,816	-0,272
0,50	0,20	0,30	1	1,000	0,250	0,429	0,250	0,429	1,714	-0,500	-0,655	-0,327
0,50	0,30	0,20	1	1,000	0,429	0,250	0,429	0,250	0,583	-0,655	-0,500	-0,327
0,50	0,40	0,10	1	1,000	0,667	0,111	0,667	0,111	0,167	-0,816	-0,333	-0,272
0,60	0,10	0,30	1	1,500	0,111	0,429	0,074	0,286	3,857	-0,408	-0,802	-0,218
0,60	0,20	0,20	1	1,500	0,250	0,250	0,167	0,167	1,000	-0,612	-0,612	-0,250
0,60	0,30	0,10	1	1,500	0,429	0,111	0,286	0,074	0,259	-0,802	-0,408	-0,218
0,70	0,10	0,20	1	2,333	0,111	0,250	0,048	0,107	2,250	-0,509	-0,764	-0,167
0,70	0,20	0,10	1	2,333	0,250	0,111	0,107	0,048	0,444	-0,764	-0,509	-0,167
0,80	0,10	0,10	1	4,000	0,111	0,111	0,028	0,028	1,000	-0,667	-0,667	-0,111
0,90	0,05	0,05	1	9,000	0,053	0,053	0,006	0,006	1,000	-0,688	-0,688	-0,053

4. Príklad korelácie disjunktných javov

Každý kvalitatívny znak s konečnou množinou r hodnôt generuje r disjunktných javov. Predpokladáme nenulové pravdepodobnosti ich výskytu. Na ilustráciu využijeme umelý príklad, ktorý ale značne zodpovedá skutočnosti. Ako nominálny znak uvažujme otázku, ktorá sa pravidelne opakuje pri výskumoch verejnej mienky – a síce koho by ste volili vo voľbách

do NR SR. Na ilustratívne účely sme zostrojili simulačný súbor, kde sa vyskytujú iba strany: KDH, ĽS-HZDS, SDKÚ, SNS, SMER a „strany“ nezúčastnil by som sa a neviem. Frekvenčná tabuľka zo simulačného súboru dáva výsledky:

Ktorú stranu by ste volili do NR SR?

		Frekvencie	Percent	Valid Percent	Cumulative Percent
Valid	KDH	158	5,7	5,7	5,7
	ĽS-HZDS	153	5,5	5,5	11,3
	SDKÚ	209	7,6	7,6	18,8
	SNS	223	8,1	8,1	26,9
	SMER	851	30,8	30,8	57,7
	SMK	179	6,5	6,5	64,2
	nezúčastnil by sa	520	18,8	18,8	83,0
	neviem	469	17,0	17,0	100,0
	Total	2762	100,0	100,0	

Skúmaný nominálny znak má v našom umelom príklade $r=8$ hodnôt a preto vytvoríme 8 indikátorových premenných. Ich názvy budú reprezentovať „hodnoty“ kvalitatívneho znaku a preto ich nazveme rovnako ako „hodnoty“ znaku.

Výsledky sú v korelačnej matici:

Correlations

		kdh	hzds	sdku	sns	smer	smk	nezucast	neviem
kdh	Pearson Correlation	1	-,060(**)	-,070(**)	-,073(**)	-,164(**)	-,065(**)	-,119(**)	-,111(**)
	Sig. (2-tailed)		,002	,000	,000	,000	,001	,000	,000
	N	2762	2762	2762	2762	2762	2762	2762	2762
hzds	Pearson Correlation	-,060(**)	1	-,069(**)	-,072(**)	-,162(**)	-,064(**)	-,117(**)	-,110(**)
	Sig. (2-tailed)	,002		,000	,000	,000	,001	,000	,000
	N	2762	2762	2762	2762	2762	2762	2762	2762
sdku	Pearson Correlation	-,070(**)	-,069(**)	1	-,085(**)	-,191(**)	-,075(**)	-,138(**)	-,129(**)
	Sig. (2-tailed)	,000	,000		,000	,000	,000	,000	,000
	N	2762	2762	2762	2762	2762	2762	2762	2762
sns	Pearson Correlation	-,073(**)	-,072(**)	-,085(**)	1	-,198(**)	-,078(**)	-,143(**)	-,134(**)
	Sig. (2-tailed)	,000	,000	,000		,000	,000	,000	,000
	N	2762	2762	2762	2762	2762	2762	2762	2762
smer	Pearson Correlation	-,164(**)	-,162(**)	-,191(**)	-,198(**)	1	-,176(**)	-,321(**)	-,302(**)
	Sig. (2-tailed)	,000	,000	,000	,000		,000	,000	,000
	N	2762	2762	2762	2762	2762	2762	2762	2762
smk	Pearson Correlation	-,065(**)	-,064(**)	-,075(**)	-,078(**)	-,176(**)	1	-,127(**)	-,119(**)
	Sig. (2-tailed)	,001	,001	,000	,000	,000		,000	,000
	N	2762	2762	2762	2762	2762	2762	2762	2762
nezucast	Pearson Correlation	-,119(**)	-,117(**)	-,138(**)	-,143(**)	-,321(**)	-,127(**)	1	-,218(**)
	Sig. (2-tailed)	,000	,000	,000	,000	,000	,000		,000
	N	2762	2762	2762	2762	2762	2762	2762	2762
neviem	Pearson Correlation	-,111(**)	-,110(**)	-,129(**)	-,134(**)	-,302(**)	-,119(**)	-,218(**)	1
	Sig. (2-tailed)	,000	,000	,000	,000	,000	,000	,000	
	N	2762	2762	2762	2762	2762	2762	2762	2762

** Correlation is significant at the 0.01 level (2-tailed).

Interpretáciou korelačnej matice vzhľadom ku meritu skúmaných javov sa v tomto príspevku nezaobráame. Na doplnenie uvádzame kontingenčnú tabuľku vybraného vzťahu SMER s SDKÚ:

smer * sdu Crosstabulation

			sdku		Total
			0	1	0
smer 0	Count	1702	209	1911	
	% within smer	89,1%	10,9%	100,0%	
	% within sdku	66,7%	100,0%	69,2%	
	Adjusted Residual	-10,0	10,0		
1	Count	851	0	851	
	% within smer	100,0%	,0%	100,0%	
	% within sdku	33,3%	,0%	30,8%	
	Adjusted Residual	10,0	-10,0		
Total	Count	2553	209	2762	
	% within smer	92,4%	7,6%	100,0%	
	% within sdku	100,0%	100,0%	100,0%	

Symmetric Measures

		Value	Asymp. Std. Error(a)	Approx. T(b)	Approx. Sig.	Exact Sig.
Nominal by Nominal	Contingency Coefficient	,188			,000	,000
Interval by Interval	Pearson's R	-,191	,007	-10,219	,000(c)	,000
Ordinal by Ordinal	Spearman Correlation	-,191	,007	-10,219	,000(c)	,000
N of Valid Cases		2762				

a. Not assuming the null hypothesis.

b. Using the asymptotic standard error assuming the null hypothesis.

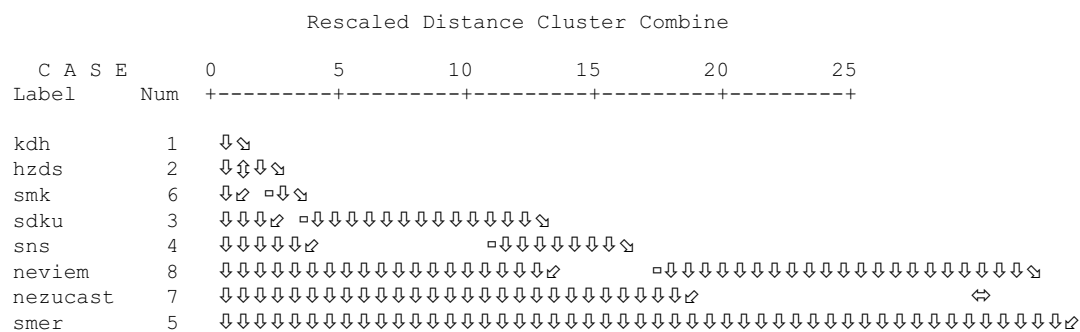
c. Based on normal approximation.

Je zrejmé, že skúmané vzťahy môžeme doplniť analýzou kontingenčných tabuliek. To ale nie je účelom tohoto článku.

Na doplnenie uvedieme ilustratívny dendrogram. Ako mieru blízkosti sme zvolili Pearsonov korelačný koeficient.

* * * * * H I E R A R C H I C A L C L U S T E R A N A L Y S I S

Dendrogram using Average Linkage (Between Groups)



5. Záver

Vlastnosti korelácií medzi disjunktnými a komplementárnymi javmi sú prirodzeným dôsledkom vlastnosti disjunktnosti, resp. v špeciálnom prípade ich komplementárnosti. Okrem vlastnosti zápornosti, môžeme uvedené vzťahy využiť pri špecifikácii modelových vzťahov určených kvalitatívnym znakom.

6. Literatúra

[1]Chajdiak, J.(2003): Štatistika jednoducho. Statis Bratislava 2003, ISBN 808565928-X.

- [2]Chajdiak J. (2005): Štatistické úlohy a ich riešenie v Exceli. STATIS, Bratislava, ISBN 80-85659-39-5.
- [3]Kanderová, M. – Úradníček, V.(2007): Štatistika a pravdepodobnosť pre ekonómov. 1. časť. OZ Financ, Banská Bystrica 2007, ISBN 978-80-969535-5-4.
- [4]Kanderová, M. – Úradníček, V.(2007): Štatistika a pravdepodobnosť pre ekonómov. 2. časť. OZ Financ, Banská Bystrica 2007, ISBN 987-80-696535-6-1.
- [5]Luha J.(2003): Skúmanie súboru kvalitatívnych dát. EKOMSTAT 2003. SŠDS Bratislava 2003.
- [6]Luha J.(2003): Matematickoštatistické aspekty spracovania dotazníkových výskumov. Štatistické metódy vo vedecko-výskumnej práci 2003, SŠDS, Bratislava 2003. ISBN 80- 88946-32-8.
- [7]Luha J.(2005): Viacrozmerné štatistické metódy analýzy kvalitatívnych znakov. *EKOMSTAT 2005*, Štatistické metódy v praxi.SŠDS Trenčianske Teplice 22.–27. 5. 2005.
- [8]Luha J.(2007): Kvótový výber. FORUM STATISTICUM SLOVACUM 1/2007. SŠDS Bratislava 2007. ISSN 1336-7420.
- [9]Luha J.(2008): Korelácie vnorených javov. FORUM STATISTICUM SLOVACUM 5/2008. SŠDS Bratislava 2008. ISSN 1336-7420.
- [10]Pecáková I.: Statistika v terénnych průzkumech. Proffessional Publishing, Praha 2008. ISBN 978-80-86946-74-0.
- [11]Řehák J., Řeháková B. (1986): Analýza kategorizovaných dat v sociologii. Academia Praha 1986.
- [12]Řezanková A.(2007): Analýza dat z dotazníkových šetření. Proffessional Publishing, Praha 2007. ISBN 978-80-86946-49-8.
- [13]Stankovičová I., Vojtková M.(2007): Viacrozmerné štatistické metódy s aplikáciami. IURA EDITION, Bratislava 2007, ISBN 978-80-8078-152-1.

Adresa autora:

RNDr. Ján Luha, CSc.
Jan.Luha@statistics.sk

Grubbsov test a hypotézy o obrobenom povrchu Test of Grubbs and Hypotheses on the Cut Surface

Anna Macurová, Dušan Macura

Abstract: The paper deals with the analysis and use of data to solve problems. This paper is designed to equip technically oriented readers with the basic statistical skills they need to practice their professions. The hypothesis is presented with testing experimental process in the next text.

Key words: Surface, Hypotheses, Testing of the Hypotheses.

Kľúčové slová: Povrch, hypotézy, testovanie hypotéz.

1. Úvod

Význam štatistiky v inžinierskej praxi z hľadiska zvyšovania kvality výroby je nesporný. V každom reálnom výrobnom procese existuje variabilita hodnôt z rôznych hľadísk. V modelovej situácii, za ktorú budeme považovať experiment, vyhodnotíme získané hodnoty zriedkavejšie používaným kritériom, závery formulované v experimente budeme považovať za hypotézy a pomocou vybraných štatistických metód budú navrhované výsledky experimentov. Je známe, že testovanie hypotéz je proces, v ktorom na základe výberových údajov odhadujeme s určitou pravdepodobnosťou, či tvrdenie o niektorom z parametrov základného súboru je, alebo nie je pravdivé. Každá úloha testovania hypotéz je formulovaná tak, že proti sebe stoja dve hypotézy a to nulová hypotéza (často jednoduchá), ktorú budeme označovať H_0 a alternatívna hypotéza (najčastejšie zložená) označovaná H_1 .

2. Test extrémnych hodnôt (Grubbsov test)

Výskum obrobených povrchov patrí medzi základné problémy, ktorých riešenie zdokonaľuje teóriu výrobných technológií. Štatistika je súčasťou rozboru experimentov aj vo výrobných technológiách. V nasledujúcom konkrétnom experimente sa sústredíme na štatistické spracovanie vybraných údajov z hľadiska extrémnej hodnoty v získanom súbore experimentálnych hodnôt. Drsnosť obrobeného povrchu sa dá považovať za náhodnú premennú s normálnym rozdelením $N(\mu, \sigma^2)$. Nech $x_1, x_2, \dots, x_8$ je náhodný výber z normálneho rozdelenia $N(\mu, \sigma^2)$

$$4,8 \ 7,3 \ 5,4 \ 6,2 \ 7,1 \ 5,1 \ 9,8 \quad (1)$$

hodnôt drsnosti obrobeného povrchu (sústružením, v mikrometroch). Hodnota 9,8 vzbudila podozrenie, že je extrémne veľká. Otestujeme túto hypotézu na hladine významnosti $\alpha = 0,05$.

Testujeme hypotézu

$$H_0 : x_{\max} = 9,8 \text{ nie je extrémne veľká hodnota}$$

proti

$$H_1 : x_{\max} = 9,8 \text{ je extrémne veľká hodnota} \quad (2)$$

Zo súboru hodnôt (1) vyplývajú hodnoty $n = 8$, $x_{\max} = 9,8$, $\bar{x} = 6,4375$, $s_x = 1,6256$ do nasledujúcej testovacej štatistiky

$$T_{max} = \frac{x_{max} - \bar{x}}{s_x} \sqrt{\frac{n}{n-1}}. \quad (3)$$

Po výpočte dostaneme $T_{max} = 2,2112$.

Ak

$$T_{max} > T_{\alpha}(n), \quad (4)$$

tak H_0 zamietame, x_{max} je extrémne veľká hodnota, $T_{\alpha}(n)$ je kritická hodnota pre Grubbsov test, ktorú určíme zo štatistických tabuliek [1]. Pre $\alpha = 0,05$ je $T_{0,05}(8) = 2,172$.

Na základe vzťahu (4) je

$$T_{0,05}(8) = 2,172 < T_{max} = 2,2112.$$

Hypotézu H_0 z tvrdenia (2) na hladine významnosti $\alpha = 0,05$ zamietame a prijímame hypotézu H_1 z (2). Hodnota 9,8 je extrémne veľká a môžeme ju zo súboru (1) vylúčiť.

3. Testovanie experimentálnych hodnôt získaných dvoma spôsobmi

12 meraní najväčšej nerovnosti obrobeného povrchu sa uskutočnilo výpočtom z diferenciálnej rovnice (premenné v nasledujúcej diferenciálnej rovnici zodpovedajú označovaniu príslušných odborných pojmov vo výrobných technológiách)

$$2fRz + 4Rz^2Rz' - f^2Rz' = 0, \quad (5)$$

Rz je nerovnosť (drsnosť) obrobeného povrchu, f je posuv nástroja, $Rz' = \frac{dRz}{df}$, riešením (5)

je funkcia $Rz = \frac{c}{8} + \sqrt{\frac{c^2}{64} - \frac{f^2}{4}}$, $c \in R$, $\frac{c}{8} = r_{\varepsilon}$ je polomer hrotu nástroja, pričom platí $r_{\varepsilon}^2 - \frac{f^2}{4} \geq 0$.

Získané hodnoty sú označené y_i a novým digitálnym meracím prístrojom na meranie nerovnosti obrobeného povrchu (drsnosti) hodnoty označené z_i . Zistíme, či je medzi získanými hodnotami podstatný rozdiel.

Výsledky meraní sú nezávislé a zapíšeme ich do tabuľky, zistené hodnoty sú párové, (y_i, z_i) , $i = 1, 2, \dots, 12$. Urobíme rozdiely $x_i = y_i - z_i$, $i = 1, 2, \dots, 12$, ktoré vyjadrujú veľkosť systematickej chyby.

Tabuľka 1: Hodnoty najväčšej nerovnosti obrobeného povrchu

y_i	28,2 2	33,9 5	38,2 5	42,5 2	37,6 2	36,8 4	36,1 2	35,1 1	34,4 5	52,8 3	57,9 0	51,5 2
z_i	28,2 7	33,9 9	38,2 0	42,4 2	37,6 4	36,8 5	36,2 1	35,2 0	34,4 0	52,8 6	57,8 8	55,5 3
x_i	-0,05	-0,04	0,05	0,10	-0,02	-0,01	-0,09	-0,09	0,05	-0,03	0,02	-0,01

Vypočítame výberové charakteristiky

$$\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i = -0,01, \quad s_x = \sqrt{\frac{1}{n-1} \left(\sum_{i=1}^n x_i^2 - n\bar{x}^2 \right)} = 0,0572,$$

$$\bar{y} = 40,7875, \quad s_y = 8,9725, \quad \bar{z} = 40,7875.$$

Použijeme párový test. Ak medzi výsledkami získanými rôznymi prístrojmi nie je rozdiel, je $\Delta = 0$.

Testujeme hypotézu o strednej hodnote

$$H_0 : \mu_1 - \mu_2 = \Delta \text{ proti } H_1 : \mu_1 - \mu_2 \neq \Delta .$$

Testovacia štatistika nadobúda hodnotu

$$T = \frac{\bar{x} - \Delta}{s_x} \sqrt{n} = \frac{-0,01 - 0}{0,0572} \sqrt{12} = -0,6056 .$$

V štatistických tabuľkách nájdeme $t_\alpha(n-1) = t_{0,05}(11) = 2,201$ [2].

Hypotézu H_0 na hladine významnosti $\alpha = 0,05$ nezamietame. Medzi hodnotami získanými výpočtom a meracím prístrojom nie je podstatný rozdiel.

4. Záver

Pri každom testovacom postupe sa dopúšťame dvoch druhov chýb. Ak je H_0 platná, môžeme vplyvom náhody dostať výsledok, pri ktorom hodnota testovacej štatistiky T dostane do kritického oboru (oboru zamietnutia) W . Podľa vopred stanoveného testovacieho postupu zamietame H_0 a prijímame H_1 , čím sa dopustíme chyby 1. druhu. Pravdepodobnosť chyby 1. druhu nazývame hladina významnosti testu a označujeme písmenom α [3].

Ak je platná alternatívna hypotéza H_1 , môžeme vplyvom náhody dostať výsledok, že hodnota testovacej štatistiky T sa dostane do oboru prijatia V . Na základe stanoveného testovacieho kritéria nezamietame H_0 , čím sa dopustíme chyby 2. druhu, ktorá sa označuje β . Pri vopred stanovenej hodnote α sa snažíme minimalizovať β . Optimalizácia hodnôt patriacich do oboru zamietnutia W nie je jednoduché, niekedy je vhodné určiť tieto hodnoty ako extrémne hodnoty T .

5. Literatúra

- [1] MACUROVÁ, A. – HREHOVÁ, S. 2008. Základy štatistiky pre technikov. TU Košice, FVT Prešov, 2008. 120 s. ISBN 978-80-8073-913-3.
- [2] MACURA, D. 2003. Obyčajné diferenciálne rovnice. 1. vyd. FHPV PU, Prešov, 2003. 79 s. ISBN 80-8068-175-9.
- [3] POTOCKÝ, R. a kol. 1991. Zbierka úloh z pravdepodobnosti a matematickej štatistiky. Alfa, Bratislava, 1991. 392 s. ISBN 80-05-00524-5.
- [4] VASILKO, K. a kol., 1991. Technológia obrábania a montáže 1. vyd. Bratislava. Alfa, 1991, 494 s.

Príspevok je zostavený v rámci riešenia projektu 1/4004/2007.

Adresa autora (-ov):

Anna Macurová, PaedDr., PhD.
FVT TU, Bayerova 1
080 01 Prešov
anna.macurova @tuke.sk

Dušan Macura, Mgr., PhD.
PU FPHV
Ul. 17. Novembra
080 01 Prešov
macura @unipo.sk

Kointegrační analýza modelu inflace v České republice

Cointegration Analysis of the Inflation Model in the Czech Republic

Jiří Neubauer

Abstract: The article deals with the question of modeling multidimensional non-stationary cointegrated processes. The multidimensional process is called cointegrated if there exists any linear combination of its one-dimensional components which is stationary. For instance this property can be found in some series of economic indices which are predominantly non-stationary. This contribution is focused on cointegration analysis of multidimensional time series of Czech macroeconomic indices and on verification of some relations describing inflation in the Czech Republic.

Key words: vector autoregressive process, cointegration, inflation.

Klíčové slová: vektorový autoregresní proces, kointegrace, inflace.

1. Úvod

Kointegrační analýza, jako metoda modelování vícerozměrných nestacionárních procesů, nachází uplatnění převážně v oblasti makroekonomických časových řad. Detailní popis této metody je možné najít například v [Hamilton, 1994], [Johansen, 1995], [Lütkepohl, 2007]. Článek je věnován kointegrační analýze čtvrtletních makroekonomických ukazatelů České republiky a pomocí testů kointegrace jsou ověřovány některé vztahy popisující vývoj inflace. Základní model inflace vychází z článku [Kim, 2001].

Uvedeme nyní v krátkosti některé základní definice a nezbytná tvrzení.

Definice 1. Necht' ε_t je n -rozměrný bílý šum s nulovou střední hodnotou a varianční maticí Ω . Stochastický proces Y_t splňující $Y_t - EY_t = \sum_{i=1}^{\infty} C_i \varepsilon_{t-i}$ se nazývá $I(0)$ proces, jestliže $C = \sum_{i=0}^{\infty} C_i \neq 0$ (EY_t značí střední hodnotu procesu Y_t).

Definice 2. O stochastickém procesu Y_t řekneme, že je *integrováný řádu d* a značíme jej $I(d)$, $d = 1, 2, \dots$, jestliže $\Delta^d(Y_t - EY_t)$ je $I(0)$ proces (Δ je diferenční operátor, $\Delta Y_t = Y_t - Y_{t-1}$).

Definice 3. Vícerozměrný proces Y_t nazveme n -rozměrným autoregresním procesem $VAR(p)$, jestliže platí

$$Y_t = \Phi_1 Y_{t-1} + \Phi_2 Y_{t-2} + \dots + \Phi_p Y_{t-p} + \Lambda D_t + \varepsilon_t \text{ pro } t = 1, 2, \dots, T, \quad (1)$$

pro dané počáteční hodnoty $Y_{-p+1}, \dots, Y_0$, kde ε_t je n -rozměrný bílý šum s normálním rozdělením $N(0, \Omega)$, $\Phi_1, \dots, \Phi_p$ jsou matice koeficientů (řádu n), Λ je matice koeficientů deterministického členu D_t , který může obsahovat konstantu, lineární člen, sezónní vlivy a další regresory, které jsou nestochastické. Matice Λ je typu $n \times s$, matice D_t je typu $s \times 1$. Autoregresní proces $VAR(p)$ definovaný rovnicí (1) je možné po jednoduchých úpravách zapsat ve tvaru

$$\Delta Y_t = \Pi Y_{t-1} + \sum_{i=1}^{p-1} \Gamma_i \Delta Y_{t-i} + \Lambda D_t + \varepsilon_t \text{ pro } t = 1, 2, \dots, T, \quad (2)$$

kde $\Pi = \sum_{i=1}^p \Phi_i - \mathbf{I}$, $\Gamma_i = -\sum_{j=i+1}^p \Phi_j$. Tento tvar vyjádření autoregresního procesu se nazývá ECM („error correction model“), používá se pro popis a analýzu kointegrovaných procesů resp. časových řad.

Hlavní myšlenku kointegrace ukážeme na 2 jednorozměrných procesech typu $I(1)$. Řekneme, že procesy X_t a Y_t jsou kointegrované, jestliže existuje taková lineární kombinace $aX_t + bY_t$, která je stacionární. Uvedeme nyní obecnou definici kointegrovaného procesu.

Definice 4. Necht' $\mathbf{Y}_t$ je n -rozměrný integrovaný proces řádu 1. Tento proces nazveme *kointegrovaný s kointegračním vektorem* β ($\beta \in R^n, \beta \neq \mathbf{0}$), jestliže při vhodné volbě počátečních hodnot je proces $\beta' \mathbf{Y}_t$ stacionární.

V případě dvourozměrného procesu může existovat pouze jeden kointegrační vektor. Pokud je $n > 2$, mohou existovat dva vektory β_1, β_2 takové, že $\beta_1' \mathbf{Y}_t$ i $\beta_2' \mathbf{Y}_t$ jsou stacionární a přitom jsou β_1, β_2 lineárně nezávislé. Obecně může existovat $r < n$ lineárně nezávislých kointegračních vektorů. Jejich maximální počet budeme označovat jako *řád (hodnota) kointegrace* a vektorový prostor, který tyto vektory generují, jako *kointegrační prostor*.

2. Maximálně věrohodné odhady kointegračních vektorů a testy kointegrace

Grangerova věta (viz např. [Johansen, 1995]) udává nutnou a postačující podmínku, aby autoregresní VAR(p) proces byl $I(1)$ a kointegrovaný. Podle hodnoty matice Π v jeho ECM reprezentaci se definuje tzv. $H(r)$ model procesu $I(1)$.

Definice 5. $H(r)$ model procesu $I(1)$ je definován jako model VAR(p) splňující

$$\Pi = \alpha\beta',$$

kde α a β jsou $n \times r$ matice. ECM reprezentace tohoto modelu má tvar

$$\Delta \mathbf{Y}_t = \alpha\beta' \mathbf{Y}_{t-1} + \sum_{i=1}^{p-1} \Gamma_i \Delta \mathbf{Y}_{t-i} + \Lambda \mathbf{D}_t + \varepsilon_t \text{ pro } t = 1, 2, \dots, T,$$

kde $\alpha, \beta, \Gamma_1, \dots, \Gamma_{p-1}, \Lambda, \Omega$ jsou parametry, na které nejsou kladena žádná omezení.

Za předpokladu platnosti hypotézy

$$H(r) : \Pi = \alpha\beta'$$

získáme maximálně věrohodný odhad parametru β následujícím postupem (viz [Johansen, 1995], [Hamilton, 1994]). Nejprve vyřešíme rovnici¹

$$|\lambda \mathbf{S}_{11} - \mathbf{S}_{10} \mathbf{S}_{00}^{-1} \mathbf{S}_{01}| = 0$$

pro vlastní čísla $1 > \lambda_1 > \dots > \lambda_n$ a vlastní vektory $\mathbf{V} = (\mathbf{v}_1, \dots, \mathbf{v}_n)$, které jsou normalizovány vztahem $\mathbf{V}' \mathbf{S}_{11} \mathbf{V} = \mathbf{I}$. Kointegrační vztahy odhadneme z

$$\hat{\beta} = (\mathbf{v}_1, \dots, \mathbf{v}_r),$$

maximalizovaná hodnota věrohodnostní funkce je dána vztahem

$$L_{\max}^{-2/T} = |\mathbf{S}_{00}| \prod_{i=1}^r (1 - \lambda_i). \quad (3)$$

¹ Konstrukce matic $\mathbf{S}_{00}, \dots, \mathbf{S}_{11}$ je podrobně popsána v [Johansen, 1995]

Test věrohodnostním poměrem („LR-test“) $Q(H(r) | H(n))$ modelu $H(r)$ v modelu $H(n)$ dvou výrazů (3) pro r a n , tedy

$$Q(H(r) | H(n))^{\frac{-2}{T}} = \frac{|S_{00}| \prod_{i=1}^r (1 - \lambda_i)}{|S_{00}| \prod_{i=1}^n (1 - \lambda_i)}.$$

Po zlogaritmování obdržíme tzv. „TRACE“ statistiku

$$-2 \ln Q(H(r) | H(n)) = -T \sum_{i=r+1}^n \ln(1 - \lambda_i). \quad (4)$$

Statistika pro testování modelu $H(r)$ v $H(r+1)$ (tzv. „MAX“ statistika) je dána vztahem

$$-2 \ln Q(H(r) | H(r+1)) = -T \ln(1 - \lambda_{r+1}). \quad (5)$$

Předpokládejme, že hodnota matice Π je rovna r . Asymptotické rozdělení statistik (4) a (5) závisí na existenci deterministického členu v modelu.

Saikkonen a Lütkepohl (viz Lütkepohl, 2007) navrhli test kointegrace, založený na tom, že nejprve je v procesu odhadnut deterministický člen D_t , ten je odečten od pozorování a na vzniklý proces je použit Johansenův LR-test. Oba typy testů budou použity při kointegrační analýze.

3. Analýza modelu inflace pro Českou republiku

Při rozboru modelu inflace budeme vycházet z článku [Kim, 2001], který se zabývá analýzou inflace v Polsku. Pokusíme se daný model aplikovat na makroekonomické ukazatele České republiky. V modelu vystupují následující veličiny (jedná se o čtvrtletní hodnoty 1996–2005):

- y_t - reálný hrubý domácí produkt České republiky
- m_t - peněžní agregát M2 České republiky
- p_t^F - zahraniční měnová hladina (index spotřebitelských cen EU)
- w_t - nominální mzdy v České republice
- p_t^D - domácí cenová hladina (index spotřebitelských cen České republiky)
- e_t - směnný kurz (české koruny a eura)

Pozn.: Zdrojem dat jsou internetové stránky České národní banky (www.cnb.cz) a EUROSTAT (epp.eurostat.ec.europa.eu), všechny uvedené časové řady jsou zlogaritmované.

Ekonomická teorie rozlišuje tři druhy inflace: monetární, mzdovou a importovanou. Lze psát

$$p_t^D = f(p_t^F, m_t - p_t^D, w_t - p_t^D, e_t, y_t - p_t^D),$$

kde $m_t - p_t^D$ je reálná peněžní zásoba, $w_t - p_t^D$ jsou reálné mzdy a $y_t - p_t^D$ je reálný hrubý domácí produkt. Monetární inflace se objevuje v případě, že růst peněžní nabídky převyšuje růst ekonomiky. Za předpokladu homogenity reálných peněz a reálného výstupu lze reálnou peněžní zásobu vyjádřit ve tvaru

$$m_t - p_t^D = b + (y_t - p_t^D). \quad (6)$$

K mzdové inflaci dochází, když růst reálných mezd je větší než úroveň mezd, která odpovídá růstu produktivity dané země. Pro reálné mzdy tedy můžeme psát ²

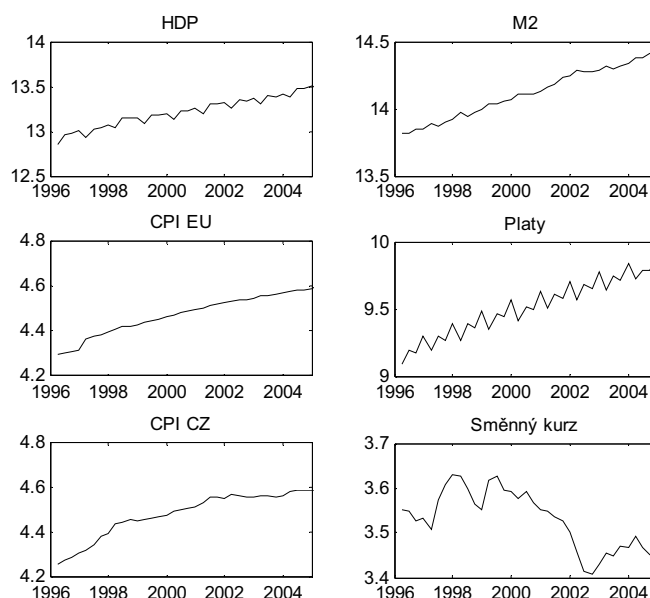
²Daná rovnice předpokládá, že vliv nezaměstnanosti na mzdy není významný (viz [Kim, 2001]).

$$w_t - p_t^D = a + (y_t - p_t^D). \quad (7)$$

Importovaná inflace je výsledkem nerovnosti cenových hladin jednotlivých národních ekonomik. Za platnosti parity kupní síly můžeme psát

$$p_t^D = e_t + p_t^F. \quad (8)$$

Podrobnější odvození rovnic lze nalézt v [Kim, 2001].



Obrázek 1: Přirozené logaritmy časových řad

Pomocí testů kointegrace nyní zkusíme ověřit platnost daných vztahů (6), (7) a (8) pro Českou republiku. Nejprve budeme analyzovat dané vztahy jednotlivě a poté se pokusíme sestavit celkový model.

Označme nejprve $\mathbf{Y}_1 = (m_t - p_t^D, y_t - p_t^D)'$. Pomocí testů jednotkových kořenů je možné ukázat, že jednotlivé složky této dvourozměrné časové řady jsou integrované řádu 1. Tuto řadu je možné popsat pomocí modelu VAR(5) s konstantou a sezónní složkou. Pomocí zmíněných testů kointegrace ověříme, zda jsou tyto řady kointegrované. Tabulka 1 obsahuje hodnoty TRACE a MAX statistiky, dále test Saikkonena a Lütkepohla (S&L) s odpovídajícím p -hodnotami.

Tabulka 1: Testy kointegrace – monetární inflace

r	TRACE	p -hodnota	MAX	p -hodnota	S&L	p -hodnota
0	21,033	0,0056	21,031	0,0028	16,23	0,0094
1	0,002	0,9630	0,002	0,002	0,00	0,9834

Nulovou hypotézu, že $r = 0$ (neexistuje kointegrace) můžeme zamítnout, hypotézu, že $r = 1$ již zamítnout nelze. Předpokládejme tedy, že daná časová řada je kointegrovaná s jedním kointegračním vektorem, jehož odhad je $\hat{\boldsymbol{\beta}} = (1; -0,78810)'$. Nyní můžeme provést test lineární hypotézy o kointegračním vektoru $\mathbf{R} \cdot \text{vec}(\boldsymbol{\beta}) = \mathbf{q}$, v našem případě je $\mathbf{R} = (1,1)$, $\boldsymbol{\beta} = (\beta_1, \beta_2)'$ a $\mathbf{q} = 0$, který má asymptoticky χ^2 rozdělení. Testová statistika (založená na věrohodnostním podílu) má hodnotu 1,886 s p -hodnotou 0,1696. Danou hypotézu tedy není

možné zamítnout, můžeme tedy předpokládat, že vztah vyjadřující monetární inflaci pro Českou republiku platí.

Označme nyní $Y_2 = (w_t - p_t^D, y_t - p_t^D)'$. Testy jednotkových kořenů opět ukázaly, že složky integrované řádu 1, danou časovou řadu lze popsat VAR(3) modelem s konstantou a sezónní složkou. Výsledky kointegrační analýzy jsou shrnuty v tabulce 2.

Tabulka 2: Testy kointegrace – mzdová inflace

r	TRACE	p -hodnota	MAX	p -hodnota	S&L	p -hodnota
0	25,444	0,0009	25,119	0,0004	28,53	0,0000
1	0,325	0,5686	0,325	0,5686	0,55	0,5157

I v tomto případě můžeme předpokládat existenci jednoho kointegračního vektoru, jehož odhad je $\hat{\beta} = (1; -0,92739)'$. Provedeme-li stejný test jako v případě monetární inflace, obdržíme hodnotu testové statistiky 0,257 s odpovídající p -hodnotou 0,6117, hypotézu tedy není možné zamítnout a opět je možné předpokládat, že vztah popisující mzdovou inflaci platí.

Pro test posledního vztahu popisující importovanou inflaci označme $Y_3 = (p_t^D, e_t + p_t^F)'$. Složky jsou opět integrované, tuto dvourozměrnou časovou řadu můžeme popsat pomocí modelu VAR(1) s konstantou. Výsledky testů kointegrace jsou shrnuty v tabulce 3.

Tabulka 3: Testy kointegrace – importovaná inflace

r	TRACE	p -hodnota	MAX	p -hodnota	S&L	p -hodnota
0	23,597	0,0019	19,208	0,0063	19,94	0,0018
1	4,390	0,0362	4,390	0,0362	0,88	0,3983

Na hladině významnosti 0,05 nejsou výsledky testů zcela jednoznačné, na hladině významnosti 0,01 můžeme předpokládat existenci jednoho kointegračního vektoru, jehož odhad je $\hat{\beta} = (1; -0,57271)'$. Test lineární hypotézy o daném kointegračním vektoru má hodnotu 2,309 s odpovídající p -hodnotou 0,1286. Ani v tomto případě není možné hypotézu o platnosti vztahu (8) zamítnout.

V následující části vytvoříme větší celkový model. Označme

$$Y_t = (p_t^D, e_t + p_t^F, w_t - p_t^D, m_t - p_t^D, y_t - p_t^D)'$$

Tuto časovou řadu, která má 5 složek, lze popsat pomocí modelu VAR(5) s konstantou a sezónní složkou. Výsledky testů kointegrace shrnuje tabulka 4.

Tabulka 4: Testy kointegrace – celkový model

r	TRACE	p -hodnota	MAX	p -hodnota	S&L	p -hodnota
0	372,31	0,0000	219,15	0,0000	158,26	0,0000
1	153,16	0,0000	97,814	0,0000	85,62	0,0000
2	55,350	0,0000	40,671	0,0000	41,29	0,0001
3	14,679	0,0648	11,633	0,1259	15,19	0,0148
4	3,046	0,0810	3,046	0,0810	6,48	0,0129

Na hladině významnosti 0,05 nejsou výsledky zcela jednoznačné, nicméně na hladině 0,01 můžeme předpokládat existenci 3 kointegračních vektorů. Na základě vztahů (6), (7) a (8) by měla matice kointegračních vektorů mít tvar

$$\beta = \begin{pmatrix} 0 & 0 & 1 \\ 0 & 0 & -1 \\ 0 & 1 & 0 \\ 1 & 0 & 0 \\ -1 & -1 & 0 \end{pmatrix}.$$

To je možné ověřit pomocí již použitého testu lineární hypotézy. Testová statistika má hodnotu 223,404, odpovídající p -hodnota se blíží nule. Hypotézu o platnosti všech tří vztahů popisující inflaci zamítáme.

4. Závěr

Jak vyplývá z jednotlivých výsledků statistické analýzy studovaných časových řad makroekonomických ukazatelů, vztahy (6), (7) a (8) popisující monetární, mzdovou a importovanou inflaci nejsou zcela v souladu s hodnotami těchto ukazatelů pro Českou republiku. Jsou-li tyto vztahy analyzovány odděleně, použité testy jejich platnost nevyvrátí. V celkovém modelu, jenž zahrnul všechny použité časové řady, byly sice identifikovány 3 kointegrační vektory, které mohou odpovídat nějakým dlouhodobým ekonomickým vztahům, nicméně test zamítl, že by se mělo jednat právě uvedené rovnice.

5. Literatura

- [1] HAMILTON, J., D. 1994. Time Series Analysis. Princeton: Princeton University Press, 1994, ISBN 0-691-04289-6.
- [2] KIM, B. Y. Determinants of Inflation in Poland: A Structural Cointegration Approach. Bank of Finland. Institute for Economies in Transition, BOTFIT, 2001.
- [3] JOHANSEN, S. 1995. Likelihood-based Inference in Cointegrated Vector Auto-regressive Models. Oxford: Oxford University Press, 1995, ISBN 0-19-877450-7.
- [4] LÜTKEPOHL, H. New Introduction to Multiple Time Series Analysis. Berlin: Springer-Verlag, 2007, ISBN 978-3-540-26239-8.

Adresa autora:

Mgr. Jiří Neubauer, Ph.D.
Katedra ekonomiky a managementu Univerzity obrany v Brně
Kounicova 65
612 00 Brno
Česká republika
Jiri.Neubauer@unob.cz

Faktory podnikatelského prostředí a jejich klasifikační síla¹

Factors of entrepreneurial environment and their power of classification

Jakub Odehnal, Marek Sedláčik, Jaroslav Michálek

Abstract: The receiver operating characteristic curves (ROC) were used to compare the power of classification of the entrepreneurial environment factors. These factors were created as the result of the factor analysis. Variables were obtained from the set of database Regio of Statistical Office of the European Communities and describe 85 European NUTS 2 regions represent 6 European countries (Czech Republic, Germany, Poland, Austria, Slovakia, Hungary).

Key words: ROC curves, entrepreneurial environment, multivariate statistical methods

Klíčová slova: ROC křivky, podnikatelské prostředí, mnohorozměrné statistické metody

1. Úvod

Tendence k dlouhodobému vyrovnávání socioekonomických rozdílů patří mezi základní priority evropské hospodářské politiky, zejména však mezi klíčové cíle regionální politiky EU. Regiony EU prostřednictvím podnikatelského prostředí vytváří podmínky pro vstup nových podnikatelů, investorů do regionu, jejichž podnikatelská činnost generuje nová pracovní místa, snižuje nezaměstnanost na regionálním trhu práce a vytváří tak podmínky pro regionální ekonomický růst a tím tak naplňuje základní cíle evropské hospodářské politiky. Kvalita podnikatelského prostředí regiony významně odlišuje, čímž vytváří specifické regionální podmínky pro podnikatele a napomáhá tak při rozhodování o přístupu investic do jednotlivých regionů. Hodnocení podnikatelského prostředí vybraných regionů EU pak poskytuje informace o regionech využitelné pro subjekty hospodářské politiky na úrovni národní i regionální.

Mezi významné projekty hodnotící podnikatelské prostředí jednotlivých zemí řadíme zejména publikace Světové Banky jako výstupy projektu Doing Business [9]. Výsledky regionálního hodnocení podnikatelského prostředí vybraných regionů EU můžeme nalézt v [10] prezentujících meziregionální odlišnosti podnikatelského prostředí krajů ČR.

Příspěvek představuje výsledky faktorového hodnocení 85 regionů EU charakterizujících 6 vybraných států EU (ČR, Maďarsko, Německo, Polsko, Slovensko, Rakousko). 21 proměnných bylo prostřednictvím faktorové analýzy redukováno do 4 faktorů podnikatelského prostředí pomocí kterých byla vytvořena užitím mnohorozměrných metod regionální klasifikace. Významnost jednotlivých faktorů pro klasifikaci byla posouzena pomocí ROC křivek a pomocí plochy pod ROC křivkou označovanou jako AUC (Area Under Curve).

2. Kvalita podnikatelského prostředí měřena na úrovni vybraných států EU

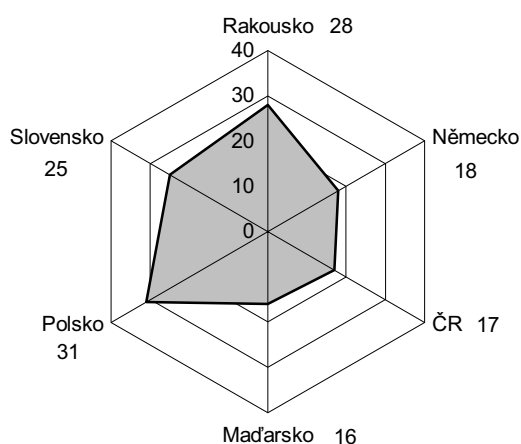
Hodnocení podnikatelského prostředí u vybraných 6 států EU vychází z výsledků šetření publikovaných v [9]. Sledováním 10 zejména regulačních ukazatelů (vznik a zánik podniku, udělování povolení, vynutitelnost plnění smluv, ochrana investorů, registrace vlastnictví, získávání úvěrů, podmínky přijímání a propouštění pracovníků, zahraniční

¹ Příspěvek vznikl za podpory výzkumného záměru MSM0021622418 a VZ04-FEM-K101.

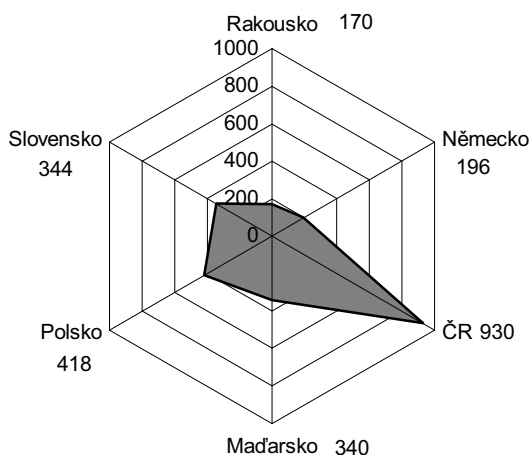
obchodování, platba daní) studie hodnotí zejména časovou a nákladovou náročnost plnění sledovaného ukazatele v dané zemi.

Z námi vybraných 6 zemí (ČR, Maďarsko, Německo, Polsko, Slovensko, Rakousko) je na nejvyšším místě sestavovaného žebříčku [9] Německo (20. místo), následováno Rakouskem (25. místo). Ze sledovaných zemí Visegradské čtyřky nejlepšího umístění dosahuje Slovensko (32. místo), dále Maďarsko (45. místo), ČR (56. místo) a s odstupem Polsko (74. místo). V celkovém hodnocení je tedy dobře patrný rozdíl mezi tradičními zeměmi EU a postavením stále ještě nových členských zemí EU.

Srovnáním dílčích charakteristik u vybraných zemí můžeme pozorovat ukazatele, u kterých vybrané země dosahují podobných hodnot a které obě skupiny zemí v celkovém pořadí sblížují (obr. 1), ale i ukazatele, u kterých existují významné rozdíly, které dané země výrazně rozlišují (obr. 2).



Obrázek 1: Časová náročnost zahájení podnikání (dny)



Obrázek 2: Časová náročnost plateb daní (hodin za rok)

Podmínky zahájení podnikání hodnoceny počtem kalendářních dnů nutných k absolvování nutných administrativních záležitostí zobrazuje obrázek 1. Ze sledovaných zemí je nejnižší časová náročnost zřejmá u Maďarska 16 dní, nejvyšší u Polska 31 dní. V případě této dílčí charakteristiky podnikatelského prostředí tak u vybraných zemí EU nepozorujeme významné rozdíly, které by výrazně odlišovaly tradiční země EU zastoupené Německem a Rakouskem od nových členských zemí Visegradské čtyřky.

V případě charakteristiky platby daní (obrázek 2), jejíž výše, administrativní a časová náročnost přímo ovlivňuje rozhodování podnikatele o realizaci podnikatelských aktivit v dané zemi, můžeme pozorovat nejvyšší hodnotu u ČR, kde je časová náročnost přípravy a platby daní 930 hodin. Srovnáním s ostatními vybranými zeměmi nejnižší časovou náročnost platby daní pozorujeme v případě Rakouska (170 hodin) a Německa (196 hodin).

Meziroční srovnání umístění vybraných zemí v pořadí dle kvality podnikatelského prostředí vypovídá o změně podnikatelského prostředí, jehož rostoucí kvalita vedla v případě Německa, Maďarska a Slovenska ke zlepšení pozic v publikovaném pořadí zemí. Opačný případ nastal u Polska a ČR, kdy došlo ke zhoršení výsledné pozice ze 74. místa na 75. místo, respektive z 52. místa na 56. místo.

3. Faktory podnikatelského prostředí vybraných regionů a jejich klasifikace

Regionální podnikatelské prostředí je charakterizováno pomocí faktorů ovlivňujících lokalizaci ekonomických aktivit v dané oblasti. Lokalizační faktory tak regiony odlišují a tím poskytují informace o podnikatelském prostředí a jeho kvalitě. Mezi hlavní lokalizační faktory ovlivňující investory ve volbě vhodné lokality dle [1] řadíme makroekonomickou a politickou stabilitu, dostatek kvalifikovaných pracovních sil, kvalitní infrastrukturu, blízkost vědecko výzkumné základny dále inovační potenciál, situaci na regionálním trhu práce.

Data pro faktorové hodnocení regionálního podnikatelského prostředí 85 regionů EU byla získána z Regional Statistics Yearbook [3] a databáze Regio [4]. Výsledkem faktorové analýzy [7] jsou 4 faktory podnikatelského prostředí charakterizující vybrané regiony EU prostřednictvím vypočteného skóre.

Faktor kvality pracovních sil a inovací (1)

Faktor trhu práce (2)

Faktor ekonomické výkonnosti (3)

Faktor infrastruktury (4)

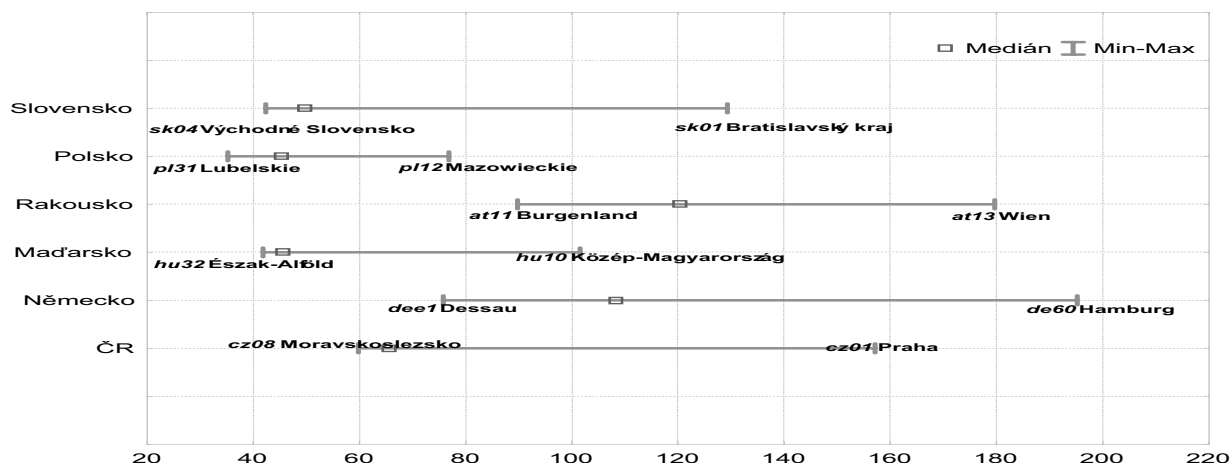
Regionální faktory podnikatelského prostředí jsou tvořeny z 21 proměnných s přihlédnutím na významnost faktorových zátěží u každé proměnné. Faktor kvality pracovních sil a inovací je tak tvořen 9 proměnnými: počet ICT patentů na mil. obyvatel, počet biotechnologických patentů na mil. obyvatel, zaměstnanost v technologicky a vědecky náročných oborech, počet High-tech patentů na mil. obyvatel, celkové výdaje na vědu a výzkum v regionech jako % HDP, procento zaměstnaných výzkumných pracovníků v regionu, zaměstnanci vědy a výzkumu jako % pracovní síly, poskytování IT služeb v regionu, podíl pracovní síly s terciálním vzděláním na celkové pracovní cíle. Faktor trhu práce 5 proměnnými: dlouhodobá nezaměstnanost, celková změna počtu obyvatel, čistá migrace, procentní vyjádření HDP na obyvatele k průměru EU, míra zaměstnanosti v regionech, míra nezaměstnanosti v regionech. Faktor ekonomické výkonnosti 5 proměnnými: HDP na obyvatele v běžných cenách, příjem domácností, produktivita práce (HDP na zaměstnance), produktivita služeb a Faktor infrastruktury proměnnou charakterizující hustotu dopravních cest v regionech. Výsledky faktorové analýzy budou dále použity ke klasifikaci regionů pomocí shlukové analýzy.

Mnohorozměrná klasifikace sledovaných regionů byla vytvořena užitím shlukové analýzy [2, 5, 7], kdy jako objekty byly zvoleny vybrané regiony EU a jako proměnné vypočtené faktorové skóre pro dílčí faktory podnikatelského prostředí. Výsledky hierarchického shlukování (wardova metoda, druhá mocnina euklidovské vzdálenosti) zobrazené na obrázku 3 a mapou podle [7] na obr. 4 ukazují na přirozenou klasifikaci regionů EU do 2 skupin odlišujících se podnikatelským prostředím a jeho kvalitou. První skupina je tvořena 16 regiony Polska s označením (pl11 – pl63), 7 regiony ČR (cz02 – cz08), 6 regiony Maďarska (hu21 – hu33) a 3 regiony Slovenska (sk02 – sk04). Druhá skupina je tvořena 41 regiony Německa (de11 – deg0), 9 regiony Rakouska (at11 – at34), 1 regionem ČR (cz01), 1 regionem Slovenska (sk01) a 1 regionem Maďarska (hu10).

Výsledky klasifikace poslouží jako základ pro zjištění významnosti jednotlivých faktorů podnikatelského prostředí pro danou klasifikaci. Zmíněná významnost sledovaných faktorů bude posuzována prostřednictvím tvaru ROC křivek (např. v [8]) a výpočtu plochy AUC (např. v [8]) umožňující srovnání faktorů podnikatelského prostředí dle kvality použitého klasifikačního kritéria. Vysoké schopnosti správné klasifikace objektů a tedy vysoké klasifikační síle odpovídá tvar ROC křivky, kdy křivka nejprve strmě roste a dále je téměř konstantní (tj. odpovídající AUC se blíží k hodnotě 1). Opačné situaci, tedy nízké schopnosti klasifikace, odpovídá tvar ROC křivky blízký se diagonále.

Mezi 2 skupinami vybraných regionů tradičních zemí EU a sledovaných regionů zemí Visegradské čtyřky můžeme dle tvaru odhadnutých ROC křivek (viz obr. 4) usuzovat vysokou klasifikační sílu zejména u faktoru ekonomické výkonnosti. Obdobně i odhadem plochy pod ROC křivkou (AUC) pozorujeme nejvyšší kvalitu klasifikačního kritéria v případě faktoru 3 podnikatelského prostředí. Prokázána je tak výrazná odlišnost sledovaných skupin regionů zejména u charakteristik makroekonomického výstupu, které daný faktor sytí.

Následující obrázek zobrazuje proměnnou HDP na obyvatele v PPS jako proměnnou faktoru ekonomické výkonnosti, která demonstruje zjištěné rozdíly v ekonomické úrovni sledovaných regionů EU. Rozdíly tak mezi regiony jednotlivých zemí můžeme posuzovat prostřednictvím maximální respektive minimální hodnoty ukazatele a pomocí mediánu.



Obrázek 6: HDP na obyvatele v PPS, NUTS2, % průměru EU

Kvalita podnikatelského prostředí odrážející se v tomto ekonomickém ukazateli ukazuje na rozdílné úrovně regionálního podnikatelského prostředí u obou skupin sledovaných zemí. Vysoká úroveň je zejména v oblastech tradičních průmyslových center a metropolitních regionů. Regionální hospodářská politika EU by tak měla být koncentrována zejména do problémových oblastí, jejichž regionální podnikatelské prostředí je oproti tradičním zemím EU méně vyspělé, což se odráží v ekonomické úrovni regionů a v životní úrovni obyvatel. Konvergenci regionů jako jeden z cílů regionální politiky je tak možno dosáhnout prostřednictvím zlepšení podmínek pro podnikání a celkového zvýšení konkurenceschopnosti regionů.

5. Závěr

Příspěvek představuje použití ROC křivek ke zjištění klasifikační síly faktorů podnikatelského prostředí. Vybraných 85 regionů EU bylo prostřednictvím zkonstruovaných faktorů a užitím shlukové analýzy klasifikováno do dvou skupin odlišujících se kvalitou podnikatelského prostředí. První skupina regionů obsahuje regiony tradičních zemí EU v textu zastoupené Německem a Rakouskem. Druhá skupina obsahuje regiony zemí Visegradské čtyřky s výjimkou metropolitních regionů Prahy, Bratislavy a Budapeště. Dle tvaru znázorněných ROC křivek a odhadem plochy pod křivkou AUC je možné pozorovat nejvyšší odlišnosti mezi skupinami regionů a tedy nejvyšší klasifikační sílu u faktoru ekonomické výkonnosti, který charakterizuje obecné ekonomické prostředí sledovaných regionů.

6. Literatura

- [1] BLAŽEK, J., UHLÍŘ, D.: Teorie regionálního rozvoje. UK Praha, Nakladatelství Karolinum 2002, ISBN 80-246-0384-5
- [2] HEBÁK, P.: Vícerozměrné statistické metody. Vyd. 1. Praha: Informatorium, 2004. 239 s.
- [3] EUROPEAN COMMISSION, REGIONAL YEARBOOK 2007, Luxembourg: Office for Official Publications of the European Communities 2007, ISBN 978-92-79-05077-0.
- [4] EUROSTAT: Regional statistics 2008.
http://epp.eurostat.ec.europa.eu/portal/page?_pageid=1996,45323734&_dad=portal&_schema=PORTAL&screen=welcomeref&open=/&product=EU_MASTER_regions&depth=2.
- [5] LUKASOVÁ, A., ŠARMANOVÁ, J. Metody shlukové analýzy. Vyd. 1. Praha: SNTL, 1985. 210 s.
- [6] MICHÁLEK, J., SEDLACÍK, M., DOUDOVÁ, L.: A Comparison of Two Parametric ROC Curves Estimators in Binormal Model, 23rd International Conference Mathematical Methods in Economics 2005, Hradec Králové, Czech Republic, 14. - 16.9. 2005.
- [7] ODEHNAL, J., MICHÁLEK, J.: Hodnocení konkurenceschopnosti vybraných regionů EU, v recenzním řízení.
- [8] PEPE, M. S. The Statistical Evaluation of Medical Tests for Classification and Prediction. Oxford University Press, New York, 2004. ISBN 0-19-856582-8.
- [9] SVĚTOVÁ BANKA: Doing Business in 2008 – Removing Obstacles to Growth. The World Bank, the International Finance Corporation and Oxford University Press copublication, Washington, 2008
- [10] VITURKA, M. a kol.: Regionální vyhodnocení kvality podnikatelského prostředí v ČR. 1.vyd. Brno : Masarykova univerzita, Ekonomicko-správní fakulta, 2003. 141 s. ISBN 80-210-3304-5

Adresy autorů:

Ing. Jakub Odehnal
KAMI
Lipová 41a
ESF MU
602 00 Brno
odehnal@mail.muni.cz

Mgr. Marek Sedláčik, Ph.D.
Katedra Ekonometrie
Kounicova 65
Univerzita obrany
612 00 Brno
marek.sedlacik@unob.cz

Doc. RNDr. Jaroslav Michálek, CSc.
Ústav matematiky
Fakulta strojního inženýrství
Vysoké učení technické v Brně
Technická 2896/2,
616 69 Brno
michalek@fme.vutbr.cz

Pokus o nájdenie spravodlivého dôchodkového systému **An attempt to find the fair pension system**

Karol Pastor

Abstract: The paper recalls the non-sustainability of the present pay-as-you-go pension system. In its present form it is unfair and decreases the birth rate, as well, because it overlooks the child-raising as an investment into the system. The solution can go in two directions: either increase the premium or to lower the parents financial contributions into the system. The paper suggests to increase the basic rate of contributions in 12 % of crude salaries and to deduct 3% for each raised child as a parental bonus.

Key words: Pension system, Pay-as-you-go, Population ageing, Parental bonus, Reduction of contributions.

Kľúčové slová: dôchodkový systém, priebežný pilier, populačné starnutie, rodičovský bonus, zníženie odvodového zaťaženia.

1. Úvod

Článok vychádza z všeobecne známeho faktu, že totiž vďaka populačnému vývoju na Slovensku a celej Európe je priebežný dôchodkový systém v dnešnej podobe neudržateľný. Ukážeme, že je neudržateľný i po zavedení II. piliera, nespravodlivý a protipopulačne pôsobiaci. Predložený článok sa zaoberá ospravedlivením tohoto systému, t.j. hľadá spôsob, ako nespravodlivý systém zmeniť na spravodlivý.

2. Neudržateľnosť priebežného piliera

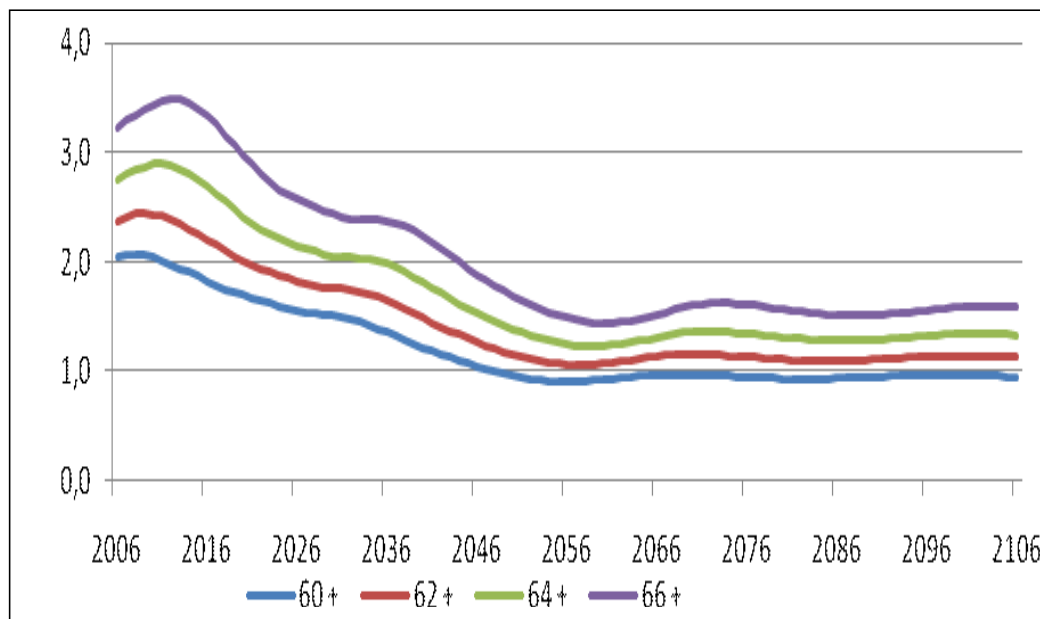
Pod tlakom obáv z dôsledkov populačného starnutia sa u nás v posledných rokoch objavilo viacero prognóz populačného vývoja Slovenska. Vychádzajú z rôznych predpokladoch o plodnosti, migrácii a pod. Nebudeme posudzovať, ktorý nich je najpravdepodobnejší. Dobrý dôchodkový systém by totiž aj tak nemal príliš závisieť od populačného vývoja. Pre ilustráciu spomeňme aspoň prognózu Výskumného demografického centra z r. 2004, publikovanú na internete [8].

Dobрым ukazovateľom zmeny vekovej štruktúry populácie je napr. index starnutia, ktorý dáva do pomeru poreprodukčnú a predreprodukčnú zložku ([7], s. 76) a ktorý v pomeroch SR do budúcnosti vykazuje trvalý nárast pri všetkých prognózach. Z hľadiska našej témy je však dôležitejší index ekonomickej závislosti starých ľudí (old age dependency ratio, napr. [7], s. 75). Možno ho definovať rôzne. Pre ilustráciu, podiel počtu osôb 65-ročných a starších a osôb 18-64 ročných, ktorý podľa údajov [8] bol v r. 2005 iba 0,17, r. 2025 bude 0,37 a roku 2050 bude 0,56. To znamená, že kým r. 2005 na jedného dôchodcu pripadalo 5,6 ľudí v produktívnom veku, r. 2025 to bude iba 2,7 ľudí a r. 2050 iba 1,8 ľudí. Ak sa nezmení výška odvodov, mali by dnešní päťdesiatnici počítať iba s polovičnými dôchodkami ako majú dnešní dôchodcovia, a dnešní absolventi VŠ iba s tretinovými (počítané v dnešných cenách). Dodajme, že kým pre rok 2050 uvedené čísla sú hypotetické, lebo pôrodnosť sa môže zmeniť, čísla pre rok 2025 sú už viac-menej isté. Všetci tí, ktorí r.2025 budú mať 18 rokov a viac, sú už na svete.

Uvedený výpočet je len orientačný, pretože z hľadiska výšky dôchodkov v priebežnom pilieri je dôležitý počet pracujúcich, t.j. prispievateľov do systému, a ten môže byť odlišný. Závisí o.i. na dôchodkovom veku, zamestnanosti žien, pracujúcich dôchodcoch, na miere nezamestnanosti. Pri zamestnanosti 70 %, čo zodpovedá odporúčaniam EU [1], [2], by pripadalo asi 2,6 prispievateľov na jedného dôchodcu. Podobný podiel je v súčasnosti na

Slovensku v čase, keď sa do systému odvádza 18 % z hrubej mzdy a dôchodky predstavujú 44 % priemernej hrubej mzdy, čo znamená, že na jedného dôchodcu pripadá 2,45 prispievateľov do systému. Ako možno vyčítať napr. z Grafu 1 prevzatom z práce [4], to potrvá ešte asi 10 rokov. Potom sa však tento podiel (vďaka starnutiu obyvateľstva) radikálne zmení. Ak náš dôchodkový systém nezareaguje včas, aj pri optimistických variantoch by dnešní mladí mali rátať v I. pilieri iba s asi polovičným dôchodkom toho dnešného (alebo s dvojnásobnými odvodmi).

Graf 1. Podiel prispievateľov do Sociálnej poisťovne na jedného dôchodcu v SR pre rôzny dôchodkový vek pri zachovaní súčasných parametrov (plodnosť, úmrtnosť, zamestnanosť).



Zdroj [4].

3. Nespravodlivý a protipopulačný dôchodkový systém

Na obdobie života, keď sa človek už nevládze sám užiť, sa možno pripraviť dvomi spôsobmi: vychovať deti, ktoré sa oňho v zhode s prirodzeným zákonom postarajú, alebo nasporiť si peniaze, za ktoré si potom takúto starostlivosť kúpi.

Tieto dve riešenia si navzájom konkurujú. Výchova detí nielenže znižuje terajší a budúci pracovný príjem rodičov, ale aj ich budúci dôchodok. Súčasný dôchodkový systém (a treba povedať, že I., II. aj III. pilier) sú výhodnejšie pre tých, ktorí nemajú deti. Inými slovami, súčasný dôchodkový systém je protipopulačný, lebo motivuje radšej sporiť ako mať deti. Súčasný populačný vývoj je potom už len logickým dôsledkom nastavených ekonomických podmienok.

Nebudeme tu rozoberať otázku II. piliera. Diverzifikuje riziko a prenáša časť zodpovednosti za zabezpečenie v starobe na samotného občana, čo je správne. Ak by sa vyriešila je miera štátnych záruk za úspory v DSS, tak nech príspevkovo definované povinné sporenie (II. pilier) zostane, nech je vlastníctvom sporiteľa a zostane „zásluhové“. Je správne oddeliť zásluhové od solidárneho, príspevkovo definované od dávkovo definovaného.

Pretože II. pilier je zásluhový, solidárnosť by mala byť vlastnosťou I. piliera. Priebežný dôchodkový systém (I. pilier, PAYG, „pay-as-you-go“) by mal mať charakter poistenia. Do I. piliera prispieva nielen ten, kto zarába (a odvádza peniaze), ale aj ten, kto „dodáva“

prispievateľov, čiže vychováva deti. Viacdetní rodičia teda do I. piliera prispievajú viac ako ostatní, no spravidla dostávajú dokonca menej ako ostatní. Bez zveličovania možno povedať, že tí, ktorí vychovávajú svoje deti, pracujú na bezdetných, čiže súčasný systém okráda viacdetných rodičov. Našťastie, nespravodlivosť, ktorú tento systém v sebe nesie, možno podstatne zredukovať.

4. Možnosti ospravodlivenia dôchodkového systému

Spôsobov, ako zmeniť nespravodlivý dôchodkový systém na spravodlivejší, je viacero. Viacdetní rodičia by mali dostávať vyšší dôchodok, alebo menej prispievať, alebo oboje. Pravda, existujú aj pokusy kompenzovať rodičom ich vklad do dôchodkového systému iným spôsobom.

Prvú možnosť ponúka zaujímavý a pomerne dobre rozpracovaný návrh V. Palka na zavedenie tzv. rodičovského piliera, podľa ktorého výška vyplácaného dôchodku súvisí s počtom vychovaných detí a to tak, že časť odvodov pracujúceho dieťaťa by štát prevádzal priamo na účet jeho rodiča [5], [6]. Do tejto kategórie patria aj tie riešenia, ktoré počítajú so skorším odchodom do dôchodku pre ženy podľa počtu vychovaných detí, alebo rôzne formy „refundácie“ výdavkov na výchovu, ktoré siahajú do výšky rádovo 1 milión korún.

Návrh, ktorý r. 2004 rozpracovali Hyzl a kol. [3], je akousi kombináciou oboch prístupov. Autori navrhujú, aby dôchodok z I. piliera poberali len tí, ktorí vychovali aspoň jedno dieťa. Počas výchovy detí neplatia príspevok do II. piliera. Ak majú dosť detí, stačí im to na slušný dôchodok.

Najjednoduchší a administratívne nenáročný sa zdá systém rodičovských bonusov pri odvodoch do I. piliera. Pracujúci, ktorý vychováva deti, má znížené odvody o určitú pevnú čiastku alebo percentá z hrubej mzdy. Takýto systém platil pred pár rokmi aj na Slovensku. Vtedy to bolo o 0,5 % za každé vyživované dieťa, čo je príliš málo. Po prvé, neodstraňuje to nespravodlivosť, len ju nepatrne znižuje, len na zalepenie očí. Po druhé, 0,5 % z hrubej mzdy 20 000.- je 100 korún, po zdanení 82, čo je tak na päť pív. Aby to pôsobilo aspoň trochu propopulačne, musí to byť cítiť, musí byť čo závidieť. Novší Brockov návrh (SME 6.8.2008) je 1%, čo je tiež primálo, ale už trochu lepšie. Verejnosť to ešte stále vníma ako zvýhodnenie rodičov, hoci ide iba o malé zmiernenie znevýhodnenia.

Tento článok uprednostňuje riešenia, kde od počtu detí nezávisí výška dôchodku, ale výška odvodov, teda tie, ktoré občan pocíti už teraz, nie až o 40 rokov. Ak „odmena“ má motivovať, musí byť vidieť už v tom čase, keď je ešte možné si ju zaslúžiť, nie až dodatočne. Po druhé, treba taký bonus, ktorý cítiť. Podľa našich predbežných výpočtov by to malo byť aspoň 5 % z hrubej mzdy za každé vyživované dieťa. Po tretie, pokiaľ „odmena ex post“ sa vzťahuje len na pracujúce deti, je veľmi slabo determinovaná, nezaručuje dostatočnú predvídateľnosť, a má charakter skôr lotérie (mladý rodič nevie, koľko z jeho (budúcich) detí sa dožije dospelosti, zostane na Slovensku, koľko budú zarábať, či neostanú pri deťoch, nevie, koľkokrát sa ešte bude meniť dôchodkový systém). Výhodné je mať veľa detí a málo vnukov.

Po štvrté „odmena ex post“ je administratívne náročná. Zisťovať po 50 rokoch, či niekto dieťa vychoval, alebo iba splodil (alebo ani to nie), najmä ak to bolo v cudzine, je takmer nemožné. Čo ak ho vychovával iba 5 rokov a potom odišiel od rodiny - koľko má dostať? „Ex post“ sa spravodlivé riešenie asi nájsť nedá. Naproti tomu systém odvodových úľav (dnes vychovávaš - dnes máš úľavu) nevyžaduje prakticky žiadne nové administratívne úkony. Všetko čo treba sa zisťuje už pri daňovom priznaní a je ľahko overiteľné. Zostáva už len určiť výšku spravodlivého rodičovského bonusu.

5. Orientačný výpočet výšky odvodového bonusu

V tejto časti budeme vychádzať z modelu, v ktorom sa výchova detí počíta ako investícia do priebežného I. piliera, čo sa zohľadňuje nie pri určovaní výšky dôchodku, ale

odvodov. Zníženie odvodov (rodičovský bonus) nemá povahu refundácie nákladov na výchovu dieťa (rodičia nepredávajú svoje deti), ale je určené cenou dieťaťa pre dôchodkový systém.

Výpočet výšky rodičovského bonusu nie je jednoznačný. Pre orientáciu môžeme vychádzať z nasledovnej úvahy: Nech A a B sú dva takmer rovnaké štáty, ktoré sa líšia iba v tom, že v B je dvakrát vyššia pôrodnosť ako v A. V oboch štátoch je príspevok do I. piliera rovnaký, 9 % hrubej mzdy. Nech v štáte A je tzv. stacionárna populácia - ľudí neubúda, ani nepribúda, čistá miera reprodukcie je 1, čo znamená že každý má priemerne 2,1 dieťaťa (pre jednoduchosť ďalej počítajme s celými číslami, t.j. iba s dvomi deťmi). V štáte B je čistá miera reprodukcie 2, čo znamená, že za 40 rokov od nástupu do zamestnania do odchodu do penzie sa počet obyvateľov zvýši asi 3 krát ($2^{40/25} = 2^{1.6} = 3,03$, kde 25 je dĺžka generácie a 40 je dĺžka pracovnej kariéry). Znamená to, že aj počet prispievateľov je trikrát vyšší a čerstvý dôchodca v štáte B môže počítať s 3-krát vyšším dôchodkom ako v štáte A. Je to akoby odmena za to, že miesto dvoch mal štyri deti.

Teda aj obrátene, na to, aby rodič štyroch detí štátu B mal rovnaký dôchodok ako rodič dvoch detí štátu A, stačí, aby prispieval tretinovým odvodom. Keďže štvordetný rodič sa počas 40-ročnej pracovnej kariéry stará priemerne o 2 deti, jedno dieťa navyše by predstavovalo zníženie odvodov z 9% na 3%. Závery pre iné počty detí možno robiť extrapoláciou. Pretože počet obyvateľov jej nelineárnou funkciou čistej miery reprodukcie, extrapolácie možno robiť rôznym spôsobom. Možno tak napr. vyvodiť, že bezdetný by mal prispievať sumou trikrát vyššou, t.j. 27 %, trojdetný sumou 1%, štvor- a viacdetný by neprispieval nič. Tak by to bolo aj motivačné.

Na druhej strane, takýto výsledok príliš revolučný, schematický (vyššie uvedené výpočty závisia od čistej miery reprodukcie v štáte, nie od individuálnej plodnosti) a neúnosný z ekonomického i administratívneho hľadiska. Viedol by k tomu, že bezdetní by chodili pracovať do zahraničia. Výpočet je teda len orientačný, ukazuje však, že minulosti platný bonus 0,5 % za každé dieťa (alebo dnešný bonus 0 %) má ďaleko od spravodlivosti.

Z uvedených dôvodov treba uprednostniť nejaký únosný kompromis, napr. zníženie odvodov o 3 % za každé vyživované dieťa. To samozrejme povedie v prípade príjmov v Sociálnej poisťovni. Tá musí čeliť ďalším dvom problémom: schodku spôsobenom odchodom časti sporiteľov do II. piliera a už spomínanému zníženiu počtu prispievateľov vďaka demografickému starnutiu. Preto tak či tak najneskôr o 10 rokov bude treba odvody zvýšiť a celý systém doladovať podľa konkrétnej situácie.

Na záver uvedme jeden konkrétny návrh. Odvody do II. piliera 9%, do I. piliera základná sadzba 12 % (neskôr možno i viac) a zníženie za každé vyživované dieťa o 3 % hrubej mzdy. V súčasnosti na dôchodok prispieva do Sociálnej poisťovne zamestnanec 4% a zamestnávateľ 14 % hrubej mzdy. Po úprave by prispieval zamestnanec 7 % a zamestnávateľ 14 %, zľavy sa týkajú prednostne tej časti, ktorú odvádza zamestnanec. Úprava by mala vstúpiť do platnosti v blízkej dobe, nie neskôr ako po r. 2012, keď za podiel prispievateľov na poberateľov začne zhoršovať. Možno ju kombinovať so zvýšením dôchodkového veku (napr. na 65 rokov).

Samozrejme, varianty a vylepšenia sú možné, prechodné opatrenia potrebné. V každom prípade, nech už sa ujme čokoľvek, je potrebné, aby sa o problémy konečne začalo vážne diskutovať.

6. Literatúra

- [1] EURÓPSKA KOMISIA 2005. Zelená kniha „Nová medzigeneračná solidarita ako odpoveď na demografické zmeny“ Oznámenie komisie KOM(2005) 94. Brusel 16.3.2005.
http://ec.europa.eu/employment_social/news/2005/mar/comm2005-94_sk.pdf

- [2]EURÓPSKA KOMISIA 2006. Demografická budúcnosť Európy - pretvorme výzvy na príležitosť. Oznámenie komisie KOM(2006) 571. Brusel 12.10.2006.
http://ec.europa.eu/employment_social/news/2006/oct/demography_sk.pdf
- [3]HYZL, J., RUSNOK, J., ŘEZNÍČEK, T., KULHAVÝ, M. 2004. Penzijní reforma pro Českou republiku. Praha 2004. www.ing.cz/cz/o_ing/INGnavrhpenzijnireformy.pdf
- [4]MANIAČKOVÁ, J. 2008. Prognózy hospodárenia dôchodkového systému SR. Bakalárska práca. FMFI UK Bratislava 2008.
- [5]PALKO, V. 2007. Prečo budeme vymierať? In: Impulz 3, č.4, 2007, s. 32-39.
- [6]PALKO, V., KRAJNIAK, M. 2007. Ekonomika konzervatívnej solidarity. Pracovný text, Bratislava 2007. www.prave-spektrum.sk
- [7]VAŇO, B., JURČOVÁ, D., MÉSZÁROS, J. 1997. Základy demografie. OZ Sociálna práca, Bratislava 2003.
- [8]VÝSKUMNÉ DEMOGRAFICKÉ CENTRUM 2004. Prognóza obyvateľstva SR do r. 2050. www.infostat.sk/vdc/pdf/prognoza_web/slov/nuts1/SR.pdf

Adresa autora:

Doc. RNDr. Karol Pastor, CSc.

KAMŠ FMFI UK

Mlynská Dolina, 842 48 Bratislava

pastor@fmph.uniba.sk

**Stanovení vrcholu a dna hospodářského cyklu ČR pomocí
neparametrického odhadu derivace trendu**
**Through and Peak Determination of the CR Business Cycles using
nonparametric estimat of the trend derivation**

Jitka Poměnková

Abstract: Presented paper deals with through and peak determination of the CR business cycle using statistical detrending technique, namely nonparametric Gasser-Müller estimate. In the first step, analysis of time trend and expert estimate of potential turning points are done. Consequently, an estimate of the first derivation of given time series is done followed by turning points determination, i.e. troughs and peaks. Then, according to the result of the second derivation time trend estimate, analysis of potential turning point (trough or peak) are done. Finally, by comparison of obtained results, determination of troughs and peaks of the CR business cycle is concluded.

Key words: business cycle, Gasser-Müller estimate, trough, peak

Klíčová slova hospodářský cyklus, Gasser-Müllerův odhad, dno, vrchol

1. Úvod

Hrubý domácí produkt (HDP) je považován za základní ukazatel poskytující informace o ekonomické úrovni a výkonnosti země. Zachycuje produkci vytvořenou výrobními faktory na území daného státu. Při sledování vývoje HDP v čase jsou zřejmé dvě jeho základní charakteristiky, a to dlouhodobě rostoucí úroveň HDP a jeho krátkodobé fluktuace projevující se projevují cyklickým vývojem meziročních temp růstu. V dlouhém období se krátkodobé fluktuace ekonomické výkonnosti zdají být zanedbatelné. Ovlivňují však významně každodenní rozhodování jednotlivých ekonomických subjektů a nositelů hospodářské politiky. Fluktuace ekonomické aktivity obvykle nelze objasnit jedinou příčinou, jde spíše o souhrn různých vlivů, a to exogenních i endogenních. Jak definují Dornbusch a Fischer (1994) hospodářský cyklus představuje více nebo méně pravidelné střídání expanze a recese ekonomické aktivity kolem dráhy trendového růstu. Na vrcholu cyklu je ekonomická aktivita vzhledem k trendu vysoká, v sedle je dosaženo spodního bodu ekonomické aktivity.

Hospodářské cykly jsou charakterizovány jako opakované sledy výrazných expanzí a recesí celkové ekonomické aktivity, které jsou ohraničeny body zvratu. Recese je období mezi vrcholem a sedlem (dnem) aktivity a expanze je období mezi sedlem (dnem) a vrcholem. Recesi, resp. expanzi, charakterizuje významný pokles, rep. růst, úrovně agregátní ekonomické aktivity přesahující několik měsíců. Celková doba mezi dvěma vrcholy určuje délku cyklu, která se pohybuje mezi jedním rokem až více než deseti lety. Vrcholy a sedla cyklu jsou identifikovány podle bodů zvratu spektra ukazatelů (Kadeřábková, 2003).

Podle Canovy (1999) můžeme provést dělení soudobých filtrační techniky na statistické a ekonomické. Mezi statistické techniky lze zařadit techniku prvních diferencí nebo eliminaci trendu předpokládající přímou nepozorovatelnost trendové nebo cyklické složky, k jejíž identifikaci však používají rozdílné statistické předpoklady. V případě statistického přístupu předpokládáme, že použitá data byla předem sezónně očištěna nebo, že sezónní a cyklická složka jsou považovány za jednu složku a že nepravidelné (vysokofrekvenční) fluktuace mají malý význam.

Předkládaný příspěvek se zabývá detekcí dna a vrcholu hospodářského cyklu pomocí statistické filtrační techniky, a to neparametrického jádrového odhadu. Zvoleným typem odhadu je konvoluční typ Gasser-Müllerova odhadu z důvodu možnosti odhadování i derivací trendu vývoje. Nejprve je odhadnut trendu vývoje, který je následován expertním odhadem potenciálních bodů zvrátů (dna a vrcholy). Poté je proveden odhad první derivace trendu vývoje a jsou stanoveny potenciální extrémny, tedy dna a vrcholy. Následně je odhadnuta druhé derivace trendu vývoje a analyzováno, zda potenciální extrémny jsou maxima (vrcholy) nebo minima (dna). Komparací dosažených výsledků jsou v závěru konstatována období dna a vrcholu hospodářského cyklu ČR.

2. Metodika

Nechť jsou hodnoty x pevně voleny experimentátorem a hodnoty Y , které mohou být reálné nebo simulované, jsou závislé na hodnotách x . Regresní rovnice popisující vztah dvojice proměnných (x_i, Y_i) , $i=1, \dots, n$ můžeme zapsat jako $Y_i = m(x_i) + \varepsilon_i$, $i=1, \dots, n$ s neznámou regresní funkcí m a chybou pozorování ε_i , pro kterou platí $E(\varepsilon_i)=0$, $D(\varepsilon_i)=\sigma^2$, $i=1, \dots, n$. Myšlenka jádrového vyhlazování je najít vhodnou aproximaci $\hat{m}$ odhadu funkce m , přičemž předpokládáme ekvidistantní rozdělení bodů plánu x_i intervalu $[0,1]$.

Předpokládejme, že jsou nezáporná celá čísla, $0 \leq v < k$. Obvykle, jádro K je kompaktně nesená reálná funkce splňující podmínku $K \in S_{v,k}^0$, kde $S_{v,k}^0$ je třída reálných funkcí:

$$S_{v,k}^0 = \left\{ K \in Lip[-1,1], \text{nosič}(K) \subseteq [-1,1] \right. \\ \left. \int_{-1}^1 x^j K(x) dx = \begin{cases} 0 & 0 \leq j < k, j \neq v \\ (-1)^v v! & j = v \\ \beta_k \neq 0 & j = k \end{cases} \right\} \quad (2)$$

říkáme, že K je jádro řádu (v, k) .

Hladkost jádrové funkce je vyjádřena následovně. Označme pro libovolné celé číslo $\mu \geq 1$

$$S_{v,k}^\mu = \left\{ K \mid K \in S_{v,k}^0 \cap C^\mu[-1,1], K^{(j)} = K^{(-j)} = 0, j = 0, 1, \dots, \mu-1 \right\}. \quad (3)$$

Tato jádra nazýváme hladká jádra. Podrobněji viz Horová (2002).

Obecně můžeme jádrová odhad zapsat $\hat{m}(x) = SY = \sum_{i=1}^n W_i(x;h)Y_i$, kde $W_i(x;h)$ označuje váhovou funkci. Váhové funkce, které lze zapsat do vyhlazovací matice $S=(s_{ij})=(W_i(x_j;h))$, $i,j=1, \dots, n$, závisí na h, i, x a K . Poznamenejme, že h označuje šířku vyhlazovacího okna a že vyhlazovací matice S nezávisí na Y .

Gasser-Müllerův, tzv. konvoluční, typ jádrového odhadu funkce m je definován následovně

$$\hat{m}^{(v)}(x) = \sum_{i=1}^n Y_i \frac{1}{h^{v+1}} \int_{s_{i-1}}^{s_i} K\left(\frac{x-u}{h}\right) du. \quad (4)$$

Body plánu $x_i \in [0,1]$, $i=1, \dots, n$, jsou vzestupně uspořádané a pro body s_i , $i=0, \dots, n$ platí $s_0=0$, $s_i=(x_{i+1}+x_i)/2$, $i=1, \dots, n-1$ a $s_n=1$. Pro prvky vyhlazovací matice $S=(s_{ij})=(W_i(x_j;h))$, $i,j=1, \dots, n$, v bodě plánu x_j , s šířkou vyhlazovacího okna h pro Gasser-Müllerův odhad platí vztah

$$W_i(x_j;h) = \frac{1}{h^{v+1}} \int_{s_{i-1}}^{s_i} K\left(\frac{x_j-u}{h}\right) du. \quad (5)$$

Výhodou využití Gasser-Mullerova odhadu je možnost odhadovat nejen funkci, ale také jejich derivací. Pro odhad optimální šířky vyhlazovacího okna je ožné použít metodu zobecněného křížového ověřování. Při odhadu funkce můžeme toto kritérium formulovat následovně:

$$\hat{h}_{opt,GCV} = \arg \min_{h \in H} GCV(h), \quad (6)$$

kde H označuje množinu, na které hledáme minimum zobecněné funkce křížového ověřování. Zpravidla postačí $H_n = [a_k n^{-1/(2k+1)}, b_k n^{-1/(2k+1)}]$ pro nějaké $0 < a_k < b_k < \infty$. Pro funkci GCV platí

$$GCV(h) = \frac{1}{n} \sum_{i=1}^n \left\{ \frac{Y_i - \hat{m}(x_i)}{1 - tr(S)/n} \right\}^2, \quad (7)$$

kde $tr(S)$ označuje součet diagonálních prvků vyhlazovací matice definované výše.

Odhad šířky vyhlazovacího okna v případě odhadu první derivace funkce m získáme ze vztahu

$$CV^{(1)}(h) = \frac{1}{n-1} \sum_{i=1}^{n-1} \{ \Delta_i^{(1)} - \hat{m}_{-(i,i+1)}^{(1)}(x_i^{(1)}) \}^2, \quad (8)$$

kde $x_i^{(0)} = x_i, \Delta_i^{(0)} = Y_i, x_i^{(1)} = (x_{i+1}^{(0)} - x_i^{(0)})/2, \Delta_i^{(1)} = (\Delta_{i+1}^{(0)} - \Delta_i^{(0)})/(x_{i+1}^{(0)} - x_i^{(0)}), i = 1, \dots, n$

a $\hat{m}_{-(i,i+1)}^{(1)}(x_i^{(1)})$ je jádrový Gasser-Müllerův odhad v bodě $x_i^{(1)}$ konstruovaný na datech $(x_1, Y_1), \dots, (x_{i-1}, Y_{i-1}), (x_{i+2}, Y_{i+2}), \dots, (x_n, Y_n)$. Číslo $\hat{h}_{opt,CV^{(1)}}$, které minimalizuje $CV^{(1)}(h)$, je odhadem optimálního vyhlazovacího parametru pro odhad 1. derivace:

$$\hat{h}_{opt,CV^{(1)}} = \arg \min_{h \in \hat{H}_n} CV^{(1)}(h), K \in S_{1,k}^0, \quad (9)$$

kde zpravidla $\hat{H}_n = [1/n, 2]$. Lze ukázat (Poměnková, 2005), že platí

$$h_{opt,v,k}^* = \left[\frac{(2v+1)k}{k-v} \right]^{1/(2k+1)} h_{opt,0,k}^* \quad \text{pro } v, k \text{ sudé a } h_{opt,v,k}^* = \left[\frac{(2v+1)(k-1)}{3(k-v)} \right]^{1/(2k+1)} h_{opt,1,k}^* \quad \text{pro } v, k \text{ liché.}$$

Postačí tedy určit odhad hodnoty h pro $v=0$ (sudé) a $v=1$ (liché) podle toho, zda počítáme odhad sudé nebo liché derivace funkce m .

3. Empirická část

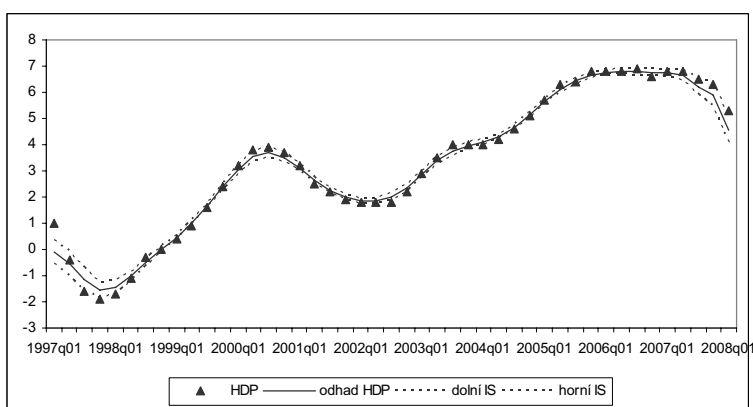
Pro emipirickou analýzu byly zvoleny meziroční procentní změny čtvrtletních absolutních hodnot HDP České republiky v konstantních cenách, sezónně očištěné v období 1997/1 – 2008/1. V dalším textu budou použité hodnoty označovány za HDP. Pro vstupní data byly provedeny neparametrické odhady a to vývoje trendu HDP (Obrázek 1), vývoje první derivace HDP (Obrázek 2) a vývoje druhé derivace HDP (Obrázek 3). Při konstrukci bylo využito optimalizace vyhlazovacího parametru šířky vyhlazovacího okna pomocí metody zobecněného křížového ověřování a optimalizace výběru vhodné jádrové funkce (Poměnková, 2005). Použité parametry uvádí následující tabulka.

Tabulka 1: Parametry příslušné použitým neparametrickým odhadům.

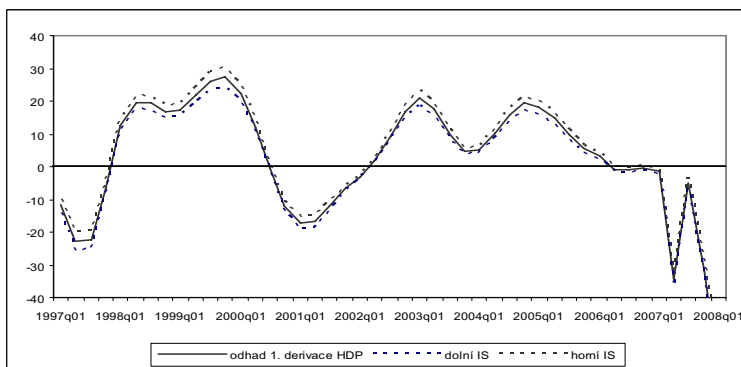
Odhadovaný trend	Šířka vyhlazovacího okna	Typ jádrové funkce
Vývoj HDP	$h=0,056$	$S_{0,2}^1$
Vývoj 1. derivace HDP	$h=0,067$	$S_{1,3}^1$
Vývoj 2. derivace HDP	$h=0,767$	$S_{2,4}^1$

Před aplikací Gasser-Mülerova odhadu byla data testována na nezávislost Kendalovým τ testem náhodnosti, přičemž nulová hypotéza o nezávislosti byla zamítnuta. Proto bylo přistoupeno k odhadu šířky vyhlazovacího okna i dalšími způsoby (metody minimalizující AMSE, Härdle, 1990), přičemž zjištěné výsledky lze považovat za shodné s výsledky zjištěné metodou zobecněného křížového ověřování. Lze se tedy domnívat, že pro odhady funkcí i derivací funkcí můžeme hodnoty šířek vyhlazovacího okna, jak je uvádí Tabulka č. 1, považovat za vyhovující.

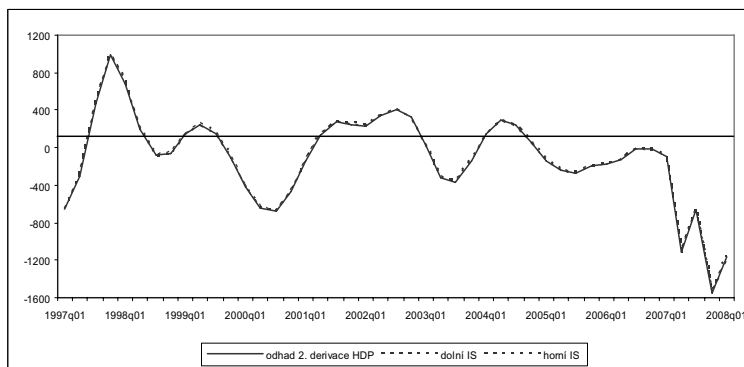
Analýza dat proběhla v následujících krocích. Nejprve byl proveden odhad vývoje HDP v České republice pomocí Gasser-Mülerova odhadu a expertně odhadnuty možné okamžiky dna a vrcholu hospodářského cyklu. Ve druhém kroku byl proveden Gasser-Mülerův odhad první derivace vývoje HDP a stanoveny stacionární body. Tj. body, ve kterých je první derivace nulová, a tedy ve kterých lze očekávat extrém. V posledním kroku byl proveden Gasser-Mülerův odhad druhé derivace vývoje HDP a posouzeno, zda v nulových bodech první derivace lze očekávat nastává maximum nebo minimum.



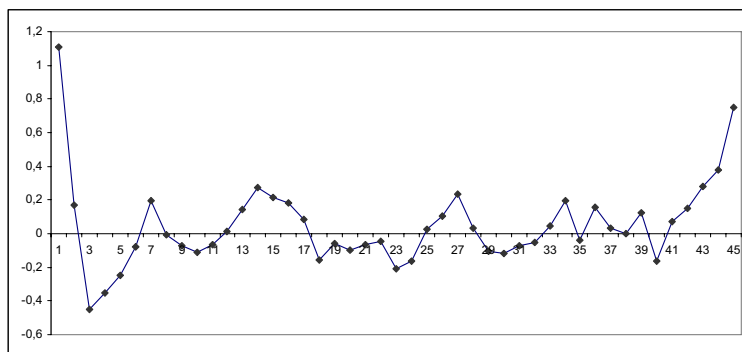
Obrázek 1: Gasser-Mülerův odhad vývoje HDP v ČR



Obrázek 2: Gasser-Mülerův odhad 1. derivace vývoje HDP v ČR



Obrázek 3: Gasser-Mülerův odhad 2. derivace vývoje HDP v ČR



Obrázek 4: Rezidua po provedení Gasser-Mülerova odhadu vývoje HDP v ČR

Na základě rozboru výsledků Gasser-Mülerova odhadu první derivace vývoje HDP a odpovídajícího intervalu spolehlivosti pro tento odhad (Obrázek 1) lze očekávat čtyři možná období pro výskyt extrému: 1997/Q4 – 1998/Q1, 2000/Q2 – 2000/Q4, 2002/Q1 – 2002/Q2 a 2006/Q1 – 2006/Q3. Porovnáme-li zjištěná období se skutečnými hodnotami HDP, pak v prvním období dosahuje HDP hodnoty minima v 1997/Q4, maxima v 2000/Q2, minima v 2002/Q1 nebo 2002/Q2 (hodnoty HDP jsou v těchto obdobích stejné, menší než v předchozích obdobích 2001/Q4 a následujících obdobích v 2002/Q3) a maxima v 2006/Q3. Doplníme-li analýzu o výsledky Gasser-Mülerova odhadu druhé derivace vývoje HDP, vidíme, že ve stanovených obdobích nastává minimum v období 1997/Q4 – 1998/Q1, maximum v období 2000/Q2 – 2000/Q4, minimum v období 2002/Q1 – 2002/Q2 a maximum v období 2006/Q1 – 2006/Q3. Výsledky odhadu trendu vývoje HDP v České republice a expertním odhadem stanovená minima (dna) a maxima (vrcholy) lze tedy pokládat za potvrzená odhadem první i druhé derivace vývoje HDP.

Provedeme-li analýzu reziduí, pak na základě Box-Piercova testu založeném na prvních 15 autokorelacích nelze zamítnout nulovou hypotézu o náhodnosti reziduí ($p\text{-value} = 0,8133 > 0,05$). Z Obrázku 4. je patrné, že v krajních hodnotách dosahují rezidua vyšších hodnot, než v průběhu posuzovaného období. Toto lze přičítat tzv. hraničním efektům, které se mohou vyskytnout v souvislosti s Gasser-Mülerovým odhadem. Korekcí hraničních efektů se však nebudeme zabývat, neboť v konkrétním případě nemají hraniční efekty zásadní vliv na detekci minim a maxim vývoje trendové funkce v posuzovaném období. Potenciálním rizikem by sice mohl být hraniční efekt jádrového odhadu na počátku posuzovaného období, ale podíváme-li se na vývoj úvodních čtyř hodnot, vyvíjejí se klesajícím způsobem, přičemž odhad trendu vývoje dosahuje v případě prvních dvou hodnot nižší úroveň. Lze se tedy domnívat, že korekce hraničního efektu v úvodu posuzovaného období by trendovou funkci upravila zvýšením hodnoty odhadu (Poměnková, 2004).

4. Závěr

Předkládaný příspěvek se zabýval detekcí dna a vrcholu hospodářského růstu ČR v období 1997/1 – 2008/1 pomocí neparametrického Gasser-Müllerova odhadu s využitím odhadu derivací. V prvním kroku byl odhadnut trendu vývoje hospodářského cyklu a expertním odhadem byly stanoveny možné body zvratu - dna a vrcholy. Poté byl proveden odhad první derivace trendu vývoje a stanoveny potenciální extrémy - dna a vrcholy. Následně analýzou výsledků druhé derivace trendu vývoje bylo zkoumáno, zda potenciální extrémy jsou extrémy, a to maxima (vrcholy) nebo minima (dna).

Komparací dosažených výsledků bylo zjištěno, že při použití neparametrického Gasser-Müllerova odhadu nastávají ve stanoveném období 1997/1 – 2008/1 dvě minima (dna) a dvě maxima (vrcholy) ve vývoji ekonomiky ČR. Jmenovitě, v období 1997/Q4 – 1998/Q1

minimum, v období 2000/Q2 - 2000/Q4 maximum, v období 2002/Q1 – 2002/Q2 minimum a v období 2006/Q1 – 2006/Q3 maximum.

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“.

5. Literatura

- [1]CANOVA, F. 1999. Does Detrending Matter for the Determination of the Reference Cycle nad Selection of Turniny Points?, *The Economic Journal*, Vol. 109, No. 452 (Jan., 1999), pp. 126-150.
- [2]DORNBUSCH, R., FISCHER, S. 1994. Makroekonomie, šesté vydání, Praha SNP a Nadace Economics 1994, Praha, 1994, pp.602, ISBN 80-04-25 556-6
- [3]HÄRDLE, W. 1990. Applied nonparametric regression. Cambridge University Press, Cambridge, 1990
- [4]KADEŘÁBKOVÁ, A. 2003. Základy makroekonomické analýzy, Linde, Praha, 2003, 176 s., ISBN 80-86131-36-X
- [5]POMĚNKOVÁ, J. 2004. Edge effects of the Gasser-Muller estimator. In Summer School Datastat03, Proceedings, Folia Fac.Sci.Nat.Univ.Masaryk.Brunensis,Mathematica 15, 2004. 15. vyd. Brno: MU Brno, 2004, s. 307-314. ISBN 80-210-3564-1.
- [6]HOROVÁ, I. 2002. Optimization problems Connected with Kernel estimates, Signal processing, Communications and Computer Science. 2002 by World Scientific and Engineering Socitey Press, pp. 339 - 334.
- [7]POMĚNKOVÁ, J. 2005. Některé aspekty vyhlazování regresní funkce, PhD-thesis, Ostrava, 2005.

Adresa autora:

Jitka Poměnková, Ph.D., RNDr.
Ústav financí PEF MZLU v Brně
Zemědělská 1
Česká republika
613 00 Brno
pomenka@mendelu.cz

Dopad finančních prostředků plynoucích z Evropské unie na ekonomiku České republiky

Impact of financial transfers from the European Union on the Czech Republic economy

Jan Přenosil, Jitka Poměnková

Abstract:

Membership of Czech Republic and other Central European states is a political, economical and social issue. Considerable inflows of money spend on pre-accession aid, structural policy and common agriculture policy are helping the structural changes in the transitional economy. Efficiency of usage of these sources is important due to possible influence of the system in financial aid in 2007-2013 period and next financial period. Experience in new member states can help with more efficient allocation.

Key words: Correlation analysis, European Union, Economic and Social Cohesion, Common Agriculture Policy, Pre-accession Aid

Klíčová slova: Korelační analýza, Evropská unie, Politika ekonomické a sociální soudržnosti, Společná zemědělská politika, Předvstupní pomoc

1. Úvod

Ačkoliv je členství České republiky i dalších států střední a východní Evropy vnímáno veřejností jako poměrně bezproblémový krok ukončený v roce 2004 plynou novým členským státům z členství jak některá omezující pravidla, tak poměrně významné finanční toky. Jak předvstupní pomoc (převážně fondy PHARE, ISPA, SAPARD), tak i strukturální pomoc (ERDF, ESF) i fondy společné zemědělské politiky (EAGGF, v novém období i EAFRD) společně dosahují objemu procent HDP přijímajícího státu. Cílem těchto prostředků je napomoci dosáhnout cíle Společenství/Unie – viz článek 2 Smlouvy o založení Evropských společenství (dále jen SES)

Článek 2¹

Posláním Společenství je vytvořením společného trhu a hospodářské a měnové unie a prováděním společných politik nebo činností uvedených v člancích 3 a 4 podporovat harmonický, vyvážený a udržitelný rozvoj hospodářských činností, vysokou úroveň zaměstnanosti a sociální ochrany, rovné zacházení pro muže a ženy, trvalý a neinflační růst, vysoký stupeň konkurenceschopnosti a konvergence hospodářské výkonnosti, vysokou úroveň ochrany a zlepšování kvality životního prostředí, zvyšování životní úrovně a kvality života, hospodářskou a sociální soudržnost a solidaritu mezi členskými státy.

Konkrétně směřují jednotlivé nástroje k naplnění jejich vlastního cíle. U předvstupní pomoci se jednalo o pomoc při dosahování standardu Společenství primárně v oblasti ochrany životního prostředí, infrastruktury, rozvoje institucí a přípravy na politiky Společenství po plném zapojení se po vstupu.

Cíle strukturální politiky je specifikován v článku 158 SES

HOSPODÁŘSKÁ A SOCIÁLNÍ SOUDRŽNOST

Článek 158 (bývalý článek 130a)

¹ Citovaná část smlouvy zvýrazněna italikou

Společenství za účelem podpory harmonického vývoje rozvíjí a prosazuje svou činnost vedoucí k posilování hospodářské a sociální soudržnosti.

Společenství se především zaměří na snižování rozdílů mezi úrovní rozvoje různých regionů a na snížení zaostalosti nejvíce znevýhodněných regionů nebo ostrovů, včetně venkovských oblastí.

Podobně i v případě společné zemědělské politiky lze citovat článek 33 SES:

Článek 33 (bývalý článek 39)

1. Cílem společné zemědělské politiky je:

- a) zvýšit produktivitu zemědělství podporou technického pokroku a zajišťováním racionálního rozvoje zemědělské výroby a optimálního využití výrobních činitelů, zejména pracovní síly*
- b) zajistit tak odpovídající životní úroveň zemědělského obyvatelstva, a to zejména zvýšením individuálních příjmů osob zaměstnaných v zemědělství;*
- c) stabilizovat trhy;*
- d) zajistit plynulé zásobování;*
- e) zajistit spotřebitelům dodávky za rozumné ceny.*

2. Při vypracovávání společné zemědělské politiky a zvláštních metod, které může zahrnovat, se bude přihlížet:

- a) ke zvláštní povaze zemědělské činnosti, vyplývající ze sociální struktury v zemědělství a ze strukturálních a přírodních rozdílů mezi různými zemědělskými oblastmi;*
- b) k nutnosti provádět vhodné úpravy postupně;*
- c) ke skutečnosti, že v členských státech zemědělství představuje odvětví, které je těsně spjato s celým hospodářstvím.*

I při letním přehledu cílů jednotlivých nástrojů je zřejmé propojení s celkovým hospodářským rozvojem ekonomiky přijímající země. Tento však není jediným cílem – dále se snažíme o podporu dlouhodobě udržitelného rozvoje, stabilizaci ekonomického rozvoje a rovnoměrný rozvoj regionů, rovnost příležitostí atp.

Při určité míře zjednodušení však lze ekonomický rozvoj – v článku měřený růstem HDP (ve stálých cenách) – použít jako indikátor dopadu vynaložených prostředků. Při růstu finančních transferů očekáváme, vzhledem k proklamovaným cílům, i růst HDP. Druhým ukazatelem využitým v článku je míra nezaměstnanosti. I tento ukazatel lze ve středně a dlouhodobém horizontu využít k hodnocení dopadu vynaložených prostředků – s růstem prostředků očekáváme stabilizaci, či dokonce pokles nezaměstnanosti.² Posledním uvažovaným ukazatelem je vývoj podílu zemědělství, rybářství a lesnictví na celkové přidané hodnotě. Očekávaný vztah odpovídá dvěma trendům ve vývoji sektoru ve starých členských zemích – snižování podílů na zaměstnanosti v rámci celého hospodářství a obdobného růstu produktivity jako v ostatních sektorech. Celkový vztah je tedy stabilizace, či mírné oslabení sektoru v souvislosti s využitím prostředků na Společnou zemědělskou politiku.³

Při vlastním hodnocení dopadu je vhodné uvažovat i možnost zpoždění pozitivního či negativního dopadu, neboť prostředky nemají působit pouze jako fiskální stimul, ale napomoci i strukturální změně. Zpoždění pravděpodobně nepřesahuje několik čtvrtletí/let (vzhledem k typickému dvouletému období pro realizaci projektů a jejich zpětnému

² Krátkodobý vztah je vzhledem ke strukturálním změnám složitější.

³ Prostředky na SZP pravděpodobně umožňují rychlejší transformaci sektoru a zpomalují negativní sociální dopad transformace.

proplácení). Nicméně vzhledem k dostupnosti dat (období 2000 až 2007 – čtvrtletní či roční údaje) není v tomto článku se zpožděním uvažováno.

2. Metodika

Analýzy dopadu finančních prostředků plynoucích z EU na ekonomiku České republiky je provedena s využitím korelační analýzy, a to výběrového korelačního koeficientu a výběrového parciálního korelačního koeficientu. Jako proměnné jsou zvoleny hodnoty hrubého domácího produktu České republiky (HDP ve stálých cenách), nezaměstnanost (N), podíl primárního sektoru (PPS) a prostředky plynoucí z EU do České republiky v eurech (P). Podrobněji bude o charakteru dat pojednáno v empirické části příspěvku. Zkoumaným obdobím je rozmezí let 2000 – 2007.

V souvislosti se stanoveným cílem byly formulovány následující slovní hypotézy:

- Nárůst finančních transferů z Evropské unie zvyšuje úroveň HDP,
- Nárůst finančních prostředků z Evropské unie neposiluje růst nezaměstnanosti (tj. nevýznamná nebo negativní korelace),
- Nárůst finančních prostředků z Evropské unie neposiluje podíl primárního sektoru na celkové přidané hodnotě ekonomiky (tj. nevýznamná nebo negativní korelace).

V první fázi bude využito výběrového korelačního koeficientu, který poukazuje na přímou nebo nepřímou závislost. V další fázi bude využit výběrový parciální korelační koeficient, který vypovídá o závislosti mezi dvěma proměnnými při vyloučení vlivu vybraných proměnných (Hebák, J., Hustopecký, J., Malá, J.; 2005 str. 154). Jak Hebák, J., Hustopecký, J., Malá, J. (2005) uvádí, zkušenosti s používáním vícerozměrných statistických metod průkazně ukazují, že parciální korelační koeficienty mohou být užitečnými mírami síly lineární závislosti ve výše uvedeném smyslu i v případech, ve kterých výběrová data pocházejí z jiného než vícerozměrného rozložení. Dalším způsobem výpočtu je využití robustního odhadu korelačního koeficientu (HEBÁK, J., HUSTOPECKÝ, J., MALÁ, J.(2005)). Výpočet tohoto koeficientu je však poměrně náročný a vzhledem k podpurnému charakteru statistického zpracování údajů bylo prozatím od tohoto upuštěno.

Výběrový korelační koeficient je vypočten na základě vztahu

$$r_{XY} = \frac{s_{XY}}{s_X s_Y}, s_X^2 > 0, s_Y^2 > 0. \quad (1)$$

Pro výpočet výběrového parciálního korelačního koeficientu je využit vztah

$$r_{YZ.X} = \frac{r_{YZ} - R_{Y.X}^T R^{-1} R_{Z.X}}{\sqrt{(R_{Y.X}^T R^{-1} R_{Y.X}) \cdot (R_{Z.X}^T R^{-1} R_{Z.X})}}, \quad R_{Y.X} = \begin{pmatrix} r_{Y.X_1} \\ \vdots \\ r_{Y.X_k} \end{pmatrix}, \quad R = \begin{pmatrix} 1 & r_{X_1 X_2} & \cdots & r_{X_1 X_k} \\ \vdots & & \ddots & \\ r_{X_k X_1} & & \cdots & 1 \end{pmatrix}, \quad (1)$$

kde R regulární, Y, Z jsou náhodné veličiny, $X = (X_1, \dots, X_k)$ je náhodný vektor

a $r_{YZ}, r_{YX_i}, r_{X_i X_j}, i, j = 1, \dots, k$ jsou příslušné výběrové korelační koeficienty. Zjištěné korelační koeficienty budou testovány na statistickou významnost pomocí t-testu (Anděl, 1978).

3. Empirická analýza

Pro empirickou analýzu bylo využito hodnot hrubého domácího produktu České republiky (HDP, meziroční změna ve stálých cenách na obyvatele), nezaměstnanosti (N, meziroční změna vyjádřena v lidech), podíl sektoru zemědělství, rybářství a lesnictví na

celkové přidané hodnotě (PPS, meziroční změna podílu sektoru na přidané hodnotě) a prostředky plynoucí z Evropské unie do České republiky (P, meziroční změna platby z EU (v rámci transferů jsou uvažovány stropy na platby v rámci rozpočtu EU v období 200-2003 položka předvstupní pomoc, v období 2004-2007 položky pokrývající strukturální fondy (ERDF, ESF), kohezní fond (CF), prostředky na Společnou zemědělskou politiku i rozvoj venkova (část záruční i rozvojová EAGGF v období 2004-2006, v období 2007-2013 i EAFRD). V případě HDP, N, PPS bylo využito čtvrtletních a ročních hodnot, v případě P byly dostupné pouze roční hodnoty. Dosažitelným obdobím pro zvolené proměnné bylo rozmezí let 2000-2007. Zdrojem dat pro HDP, N, PPS byl Český statistický úřad (Český statistický úřad, 2008). Pro finanční transfery data zveřejňovaná Evropskou komisí týkající se rozpočtu EU (Evropská komise, 2008).

V první fázi byl proveden výpočet výběrových korelačních koeficientů pro čtvrtletní hodnoty HDP, N a PPS. Pro posouzení případného vlivu třetí proměnné při korelaci byly dále vypočteny výběrové parciální korelace a tak zjištěna čistá korelace mezi uvažovanými proměnnými. Výsledné hodnoty uvádí tabulka 1, kde číslo v závorce (p-value) označuje hladinu významnosti α .

Tabulka 1: Výběrové korelační koeficienty a parciální korelační koeficienty (n=28)

	Výběrový korelační koeficient			Výběrový parciální korelační koeficient		
	HDP	N	PPS	HDP	N	PPS
HDP	1			1		
N	0,525 (0,0035)	1		0,5446	1	
PPS	0,0882 (0,6490)	-0,1389 (0,4724)	1	0,1912	-0,2185	1

Z uvedené tabulky vidíme, že HDP a nezaměstnanost, mezi kterými nastává statisticky významná korelace, nejsou primárním sektorem nijak významně ovlivňovány. Jak klasická tak parciální korelace vykazují velmi blízkých hodnot. V případě korelace mezi nezaměstnaností a primárním sektorem, stejně jako HDP a primárním sektorem, se hodnoty parciální korelace ve vztahu ke klasické korelaci zvýšily, nikoliv však významným způsobem. I v tomto případě je proto budeme považovat za přibližně stejné. Můžeme říci, že vztah mezi HDP a PPS, stejně jako N a PPS vykazuje nekorelovanost.

V následujícím kroku byl proveden výpočet pro roční hodnoty s přidáním proměnné reprezentující prostředky plynoucí z EU (P). Vzhledem k faktu, že dostupné byly pouze roční hodnoty od roku 2000, je posuzovaný soubor malý, $n = 7$ hodnot. Zjištěné výsledky výběrové korelace a výběrové parciální korelace uvádí tabulka 2. Porovnáme-li vypočtené hodnoty výběrové parciální korelace pro rozsah souboru $n = 28$ (tabulka 1) a $n=7$ (tabulka 2), vidíme, že přestože je rozsah souboru výrazně menší, lze vypočtené hodnoty korelací považovat za přibližně stejné. A tedy můžeme provést analýzu ve vztahu k prostředkům plynoucím z EU.

Tabulka 2: Výběrové korelační koeficienty a parciální korelační koeficienty (n=7)

	Výběrový korelační koeficient				Výběrový parciální korelační koeficient			
	HDP	N	PPS	P	HDP	N	PPS	P
HDP	1				1			
N	0,6082 (0,1473)	1			0,6409	1		
PPS	0,0841 (0,8577)	-0,2070 (0,6560)	1		0,2419	-0,2682	1	
P	-0,0044 (0,9925)	0,1539 (0,7418)	-0,5011 (0,2520)	1	0,0075	0,0407	-0,4728	1

Pro zjednodušení znovu uvádíme slovní hypotézy formulované v souladu se stanoveným cílem:

- Nárůst finančních transferů z Evropské unie zvyšuje úroveň HDP,
- Nárůst finančních prostředků z Evropské unie neposiluje růst nezaměstnanosti (tj. nevýznamná nebo negativní korelace),
- Nárůst finančních prostředků z Evropské unie neposiluje podíl primárního sektoru na celkové přidané hodnotě ekonomiky (tj. nevýznamná nebo negativní korelace).

V případě vzájemného vlivu mezi prostředky plynoucími z EU (P) a primárním sektorem (PPS) byla výběrovou korelací zjištěna střední negativní závislost ($r = -0,5011$). Při porovnání s výběrovou parciální korelací ($r = -0,4728$) vidíme, že ani nezaměstnanost ani HDP neovlivňuje korelaci posuzovaných hodnot. Přestože je rozsah souboru malý, můžeme se domnívat, že detekovanou střední negativní závislost lze považovat za podporující naši teorii.

V případě vzájemného vztahu mezi prostředky plynoucími z EU (P) a nezaměstnaností (N) byla analýzou detekována statistická nevýznamnost ($r = 0,1539$). Pokud provedeme očištění posuzovaných proměnných o vliv HDP a PPS, pak čistá korelace mezi P a N poklesne ($r = 0,0407$). Tento fakt plyne mimo jiné z hodnoty korelace mezi P a PPS, která byla výše identifikována jako střední negativní závislost. Ze zjištěných výsledků tedy můžeme konstatovat, že čistá korelace mezi prostředky plynoucími z EU a nezaměstnaností je v posuzovaném období je nevýznamná. Na základě toho výsledku se proto autoři přiklání k závěru, že prostředky plynoucí z EU do ČR nemají zásadní vliv na nezaměstnanost. Poznamenejme, že vzhledem k rozsahu datového souboru a k charakteru dat, kterými byly roční hodnoty, nebylo uvažováno žádné zpoždění.

Poslední posuzovaný vzájemný vztah je mezi prostředky plynoucími z EU a HDP. Z výsledku výběrové korelace vyplývá, že tento vztah je statisticky nevýznamný ($r = -0,0044$). Vzhledem k faktu, že HDP významně koreluje s nezaměstnaností ($r = 0,6082$, tabulka 2) a že tento fakt byl potvrzen i na čtvrtletních hodnotách ($r = 0,5250$, tabulka 1), budeme přihlížet k výsledku čisté korelace, za kterou je považovaná výběrová parciální korelace mezi P a HDP očištěná o vliv N a PPS. Ta ovšem také dosáhla statisticky nevýznamné hodnoty ($r = 0,0075$). Lze proto usuzovat, že mezi prostředky plynoucími z EU a HDP v ČR nebyla prokázána závislost.

4. Závěr

Provedená analýza prokázala nevýznamnou statistickou závislost mezi prostředky proudícími z Evropské unie vynakládanými na politiku hospodářské a sociální soudržnosti a společnou zemědělskou politiku. Závěr, který je v rozporu s hypotézou položenou na začátku článku. Korelace $-0,0044$ na hladině významnosti $0,9925$ poukazuje i přes velmi omezený zdroj dat na minimální dopad těchto prostředků na ekonomický růst. Toto dílčí zjištění napovídá nedostatečné efektivitě při využívání prostředků plynoucích z Evropské unie. Důvody pro tuto neefektivitu lze očekávat z několika důvodů:

- a) sledování několika cílů
- b) nedostatečný objem prostředků vynakládaných z rozpočtu EU
- c) nedokonalá administrace prostředků
- d) příliš široké zacílení prostředků
- e) nedostatečný tlak na snižování finanční náročnosti projektů podporovaných z Evropské unie

Ad a) Už v úvodu článku byla zmíněno vícecílové zaměření prostředků poskytovaných z rozpočtu EU. Ačkoliv cíle pokrývají kromě ekonomického rozvoje i ochranu životního prostředí, či rovnost příležitostí jedná se, alespoň z dlouhodobého pohledu, o souřadné cíle. Důvodem neprokázání závislosti tedy může být nedostatečná délka zkoumané časové řady.

Ad b) Celkový objem transferů z evropské unie je ve srovnání s objemem HDP či státního rozpočtu poměrně malý (objem HPD ve stálých cenách pro rok 2007 je 2993 mld. Kč versus 1,867 mld. Eur). Přesto je nutné si uvědomit povahu prostředků určených ať už na investice, či spolufinancování rozvojových programů dofinancovávaných z národní úrovně. Při srovnání s prostředky státního rozpočtu na čisté investice dosahuje objem finančních prostředků z EU takřka třetiny.

Ad c) Vlastní administrace prostředků je stále více předávaná národní (či regionální) administrativě. Lze očekávat, že s postupem času se zlepšuje i úroveň administrace prostředků. I v tomto bodě je však nutné upozornit na příliš krátkou časovou řadu. Tu se pozitivní dopad zlepšení administrace asi neprojeví.

Ad d) Vlastní zacílení je řešeno na úrovni programů. Tzv. princip programování zajišťuje vyjednání rámcového cíle a měřitelných indikátorů s Evropskou komisí na počátku programu. Vlastní výběr projektu v souladu s cíli programu leží na národní administrativě. Vlastní hodnocení naplnění vztahu program-projekt leží mimo obsah tohoto článku, přesto určitá míra flexibility je při sedmiletém plánování nutná.

Ad e) Finanční omezení celkové výše rozpočtu v rámci řízení projektů dnes existují spíše technického rázu – typicky tabulkové ukazatele ceny jednotlivých položek. Tu se projevují dvě nečnosti současného řízení – vysoká averze k riziku a nepropojenost mezi cíly projektu a jejich finanční náročností. Současná regulace uplatňuje spíše „tradiční ověřené řešení“ oproti potenciálně riskantnější variantě. Je otázkou nakolik je řešitelné na úrovni veřejných programů silnější zapojení motivačních prvků.

Pro vlastní závěrečné hodnocení je nutné upozornit, že rozsah dat použitých v článku je omezený a celková délka časové řady je poměrně krátká. Tato skutečnost ovlivňuje validitu jednotlivých tvrzení, bohužel nejsou v současnosti dostupná data v delším časovém horizontu.

5. Literatura

- [1] ANDĚL, J. *Matematická statistika*. Praha : SNTL/ALFA, 1978.
- [2] HEBÁK, J, HUSTOPECKÝ, J, MALÁ, J. *Vícerozměrné statistické metody (2)*. Praha : Informatorium, 2005. 200 s. ISBN 80-7333-036-9.
- [3] Český statistický úřad. *HRUBÝ DOMÁCÍ PRODUKT : ČASOVÉ ŘADY UKAZATELŮ ČTVRTLETNÍCH ÚČTŮ* [online]. [2008] [cit. 2008-09-05]. Dostupný z WWW: <http://www.czso.cz/csu/redakce.nsf/i/hdp_cr>.
- [4] Evropská komise. *EU budget 2007 - Financial Report : Detailed data 2000-2007* [online]. 2008 , June 26, 2008 [cit. 2008-09-03]. Dostupný z WWW: <http://ec.europa.eu/budget/library/publications/fin_reports/fin_report_07_data.xls>.
- [5] *SMLOUVA O ZALOŽENÍ EVROPSKÉHO SPOLEČENSTVÍ* [online]. [2002] [cit. 2008-08-05]. Dostupný z WWW: <http://www.euroskop.cz/gallery/2/756-smlouva_o_es_nice.pdf>.

Adresa autora (-ov):

Jan Přenosil, Ing., Mgr.
Ústav financí PEF MZLU v Brně
Zemědělská 1
Česká republika
613 00 Brno
prenosil@mendelu.cz

Jitka Poměnková, Ph.D., RNDr.
Ústav financí PEF MZLU v Brně
Zemědělská 1
Česká republika
613 00 Brno
pomenka@mendelu.cz

**Vybrané metody eliminace trendu a identifikace cyklické složky
v makroekonomických časových řadách**
**Selected Detrending Techniques and Methods of Cyclical Component
Identificaion in the Macroeconomic Time Series**

Petr Rozmahel

Abstract: The aim of the paper is to describe and compare selected detrending techniques used in the analysis of classical and growth business cycles. The text focuses on the different characteristics in the output cyclical component series identified by the selected detrending techniques. First order differencing, Hodrick-Prescott filter and Baxter-King Band-Pass filters were used in the analysis. The results show the different volatility and autocorrelation in the output time series produced by the filters. Such differences may influence negatively the final results of the business cycles analysis.

Key words: Business Cycle, Band-Pass Filter, Hodrick-Prescott Filter, First Order Differencing

Klíčové slova: hospodářský cyklus, pásmový filtr, Hodrick-Prescottův filtr, Logartmická difference prvního řádu.

1. Úvod

Techniky eliminace trendu resp. identifikace cyklické složky (tzv. detrendovací či filtrační techniky) představují v současnosti neodmyslitelnou součást analýzy makroekonomických časových řad a zejména pak analýzy hospodářského cyklu. Jedná se zejména o problematiku datování hospodářského cyklu (tj. určení bodů zvratu a jednotlivých fází cyklu) a dále o měření sladění hospodářských cyklů, jež své uplatnění nachází např. při posuzování připravenosti ekonomik kandidátských zemí na vstup do eurozóny. Ekonomická teorie rozlišuje mezi dvěma základními pojetími technické identifikace hospodářského cyklu. Klasické cykly jsou vnímány jako cyklické střídání fází absolutního reálného poklesu a růstu zvoleného indikátoru aproximující hospodářský cyklus (Burns-Mitchell, 1946). Pojetí růstového hospodářského cyklu rozšířené zejména ve studiích obsahujících korelační analýzu cyklů prezentuje hospodářský cyklus jako cyklické fluktuace proměnné (resp. cyklické složky časové řady) okolo svého trendu (Lucas, 1977). Cílem příspěvku je popsat a porovnat vybrané detrendovací techniky používané v analýze klasických i růstových hospodářských cyklů a poukázat na možné odlišnosti v generovaném výstupu cyklické složky proměnné – resp. identifikovaného cyklu, které pak mohou ovlivnit i konečné výsledky procesu identifikace či měření sladění hospodářských cyklů. Objektem porovnání jsou techniky logaritmické difference prvního řádu (FOD), Hodrick-Prescottův Filtr (HP) a pásmový filtr spektrální domény (BK_BP) v modifikaci dle Baxter-King (1999), které patří mezi nejčastěji užívané detrendovací techniky ve studiích obsahujících analýzu hospodářských cyklů.

2. Popis vybraných detrendovacích technik

a) Difference prvního řádu (FOD). Mezi neužívanější techniky eliminace trendu v makroekonomických časových řadách – paradoxně jak ve studiích klasického tak růstového cyklu – patří diferencování časových řad, konkrétně difference prvního řádu (FOD) hodnot logaritmů původní časové řady. Canova (1998, 1999) charakterizuje tuto techniku jako

vhodný filtr produkující stacionární časovou řadu, za předpokladu, že trendová komponenta představuje proces prosté náhodné procházky a cyklická složka je stacionární, přičemž tyto dvě složky jsou nekorelované. Navíc se předpokládá jednotkový kořen řady y_t , zcela díky systematické složce y_{t-1} , pro:

$$y_t = y_{t-1} + \varepsilon_t \quad (1)$$

kde trend je definován jako $g_t = y_{t-1}$ a kde $\hat{c}_t = y_t - y_{t-1}$.

b) Hodrick–Prescottův filtr. Optimální filtrační technikou umožňující flexibilně extrahovat trend v časové řadě, který se stochasticky vyvíjí v čase, se zdá být Hodrick–Prescottův filtr. Tato filtrační technika umožňuje přizpůsobit se preferencím výzkumníka dle předpokladu o charakteru trendu a hospodářském cyklu. Spolu s možností extrahovat trend, který se stochasticky posunuje v čase, přičemž je nekorelovaný s cyklickou komponentou umocnil jeho oblibu a široké využití zejména rozvoj literatury reálného hospodářského cyklu. K vyrovnání hodnot časové řady $y_t = g_t + c_t$ dochází prostřednictvím řešení problému minimalizace druhé difference součtu čtverců trendové komponenty g_t .

$$\min_{\{g_t\}_{t=1}^T} \sum_{t=1}^T c_t^2 + \lambda \sum_{t=1}^T [(g_{t+1} - g_t) - (g_t - g_{t-1})]^2 \quad \lambda > 0 \quad (2)$$

Míra hladkosti trendu je ovlivněna volitelným parametrem λ , který penalizuje nedostatečnou hladkost trendu a reguluje tak výsledné fluktuace cyklické c_t od trendu. S růstem λ se rozsah fluktuací kolem trendu zvyšuje a trend g_t se tak stává hladším. V případě, že $\lambda \rightarrow \infty$, tedy že hodnota parametru λ se bude blížit nekonečnu, stává se trend lineární přímkou. Pokud zvolíme $\lambda=0$, pak trendová složka a původní řada splynou, tj. $y_t = g_t$. Volba parametru λ je v literatuře hospodářského cyklu předmětem rozsáhlé diskuse. Dle práce Hodricka a Prescottta

(1980) je parametr roven $\lambda = \frac{\sigma_g^2}{\sigma_c^2}$, kde σ_g^2 a σ_c^2 jsou rozptyly změn trendu a cyklu. S ohledem

na uvedenou podmínku autoři navrhuji optimální hodnotu $\lambda=1600$ pro čtvrtletní data. Nelson–Plosser (1982) navrhuji hodnotu λ v intervalu $[1/6; 1]$ pro většinu zkoumaných makroekonomických časových řad. To implikuje, že většina variability, kterou Hodrick a Prescottt přisuzují cyklické komponentě, je dle Nelsona a Plossera ve skutečnosti částí trendu.

c) Filtry spektrální (frekvenční) domény. Zatímco předchozí techniky eliminace trendu resp. identifikace cyklu v časových řadách pracují v tzv. časové doméně, následující technika je založena na pojetí časových řad ve spektrální popř. frekvenční doméně. Spektrální analýza časových řad je založena na transformaci časové řady na směsici sinusových a kosinusových křivek (na základě tzv. Fourierovy transformace) o různých frekvencích a amplitudách. Metody spektrální analýzy umožňují výzkumníkovi získat obraz o spektru časové řady tedy informaci o zastoupení jednotlivých frekvencí v časové řadě. (Cipra, 1986)

Fabio Canova (1998) ve svém diskusním článku o hodnocení vlivu jednotlivých detrendovacích technik na kvalitu identifikace a určení hospodářského cyklu popisuje princip frekvenčních filtrů za předpokladů, že cyklická a trendová složka časové řady jsou nezávislé a že trendová komponenta má největší sílu v nízkých frekvencích spektra, přičemž dále od nuly síla trendové složky klesá velmi rychle. Vstupní předpoklad navíc nevyžaduje, aby trend časové řady $y_t = g_t + c_t + \varepsilon_t$ byl definován jako deterministický či stochastický, neboť filtr připouští změny ve vývoji trendu v čase, pokud nejsou příliš časté. Trendová a následně i nepravidelná a cyklická složka mohou být v časové řadě odhaleny pomocí následující procedury:

$$a(\omega)F_y(\omega) = F_g(\omega) \quad (3)$$

kde $a(\omega)$ je nízko-pásmový filtr (low pass), $F_y(\omega)$ a $F_g(\omega)$ jsou Fourierovy transformace¹ y_t a g_t . Polynom $a(l)$ v časové doméně vzniká inverzní Fourierovou transformací $a(\omega)$ ve formě

$$a(l) = \frac{\sin(\omega_2 l) - \sin(\omega_1 l)}{\pi l} \quad (4)$$

kde ω_1 a ω_2 jsou horní a dolní limity frekvenčního pásma, kde systematická složka vykazuje veškerou svoji sílu. Odhad cyklické složky je potom $(1-a(l))y_t$. Klíčovým bodem procedury je vhodný výběr horních a dolních limitů frekvenčního pásma. Za předpokladu, že hospodářský cyklus nepřesahuje délky 30 čtvrtletí, jsou stanoveny $\omega_1=0$ a $\omega_2=\pi/15$. Takto definovaný filtr identifikuje cykly s periodicitou ne vyšší než 30 čtvrtletí, přičemž vyšší frekvence jsou přisuzovány trendové komponentě, která je takto z časové řady eliminována. Definovaný proces však stále ponechává příliš velké množství nežádoucí variability vysokých frekvencí, které nemusí nutně znamenat fluktuace přisuzované hospodářskému cyklu. Proto je dále vytvořen filtr s cílem eliminovat nepravidelnou složku ε_t časové řady, která je identifikována s předpokladem, že její největší síla je ve vysokofrekvenčním pásmu spektra s periodicitou nižší než 6 čtvrtletí. Výsledná specifikace $a(\omega) = \omega \in [0, \pi/15] \cup [\pi/3, \pi]$ propouští cyklické komponenty časové řady s nižší periodicitou než 30 čtvrtletí a zároveň vyšší periodicitou než 6 čtvrtletí. Maximální délka takto identifikovaného cyklu je tedy necelých 8 let a minimální délka je 1,5 roku.

Pravděpodobně nejrozsáhlejší modifikací spektrálních filtračních technik s cílem identifikovat komponentu hospodářského cyklu v časové řadě indikátoru agregátní ekonomické aktivity je v současné literatuře Baxter-Kingův pásmový filtr. Autoři Marianne Baxter a Robert King ve své článku z roku 1999 vyvinuli pásmový filtr, který dle vstupních předpokladů považují za nejlepší aproximaci ideálního pásmového filtru. Autoři kritizují ve své práci dosavadní detrendovací metody, které byly aplikovány s cílem generovat stacionární časovou řadu, přičemž nezohledňují vstupní požadavky na statistické vlastnosti hospodářského cyklu. Rozšíření „mechanické aplikace“ detrendovacích technik typu klouzavých průměrů, odstranění lineárních či exponenciálních trendů, FOD, HP filtrů apod. pramenilo z opomenutí výchozí definice hospodářského cyklu. Autoři se snaží vyvinout ideální pásmový filtr, který povede k dekompozici časové řady dle vstupních dispozic výzkumníka. Výchozí záměr autorů je založen na konstrukci lineárního filtru, který eliminuje nízkofrekvenční trendovou komponentu a nepravidelnou složku zahrnující naopak vysokofrekvenční fluktuace. Výsledkem je identifikovaná komponenta hospodářského cyklu.

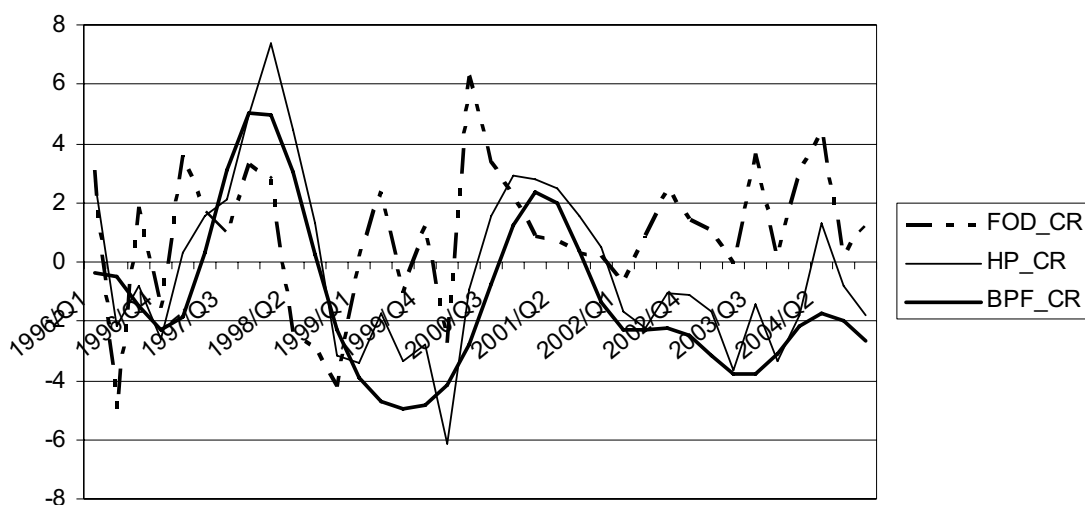
Vstupní podmínkou je zde pojetí hospodářského cyklu dle definice Burnse a Mitchella (1946), kteří definovali hospodářský cyklus jako fluktuace agregátní ekonomické aktivity národa ne kratší než 1,5 roku a ne delší než 8 let. Definování hospodářského cyklu vede ke konkrétnímu dvoustrannému klouzavému průměru (lineárnímu filtru), který získává podobu pásmového filtru (kombinace „low-pass“ a „high-pass“ filtrů – viz výše popsaná procedura dle Canovy) propouštějící komponenty časové řady s periodickými fluktuacemi mezi 6-32 čtvrtletími, přičemž odstraňuje fluktuace o vyšší i nižší frekvenci. Jelikož takto definovaný filtr pracuje na bázi klouzavých průměrů zkombinovaných s metodami spektrální analýzy, je nutno určit vhodnou délku klouzavého úhrnu K . Autoři zjišťují že vhodnou délkou klouzavé řady je $K=12$, což vede k odříznutí 12 pozorování (o čtvrtletní frekvenci) na začátku i konci původní časové řady. Délka klouzavého úhrnu ovlivňuje kvalitu výsledné dekompozice, přičemž růst klouzavého úhrnu K vede k lepší aproximaci ideálního filtru, ale větší ztrátě dat. Uvedená vlastnost je problémem zejména u krátkých dostupných časových řad, což je případ kandidátských zemí střední a východní Evropy připravující se na vstup do eurozóny, neboť aplikace filtru vyžaduje odříznutí 6 let vstupní časové řady, což sníží vypovídací schopnost

¹ K bližšímu zkoumání principů Fourierovy transformace lze odkázat např. na Rektoryse (2000)

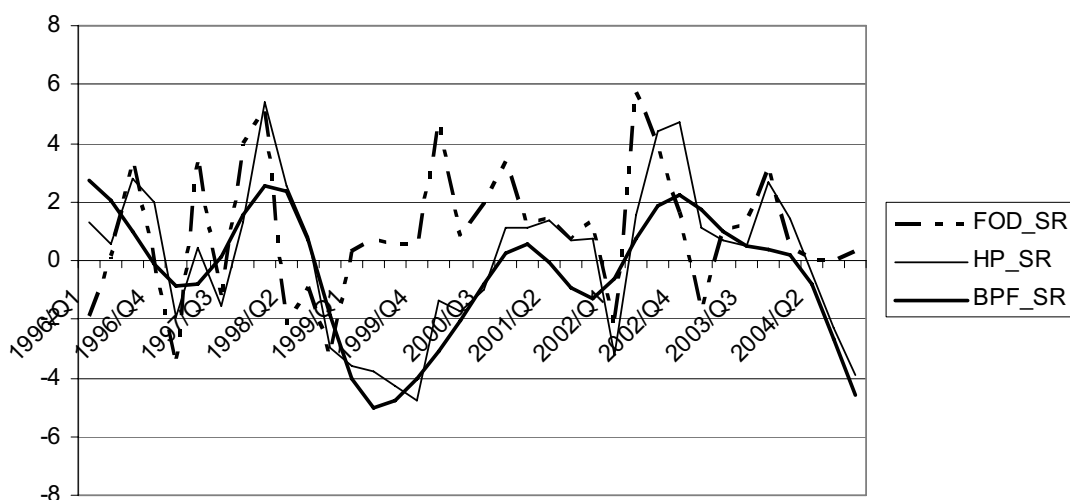
např. o naměřené podobnosti takto identifikovaných cyklů srovnávaných zemí v transformačním resp. posttransformačním období.

3. Aplikace a srovnání vybraných detrendovacích technik

Podobně jako v duchu studie Baxter–King (1999) následující část příspěvku porovnává vybrané statistiky generovaného výstupu filtračních technik v současnosti všeobecně užívaných k analýze hospodářských cyklů: logaritmická diference prvního řádu (FOD), Hodrick-Prescottův filtr (HP, $\lambda=1600$) a vlastní Baxter-Kingův pásmový filtr BK_BP (6,32,12). Vstupní data analýzy tvořily makroekonomické časové řady průmyslové produkce (ve čtvrtletních frekvencích) české a slovenské ekonomiky v období 1993-2007. S ohledem na zvolený klouzavý úhrn pásmového filtru (12 čtvrtletí), je tříletý úhrn na každé straně časové řady oříznut. Zdrojem dat je finanční statistika mezinárodního měnového fondu (IFS IMF).



Obrázek 1: Identifikace cyklu průmyslové produkce ČR v období 1993-2007; zdroj: IMF, vlastní výpočty



Obrázek 2: Identifikace cyklu průmyslové produkce SR v období 1993-2007; zdroj: IMF, vlastní výpočty

Tabulka 1: Statistické charakteristiky cyklu průmyslové produkce ČR identifikované vybranými filtračními technikami; zdroj: vlastní výpočty

	Volatilita Směr.odchylka	koeficient autokorelace		
		1.	2.	3.
FOD_CR	2,3804	0,0013	0,1384	-0,1911
HP_CR	2,8028	0,6574	0,3622	0,0335
BPF_CR	2,5854	0,8687	0,5459	0,1590

Tabulka 2: Statistické charakteristiky cyklu průmyslové produkce SR identifikované vybranými filtračními technikami; zdroj: vlastní výpočty

	Volatilita Směr.odchylka	koeficient autokorelace		
		1.	2.	3.
FOD_CR	2,2687	-0,3556	-0,1216	-0,0894
HP_CR	2,5426	0,9121	0,8257	0,7410
BPF_CR	2,1428	0,7918	0,4543	0,1199

Ze srovnání volatility hospodářských cyklů generovaných jednotlivými technikami vyplývá rozdíl mezi technikou FOD a HP a BK_BP filtry. Technika FOD vykazuje stabilně nízkou směrodatnou odchylku, přičemž frekvence změn jsou oproti relativně hladkým cyklům generovaným HP a BK_BP vyšší. Výsledkem je tedy časová řada s vyšší frekvencí bodů zvratu, což je důsledkem faktu, že technika FOD přesouvá váhu a zvýrazňuje vysokofrekvenční komponenty původní časové řady. HP pracuje jako „high-pass“ filtr, což znamená, že ponechává v časové řadě vysokofrekvenční volatilitu, která je BK_BP filtrem odstraněna. HP produkuje jen mírně vyšší volatilitu v porovnání s BK_BP filtrem neboť makroekonomické časové řady typu HDP nemají mnoho síly ve vysokých frekvencích. Také závislost v rámci zkoumaných řad měřena koeficientem autokorelace je nižší při použití techniky FOD.

4. Závěr

Obecně lze tedy konstatovat, že zkoumání vstupních makroekonomických časových řad pomocí technik HP a BK_BP filtrů přináší podobný dojem o povaze hospodářského cyklu, i když míra podobnosti je závislá na volitelných parametrech λ v případě HP a klouzavém úhrnu K u BK_BP. Technika FOD poskytuje odlišný pohled na základní charakteristiky hospodářských cyklů. FOD produkuje časové řady s nižší volatilitou než BK_BP a HP, což je přímý následek faktu, že FOD snižuje váhu nízkých frekvencí oproti alternativním filtrům. Ze stejného důvodu FOD generuje časové řady, které vykazují mnohem nižší závislost (autokorelaci), než ostatní a také výsledná vzájemná korelace jednotlivých časových řad indikátorů s HDP je daleko nižší. Uvedené výsledky jsou v souladu se studií Baxter-King (1999), kde autoři závěrem konstatují, že techniky FOD podobně jako techniky odrážení lineárního trendu nejsou vhodné pro analýzu hospodářských cyklů. Nadřazené jsou v tomto ohledu techniky HP a BK_BP filtru, přičemž preferují BK pásmový filtr. Důvodem je diskuse o vlivu rozdílného stanovení parametru λ a rapidní změna vah jednotlivých frekvencí na koncích časové řady při použití HP filtru, což může vést k podstatným zkreslením výsledných cyklických pozorování. BK_BP oproti tomu obsahuje méně bodů, které by mohly být v tomto ohledu potenciálním terčem kritiky a poskytuje bližší aproximaci ideálního filtru.

Výsledky analýzy poukazují na možný problematický bod při analýze hospodářských cyklů. Volba detrendovací techniky včetně souvisejících parametrů (koeficient λ v případě Hodrick-Prescottova fitu a délka klouzavého úhrnu K u pásmových filtrů) mohou mít odlišný dopad na celkové výsledky např. měření sladěnosti hospodářských cyklů či identifikace bodů zvratu cyklu.

5. Literatura

- [1]BAXTER, M., KING, R.G. 1999. Measuring Business Cycles: Approximate Band-Pass Filters for Economic Time Series. In: Review of Economics and Statistics, vol. 81, č.4, 1999, s. 575-593.
- [2]BURNS, A.F., MITCHELL, W.C. 1946. Measuring Business Cycles: Vol. 2 of Studies in Business Cycles. New York: NBER, 1946.
- [3]CANOVA, F. 1998. Detrending and Business Cycle Facts. In: Journal of Monetary Economics, vol. 41, 1998, s. 533-540.
- [4]CANOVA, F. 1999. Does Detrending Matter for the Determination of the Reference Cycle and the Selection of Turning points? In: Economic Journal, vol.109, č. 452, 1999,p. 126-150.
- [5]CIPRA, T. 1986. Analýza časových řad s aplikacemi v ekonomii. Praha: SNTL-Nakladatelství technické literatury/Alfa. 1986
- [6]HODRICK, R., PRESCOTT, E.C. 1980. Post-War U.S. Business Cycles. Pittsburgh PA: Carnegie Mellon University, Working Paper, 1980.
- [7]LUCAS, R.E. Understanding Business Cycles. In BRUNNER. K., MELTZER, A.H. (eds.): Stabilisation Domestic and International Economy. Carnegie-Rochester Conference Series on Public Policy 1977, vol. 5, s.7-29.
- [8]NELSON, C. R., PLOSSER, C. I. 1982. Trends and Random Walks in Macroeconomic Time Series: Some Evidence and Implications. In: Journal of Monetary Economics, vol. 10, 1982, s. 139–167.
- [9]REKTORYS, K. 2000. Přehled užití matematiky I. 7. vyd. Praha: Prometheus, 2000, 720 s. ISBN 80-7196-179-5.

Článek byl zpracován v rámci výzkumného projektu 402/08/P494 Posouzení procesu konvergence České republiky a dalších kandidátských ekonomik k eurozoně za finanční podpory Grantové agentury ČR.

Adresa autora:

Petr Rozmahel, Ing., Ph.D.
Výzkumné centrum
PEF MZLU v Brně
Zemědělská 1
Brno 613 00
rozi@mendelu.cz

Analýza odľahlých údajov Analysis of Outliers

Milan Terek¹

Abstract: The article deals with the analysis of the outliers. The definition of the finite-sample breakdown point is given and compared for sample mean and median. Then the methods of the elimination of outliers by trimming and winsorizing is described and illustrated on examples. Finally the methods of the detection of outliers are given and illustrated.

Key words: outlier, finite-sample breakdown, trimming, winsorizing, detection of outliers

Kľúčové slová: odľahlý údaj, bod prelomu konečného výberu, zastrihnutie, winsorizácia, detekcia odľahlých údajov

1. Úvod

Odľahlé údaje sú neobvykle malé alebo neobvykle veľké hodnoty. Ich zisťovanie je často veľmi dôležité. Ak ich existenciu nezoberieme pri analýze údajov do úvahy, môžu výsledky analýzy stratiť výpovednú schopnosť. Pokúsime sa to ilustrovať na príklade. Opíšeme možnosti eliminácie odľahlých údajov a niektoré možnosti ich detekcie. Postupy budeme ilustrovať na príkladoch.

2. Citlivosť výberového priemeru a mediánu na odľahlé údaje

Ak chceme kvantifikovať citlivosť napríklad výberového priemeru na odľahlé údaje, možno využiť tzv. *bod prelomu konečného výberu (finite-sample breakdown point)*². Pre výberový priemer $\bar{X}$, bod prelomu konečného výberu je najmenší podiel pozorovaní, ktorý spôsobuje, že hodnota výberového priemeru je neodôvodnene veľká alebo malá. Inak, pre výberový priemer, bod prelomu konečného výberu je najmenší podiel z n pozorovaní, ktorý môže spôsobiť stratu jeho výpovednej schopnosti. Jediné pozorovanie môže spôsobiť, že hodnota výberového priemeru je neodôvodnene malá alebo veľká, preto jeho bod prelomu je $1/n$.

Všimnime si teraz výberový medián $X_{1/2}$ a jeho bod prelomu konečného výberu. Skoro polovica z n hodnôt môže byť neodôvodnene veľká alebo malá bez toho, že by hodnota mediánu bola neodôvodnene veľká alebo malá. Bod prelomu pre výberový medián je približne³ 0,5. Výberový priemer a výberový medián sú z hľadiska citlivosti na odľahlé údaje, najviac rozdielne.

Príklad 1. Náhodne sa vybralo 100 mužov. Odpovedali na otázku, koľko rozličných partneriek by chceli v živote mať. Výsledky triedenia odpovedí sú v tabuľke 1.

¹ Tento príspevok vznikol s príspevom grantovej agentúry VEGA v rámci projektu číslo 1/0437/08 Kvantitatívne metódy v stratégii šesť sigma a projektu číslo 1/3182/06 Zlepšovanie kvality produkcie strojárskych výrobkov pomocou štatistických metód.

² Pozri WILCOX, R. R.: *Applying Contemporary Statistical Techniques*. USA: Academic Press, 2003, s. 59.

³ Hodnota 0,5 je najväčšia možná hodnota.

Tabuľka 1: Želaný počet rozličných partneriek

Počet partneriek	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	17	40	60	100	8000
Počet odpovedí	3	51	6	1	2	6	1	3	1	1	11	1	4	1	2	2	1	1	1	1

Hodnota výberového priemeru je 85,92. Táto hodnota nemá prakticky žiadnu výpovednú schopnosť, pretože 98 % hodnôt pozorovaní je menších ako táto hodnota. Hodnota výberového mediánu je 1, a určite lepšie charakterizuje typické želania vybratých mužov. V tomto príklade vidno, že výberový priemer je výrazne ovplyvnený jedinou odľahlou hodnotou. Táto hodnota ale neovplyvnila výberový medián.

3. Možnosti eliminácie odľahlých hodnôt

Všimnime si pojem *zastrihnutý priemer* (*trimmed mean*)⁴. Zastrihnutý priemer je vypočítaný z údajov, z ktorých bola vynechaný určitý podiel najväčších a najmenších hodnôt. Špeciálnym prípadom zastrihnutého priemeru je výberový priemer, v ktorom neboli pri výpočte vynechané žiadne údaje z pôvodného výberového súboru.

Napríklad 10%-ný zastrihnutý priemer znamená, že 10% najmenších a 10% najväčších hodnôt pozorovaní bolo pri výpočte priemeru vynechaných.

Nech γ je nejaká konštanta, $0 \leq \gamma < 0,5$, $\eta = 100\gamma$ a $g = \text{ent } \eta$, kde značka *ent*⁵ znamená najväčšie celé číslo menšie alebo rovné číslu η . Potom γ -zastrihnutý priemer je priemer vypočítaný z hodnôt, z ktorých bolo g najmenších a g najväčších hodnôt pozorovaní odstránených:

$$\bar{x}_{t,\eta} = \frac{1}{n-2g} (x_{g+1} + x_{g+2} + \dots + x_{n-g})$$

Príklad 1 – pokračovanie 1. Vypočítame hodnotu 10%-ného a 20%-ného zastrihnutého priemeru z údajov v príklade 1. Výsledky sú takéto.

Hodnota 10%-ného zastrihnutého priemeru $\bar{x}_{t,10}$ je

$$\bar{x}_{t,10} = \frac{1}{80} (x_{11} + x_{12} + \dots + x_{90}) = 3,6$$

Hodnota 20%-ného zastrihnutého priemeru $\bar{x}_{t,20}$ je

$$\bar{x}_{t,20} \approx 3,02$$

⁴ Podľa WILCOX, R. R.: *Applying Contemporary Statistical Techniques*. USA: Academic Press, 2003, s. 63.

⁵ Podľa STN ISO 31 – 11. *Veličiny a jednotky. 11. časť: Matematické značky používané vo fyzikálnych vedách a v technike*. Bratislava: Slovenský ústav technickej normalizácie 1998.

Príklad ilustruje výrazné zlepšenie výpovednej schopnosti výberového priemeru. Zásadnou otázkou je, aký podiel údajov by sme mali „odstrihnúť“. Skôr ako sa budeme venovať tejto otázke, všimneme si alternatívny postup, ktorý sa nazýva *winsorizácia* (*winsorizing*).

Winsorizovaný priemer je podobný ako zastrihnutý priemer, ale „odstrihnuté“ údaje sa nahradia inými údajmi. Napríklad pri výpočte 10%-ného winsorizovaného priemeru sa 10% najväčších hodnôt „neodstrihne“, ale každá z nich sa nahradí najväčšou hodnotou, ktorá by už nebola „odstrihnutá“ pri výpočte 10%-ného odstrihnutého priemeru. Podobne, 10% najmenších hodnôt sa „neodstrihne“, ale každá z nich sa nahradí najmenšou hodnotou, ktorá by už nebola „odstrihnutá“ pri výpočte 10%-ného odstrihnutého priemeru. Potom γ -winsorizovaný priemer sa vypočíta podľa vzťahu:

$$\bar{x}_{w,\eta} = \frac{1}{n} \{ (g+1)x_{g+1} + x_{g+2} + \dots + x_{n-g-1} + (g+1)x_{n-g} \}$$

Príklad 1 – pokračovanie 2. Vypočítame hodnotu 10%-ného a 20%-ného winsorizovaného priemeru z údajov v príklade 1. Výsledky sú takéto.

Hodnota 10%-ného winsorizovaného priemeru $\bar{x}_{w,10}$ je

$$\bar{x}_{w,10} = \frac{1}{100} \{ 11 \cdot x_{11} + x_{12} + \dots + x_{89} + 11 \cdot x_{90} \} = 4,28$$

Hodnota 20%-ného winsorizovaného priemeru $\bar{x}_{w,20}$ je

$$\bar{x}_{w,20} = 4,01$$

4. Detekcia odľahlých údajov

Všimneme si stratégie detekcie odľahlých údajov. Prvá stratégia je založená na výberovom priemere a výberovom rozptyle.

Keď prijmeme predpoklad o normálnom rozdelení základného súboru, je obvyklé považovať za odľahlú hodnotu, ktorá je vzdialená od strednej hodnoty viac ako o 2,24 smerodajných odchýlok. Inak, hodnota x je považovaná za odľahlú, keď

$$\frac{|x - \mu|}{\sigma} > 2,24$$

Keď predpokladáme normálne rozdelenie základného súboru so strednou hodnotou μ a smerodajnou odchýlkou σ , pravdepodobnosť, že hodnota bude považovaná za odľahlú je 0,025. Strednú hodnotu μ a smerodajnú odchýlku σ väčšinou nepoznáme a odhadujeme ich najčastejšie pomocou výberového priemeru a výberovej smerodajnej odchýlky. Potom možno formulovať takéto rozhodovacie pravidlo:

Hodnota x je považovaná za odľahlú, keď

$$\frac{|x - \bar{x}|}{s} > 2,24 \quad (1)$$

Príklad 2. Uvažujme o takýchto hodnotách

1, 1, 1, 1, 2, 2, 2, 2, 3, 3, 3, 3, 4, 4, 4, 4, 500

Hodnota výberového priemeru je $\bar{x} = 31,76$, hodnota výberovej smerodajnej odchýlky je $s = 120,67$. Odhadneme vzdialenosť hodnoty 500 od strednej hodnoty, meranej v smerodajných odchýlkach:

$$\frac{|500 - 31,76|}{120,67} \approx 3,88$$

V tomto prípade, podľa uvedeného rozhodovacieho pravidla, hodnota 500 je odľahlá hodnota.

Príklad 3. Pridajme k údajom z príkladu 2 ešte jednu hodnotu: 1000. Potom je súbor údajov takýto

1, 1, 1, 1, 2, 2, 2, 2, 3, 3, 3, 3, 4, 4, 4, 4, 500, 1000

Hodnota výberového priemeru je $\bar{x} = 85,56$, hodnota výberovej smerodajnej odchýlky je $s = 256,49$. Odhadneme vzdialenosť hodnoty 500 od strednej hodnoty, meranej v smerodajných odchýlkach:

$$\frac{|500 - 85,56|}{256,49} \approx 1,62$$

Podľa uvedeného rozhodovacieho pravidla, hodnota 500 nie je odľahlá, aj keď z jednoduchého porovnania hodnôt je zrejmé, že ide o odľahlú hodnotu.

Posledný príklad je ilustráciou problému, ktorý je známy ako *maskovanie* (*masking*). Odľahlé hodnoty ovplyvňujú výpočet výberového priemeru aj výberovej smerodajnej odchýlky, čo môže na druhej strane maskovať ich prítomnosť, keď na ich detekciu využívame vzťah (1)⁶.

Potrebuje také rozhodovacie pravidlo na detekciu odľahlých hodnôt, ktoré samotné nie je ovplyvnené odľahlými hodnotami.

Najprv si všimneme jednu charakteristiku variability, ktorá sa nazýva *mediánová absolútna odchýlka* (*median absolute deviation*) – MAD. Na výpočet hodnoty MAD je

⁶ Pozri WILCOX, R. R.: *Applying Contemporary Statistical Techniques*. USA: Academic Press, 2003, s. 77.

nevyhnutné najprv vypočítať hodnotu $x_{1/2}$ výberového mediánu $X_{1/2}$, potom vypočítať absolútne hodnoty odchýlok hodnôt od mediánu:

$$|x_i - x_{1/2}| \quad \text{pre } i = 1, 2, \dots, n$$

MAD je medián absolútnych hodnôt odchýlok hodnôt od mediánu.

Možno ukázať, že keď máme náhodný výber z normálneho rozdelenia, potom

$$\text{MADN} = \frac{\text{MAD}}{0,6745}$$

je dobrým bodovým odhadom smerodajnej odchýlky σ základného súboru⁷.

Vráťme sa teraz k hľadaniu rozhodovacieho pravidla na detekciu odľahlých hodnôt, ktoré samotné nie je ovplyvnené odľahlými hodnotami. Jednou z možností je takéto rozhodovacie pravidlo.

Hodnota x je odľahlá, keď

$$\frac{|x - x_{1/2}|}{\text{MADN}} > 2,24 \quad (2)$$

Príklad 3 – pokračovanie 1. Posúdme odľahlosť hodnoty 500 pomocou rozhodovacieho pravidla (2). Hodnota výberového mediánu $x_{1/2} = 3$. Absolútne hodnoty odchýlok hodnôt od mediánu, usporiadané podľa veľkosti vzostupne sú:

$$0, 0, 0, 0, 1, 1, 1, 1, 1, 1, 1, 1, 2, 2, 2, 2, 497, 997$$

Potom $\text{MAD} = 1$ a $\text{MADN} \approx 1,48$.

Dosaďme za x , $x_{1/2}$ a MADN do vzťahu (2). Dostaneme

$$\frac{|500 - 3|}{1,48} \approx 335,81 > 2,24$$

Podľa rozhodovacieho pravidla (2), hodnota 500 je odľahlá.

⁷ Podľa WILCOX, R. R.: *Applying Contemporary Statistical Techniques*. USA: Academic Press, 2003, s. 73.

5. Literatúra

- [1]GARAJ, I. – JANIGA, I.: Jednostranné tolerančné medze normálneho rozdelenia s neznámou strednou hodnotou a rozptylom. One sided tolerance limits of normal distribution with unknown mean and variability. Bratislava: STU, 2005.
- [2]KANTARDZIC, M.: Data Mining. Concepts, Models, Methods, and Algorithms. USA: J. Wiley and Sons, 2003. ISBN 0-471-22852-4.
- [3]LAROSE, D. T.: Data Mining. Methods and Models. USA, J. Wiley and Sons 2006. ISBN 0-471-66656-4.
- [4]TEREK, M.: Analýza rozhodovania. Bratislava: IURA EDITION, 2007. ISBN 978-80-8078-131-6.
- [5]TEREK, M. – HRNČIAROVÁ, Ľ.: Výberové skúmanie. Bratislava: Ekonóm 2008. ISBN 978-80-225-2440-7.
- [6] WILCOX, R. R.: Applying Contemporary Statistical Techniques. USA: Academic Press, 2003. ISBN 0-12-751541-0.

Adresa autora:

Milan Terek, prof., Ing., PhD.
Ekonomická univerzita
Dolnozemska cesta 1
852 35 Bratislava
milan.terek@euba.sk

KLASIFIKAČNÍ METODY V AKUSTICKÉ EMISI - STATISTICKÝ PŘÍSTUP CLASSIFICATION METHODS IN ACOUSTIC EMISSION - STATISTICAL APPROACH

Jan Tláškal, Václav Kůs

Abstract: We present a method and a real-time algorithm which allows us to classify acoustic emission (AE) signals into the finite number of groups which correspond with the physical nature of AE sources. We use spectral parameters and non-metric distances, both independent on the signal amplitudes. We would like to distinguish detected signals emitted by different suppressed acoustic emission sources in materials (for example long-distant micro defect, corrosion, etc.) and also to separate them from common noise. We produce sets of several types of experimental signals to verify the algorithm. The model, which estimates the underlying distribution of spectral parameters originated from the frame of multimodal probability density estimates, i.e. statistical model of distribution mixtures.

Abstrakt: Prezentujeme metodu a algoritmus, který využitím statistických metod dokáže třídit signály akustické emise (AE) do skupin odpovídajícím fyzikální podstatě zdrojů (např. microtrhliny, koroze, laminární posuny, apod.) a rozlišit emisní události způsobené defektem v materiálu od různých šumů. Ke třídění používáme spektrální parametry a nemetrická zobrazení, které nezávisí na amplitudě signálu. Model popisující rozdělení extrahovaných spektrálních parametrů signálů pochází z abstraktního statistického modelu distribučních směsí, používaného na odhady multimodálního rozložení pravděpodobnosti.

Key words: Acoustic emission , Statistical methods, Distribution mixtures, EM algorithm

Klíčová slova: Akustická emise, Statistické metody, Distribuční směsi, EM algoritmus

1 Klasifikace v akustické emisi

Pracujeme na algoritmu, který pomocí statistických metod dokáže třídit signály akustické emise (AE) do skupin odpovídajícím fyzikální podstatě zdrojů, respektive rozpoznat akustickou emisi způsobenou vznikem defektu v materiálu od běžných šumů. Nepokoušíme se příliš využívat informace o fyzikální příčině vzniku akustické emise, zaměřujeme se na statistické metody třídění.

Emisní události (EU) rozumíme záznam AE piezoelektrickým snímačem na disku počítače se vzorkovací frekvencí 2 MHz. Samotný piezoelektrický snímač je bohužel citlivý na poměrně úzkém intervalu frekvencí, některé frekvence potlačuje poměrně výrazně. Navíc každé dva snímače, třebaže stejného typu, se ve svých spektrálních vlastnostech liší, což je způsobeno samotnou složitostí výrobní technologie. Díky deformaci signálu snímačem a mnohonásobným odrazům v měřicí soustavě nelze z teoretické znalosti šíření elastického vlnění učinit jednoznačný závěr o tom, o jakou EU se jedná a statistický přístup ke klasifikaci je nevyhnutelný. Záznam emisní události probíhá tak, že měřicí aparatura kontinuálně sleduje dění na snímači a jakmile přichází signál přesáhne danou prahovou hodnotu, přístroj se ve svém bufferu vrátí o 0,5 ms zpět a uloží na pevný disk 6 ms (tzn. 12500 čísel) dlouhý záznam. Pro klasifikaci používáme ještě kratší část EU délky T obsahující pouze náběh signálu bez deformujících odrazů a interferencí.

Naměříme dostatečně velkou trénovací množinu EU. Z každé emisní události vybereme D parametrů a poté mezi nimi hledáme shluky. Nově přichází signál pak zařadíme do toho z natrénovaných shluků, ke kterému bude mít největší pravděpodobnost příslušnosti.

Klasifikaci signálu AE jsme v podstatě převedli na úlohu shlukové analýzy.

Po hodnotách parametrů budeme vyžadovat, aby v rámci jednoho typu signálu neměly příliš velký rozptyl a aby tvořily pokud možno co nejvíce vzájemně oddělené shluky. Nepodaří-li se nám rozlišit dva zdroje AE, dojdeme k závěru, že vinu nese buď nesprávná volba parametrů, nebo že zdroje budí pro nás fyzikálně nerozlišitelnou AE.

2 Výběr spektrálních parametrů

Při klasifikaci se nejlépe osvědčily parametry ze spektra EU, které zachycují polohu vysokých amplitud ve spektru signálu vzhledem k frekvencím. Jestliže označíme t jako časový index emisní události z_t , pak pomocí diskrétní fourierovy transformace dostaneme každou EU ve formě kombinace harmonických složek s komplexními koeficienty Z_f ,

$$z_t = \sum_{f=1}^{T-1} Z_f e^{-2\pi i f t / T}, \quad t \in \{0, \dots, T-1\}.$$

Absolutní hodnotu $|Z_f|$ nazýváme *amplitudou harmonické složky*. Kdybychom indexu f chtěli dát fyzikální smysl frekvence v KHz, museli bychom ho vynásobit členem $\frac{2000}{1536} \text{ KHz} \cong \frac{4}{3} \text{ KHz}$.

V textu ztotožňujeme pojmy spektrum a amplitudy spektra, fáze pro nás nejsou nijak významné. Spektrum následně pro naše účely klasifikace normujeme na jedničku a vypočteme pro každý signál následující spektrální parametry [1]:

Parametr 0,33 Kvantil:

Zavedeme v souladu s analogií z teorie pravděpodobnosti parametr β -kvantil

$$Q_\beta = \min\{F \in \{0, T-1\} \mid \sum_{f=0}^F |Z_f| \geq \beta\}, \quad \beta \in (0, 1). \quad (1)$$

Parametr W :

Parametr W je druhým parametrem, který zachycuje polohu významných amplitud spektra ve smyslu frekvence

$$W = \arg \min_{l \in \{0, T-1\}} \sum_{f=0}^{T-1} |l - f| (|Z_f| - \overline{|Z_f|})^2, \quad (2)$$

kde

$$\overline{|Z_f|} = \frac{1}{T} \sum_{f=0}^{T-1} |Z_f| = \frac{1}{T}.$$

Poznamenejme v této souvislosti, že běžně dříve užívané přímé parametry ze signálu AE, například *Risetime* (doba od počátku emisní události tzn. prvního překonání šumového prahu do nabytí maxima EU) nebo *County* (počty překmitů různě nastavených prahových hladin) se pro klasifikaci zdrojů AE neosvědčily.

Vystředovaný pentest:

Motivací je zavést jakýsi referenční signál pro danou fyzikální soustavu. Pentest se zdá být logickou volbou, neboť je standardně používaným budičem AE s charakteristickým tvarem spektra. V soustavě vybudíme a zachytíme N emisních událostí délky T vybuzených pentesty na stejném místě. Označme $Z_{i,f}$, $i \in \{0, \dots, N-1\}$, $f \in \{0, \dots, T-1\}$, f -tou složku spektra i -tého pentestu, $\mathfrak{F}$ diskretní fourierovu transformaci a $\tilde{\mathfrak{F}}$ inverzní fourierovu transformaci.

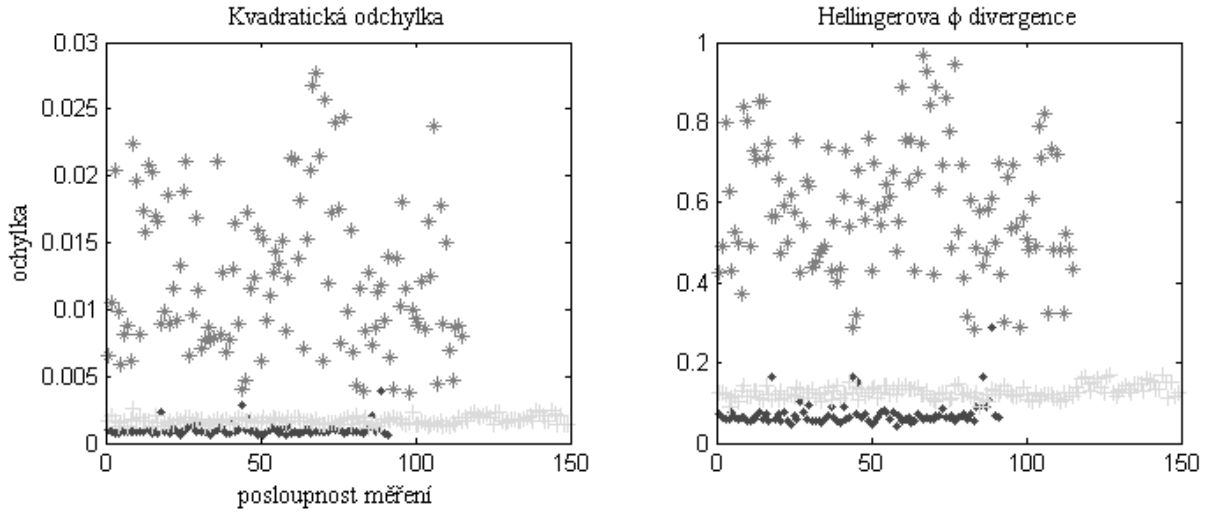
Vystředovaným pentestem nazveme vektor $\bar{P} \in \mathbb{R}^T$,

$$\bar{P} = \mathfrak{Z} \left(\frac{1}{N} \sum_{f=0}^{T-1} \sum_{i=0}^{N-1} Z_{i,f} \right).$$

Zavedeme-li nějakou formu vzdálenosti mezi spektrem $\bar{P}$ a spektry ostatních EU, získáváme další funkční spektrální parametr k již zavedeným parametrům $Q_{0.33}$ a W . Klasické metriky selhávají v porovnávání rozdílů a společných rysů mezi spektry EU, proto používáme symetrizovanou verzi tzv. blendované Hellingerovy divergence [2,3], která je dána vztahem

$$D_H(Z, \mathfrak{Z}(\bar{P})) = 4 \left(2 - \sum_{f=0}^{T-1} (\sqrt{Z_f} + \sqrt{\mathfrak{Z}(\bar{P}_f)}) \sqrt{\frac{1}{2} (Z_f + \mathfrak{Z}(\bar{P}_f))} \right).$$

Na obrázku 1 uvádíme porovnání separace hodnot vzdáleností vystředovaného pentestu od prvků trénovací množiny pomocí kvadratické odchylky a Hellingerovi ϕ -divergence pro tři typy experimentálně naměřených sad signálů akustické emise.



Obrázek 1: Separace tří typů AE pomocí kvadratické odchylky a Hellingerovi ϕ -divergence

3 Distribuční směsi

Jak bylo již řečeno, vybíráme z emisní události několik parametrů, mezi nimiž hledáme shluky. Běžné metody shlukové analýzy nám nevyhovují, protože tíhnou k vytváření kruhových shluků. Díváme se na vypočtené parametry jako na realizace náhodné veličiny X s multimodálním rozložením pravděpodobnosti, jejíž rozdělení odhadujeme pomocí modelu normální distribuční směsi [4].

Definice: Distribuční směs (DS) budeme rozumět každou konvexní kombinaci

$$p(x | \Theta) = \sum_{j \in M} \alpha_j p_j(x | \theta_j), \quad (3)$$

$$\sum_{j \in M} \alpha_j = 1,$$

kde $M = (1, \dots, M)$ je indexová množina složek směsi, $p(x | \theta)$ jsou jednoduché hustoty pravděpodobnosti určené vektorem θ a dále α jsou váhy, případně frekvenční funkce výskytu komponenty distribuční směsi. Distribuční směr zahrnuje komponenty obecně různých typů. V technické praxi se převážně používá směr komponent parametrických rozdělení stejného typu, v našem případě to bude směr normálních rozdělení. Odhadování rozdělení distribuční směsi se redukuje na odhad skupiny parametrů $\Theta = (\alpha_1, \dots, \alpha_M, \theta_1, \dots, \theta_M)$. Komponenta $p(x | \theta_j)$ má tvar

$$p(x | \theta_j) = \frac{1}{(2\pi)^{D/2} |\mathbb{C}_j|^{1/2}} \exp\left(-\frac{1}{2}(x - \mu_j)^T \mathbb{C}_j^{-1}(x - \mu_j)\right). \quad (4)$$

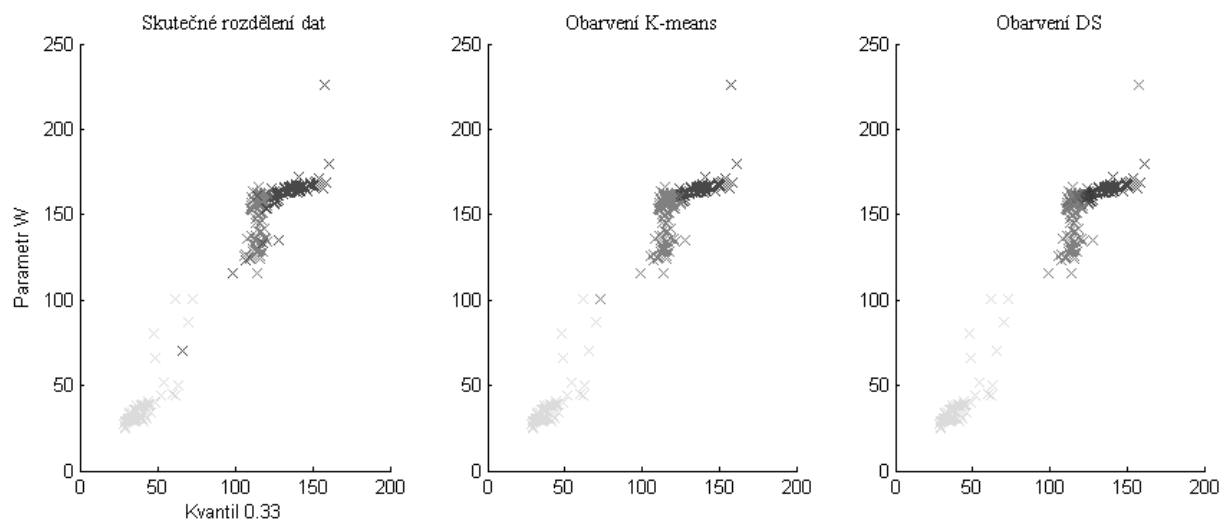
Parametr θ_j má tvar $\theta_j = (\mu_j, \mathbb{C}_j)$, kde $\mu_j \in \mathbb{R}^{D \times 1}$ značí vektor středních hodnoty a $\mathbb{C}_j \in \mathbb{R}^{D \times D}$ je kovariační matice, $x \in \mathbb{R}^{D \times 1}$ jsou vybrané parametry signálu. Parametry směsi odhadujeme *maximálně věrohodným odhadem*. Maximálně věrohodný odhad získáme iterační metodou známou jako EM algoritmus, který konverguje k lokálnímu maximu věrohodnostní funkce [5].

Tvoří-li data pouhým okem dostatečně oddělené shluky, je v podstatě lhostejné jakou metodu pro určení příslušností do shluků použijeme, všechny dosáhnou uspokojivého výsledku. Zcela jiná situace nastává pokud jsou shluky blízko u sebe, částečně se překrývají, mají různou velikost nebo data obsahují výrazné znečištění (tzv. outliers). Proložit shluky směsí normálních hustot pravděpodobnosti a tuto hustotu použít následně pro třídění je výhodné, pokud body v jednotlivých shlucích nemají rovnoměrné rozložení a tvoří přibližně elipsoidy. Takovému rozložení dat odpovídá mnoho úloh v přírodě včetně té naší. Další výhodou modelu DS je robustnost vůči znečištění a díky EM algoritmu i přímočará implementace maximálně věrohodného odhadu.

4 Clustering pomocí distribučních směsí

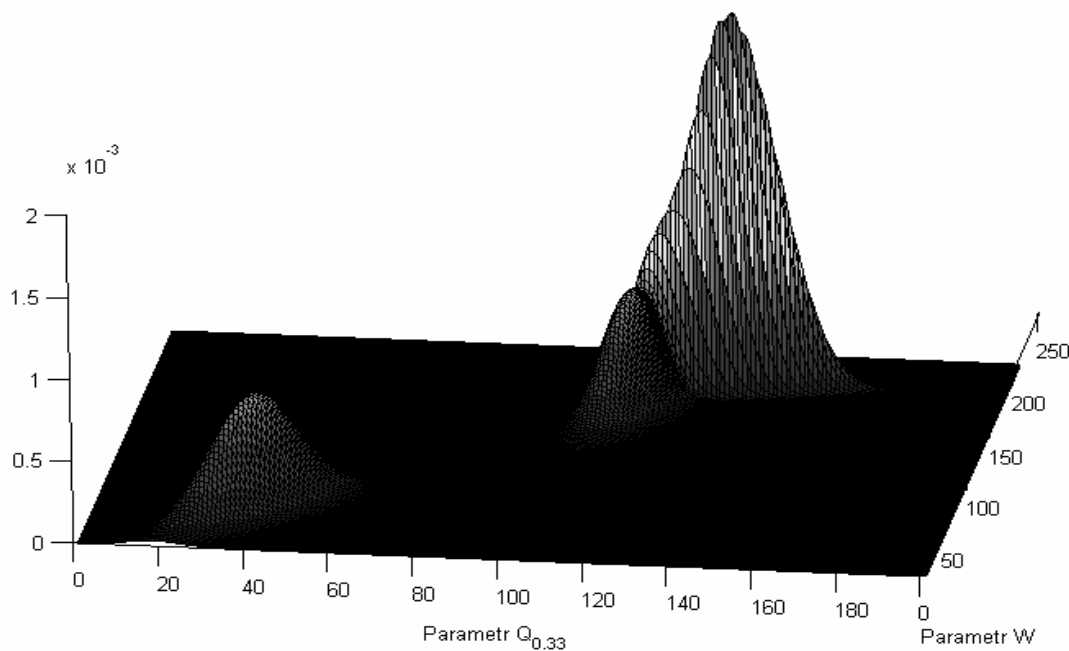
Nechť vzorek $(x_1, \dots, x_N)$ pochází z populace s předem známým počtem M oddělených shluků. Tato data budeme považovat za realizaci distribuční směsi o M komponentách. Ke každému $x_i, i \in (1, \dots, N)$, předepíšeme příslušnost $y_i, y_i \in (1, \dots, M)$ ke k -té komponentě směsi přímým vztahem:

$$y_i = \arg \max_{k \in M} \alpha_k p_k(x_i | \theta_k). \quad (5)$$



Obrázek 2: Metoda DS použitá na překrývající se data

Na obrázku 2 vlevo bylo vyneseno 300 hodnot $X = (Q_{0.33}, W)$. Data jsou legitimní, pocházejí z laboratorního měření 3 různých typů akustických signálů na tenké plechové desce. Nabízíme srovnání výsledků jednoho průběhu EM algoritmu a nejlepšího výsledku metody nejbližšího souseda se 100 náhodnými inicializacemi. Na obrázku 3 je zobrazen graf odhadnuté hustoty distribuční směsi mající maximální věrohodnost, která byla využita pro obarvení shluků v obrázku 2.



Obrázek 3: Graf hustoty distribuční směsi

5 Literatura

[1] TLÁSKAL, J. 2007. Statistické metody klasifikace signálů, Výzkumná zpráva, Katedra matematiky

FJFI ČVUT Praha.

[2] KŮS, V., 2003. Blended ϕ -divergences with examples, In: Kybernetika, vol.39/1, s.43-54.

[3] KŮS, V., Morales D. and Vajda I., 2008. Extensions of the Parametric Families of Divergences Used in Statistical Inference, In: Kybernetika, s. 1-16.

[4] MCLACHLAN G., PEEL D., 2001. Finite Mixture Models. New York: J.Wiley & Sons, Scientific, Technical and Medical Division.

[5] BISHOP, CH.,M., 2006. Pattern Recognition and Machine Learning, New York, Springer.

Jan Tláškal a Ing. Václav Kůs, PhD. - KM FJFI ČVUT Praha, Trojanova 13, 120 00 Praha 2.

Opravitelné systémy s meniacim sa počtom zamietnutých dávok Repairable systems with changing failure rate

Božena Viktorínová *

Abstract: The article deals with the discovering number of failures per day in the repairable systems and fix of average number of failures per system by means of model NHPP (Nonhomogenous Poisson process with Weibull intensity (1)) and by means of relationship (4).

Key words: Weibull intensity, repairable system, failure rate

Kľúčové slová: Weibullova intenzita, opraviteľný systém, zamietnutá dávka

1. Úvod

Pri kontrole kvality výrobkov ktorá prebieha niekoľko dní je zaujímavé zistiť počet vadných výrobkov pripadajúcich na jeden deň v opraviteľných systémoch a stanoviť priemerný počet vadných výrobkov na jeden systém a to pomocou modelu NHPP (Nehomogénny Poissonov proces s Weibullovou intenzitou) (1), a pomocou vzťahu (4). Článok sa zaoberá touto problematikou, ktorá je popísaná v [1].

2. Popis uvedenej problematiky

Nazvime systémami dávky (množiny produktov) ktoré prichádzajú ku kontrole a v ktorých sa môžu vyskytnúť nezhodné produkty. Tieto systémy (každý jeden) sa kontrolujú rôzne dlhý čas (niekoľko dní). Všetky systémy sa však začnú kontrolovať od toho istého dňa a kontrola systémov prebieha súbežne. Ak sa v systéme nájde nezhodný produkt, tento sa opraví alebo nahradí dobrým a systém sa znovu skontroluje. Preto hovoríme o opraviteľných systémoch. Ak predpokladáme, že s rastúcim počtom dní kontroly systémov počet zistených nezhodných produktov pripadajúci na jeden deň klesá, na popis tohto javu môžeme použiť NHPP model ,ktorý môžeme vyjadriť nasledovne:

$$r(t) = \lambda \cdot b \cdot t^{(b-1)} \quad (1)$$

kde λ a b sú neznáme parametre, ktoré pre tento prípad boli odhadnuté metódou podmienenej maximálnej vierohodnosti [1]:

$$\hat{b} = \frac{\sum_{q=1}^K N_q}{\sum_{q=1}^K \sum_{i=1}^{N_q} \ln \frac{T_q}{X_{iq}}} \quad (2)$$

$$\hat{\lambda} = \frac{\sum_{q=1}^K N_q}{\sum_{q=1}^K T_q^{\hat{b}}} \quad (3)$$

kde K je celkový počet testovaných systémov (u nás $K = 20$)
 q je poradie (číslo) systému (1,2,...K)

* Tento príspevok vznikol s príspevom grantovej agentúry VEGA, v rámci projektu číslo 1/0437/08: Kvantitatívne metódy v stratégii šesť sigma

N_q je počet zamietnutých produktov v určitom testovacom čase, ktoré sa vyskytujú v q-tom systéme

T_q je dĺžka trvania testu q-teho systému

X_{iq} je čas v ktorom sa vyskytol i-ty nezhodný produkt v q-tom systéme

$\ln$ je prirodzený logaritmus

V našom príspevku budeme riešiť dva problémy. 1. Pomocou rovnice (1) odhadneme počet zamietnutí (alebo počet nezhodných produktov) pripadajúci na jeden deň. A určíme aj priemerný počet zamietnutí pripadajúci na jeden systém a to pomocou vzťahu

$$\frac{n}{K} = c \text{ zamietnutí / systém} \quad (4)$$

kde K je počet systémov a n je počet zamietnutí. Na lepšiu ilustráciu vyššie uvedenej problematiky uvidíme príklad.

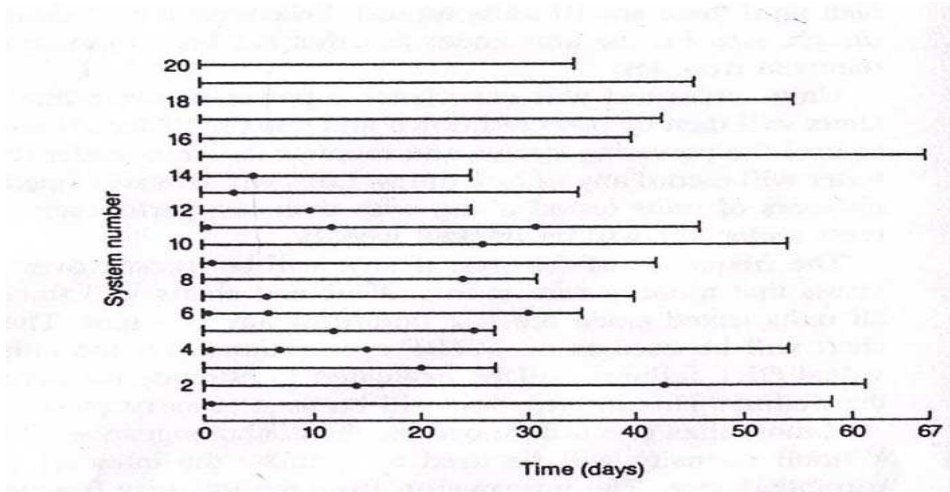
3. Príklad

Vychádzajme z nasledujúcej tabuľky údajov

Tabuľka 1: Vstupné údaje príkladu

1. Číslo systému q	2. Čas, za ktorý sa vyskytol nezhodný produkt v systéme X_{iq} (dni)	3. Počet zamietnutí v systéme q N_q	4. Dĺžka trvania testu v systéme q T_q (dni)
1	0.3	1	58
2	14; 42	2	61
3	20	1	27
4	1; 7; 15	3	54
5	12; 25	2	24
6	0.3; 6; 30	3	35
7	6	1	40
8	24	1	31
9	1	1	42
10	26	1	54
11	0.3; 12; 31	3	46
12	10	1	25
13	3	1	35
14	5	1	25
15	nevyskytol	/	67
16	nevyskytol	/	40
17	nevyskytol	/	43
18	nevyskytol	/	55
19	nevyskytol	/	46
20	nevyskytol	/	31

Znázorníme si zistené skutočnosti z Tab.1 v nasledujúcom obrázku:



Obrázok 1: Časový výskyt zistení i-teho nezhodného produktu v q-tom systéme

Body na úsečkách rovnobežných s osou x-ovou sú hodnoty z 2. stĺpca Tab.1

Ďalej vypočítajme odhady parametrov $\hat{\lambda}$ a $\hat{b}$ z rovníc (2) a (3) a Tab.1:

$$\hat{b} = \frac{1 + 2 + 1 + 3 + \dots + 1}{\ln \frac{58}{0.3} + \ln \frac{61}{14} + \dots + \ln \frac{25}{5}} = \frac{22}{40.663} = 0.54 \quad (5)$$

$$\hat{\lambda} = \frac{1 + 2 + 1 + \dots + 1}{58^{0.54} + 61^{0.54} + \dots + 46^{0.54} + 31^{0.54}} = \frac{22}{149.059} = 0.15 \quad (6)$$

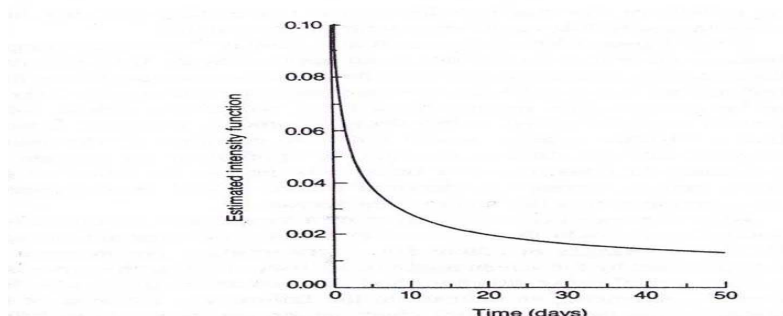
Očakávaný počet zamietnutí na jeden deň potom vypočítame tak, že do rovnice (1) dosadíme (5) a (6):

$$r(t) = 0,15 \cdot 0,54 \cdot t^{0,54-1} = 0,081 \cdot t^{-0,46} \quad (7)$$

Pre rôzne hodnoty t (u nás v Tab.1 sú označené ako X_{iq}) ktoré usporiadame podľa veľkosti od najmenej (0.3) po najväčšiu (42), vypočítame $r(t)$. Hodnoty funkcie $r(t)$ sú uvedené v 3.stĺpci Tab.2. Funkciu $r(t)$ nazveme funkciou intenzity. Napríklad ak čas za ktorý sa vyskytol nezhodný produkt bol $X_{iq} (= t)$ desať dní, potom podľa (7):

$$r(10) = 0,081 \cdot 10^{-0,46} = 0,028 \text{ nezhodných / deň}$$

Znázorníme si tieto výpočty graficky na Obr.2:



Obrázok 2: Funkcia intenzity

Pre rôzne hodnoty t (2. stĺpec Tab.2) vypočítame hodnoty funkcie intenzity $r(t)$, ktoré sa nachádzajú v 3. stĺpci Tab.2.:

Tabuľka 2: Tabuľka pomocných výpočtov

1.	2.	3.	4.	5.	6.	7.
Počet zamietnutí	$X_{iq} (= t)$	$r(t)$	$\frac{n}{K} (= y)$	$\log_{10} X_{iq}$	$\log_{10} \frac{n}{K} (= \log_{10} y)$	$y_{teoretické}$
n	dni					
1	0.3	0.141	1/20=0.05	-0.523	-1.301	0.098
2	0.3	0.141	2/20=0.1	-0.523	-1	0.098
3	0.3	0.141	3/20=0.15	-0.523	-0.824	0.098
4	1	0.081	0.2	0	-0.699	0.180
5	1	0.081	0.25	0	-0.602	0.180
6	3	0.049	0.3	0.477	-0.523	0.312
7	5	0.039	0.35	0.699	-0.456	0.403
8	6	0.035	0.4	0.778	-0.398	0.442
9	6	0.035	0.45	0.778	-0.347	0.442
10	7	0.033	0.5	0.845	-0.301	0.477
11	10	0.028	0.55	1	-0.260	0.571
12	12	0.026	0.6	1.079	-0.222	0.625
13	12	0.026	0.65	1.079	-0.187	0.625
14	14	0.024	0.7	1.146	-0.155	0.676
15	15	0.023	0.75	1.179	-0.125	0.699
16	20	0.020	0.8	1.301	-0.097	0.808
17	24	0.019	0.85	1.380	-0.071	0.885
18	25	0.0184	0.9	1.398	-0.046	0.904
19	26	0.018	0.95	1.415	-0.022	0.922
20	30	0.017	20/20=1	1.477	0	0.990
21	31	0.016	21/20=1.05	1.491	0.021	1.006
22	42	0.014	22/20=1.10	1.623	0.041	1.172

Prudký pokles funkcie intenzity $r(t)$ znamená (Obr.2), že čím neskôr sa vyskytnú nezhodné produkty, tým počet zamietnutí na jeden deň klesá. Tak napríklad, ak zo začiatku bol počet zamietnutí 0.141 / deň, po desiatich dňoch poklesol počet zamietnutí na 0.028 / deň.

2. Aby sme zistili priemerný počet zamietnutí pripadajúcich na jeden systém, budeme počítať rovnicu (4) pre $K = 20$ (je počet systémov v našom príklade) a $n = 1, 2, \dots, 22$, čo je kumulatívny počet zamietnutí v 20-tich systémoch. Podiely typu (4) zapíšeme do 4. stĺpca Tab.2, aby sme zbytočne nezväčšovali rozsah nášho príspevku.

Napríklad, ak zistíme celkovo 8 zamietnutí, potom podľa (4) bude:

$$8 / 20 = 0.4 \text{ zamietnutí na systém.}$$

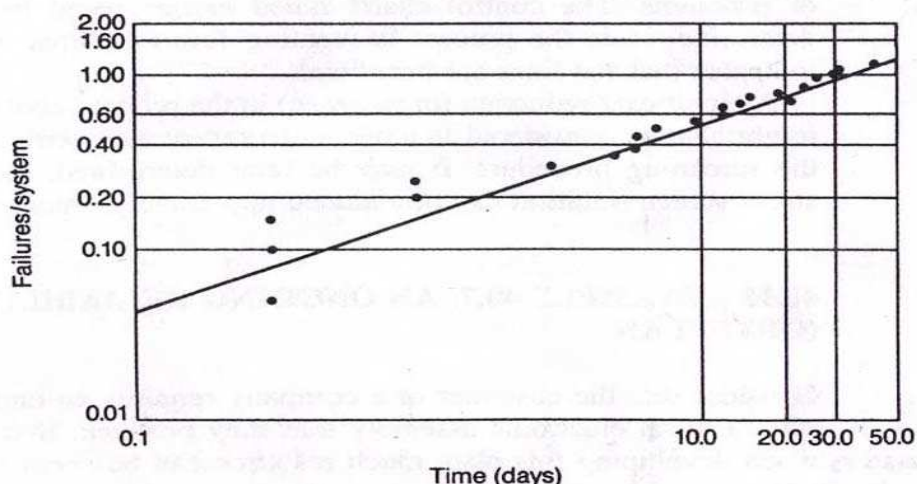
Hodnota 0.4 sa nachádza 4. stĺpci Tab.2 (pre $n = 8$) a tomu zodpovedá dĺžka trvania kontroly $X_{iq} (= t) = 6$ dní, ktorá sa nachádza v 2. stĺpci Tab.2.

Vynesme si tieto výpočty do grafu v bodoch t , s logaritmickou x-ovou a y-ovou stupnicou a to preto, lebo ak by sme tieto body vyniesli do lineárnej stupnice, teoretická krivka preložená cez tieto empirické body by bola síce priamka, ale s veľkou odchýlkou od empirických hodnôt na začiatku a na konci grafu. Preto teoretickou funkciou ktorá lepšie vystihuje zistený trend je funkcia

$$y = a \cdot x^b \quad (8)$$

s ktorou budeme pracovať.

Na os x-ovú budeme vynášať hodnoty 2. stĺpca Tab.2 a na os y-ovú hodnoty zo 4.stĺpca Tab.2.



Obrázok 3: Priemerný počet nezhodných produktov pripadajúcich na systém

Ak chceme preložiť priamku cez empirické body na Obr.3, rovnicu (8) musíme zlogaritmoviť

$$\log_{10} y = \log_{10} a + b \cdot \log_{10} x \quad (9)$$

A teda zlogaritmujeme aj hodnoty v 2. a 4. stĺpci Tab.2. Rovnica (9) je rovnicou priamky, ktorej koeficienty $\log_{10} a$ a b sme odhadli metódou najmenších štvorcov:

$$\log_{10} y = -0.744745 + 0.501358 \cdot \log_{10} x \quad (10)$$

Dosadiac do rovnice (10) piaty stĺpec Tab.2, dostaneme teoretické hodnoty $\log_{10} y$ (6.stĺ.Tab.2). Ak tieto hodnoty odlogaritmujeme, dostaneme teoretické hodnoty y , uvedené v 7.stĺ.Tab.2. Napríklad pre 1.číslo z Tab.2 by bolo: $\log_{10} y = -0.744745 + 0.501358 \cdot (-0.523) = -1.006955$. Z toho $y_{\text{teoretické}} = 1 / 10^{1.006955} = 0.0984$. Pre 6. číslo z Tab.2: $\log_{10} y = -0.744745 + 0.501358 (0.477) = -0.505597$, z čoho $y_{\text{teoretické}} = 1 / 10^{0.505597} = 0.3122$.

Keďže priamka je určená dvomi bodmi, stačí do Obr.3 vyniesť v bode $t = 0.3$ hodnotu 0.0984 a v bode $t = 3$ hodnotu 0.3122 a priamku máme cez empirické údaje preloženú. Z Obr.3 môžeme skonštatovať, že so zvyšujúcim sa počtom dní kontroly priemerný počet zamietnutí pripadajúci na jeden systém stúpa. Je to prirodzené, pretože systém je kontrolovaný vzhľadom na kvalitu dôkladnejšie.

4. Záver

V článku sme sa zaoberali zaujímavou kontrolou kvality produktov podľa [1] a to vzhľadom na dĺžku trvania kontroly rôznych systémov, ktorá môže byť ľubovoľne dlhá (vyjadrená v dňoch). Zvyčajne sa kontroluje počet náhodne vybraných produktov z určitej pevne stanovenej dávky. V tomto prípade ale kontrolujeme paralelne niekoľko systémov a to každý systém počas nie rovnakého počtu dní. Zistili sme, že pri takomto type kontroly 1. počet zistených nezhodných produktov pripadajúci na jeden deň s pribúdajúcim časom kontroly klesá a 2., že s pribúdajúcim počtom dní kontroly priemerný počet nezhodných produktov pripadajúci na jeden systém rastie.

5. Literatúra

[1] Forrest,W., Breyfogle,III.: Implementing six sigma , John Wiley&sons, inc., Toronto 2003, ISBN 978-0-471-265 72-6, str.585 - 590

Adresa autora

Božena Viktorínová, Mgr.CSc.

Dolnozemska 1

85235 Bratislava

Ekonomická univerzita, FHI, KŠ

viktorin@euba.sk

Interval spoľahlivosti pre odhad počtu nezhodných produktov v opraviteľných systémoch

Confidence interval for estimation the number of defect products in repairable systems

Božena Viktorínová *

Abstract:

The article deals with the construction of confidence intervals (according to [1]) for estimation the number of defect products in series of control systems. This control is border each only through time and so it is possible to assume that the control systems will not have the same size. At the same time it is not speaking about 100% control. If in the control system occurs one defect product, this product is being replaced with a good product and the system is returned back into the control and so the systems are repairable systems.

Key words:

Confidence intervals, defect product, tested systems, repairable systems

Kľúčové slová:

Interval spoľahlivosti, nezhodný produkt, testovanie systémov, opraviteľný systém

1. Úvod

Definujme si pojem systém ako množinu prvkov, ktoré prichádzajú ku kontrole, a ak sa kontrolou zistí nezhodný produkt, tento je nahradený dobrým a systém sa vráti späť ku kontrole. Systémy sú kontrolované v rovnako dlhých časových intervaloch (napr. každý po 1000 hodín). Ak sa počas 1000 hodín vyskytne 1 nezhodný produkt, potom za 1 hodinu je to 0,001 ($1/1000=0,001$) nezhodných produktov.

Predstavme si, že takýchto systémov budeme mať n , a každý sa bude kontrolovať 1000 hodín. Teda bude sa jednať o časovo ohraničené testy. Potom kumulovaný čas, ktorý venujeme kontrole bude $T = n \cdot 1000$ hodín.

Označme písmenom r počet zistených nezhodných produktov za čas T . Chceme zistiť, aký bude počet nezhodných produktov v n kontrolovaných systémoch vyrobených produktov, čiže aké bude ρ . Toto množstvo môžeme iba odhadnúť, nakoľko nejde o 100%-nú kontrolu produktov, a to pomocou intervalu spoľahlivosti, ktorý je daný jeho dolnou a hornou hranicou (podľa [1], str.577).

2. Metodika výpočtu intervalu spoľahlivosti

pre odhad počtu nezhodných produktov v n systémoch, ak jednotlivé systémy (dávky) prichádzajúce ku kontrole sú opraviteľné:

$$\frac{\chi^2_{2r;(1+c)/2}}{2T} \leq \rho \leq \frac{\chi^2_{2r+2;(1-c)/2}}{2T} \quad (1)$$

* Tento príspevok vznikol s príspevom grantovej agentúry VEGA, v rámci projektu číslo 1/0437/08: Kvantitatívne metódy v stratégii šesť sigma

- kde T je kumulatívny počet hodín počas ktorých testujeme všetky systémy
 χ^2 je chí kvadrát hodnota z tabuľky G, kde 2r a 2r+2 je počet stupňov voľnosti
c je zvolená pravdepodobnosť s akou je interval spoľahlivosti určený
r je počet zistených nezhodných produktov počas T hodín testu

Tabuľka 1: Kumulatívne χ^2 rozdelenie

Tab.G = Tab.1 – uverejnená v [2].

Pre výpočet intervalu spoľahlivosti môžeme zvoliť aj alternatívny postup, ktorý nám dá rovnaký výsledok ako v prípade použitia rovnice (1):

$$\frac{A_{r;(1+c)/2}}{T} \leq \rho \leq \frac{B_{r;(1+c)/2}}{T} \quad (2)$$

pričom A a B sú faktory zistené z tabuľky K ,pre r zistených nezhodných produktov vo všetkých n kontrolovaných systémoch za čas T, c je zvolená pravdepodobnosť s akou je interval spoľahlivosti určený

Tabuľka 2: Faktory Poissonovho rozdelenia

TABLE K Poisson Distribution Factors

Decimal	Poisson Distribution Confidence Factors B												
Confidence	Number of Failures (r)												
Level													
(c)	0	1	2	3	4	5	6	7	8	9	10		
.999	6.908	9.233	11.229	13.062	14.794	16.045	18.062	19.626	21.156	22.657	24.134	.001	
.99	4.605	6.638	8.406	10.045	11.604	13.108	14.571	16.000	17.403	18.783	20.145	.01	
.95	2.996	4.744	6.296	7.754	9.154	10.513	11.842	13.148	14.435	15.705	16.962	.05	
.90	2.303	3.890	5.322	6.681	7.994	9.275	10.532	11.771	12.995	14.206	15.407	.10	
.85	1.897	3.372	4.723	6.014	7.267	8.495	9.703	10.896	12.078	13.249	14.411	.15	
.80	1.609	2.994	4.279	5.515	6.721	7.906	9.075	10.232	11.380	12.519	13.651	.20	
.75	1.386	2.693	3.920	5.109	6.274	7.423	8.558	9.684	10.802	11.914	13.020	.25	
.70	1.204	2.439	3.616	4.762	5.890	7.006	8.111	9.209	10.301	11.387	12.470	.30	
.65	1.050	2.219	3.347	4.455	5.549	6.633	7.710	8.782	9.850	10.913	11.974	.35	
.60	0.916	2.022	3.105	4.175	5.237	6.292	7.343	8.390	9.434	10.476	11.515	.40	
.55	0.798	1.844	2.883	3.916	4.946	5.973	7.000	8.021	9.043	10.064	11.083	.45	
.50	0.693	1.678	2.674	3.672	4.671	5.670	6.670	7.669	8.669	9.669	10.668	.50	
.45	0.598	1.523	2.476	3.438	4.406	5.378	6.352	7.328	8.305	9.284	10.264	.55	
.40	0.511	1.376	2.285	3.211	4.148	5.091	6.039	6.991	7.947	8.904	9.864	.60	
.35	0.431	1.235	2.099	2.988	3.892	4.806	5.727	6.655	7.587	8.523	9.462	.65	
.30	0.357	1.097	1.914	2.764	3.634	4.517	5.411	6.312	7.220	8.133	9.050	.70	
.25	0.288	0.961	1.727	2.535	3.369	4.219	5.083	5.956	6.838	7.726	8.620	.75	
.20	0.223	0.824	1.535	2.297	3.090	3.904	4.734	5.576	6.428	7.289	8.157	.80	
.15	0.162	0.683	1.331	2.039	2.785	3.557	4.348	5.154	5.973	6.802	7.639	.85	
.10	0.105	0.532	1.102	1.745	2.432	3.152	3.895	4.656	5.432	6.221	7.021	.90	
.05	0.051	0.355	0.818	1.366	1.970	2.613	3.285	3.981	4.695	5.425	6.169	.95	
.01	0.010	0.149	0.436	0.823	1.279	1.786	2.330	2.906	3.508	4.130	4.771	.99	
.001	0.001	0.045	0.191	0.429	0.740	1.107	1.521	1.971	2.453	2.961	3.492	.999	
	1	2	3	4	5	6	7	8	9	10	11	Deci	

	mal
Number of Failures (r)	Conf
	.
Poisson Distribution Confidence Factors A	Leve
	l
	(c)

Applications of Table K

Total test time: $T = B_{rx}/p_0$ for $p \leq p_0$ where p_0 is a failure rate criterion, r is allowed number of test failures, and c is a confidential factor.

Confidential interval statements (time-terminated test); $p \leq B_{rx}/T$ $p \geq A_{rx}/T$

Examples: 1 failure test for a 0,0001 failures/hour criterion (i.e., 10,513/10,000)

95 % confidence test: Total test time = 47,440 hr (i.e., 4744/0,001)

5 failures in a total of 10,000 hours

95 % confident: $p \leq 0,0010513$ failures/hour (i.e., 10,513/10,000)

95 % confident: $p \geq 0,0001970$ failures/hour (i.e., 1.970/10,000)

90 % confidence: $0,0001970 \leq p \leq 0,0010513$

Príklad 1 : Majme 10 systémov (o ich rozsahu sa nehovorí) o ktorých vieme, že každý jeden je kontrolovaný 1000 hodín. Kumulatívny čas kontroly (testu) teda bude $T = 10 \cdot 1000 = 10000$ hodín. Zistili sme, že počas 10000 hodín testu sa vyskytlo $r = 9$ nezhodných produktov. Máme nájsť 90%-ný interval spoľahlivosti pre počet nezhodných produktov ρ v týchto kontrolovaných systémoch.

Riešenie 1 : Zistené skutočnosti dosadíme do rovnice (1):

$$\frac{\chi^2_{2.9;(1+0.9)/2}}{2 \cdot 10 \cdot 1000} \leq \rho \leq \frac{\chi^2_{(2.9+2);(1-0.9)/2}}{2 \cdot 10 \cdot 1000}$$

Ak $\chi^2_{18;0.95} = 9,39$ a $\chi^2_{20;0.05} = 31,41$ sú hodnoty nájdené v Tab.G, potom

$$\frac{9,39}{20000} \leq \rho \leq \frac{31,41}{20000},$$

z čoho

$$0,0004695 \leq \rho \leq 0,0015705.$$

Teda môžeme povedať, že v 10 systémoch, ktoré sme kontrolovali 10 000 hodín, sa s 90%-nou pravdepodobnosťou môže vyskytnúť od 4,695 do 15,705 kusov nezhodných produktov, pričom my sme ich našli 9.

Riešenie 2 : Ten istý výsledok dostaneme, ak hodnoty z príkladu 1 dosadíme do rovnice (2):

$$\frac{A_{9;(1+0.9)/2}}{10 \cdot 1000} \leq \rho \leq \frac{B_{9;(1+0.9)/2}}{10 \cdot 1000}$$

Ak $A_{9;0.95} = 4,695$ a $B_{9;0.95} = 15,705$ sú hodnoty nájdené v Tab.K, potom

$$\frac{4,695}{10000} \leq \rho \leq \frac{15,705}{10000},$$

z čoho

$$0,0004695 \leq \rho \leq 0,0015705.$$

Komentár k výsledku je ten istý, ako v riešení 1.

3. Záver

Konštrukcia vyššie uvedených intervalov spoľahlivosti (1) a (2)) je zaujímavá tým, že neuvažuje rozsah kontrolovaného výberu, z ktorého sa potom usudzuje na celý základný súbor, ale berie do úvahy počet zistených nezhodných produktov (výrobkov) r v n systémoch, z ktorých každý podlieha časovo ohraničenej kontrole (testu). Teoreticky je teda možné, že kontrolované systémy majú rôzny rozsah. Z (1) a (2) sme teda schopní určiť ľubovoľný (90%, 95%, 99%-ný) interval spoľahlivosti, t.j. dolnú a hornú hranicu predpokladaného počtu nezhodných produktov v n systémoch, ktoré nepodliehajú 100%-nej kontrole.

4. Literatúra

- [1] Forrest, W., Breyfogle, III.: Implementing six sigma, John Wiley&sons, inc., Toronto, 2003, ISBN 978-0-471-265 72-6, str.577 – 578
- [2] Viktorínová, B.: Sekvenčný test pre opraviteľné systémy v závislosti na čase; Zborník konferencie: Nitrianske štatistické dni, SŠDS, Nitra, FSS 4/2008, júl 2008, ISSN 1336-7420

Adresa autora

Božena Viktorínová, Mgr., CSc.

Dolnozemska 1

85235 Bratislava

Ekonomická univerzita, FHI, KŠ

viktorin@euba.sk

Z histórie FERNSTATov **From the history of FERNSTATs**

Jozef Chajdiak, Ján Luha, Vladimír Úradníček

1. Úvod

V príspevku uvádzame okolnosti vzniku FERNSTATu a stručne základné informácie o prvých 5 ročníkoch tejto medzinárodnej konferencie aplikovanej štatistiky.

2. Vnik FERNSTATu

Myšlienku organizovať konferenciu FernStat s tematickou orientáciou na aplikovanú štatistiku vyslovil počas panelovej diskusie, konanej v rámci 13. ročníka Medzinárodného semináru Výpočtová štatistika 2. decembra 2003 vo večerných hodinách nemenovaný spoluautor tohoto príspevku (V. Ú.) v úzkom kruhu, ktorý tvorili autori tohoto príspevku (J.Ch., J.L., V.Ú.). Prvotná myšlienka bola zaznamenaná do PC ďalšieho nemenovaného (J. L.) na podnet tretieho nemenovaného (J. Ch.). Tak sa myšlienka zachovala a postupne bola rozpracovaná koncepcia konferencie, ako aj význam skratky **FernStat**: Medzinárodná konferencia aplikovanej štatistiky (**F**inancie, **E**konomika, **R**iadenie, **N**ázory). Okrem aplikovanej štatistiky boli stanovené aj tematické okruhy širšieho rámca, ktoré sa za prvých 5 ročníkov zatiaľ nemenili: **Tematické okruhy konferencie pre rok 2004 boli**: Aplikovaná štatistika, Demografická štatistika, Matematická štatistika a Štatistické riadenie kvality.

Konferenciu organizuje Slovenská štatistická a demografická spoločnosť v spolupráci s UMB Banská Bystrica – v rokoch 2004 a 2005 s Fakultou financií a po jej zlúčení s Ekonomickou fakultou je partnerskou organizáciou Ekonomická fakulta UMB v Banskej Bystrici.

Na pôde Slovenskej štatistickej a demografickej spoločnosti sa snažíme o rozvoj teórie a aplikácií štatistických metód. Dôležitým aspektom je tiež popularizácia štatistiky a poukazovanie na jej vážnosť a tiež presadzovanie vážnosti tejto vednej disciplíny. Taktiež je snaha aktivizovať aj aktivity štatistickej mládeže a umožňovať im účasť a prezentáciu na našich odborných akciách. Spoluautor (J.L.) tohto príspevku už dávno zistil, že keď sa sám nepochváli, tak sa pochvaly asi ťažko dočká a preto pri každej príležitosti poukazuje na to, že je autorom nielen znaku Slovenskej štatistickej a demografickej spoločnosti, ale v rámci snahy robiť vážne veci veselo, objavil aj „Zákon troch fernetov“, v stratke zákon TRIeF, alebo 3F. So znakom Spoločnosti sa stretávame na každom podujatí, nachádza sa aj na stránke www.ssds.sk. O ZÁKONE TRIeF, alebo 3F zasvätení vedia a bližšie sa o ňom možno dozvedieť na spoločenskej časti každej odbornej akcie SŠDS. Nuž a novší objav definuje skratku KAFE – nie je to ale ten mok, čo sa zvykne piť cez prestávku, ale **K**arpatský **f**ernet.

Obrázok 1 a 2 Hotel Lesák, Tajov a foto z rokovania konferencie



3. Chronológia FERNSTATov

Spomedzi piatich doterajších ročníkov konferencie FernStat patrí osobitné miesto druhému ročníku konferencie. Uvedený ročník vstupuje do histórie SŠDS prvým číslom nového vedeckého časopisu Slovenskej štatistickej a demografickej spoločnosti FORUM STATISTICUM SLOVACUM (*FSS 1/2005, ISSN 1336-7420*), ktoré obsahovo bolo „naplnené“ príspevkami konferencie FERNSTAT 2005. Navyiac počnúc týmto ročníkom sa zaviedla tradícia konania konferencie vždy prvý štvrtok a piatok v októbri.

Miesto konania konferencie je stabilné: **Tajov pri Banskej Bystrici, Hotel Lesák.**

Počas doterajších ročníkov konferencie FernStat sa s menšími obmenami vytvoril viac-menej stabilný organizačný a programový výbor v zložení – V. Úradníček – predseda, M. Kanderová – tajomník, J. Chajdiak, J. Luha, R. Gavliak, S. Koróny (v prvých dvoch ročníkoch). Pri organizovaní konferencie aktívne participovali aj niektorí ďalší členovia SŠDS, ako aj študenti UMB v Banskej Bystrici.

Prehľad dátumov konania jednotlivých ročníkov konferencie, počtu príspevkov publikovaných z príslušného ročníka, ako aj celkový počet autorov týchto príspevkov uvádza tabuľka 1.

Tabuľka 1 Chronológia FernStatov

	Dátum konania	Počet príspevkov	Počet autorov príspevkov
FernStat 2004	6. – 7. mája 2004	17	21
FernStat 2005	6. – 7. októbra 2005	35	45
FernStat 2006	5. – 6. októbra 2006	35	49
FernStat 2007	4. – 5. októbra 2007	28	37
FernStat 2008	2. – 3. októbra 2008	32	46

Prameň: Vlastné spracovanie

4. Záver

Medzinárodná konferencia FERNSTAT už za krátku päťročnú históriu dokázala svoju užitočnosť. Stala sa prostriedkom na stretávanie odborníkov a umožňuje publikovať hodnotné príspevky a tak rozširuje publikačné možnosti širokému okruhu odborníkov.

Ambíciou usporiadateľov konferencie je rozšírenie FernStatu aj o regionálnu prehliadku prác mladých aplikovaných štatistikov, z ktorých najlepší by sa kvalifikovali na Prehliadku prác mladých štatistikov a demografov, ktorá sa každoročne koná v decembri v rámci Medzinárodného semináru Výpočtová štatistika.

5. Literatúra

- [1] FernStat2004 – Zborník príspevkov zo slovenskej konferencie aplikovanej štatistiky. Tajov, 6. – 7. mája 2004. Bratislava : SŠDS, 2004. 92 s. ISBN 80-88946-34-4.
- [2] FORUM STATISTICUM SLOVACUM 1/2005. SŠDS, 2005. ISSN 1336-7420.
- [3] FORUM STATISTICUM SLOVACUM 4/2006. SŠDS, 2006. ISSN 1336-7420.
- [4] FORUM STATISTICUM SLOVACUM 4/2007. SŠDS, 2007. ISSN 1336-7420.
- [5] FORUM STATISTICUM SLOVACUM 6/2008. SŠDS, 2008. ISSN 1336-7420.

Adresa autora (-ov):

Ing. Vladimír Úradníček, PhD.
vladimir.uradnicek@umb.sk

RNDr. Ján Luha, CSc.
Jan.Luha@statistics.sk

Doc. Ing. Jozef Chajdiak, CSc.
chajdiak@statis.biz

OBSAH

	Úvod	1
Bod'a Martin	Approaches to volatility estimation based on historical data	2
Bod'a Martin Kanderová Mária	Overovanie efektívnosti trhového portfólia v priestore mean-variance: model CAPM	8
Boháčová Hana, Heckenbergerová Jana, Jindrová Pavla,	Branch structure of employment in particular regions of the Czech Republic	15
Čief Rastislav	Vývoj rozvodovosti na Slovensku 1997-2007	20
Durník Eduard, Kamenický Jiří	Regionální demografie v ČSÚ - prezentace a analytické využívání	25
Fiala Tomáš	Analýza růstu střední délky života v ČR metodou klouzavé lineární regrese	31
Hajduová Zuzana, Liptáková Erika, Závadský Cyril	Aplikácia metódy QFD v hutníckej výrobe	36
Hanousková Jitka, Kůs Václav	Asymptotické vlastnosti odhadů s minimální Kolmogorovskou vzdáleností	41
Horníková Adriana	Štatistické riadenie kvality – Taguchiho návrhy experimentov v SAS-e	47
Hrehová Stella	Možnosti využitia MS Excel pri operatívnom sledovaní kvality výrobného procesu	53
Hrnčiarová Ľubica, Terek Milan	Štatistické riadenie kvality. Analýza spôsobilosti procesu	57
Chajdiak Jozef, Suchánek Maroš	Štatistické porovnanie zhody analyzovaného a referenčného súboru	63
Janiga Ivan	Štatistická regulácia v plniacom procese	67
Kačerová Eva	Nezaměstnanost cizinců v České republice podle vzdělání	74
Kalina Jan	Robust GMM estimation	78
Kapounek Svatopluk	Stabilita poptávky po penězích v České republice	84
Kaščáková Alena, Nedelová Gabriela	Diferenciácia slovenských domácností vyplývajúca z ekonomickej aktivity	90
Koróny Samuel	Grilichesove konzistentné hranice parametrov Cobbovej-Douglasovej produkčnej funkcie stavebných podnikov	95
Linda Bohdan, Kubanová Jana	Situation in University Educational System in the Czech Republic	100
Luha Ján	Korelácia komplementárnych a disjunktných javov	105
Macurová Anna, Macura Dušan	Grubbsov test a hypotézy o obrobenom povrchu	114
Neubauer Jiří	Kointegrační analýza modelu inflace v České republice	117
Odehnal Jakub, Michálek Jaroslav, Sedlačík Marek	Faktory podnikatelského prostředí a jejich klasifikační síla	123

Pastor Karol	Pokus o nájdenie spravodlivého dôchodkového systému	129
Poměnková Jitka	Stanovení vrcholu a dna hospodářského cyklu ČR pomocí neparametrického odhadu derivace trendu	134
Přenosil Jan, Poměnková Jitka	Dopad finančních prostředků plynoucích z Evropské unie na ekonomiku České republiky	140
Rozmahel Petr	Vybrané metody eliminace trendu a identifikace cyklické složky v makroekonomických časových řadách	146
Terek Milan	Analýza odláhlých údajov	152
Tláškal Jan, Kůs Václav	Klasifikační metody v akustické emisi – statický přístup	158
Viktorínová Božena	Opravitel'né systémy s meniacim sa počtom zamietnutých dávok	164
Viktorínová Božena	Interval spoľahlivosti pre odhad počtu nezhodných produktov v opraviteľných systémoch	170
Chajdiak Jozef, Luha Ján, Úradníček Vladimír	Z histórie FERNSTATov	174

CONTENS

	Introduction	1
Bod'a Martin	Approaches to volatility estimation based on historical data	2
Bod'a Martin Kanderová Mária	Overovanie efektívnosti trhového portfólia v priestore mean-variance: model CAPM	8
Boháčová Hana, Heckenbergerová Jana, Jindrová Pavla,	Branch structure of employment in particular regions of the Czech Republic	15
Čief Rastislav	Divorce Rate Progress in Slovakia from 1997 to 2007	20
Durník Eduard, Kamenický Jiří	Regionální demografie v ČSÚ - prezentace a analytické využívání	25
Fiala Tomáš	The analysis of the growths of the life expectation in the Czech Republic by the method of moving linear regression	31
Hajduová Zuzana, Liptáková Erika, Závadský Cyril	Applikation of method QFD in metallurgical production	36
Hanousková Jitka, Kůs Václav	Asymptotic properties of Minimum Kolmogorov distance density estimators	41
Horníková Adriana	Statistical Quality Control – Taguchi Design of Experiments within SAS	47
Hrehová Stella	Possibilities of using MS Excel to operative control of manufacturing process quality	53
Hrnčiarová Ľubica, Terek Milan	Statistical Quality Control. Process Capability Analysis	57
Chajdiak Jozef, Suchánek Maroš	Statistical comparison deuces to analyse and reference file	63
Janiga Ivan	Statistical control in filling process	67
Kačerová Eva	Unemployment of Foreigners in the Czech Republic	74
Kalina Jan	Robust GMM estimation	78
Kapounek Svatopluk	Stabilita poptávky po penězích v České republice	84
Kašáková Alena, Nedelová Gabriela	Differentiation of Slovak Households Followed from the Activity Status	90
Koróny Samuel	Griliches' consistent bounds of Cobb-Douglas production function parameters of construction companies	95
Linda Bohdan, Kubanová Jana	Situation in University Educational System in the Czech Republic	100
Luha Ján	Correlation of complementary and disjoint events	105
Macurová Anna, Macura Dušan	Test of Grubbs and Hypotheses on the Cut Surface	114
Neubauer Jiří	Cointegration Analysis of the Inflation Model in the Czech Republic	117
Odehnal Jakub, Michálek Jaroslav, Sedlačík Marek	Factors of entrepreneurial environment and their power of classification	123

Pastor Karol	An attempt to find the fair pension system	129
Poměnková Jitka	Through and Peak Determination of the CR Business Cycles using nonparametric estimat of the trend derivation	134
Přenosil Jan, Poměnková Jitka	Impact of financial transfers from the European Union on the Czech Republic economy	140
Rozmahel Petr	Selected Detrending Techniques and Methods of Cyclical Component Identificaion in the Macroeconomic Time Series	146
Terek Milan	Analysis of Outliers	152
Tláškal Jan, Kůs Václav	Classification methods in acoustic emission – statistical approach	158
Viktorínová Božena	Repairable systems with changing failure rate	164
Viktorínová Božena	Confidence interval for estimation the number of defect products in repairable systems	170
Chajdiak Jozef, Luha Ján, Úradníček Vladimír	From the history of FERNSTATs	174

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 nevynechávajú.

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 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

02/63812565

e-mail

chajdiak@statis.biz
Jan.Luha@statistics.sk

Registráciu vykonalo

Ministerstvo kultúry Slovenskej republiky

Registračné číslo

3416/2005

Tematická skupina

B1

Dátum registrácie

22. 7. 2005

Objednávky

Slovenská štatistická a demografická
spoločnosť
Miletičova 3, 824 67 Bratislava 24
Slovenská republika
IČO: 178764
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. Mikuláš Cár, CSc.
Ing. Ján Cuper
Ing. Pavel Flák, DrSc.
Ing. Edita Holičková
Doc. RNDr. Ivan Janiga, CSc.
Ing. Anna Janusová
RNDr. PaedDr. Stanislav Katina, PhD.
Prof. RNDr. Jozef Komorník, DrSc.
RNDr. Samuel Koróny
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.
Doc. RNDr. Beata Stehlíková, 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

IV.

Číslo

6/2008

Cena výtlačku 500 SKK / 20 EUR
Ročné predplatné 1500 SKK / 60 EUR